

L'USINE NOUVELLE

SÉRIE | EEA

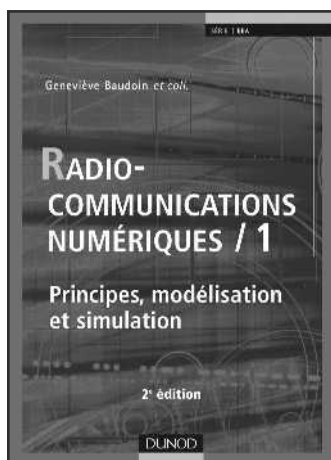
Odile Picon *et coll.*

LES ANTENNES

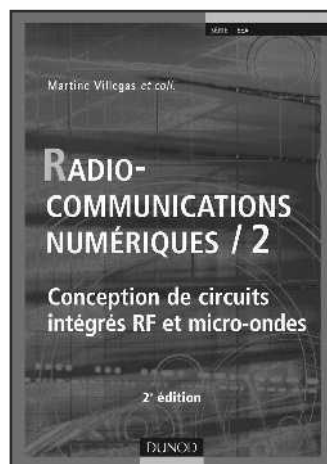
Théorie, conception et applications

Préface de Maurice Bellanger

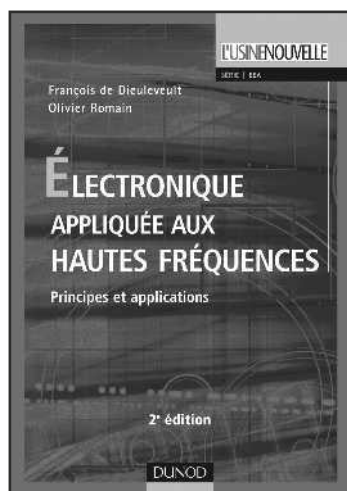
DUNOD



Martine VILLEGAS *et coll.*
Radiocommunications
numériques, Tome 1, 672 p.



Martine VILLEGAS *et coll.*
Radiocommunications
numériques, Tome 2, 368 p.



François DE DIEULEVEULT
 et Olivier ROMAIN
Électronique appliquée aux
hautes fréquences, 552 p.



Dominique PARET
RFID en ultra et super hautes
fréquences UHF-SHF, 496 p.

Odile Picon *et coll.*

LES ANTENNES

Théorie, conception et application

Préface de Maurice Bellanger

L'USINENOUVELLE

DUNOD

Consultez nos parutions sur dunod.com



Le pictogramme qui figure ci-contre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, particulièrement dans le domaine de l'édition technique et universitaire, le développement massif du photocopillage.

Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée dans les établissements



d'enseignement supérieur, provoquant une baisse brutale des achats de livres et de revues, au point que la possibilité même pour les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée. Nous rappelons donc que toute reproduction, partielle ou totale, de la présente publication est interdite sans autorisation de l'auteur, de son éditeur ou du Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris).

© Dunod, Paris, 2009
ISBN 978-2-10-054245-1

Le Code de la propriété intellectuelle n'autorisant, aux termes de l'article L. 122-5, 2° et 3° a), d'une part, que les « copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » et, d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation ou reproduction intégrale ou partielle faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (art. L. 122-4).

Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles L. 335-2 et suivants du Code de la propriété intellectuelle.

PRÉFACE

En une décennie, les radiocommunications ont apporté des changements fondamentaux dans la société, au niveau mondial. Ainsi, la téléphonie mobile a apporté à l'individu la liberté de communication et l'accès universel aux réseaux d'information. Une telle évolution a été rendue possible par les progrès de l'électronique et des techniques et moyens de traitement numérique de l'information dans les terminaux. Cependant, il ne faut pas oublier que la liaison entre ces terminaux, téléphones portables, ordinateurs, stations de base et autres infrastructures, est assurée par les ondes électromagnétiques, qui sont un point de passage obligé. C'est dire l'importance, dans les systèmes de radiocommunications et pour leur fonctionnement, des organes qui effectuent l'interface entre les moyens de traitement de l'information et les ondes qui véhiculent cette information, c'est-à-dire les antennes. Parfois spectaculaires, souvent non visibles, ces antennes, en réalité, déterminent des paramètres essentiels des communications, comme la puissance émise, la direction de rayonnement ou la portée, et elles ont un impact fort sur des caractéristiques critiques comme les débits numériques ou les dimensions des équipements.

Les antennes suscitent aussi, parfois et chez certains, des inquiétudes, comme dans le cas des antennes-relais. En effet, les ondes électromagnétiques utilisées en radiofréquences transportent de l'énergie et sont invisibles. En réalité, c'est par la présence des antennes que l'on réalise la présence des ondes et que l'on peut apprécier ou imaginer certaines de leurs caractéristiques. Indépendamment des recherches sur les effets potentiels de ces ondes sur l'environnement et sur l'individu, notamment sur la santé, il est certain que la connaissance des propriétés des antennes et la maîtrise de leurs caractéristiques sont des éléments importants pour leur acceptation par le public. En fait, il faut optimiser les antennes et leur utilisation, afin de minimiser leur impact sur l'environnement et les recherches qui visent à améliorer leurs performances techniques et à augmenter leur souplesse d'application, tout en limitant les coûts, vont contribuer au progrès des systèmes de radiocommunications de tous types et à leur exploitation harmonieuse dans la société. Dans le présent contexte, l'ouvrage d'Odile Picon et de ses collègues est une contribution importante aux progrès de l'électronique en général et des radiocommunications en particulier. D'abord, par la synthèse des connaissances de base sur la nature des ondes électromagnétiques, leur propagation et leur manipulation, qu'il met à la portée des ingénieurs et des étudiants en électronique. Ensuite, par les correspondances et les liaisons qu'il développe entre les antennes et les traitements de l'information dans les terminaux et l'importance qu'il fait apparaître pour les conceptions conjointes. Enfin, par les axes de recherche qu'il souligne et les stimulations qui vont en résulter pour les scientifiques, ingénieurs et chercheurs. Les méthodes de mesure et de modélisation se perfectionnent et la meilleure maîtrise de l'environnement électromagnétique qui en résulte est de nature à donner confiance aux utilisateurs, dans tous les domaines d'application. Bien que l'ouvrage s'adresse principalement aux professionnels de l'électronique, débutants ou confirmés, le public moins averti devrait en tirer profit également. En fait, en rendant accessibles et en clarifiant des propriétés fondamentales des antennes, en montrant les avancées dans la connaissance et la maîtrise de l'environnement électromagnétique, en faisant apparaître des perspectives de progrès, l'ouvrage contribuera à l'acceptation par le public d'une technologie dont la société moderne ne peut plus se passer.

Paris, le 26 juin 2009.

Maurice BELLANGER

Professeur au CNAM

Membre de l'Académie des technologies

TABLE DES MATIÈRES

Avant-propos	XI
Introduction	1
1 • Principe de fonctionnement des antennes	5
1.1 Le rôle des antennes	5
1.2 Différents types d'antennes	5
1.3 Le rôle de la diffraction	9
1.4 Classification des antennes en fonction du type de la source	10
2 • Principes théoriques pour l'étude des antennes	11
2.1 Équations de Maxwell	11
2.2 Potentiels électromagnétiques	21
2.3 Vecteur de Poynting	27
2.4 Théorèmes importants de l'électromagnétisme	30
3 • Calcul du rayonnement d'antennes de référence	37
3.1 Rayonnement créé par des courants	37
3.2 Diffraction par une ouverture et principe d'équivalence	51
3.3 Rayonnement des ouvertures planes	61
3.4 Différentes zones de rayonnement	68
Bibliographie	72
4 • Caractéristiques des antennes	73
4.1 Fonction caractéristique de rayonnement	73
4.2 Diagramme de rayonnement	74
4.3 Directivité	78
4.4 Gain d'une antenne	82
4.5 Facteur d'antenne	83
4.6 Polarisation d'une onde	84
4.7 Impédance d'une antenne	92
4.8 Alimentation	94
4.9 Largeur de bande	97
4.10 Bilan de liaison	99
4.11 Température de bruit d'une antenne	100
Bibliographie	101

5 • Réseaux d'antennes	103
5.1 Les réseaux d'antennes	103
5.2 Antennes multifaisceaux	115
5.3 Synthèse de réseaux	118
5.4 Alimentation des réseaux	121
Bibliographie	125
6 • Différents domaines d'utilisation des antennes	127
6.1 Gestion du spectre radiofréquence	127
6.2 Les systèmes de communications	130
6.3 Télévision et radiodiffusion FM	162
6.4 Radar	167
6.5 Télédétection	176
6.6 Radioastronomie	194
Bibliographie	196
7 • Systèmes multi-antennes	197
7.1 Éléments de traitement d'antenne	197
7.2 Diversité	213
7.3 Systèmes MIMO	220
Bibliographie	245
8 • Antennes définies selon le phénomène physique créant le rayonnement	249
8.1 Antennes boucles	249
8.2 Antennes à résonateurs	253
8.3 Antennes à ondes progressives	256
8.4 Antennes à ondes de surface et à ondes de fuite	260
8.5 Antennes Yagi-Uda	264
8.6 Antennes à réflecteurs	270
Bibliographie	275
9 • Antennes définies selon des caractéristiques systèmes	277
9.1 Différentes géométries d'antennes	277
9.2 Antennes définies par leur largeur de bande en fréquence	300
Bibliographie	319
10 • Mesures d'antennes	323
10.1 Introduction	323
10.2 Rappels sur les différentes zones de rayonnement	323
10.3 Diagramme de rayonnement et directivité	324
10.4 Conditions sur la mesure du diagramme de rayonnement	328
10.5 Gain	330
10.6 Polarisation	332
10.7 Impédance	333

10.8	Efficacité	335
	Bibliographie	336
11 •	Modélisation numérique du rayonnement des antennes	337
11.1	Introduction aux méthodes de modélisation numérique	337
11.2	Méthode intégrale associée à la méthode des moments	338
11.3	La méthode des éléments finis	350
11.4	La méthode des différences finies dans le domaine temporel (FDTD)	357
	Bibliographie	366
	Index	367

Les premières antennes sont apparues à la fin du XIX^e siècle, à une époque où les travaux sur l'électromagnétisme ont connu un développement considérable. Depuis, leur réalisation n'a cessé d'évoluer, d'abord, grâce aux progrès scientifiques de l'électromagnétisme, plus tard, sous la pression de nombreuses demandes technologiques dans des domaines d'application variés. L'essor actuel des communications impose des innovations importantes au niveau de la conception des systèmes et des antennes associées, dont les formes aujourd'hui très diverses varient beaucoup selon les utilisations : télécommunications mobiles, satellites, télévision, radio, identification, objets communicants...

Malgré cette grande diversité, toutes les antennes ont en commun de transformer un signal guidé en un signal rayonnant (ou réciproquement), dans un spectre électromagnétique relativement large allant des ondes radio aux hyperfréquences. Un principe fondamental régit leur rayonnement, celui de la diffraction des ondes.

Actuellement, la course à l'innovation concernant les systèmes de communication entraîne des études poussées dans le domaine des antennes. Dans ce contexte, les méthodes de conception, de mesures et de modélisation constituent une aide considérable. De nombreuses équipes contribuent à ce développement.

Cet ouvrage a pour objectif de fournir les connaissances nécessaires à la compréhension du fonctionnement des antennes. Nous proposons successivement :

- Des bases théoriques sur le fonctionnement des antennes.
- Un classement des nombreux types d'antennes existants en précisant, d'une part leurs caractéristiques à partir d'une modélisation, et d'autre part leurs utilisations dans les systèmes.
- Des éléments de conception applicables à chaque type. Ceux-ci sont complétés par des principes de mesure et de simulation numérique.

L'ouvrage débute par une présentation globale des antennes suivie du rappel de bases théoriques en électromagnétisme (chapitres 1 et 2). La théorie des antennes est ensuite développée en détail dans le chapitre 3. Des démonstrations y sont proposées en détail et permettront au lecteur qui le souhaite d'approfondir certains points théoriques. Le chapitre 4 s'appuie sur ces bases pour définir les caractéristiques générales du rayonnement. Celles-ci sont indispensables au dimensionnement d'une antenne dans le cadre d'une application précise. Le chapitre 5 concerne la mise en réseau des éléments rayonnants et les avantages qu'on peut en attendre.

Le chapitre 6 est consacré aux différents domaines d'applications des antennes : communications, diffusion, radars, télédétection, radioastronomie. Le lecteur pourra parcourir ce chapitre pour y découvrir toute la variété des domaines utilisant le rayonnement électromagnétique. Il pourra aussi s'attarder sur un type d'application qui le concerne plus particulièrement. Les domaines d'application ont été décrits de façon rapide en insistant sur le contexte, afin de mettre en avant les spécifications des antennes.

De nouveaux développements dans le domaine des communications ont prouvé que les performances étaient nettement améliorées par l'utilisation de systèmes multi-antennes. Le chapitre 7 en présente les avantages attendus. On y aborde les bases du traitement d'antenne et de la

diversité. Les principes du fonctionnement des systèmes MIMO (*Multi Input, Multi Output*) terminent ce chapitre.

Dans la suite, en s'appuyant sur la théorie, un grand nombre d'exemples est proposé en insistant sur les caractéristiques des différents types d'antennes. Dans le chapitre 8 les antennes sont regroupées selon les phénomènes physiques qui sont à la base de leur rayonnement (par exemple : antennes résonnantes, antennes à ondes de fuite, antennes Yagi-Uda). Le chapitre 9 est consacré aux antennes dont les caractéristiques sont imposées par des considérations vis-à-vis du système : constitution (par exemple : antennes guides, antennes miniatures) ou bien largeur de bande. Toutes les descriptions d'antennes s'attachent à donner les éléments essentiels intervenant dans leur conception. L'ensemble est largement illustré de schémas et d'abaques permettant leur dimensionnement.

Une étape de la conception consiste à valider la structure par des mesures dont les principales méthodes sont décrites dans le chapitre 10, uniquement consacré à ce point.

Dans le chapitre 11 sont présentés les principes de quelques méthodes de simulations numériques. Ces méthodes constituent une aide précieuse à la conception d'antennes complexes pour lesquelles le calcul analytique est impossible. Elles évitent aussi la réalisation de maquettes intermédiaires très coûteuses.

L'ensemble des chapitres s'organise en partant du principe général du rayonnement des antennes pour aller vers leurs diversités. Cet ouvrage peut être lu, soit de façon théorique comme base aux calculs de rayonnement, soit de façon pratique dans le but de choisir un type d'antenne et de poursuivre jusqu'à sa conception. Ces aspects complémentaires permettront aux débutants comme aux spécialistes de trouver un intérêt dans la lecture de cet ouvrage. Il contient les éléments indispensables pour des études en recherche et développement dans le domaine des antennes.

■ À propos des auteurs

Cet ouvrage est le résultat du travail d'une équipe de spécialistes en antennes, électromagnétisme, systèmes hyperfréquences et traitement du signal, tous enseignants chercheurs à l'université de Paris-Est. Les auteurs ont une large expérience de l'enseignement dans le domaine des antennes. Ils enseignent au niveau master et en école d'ingénieurs. Leur compétence est reconnue grâce à leurs nombreux travaux de recherche tant théoriques qu'appliqués, menés en partie en collaboration avec des partenaires du milieu industriel.

Laurent Cirio est diplômé de l'université de Nice Sophia-Antipolis. Maître de conférences à l'université Paris-Est Marne-la-Vallée, ses travaux portent sur la conception et la caractérisation des antennes planaires et leurs applications.

Christian Ripoll est diplômé de l'University College London, professeur associé en systèmes électroniques à l'université Paris-Est-ESIEE-Paris.

Geneviève Baudoin est diplômée de Télécom Paris, habilitée à diriger des recherches de l'université Paris-Est, professeur en traitement du signal et communications numériques et directrice de la recherche à l'université Paris-Est-ESIEE-Paris.

Jean-François Bercher est ingénieur INPG, docteur de l'université de Paris-Sud, professeur en traitement de signal à l'université Paris-Est-ESIEE-Paris.

Martine Villegas est diplômée de l'ENSEA, habilitée à diriger des recherches de l'université Paris-Est, professeur en architecture de systèmes à l'université Paris-Est-ESIEE-Paris, responsable ESIEE pour le master *Systèmes de communications hautes fréquences*.

Odile Picon, ancienne élève de l'École normale supérieure de Fontenay-aux-Roses, agrégée de physique, docteur de l'université de Paris-Sud, docteur d'état de l'université de Rennes I, est professeur en électromagnétisme et antennes à l'université Paris-Est-Marne-la-Vallée, directrice du laboratoire ESYCOM, (commun à l'UPE-MLV, à l'UPE-ESIEE-Paris et au Cnam) et responsable du master *Systèmes de communications hautes fréquences* de l'UPE-MLV.

INTRODUCTION

Les antennes sont des dispositifs utilisés pour rayonner le champ électromagnétique dans l'espace ou pour le capter. Comme nous le verrons dans cet ouvrage, il existe de nombreux types d'antennes. Il est important d'avoir une connaissance globale de leur fonctionnement lors du choix d'un dispositif rayonnant. La compréhension de ce fonctionnement aidera, d'une part à utiliser l'antenne au mieux de ses performances et d'autre part, à en réaliser une conception optimale.

Les techniques de conception et de réalisation d'antennes se sont affinées au fur et à mesure que le domaine de l'électromagnétisme s'est développé. C'est un domaine relativement récent, puisque c'est en s'appuyant sur les équations de Maxwell que tous les développements théoriques et techniques ont pu progresser. Les avancées dans ce domaine ont été rapides car touchant aux transmissions radioélectriques dont le nombre d'applications est considérable.

Dans cet ouvrage, nous présenterons la théorie liée aux antennes qui servira de base à leur conception. Nous tenterons de classer les antennes par rapport à leur principe de fonctionnement et par rapport à leur rôle dans les systèmes. Nous terminerons par des principes de modélisation numérique.

La transmission d'information

La transmission d'information s'effectue généralement grâce à une onde porteuse, caractérisée par sa fréquence. C'est une onde sinusoïdale. Sa modulation par un signal de plus basse fréquence représente l'information à transmettre. Nous n'aborderons pas ici les différents types de modulations existant qui peuvent être analogiques ou numériques.

Les ondes porteuses sont de différentes natures. Ce sont les caractéristiques du système qui permettent de choisir le type de transmission. Les ondes les plus utilisées pour la transmission sont les ondes acoustiques et les ondes électromagnétiques. Nous ne parlerons, dans cet ouvrage, que de ce dernier type.

Bien qu'une onde électromagnétique n'ait besoin d'aucun support pour se propager, il se trouve que, dans son utilisation pour la transmission d'information, elle se propage à travers un milieu. Les différents milieux dans lesquels s'effectue la propagation d'ondes électromagnétiques sont :

- les conducteurs (transmission filaire),
- la silice (transmissions par fibres optiques),
- l'air (transmissions hertziennes).

L'exemple de transmission dans les conducteurs est celui du téléphone filaire. Afin d'augmenter le débit d'informations, la téléphonie fixe utilise un autre support, constitué par la fibre optique. Dans ce cas, l'onde optique constitue la porteuse. Le troisième cas est celui de la téléphonie sans fil dont le canal de propagation est l'air. Actuellement, les applications utilisant ce canal de propagation sont nombreuses : télévision, téléphonie mobile, téléphonie fixe par liaisons hertziennes ou satellites, radar, télédétection, etc., d'où le rôle des antennes dans ces systèmes.

Rappelons ici le spectre des ondes électromagnétiques afin de bien situer les différentes utilisations de celui-ci. La figure I.1 présente une échelle par rapport aux longueurs d'ondes (λ) et une autre par rapport aux fréquences (f), sachant que, dans le vide, ces deux grandeurs sont liées par

l'intermédiaire de la vitesse de la lumière dans le vide c (cette valeur est sensiblement la même dans le vide et dans l'air) :

$$\lambda = \frac{c}{f} \quad \text{avec} \quad c = 3 \times 10^8 \text{ ms}^{-1}$$

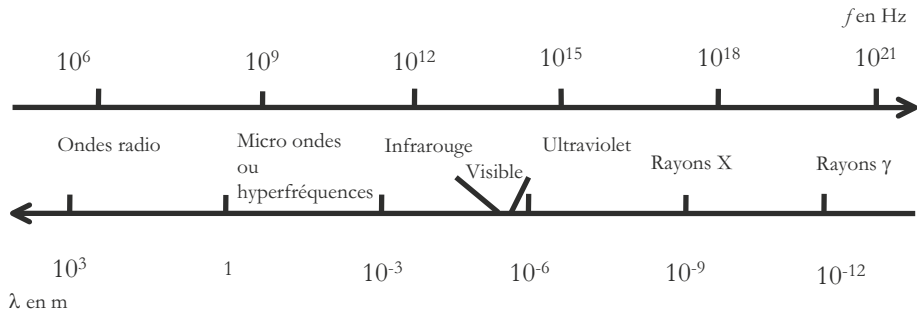


Figure I.1 – Spectre des ondes électromagnétiques.

Les ondes radiofréquences occupent la région du spectre correspondant aux basses fréquences du spectre électromagnétique. Les longueurs d'ondes correspondant aux ondes longues sont de l'ordre de 1 000 m. En montant en fréquence, on trouve les ondes courtes, puis un peu plus haut les micro-ondes ou hyperfréquences. Ensuite, pour des longueurs d'ondes de l'ordre du micromètre se trouvent les infrarouges, domaine assez étendu dont la limite haute en fréquence est le visible. Cette zone s'étend de 0,4 à 0,8 micromètre de longueur d'onde. Ensuite vient le domaine des ultra-violet. Les longueurs d'ondes de l'ordre de l'angström définissent le domaine des rayons X. Ceux-ci permettent de scruter la structure atomique en utilisant le phénomène de diffraction sur les atomes, dont la taille est l'ordre de grandeur de la longueur d'onde. Plus haut en fréquence, on trouve les rayons gamma, accompagnant les réactions nucléaires.

La partie du spectre qui concerne l'utilisation des antennes est celui qui va des ondes radio aux ondes submillimétriques ou quasi optiques qui sont intermédiaires entre les micro-ondes et l'infrarouge. Les hyperfréquences ou micro-ondes occupent la bande de fréquences comprise entre 300 MHz et 300 GHz. Cette partie du spectre est divisée en bandes de fréquences standardisées (voir chapitre 6).

La transmission d'ondes électromagnétiques dans l'air

Les ondes électromagnétiques se propagent bien dans l'atmosphère pour des fréquences basses du spectre radiofréquence jusqu'à des fréquences de l'ordre d'une vingtaine de giga-hertz. L'atténuation due à l'atmosphère se manifeste au-delà de 20 GHz. La figure I.2 donne les valeurs de l'absorption des ondes relativement aux deux composants les plus absorbants dans cette partie du spectre (vapeur d'eau et oxygène), pour 1 km d'atmosphère standard en fonction de la fréquence. On constate une atténuation globale qui augmente au fur et à mesure que la fréquence augmente. Pour certaines fréquences (environ : 22, 60, 120 et 180 GHz), l'atmosphère absorbe spécifiquement les ondes électromagnétiques. Ces phénomènes d'absorption sont dus à des interactions (résonantes ou non) des molécules avec les ondes électromagnétiques. Ils ont des explications différentes. Pour les fréquences correspondant au spectre micro-ondes, ils sont dus soit aux résonances rotationnelles, soit à la polarisation d'orientation liée aux moments dipolaires des molécules, compte tenu des énergies mises en jeu. La molécule d'eau possède un moment dipolaire électrique, alors que la molécule d'oxygène possède un moment dipolaire magnétique.

Les bandes d'absorption à 60 et 120 GHz sont dues à la résonance rotationnelle des molécules d'oxygène. À 22 et 180 GHz, un phénomène analogue apparaît pour les molécules d'eau. Les bandes sont élargies par des résonances vibrationnelles, en particulier à 60 GHz.

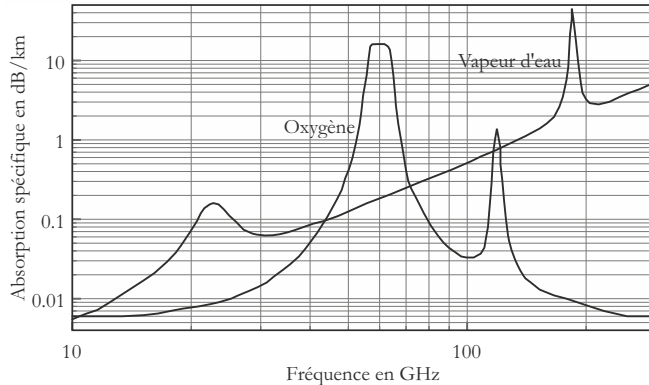


Figure 1.2 – Absorption par la vapeur d'eau et l'oxygène des ondes électromagnétiques pour 1 km d'atmosphère standard.

L'observation du diagramme d'absorption de l'oxygène et de l'eau conduit à plusieurs remarques concernant l'utilisation du spectre électromagnétique.

- Tout d'abord il faut remarquer que, d'une manière générale, l'objectif des communications est d'obtenir un débit toujours plus grand. Or, plus la fréquence d'utilisation est élevée, plus la largeur de bande utilisable est grande et donc plus le débit supporté est élevé. Cette tendance à l'augmentation de la fréquence des systèmes de communication au cours de leur évolution est incontestable.
- L'utilisation de fréquences basses (< 20 GHz) ne pose pas de problèmes particuliers d'atténuation.
- L'atténuation existant autour de 22 GHz est relativement peu importante. Elle peut cependant dégrader le signal. Pour les transmissions satellites, en particulier, pour lesquelles la distance atmosphérique parcourue est grande, les systèmes sont conçus de façon à se placer sur les bords de cette bande. C'est actuellement le cas de certains systèmes satellites japonais utilisant les bandes 20-30 GHz. Mais en général les bandes satellites utilisées sont plus basses.
- Une atténuation très forte existe à 60 GHz, qui est de l'ordre de – 16 dB/km. Il est envisagé d'utiliser cette bande pour des transmissions intrabâtiment. On profite alors de l'atténuation existant pour éviter la pollution électromagnétique d'un bâtiment à l'autre.
- Vers le haut du spectre, l'atténuation est très forte. Les fréquences utilisées dans cette partie sont celles qui intéressent les radioastronomes. C'est la raison pour laquelle les observatoires de radioastronomie sont placés dans des lieux où l'atmosphère a un rôle moindre. En effet, en altitude, la densité de la couche atmosphérique traversée est plus faible et l'air est plus sec.
- Signalons aussi que l'atténuation par les hydrométéores (pluie et neige) est un phénomène à prendre en compte pour l'évaluation de la qualité d'une liaison lors de transmissions terrestres.

1 • PRINCIPE DE FONCTIONNEMENT DES ANTENNES

1.1 Le rôle des antennes

☐ Antenne d'émission

Afin d'assurer la propagation dans l'air, il est nécessaire qu'un dispositif génère une onde rayonnée. Le rôle de l'antenne d'émission est de transformer la puissance électromagnétique guidée, issue d'un générateur en une puissance rayonnée. Dans ce sens, c'est un transducteur.

☐ Antenne de réception

De façon inverse, la puissance rayonnée peut être captée par une antenne de réception. Dans ce sens, l'antenne apparaît comme un capteur et un transformateur de puissance rayonnée en puissance électromagnétique guidée. Elle joue le même rôle qu'un télescope qui capte la lumière issue des étoiles et la transforme.

☐ Réciprocité

Dans la plupart des cas, une antenne peut être utilisée en réception ou en émission avec les mêmes propriétés rayonnantes. On dit que son fonctionnement est réciproque. Ceci est une conséquence du théorème de réciprocité qui sera démontré plus loin. Dans quelques cas exceptionnels pour lesquels les antennes comportent des matériaux non linéaires ou bien anisotropes, elles ne sont pas réciproques.

Du fait de la réciprocité des antennes, il ne sera pratiquement jamais fait de différence entre le rayonnement en émission ou en réception. Les qualités qui seront annoncées pour une antenne le seront dans les deux modes de fonctionnement, sans que cela soit précisé dans la plupart des cas.

1.2 Différents types d'antennes

Afin de comprendre comment s'effectue cette transformation entre la puissance guidée et la puissance rayonnée, nous allons présenter un certain nombre d'antennes. Elles sont classées ici selon un ordre qui suit approximativement leur chronologie d'apparition. Il n'est pas question dans ce paragraphe de présenter tous les types d'antennes, mais d'en introduire certains des plus utilisés.

En conclusion, nous aboutirons à un classement des antennes selon le type de la source rayonnante qui apparaîtra soit comme un courant électrique, soit comme une surface caractérisée par un champ électrique.

☐ Antenne dipolaire

L'antenne dipolaire est constituée de deux fils alignés, très courts et reliés chacun à deux fils parallèles et très proches constituant une ligne bifilaire (figure 1.1). En émission, cette ligne est reliée à un générateur alternatif, caractérisé par sa fréquence et son impédance interne. À la réception, la ligne bifilaire est branchée sur un récepteur.

Dans la ligne bifilaire, les courants sont de sens contraire, alors que dans le dipôle les courants sont dans le même sens. L'influence de ces deux courants s'annule dans la ligne bifilaire. Ce sont les courants variables, de même sens, qui rayonnent et créent l'onde électromagnétique dans l'espace. Étant donnée la symétrie du dipôle, le rayonnement s'effectue autour de l'axe, matérialisé par le fil. Il est isotrope dans un plan perpendiculaire à cet axe. Le rayonnement est nul dans la direction du fil. On ne peut donc pas parler d'un rayonnement isotrope.

À l'extrémité de chaque fil apparaissent des charges de signes opposées dont l'existence s'explique par la conservation de la charge. En effet, la relation de conservation suivante lie les charges au courant :

$$I = \frac{dq}{dt}$$

D'autres antennes de même type sont obtenues avec des fils rayonnants plus longs. Ces antennes de type filaires ont de nombreuses applications qui seront décrites dans les chapitres suivants. Citons rapidement les antennes pour récepteurs radio, les antennes des talkies-walkies, etc.

□ Boucle magnétique

La boucle magnétique est constituée d'un fil conducteur ayant une forme qui permet le retour du fil sur lui-même (figure 1.2). La boucle est ainsi branchée sur une ligne bifilaire reliée au générateur.

Le rayonnement, à grande distance, est maximal dans le plan de la boucle et s'effectue de façon radiale. Le courant circulant dans le fil crée un champ magnétique qui se propage. Sa variation engendre le champ électrique associé, d'où le rayonnement électromagnétique associé.

En champ lointain, les boucles magnétiques ont été très utilisées pour les récepteurs de grandes ondes radio sous forme d'un cadre sur lequel étaient enroulées plusieurs spires de fil. En champ proche, on les utilise dans tous les dispositifs RFID (identification radio fréquence). Les cartes à puce sans contact sont munies de ce type d'antenne, incluse dans le support plastique. Les détecteurs d'objets métalliques sont aussi des boucles magnétiques sensibles au champ magnétique.

□ Antenne cornet

Un dispositif très utilisé pour la propagation d'ondes guidées est le guide d'onde rectangulaire. Sa qualité de transmission est excellente. Pour cette raison, il est utilisé en haute en fréquence. Son utilisation est très répandue en hyperfréquences. Le transformateur de puissance électromagnétique guidée en puissance rayonnée est l'antenne cornet (figure 1.3). Sa forme permet de passer graduellement des dimensions du guide d'onde à l'espace libre. L'onde est ainsi naturellement projetée dans l'espace libre. C'est le même principe que le cornet acoustique.

Les transitions présentent des formes variées : linéaires, exponentielles... Le cornet sert de dispositif d'adaptation entre l'impédance du cornet et celle du vide.

De façon très naturelle, le rayonnement a lieu dans l'axe du guide d'onde. Cette antenne est plus directive que les précédentes, dans la mesure où la puissance n'est émise que dans une région de l'espace limitée.

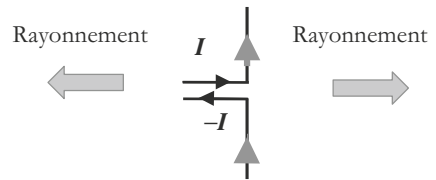


Figure 1.1 – Antenne dipolaire.

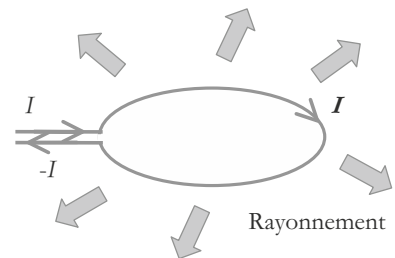


Figure 1.2 – Boucle magnétique.

Le guide d'onde est un dispositif hyperfréquence très utilisé du fait des pertes très faibles engendrées par la propagation dans celui-ci, même à hautes fréquences et de sa capacité à supporter de la puissance. Les antennes cornets qui lui sont associées sont donc aussi très utilisées comme moyen de transformation de l'onde guidée en onde rayonnée. On les retrouve, dans toutes les bandes de fréquences, dans de nombreux systèmes tels que les radars, les antennes satellites...

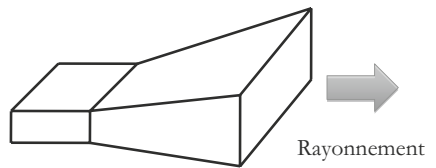


Figure 1.3 – Antenne cornet.

□ Réseau de fentes

Toujours en utilisant le guide d'onde comme dispositif de transmission, il est possible d'envisager un rayonnement dans une direction différente de l'axe du guide, en usinant des fentes dans le corps du guide (figure 1.4)

Le rayonnement s'effectue alors perpendiculairement au plan troué du guide.

Ce type de dispositif est utilisé lorsque le rayonnement doit être localisé. Par exemple, dans des tunnels, où la transmission des ondes s'effectue mal, on peut placer un réseau de fentes rayonnantes. En général la ligne est en haut du tunnel avec émission vers le bas.

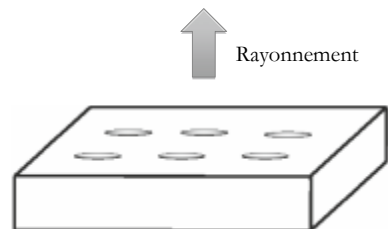


Figure 1.4 – Réseau de fentes.

□ Antenne à réflecteur parabolique

L'antenne à réflecteur est constituée de la source d'émission associée à une partie métallique réfléchissante, souvent de forme parabolique (figure 1.5)

La source, placée au foyer de la parabole envoie l'onde vers le réflecteur parabolique. Selon la propriété bien connue de la parabole, tous les rayons sont réfléchis parallèlement. Ce type d'antenne est utilisé pour viser dans une direction très précise, puisque tous les rayons passant par le foyer sortent parallèles. Par décalage de la source dans le plan focal, les rayons parallèles à la sortie du réflecteur, peuvent présenter une inclinaison par rapport à l'axe de la parabole.

Ces antennes permettent de recevoir un signal d'un satellite, placé à très grande distance. Les antennes de ce type sont très répandues pour la réception de la télévision. Leur orientation est choisie de façon à viser un satellite particulier.

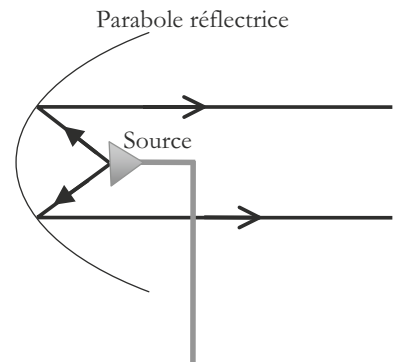


Figure 1.5 – Antenne à réflecteur parabolique.

Afin d'éviter les perturbations par la pluie ou la neige, ces antennes sont souvent recouvertes d'un radôme. C'est le cas des antennes très exposées aux conditions climatiques, utilisées pour les transmissions hertziennes. Elles sont reconnaissables par leur forme, parabolique à l'arrière et conique à l'avant du fait de la forme du radôme qui protège la source, placée au foyer.

□ Antennes de type Cassegrain

Une variante des antennes précédentes consiste à utiliser un réflecteur principal et un réflecteur secondaire, comme dans le montage Cassegrain (figure 1.6). Ce nom provient du télescope du même nom reposant sur le même principe.

Les rayons issus de la source se réfléchissent sur un premier réflecteur de forme hyperbolique, puis sur le réflecteur principal de forme parabolique. Les rayons ressortent parallèlement. La propriété de l'antenne parabolique est ainsi conservée. L'intérêt de ce type d'antenne est d'être moins sensible aux parasites provenant de l'extérieur de la parabole. De plus, les câbles reliant la source à l'électronique sont plus courts que dans les systèmes d'alimentation d'une antenne parabolique. La qualité du signal s'en trouve améliorée.

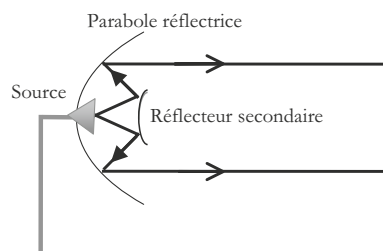


Figure 1.6 – Antenne à réflecteur parabolique.

□ Antennes plaquées

L'antenne plaquée, appelée aussi antenne patch est un type récent d'antenne dont le développement et l'utilisation sont de plus en plus fréquents. Elle est constituée d'un diélectrique, possédant un plan de masse métallique sur une face. Sur l'autre face, une gravure métallique permet de supporter des courants de surface qui créent le rayonnement électromagnétique (figure 1.7). Les courants sont amenés du générateur à l'antenne par une ligne micro ruban.

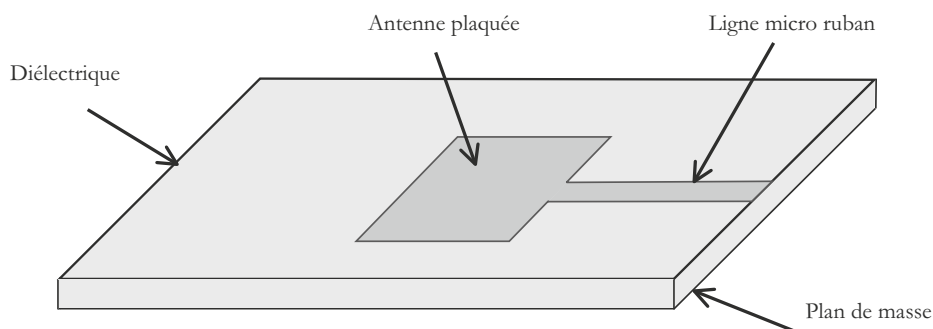


Figure 1.7 – Antenne plaquée.

Elle présente l'avantage du poids sur certaines antennes décrites précédemment. Les gravures des parties métalliques peuvent prendre des formes très variées en fonction des objectifs fixés pour la répartition du rayonnement dans l'espace. Ceci donne une grande souplesse de conception.

La géométrie des antennes plaquées peut être multipliée à l'infini. Dans le chapitre 7, les différents fonctionnements de ces antennes seront décrits.

Pour terminer avec l'introduction à ce genre d'antenne, mentionnons les antennes à alimentation par couplage (figure 1.8). L'alimentation se trouve prise en sandwich entre deux diélectriques. Un couplage électromagnétique entre l'extrémité de la ligne micro ruban et le patch, qui se trouve au-dessus du dispositif, permet d'exciter l'antenne. Le diélectrique supérieur joue alors un rôle d'écran pour la ligne d'alimentation, qui sinon pourrait éventuellement perturber le rayonnement. Il est alors intéressant de placer l'électronique au niveau de ce second diélectrique pour les antennes actives.

□ Antennes actives

Les progrès réalisés sur la fabrication des antennes plaquées, rendent possible le report d'un circuit actif sur l'antenne. L'antenne a des fonctions qui dépassent son rôle simple de transformateur d'énergie. Selon les fonctions électroniques adjointes, on obtient un dispositif complexe. On parle ainsi d'antennes intelligentes si le dispositif a une partie de contrôle et de commande.

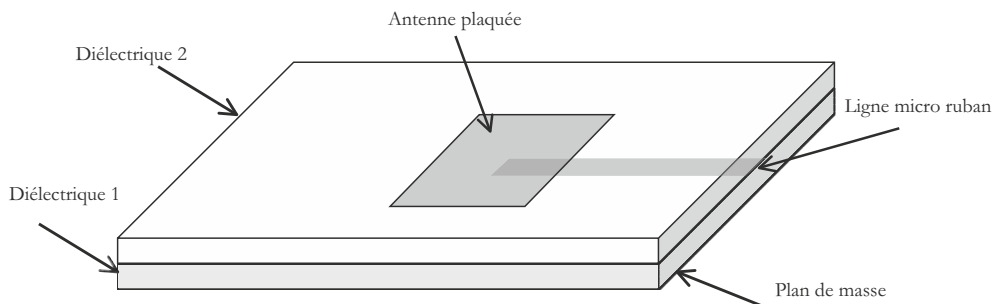


Figure 1.8 – Antenne plaquée avec alimentation à couplage électromagnétique.

Les applications des antennes actives sont très variées. Elles sont utilisées pour des tâches nécessitant :

- de la commutation,
- du déphasage dans les réseaux d'antennes,
- de l'amplification (de puissance à l'émission ou faible bruit à la réception)
- de l'agilité en fréquence, etc.

On distingue les antennes actives intégrées des antennes hybrides sur lesquelles des composants sont reportés. L'intérêt actuel porte toutefois sur les antennes intégrées, pour lesquelles l'antenne est au plus près du circuit intégré car réalisée en même temps, sur le même support.

Il existe d'autres types d'antennes qui seront vus dans la suite. Ce n'est ici qu'une introduction à leurs principes de fonctionnement.

1.3 Le rôle de la diffraction

Le principe de fonctionnement des antennes repose sur le phénomène de diffraction. Dans la suite de cet ouvrage, nous en verrons les différents aspects théoriques. Il résulte de ce point que c'est la forme du dispositif diffractant qui détermine les caractéristiques du rayonnement. Les antennes ont des dimensions qui sont du même ordre de grandeur que celui de la longueur d'onde rayonnée.

La diffraction est un phénomène lié au caractère ondulatoire d'une grandeur. En électromagnétisme, la diffraction joue sur le champ électromagnétique. L'expérience simple mettant en évidence ce phénomène est celui d'une onde plane qui interagit avec une ouverture percée dans un écran. Si la dimension du trou est nettement plus grande que la longueur d'onde, l'onde électromagnétique va passer sans perturbation apparente. Donc elle suit une trajectoire rectiligne. Si, au contraire, la dimension de l'ouverture est du même ordre de grandeur que la longueur d'onde, le phénomène de diffraction se manifeste par le fait que la répartition de la puissance s'étale autour de la direction d'incidence, après passage à travers l'écran (figure 1.9)

L'analyse de la répartition de la puissance dans le demi-plan situé après l'ouverture est complexe. Elle sera étudiée ultérieurement, grâce aux équations de Maxwell. L'interaction des bords de l'ouverture avec l'onde électromagnétique est plus forte si la dimension D est petite. Huyghens a posé comme principe que la perturbation de l'onde à la traversée de l'écran troué se manifeste de la même façon que si la surface de l'ouverture était tapissée d'une infinité de petites sources, dont la phase relative de l'une à l'autre était la même que la phase relative de l'onde incidente en chacun des points. Le champ électromagnétique après l'écran résulte de l'interférence des ondes issues des sources secondaires. Ce principe permet des calculs dans des cas simples de diffraction.

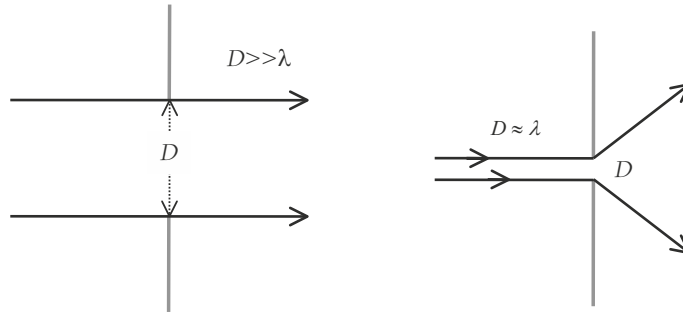


Figure 1.9 – Illustration du phénomène de diffraction.

Chaque source secondaire émet une onde sphérique, avec sa propre phase. Supposons que les sources secondaires soient toutes en phase parce que l'onde incidente est normale à la surface. Leurs interférences impliquent une répartition faisant apparaître des maxima et de minima en fonction de la direction d'observation, avec un maximum dans la direction de l'incidence. Les maxima correspondent à des directions d'interférences constructives. Si les sources sont peu nombreuses (surface d'ouverture petite), les directions d'interférences constructives sont identifiables. Au contraire plus les sources seront nombreuses et étalées dans l'espace (surface d'ouverture large), plus il est difficile de trouver des directions dans lesquelles il existe des interférences constructives, mise à part la direction correspondant au déphasage nul, c'est-à-dire la direction d'incidence. Donc, à la limite, lorsque les sources sont très nombreuses, l'interférence constructive n'a lieu que dans la direction d'incidence. On est alors dans le cas où le trou est grand par rapport à la longueur d'onde.

En s'appuyant sur le principe d'Huyghens, il est facile de concevoir que la forme de l'ouverture rayonnante va jouer un rôle important dans la répartition de la puissance dans le demi-espace concerné par ce phénomène, d'où l'importance de la géométrie de l'antenne dans sa conception.

1.4 Classification des antennes en fonction du type de la source

Il existe de nombreuses façons de classer les antennes. Nous en proposons certaines dans cet ouvrage. Une première classification repose sur le type de la source rayonnante. En effet, parmi les différentes antennes décrites dans ce chapitre, il est évident pour certaines que la source du rayonnement est constituée de courants (dipôle, boucle, antenne planaire). Ces courants sont soit *linéiques*, soit *surfaciques*.

L'interprétation du rayonnement est alors aisée : les courants variables vont créer un rayonnement dans tout l'espace selon leur répartition.

Comment alors interpréter le rayonnement d'un cornet ? Certes, il existe des courants surfaciques sur les parois du cornet qui permettraient de calculer le champ rayonné. On préfère un autre point de vue qui est celui de dire que, sur la surface de sortie du cornet, il existe un champ électrique dont la répartition impose la forme du champ rayonné dans tout l'espace. On parle alors de surface rayonnante.

On distingue alors :

- le rayonnement par les courants,
- le rayonnement par les ouvertures.

Ces deux types d'antennes seront donc étudiés séparément en s'appuyant sur des méthodes différentes.

2 • PRINCIPES THÉORIQUES POUR L'ÉTUDE DES ANTENNES

Dans ce chapitre, nous présentons les bases d'électromagnétisme relatives au fonctionnement des antennes. Nous partirons comme postulat des équations de Maxwell qui décrivent la variation spatio-temporelle du champ électromagnétique.

2.1 Équations de Maxwell

Les équations de Maxwell contiennent pratiquement toutes les informations concernant les caractéristiques du champ électromagnétique. Nous allons voir, au cours de ce chapitre, comment en extraire les propriétés. Les antennes fonctionnant dans le vide, on utilisera essentiellement les équations de Maxwell dans le vide. Cependant dans certains cas, la prise en compte du matériau constituant l'antenne est nécessaire. C'est pourquoi, après avoir présenté les équations de Maxwell dans le vide, nous décrirons rapidement les équations de Maxwell dans la matière.

2.1.1 Équations de Maxwell dans le vide

Nous supposons que le lecteur a déjà étudié précédemment l'électrostatique et la magnétostatique. Il sait donc ce que sont un champ électrique statique $\vec{E}$ et une induction magnétique statique $\vec{B}$. Nous allons étudier dans la suite non seulement la variation spatiale de ces champs mais aussi leur variation temporelle. L'association du champ électrique $\vec{E}(\vec{r}, t)$ et du champ magnétique $\vec{H}(\vec{r}, t)$, tous deux variant dans l'espace et dans le temps, définit le champ électromagnétique. Nous verrons dans un paragraphe ultérieur ce qu'on appelle champ magnétique $\vec{H}(\vec{r}, t)$. Pour l'instant, nous retiendrons que, dans le vide, c'est un vecteur colinéaire à l'induction magnétique. On se contentera donc, lorsqu'on est dans le vide d'étudier les variations de $\vec{E}(\vec{r}, t)$ et $\vec{B}(\vec{r}, t)$.

Avant de se lancer dans les équations qui décrivent le problème, définissons la notion de vide en l'électromagnétisme.

■ Notion de vide

Le vide pour l'électromagnétisme représente un espace ne contenant pas de matière mais pouvant contenir des charges électriques, sans support matériel. C'est un espace idéalisé. En première approximation, nous pouvons considérer qu'un espace rempli d'air et de charges électriques présente les propriétés du vide. Par extension, ce milieu peut aussi contenir des conducteurs dont on prend en compte uniquement les propriétés électriques.

■ L'électromagnétisme et les équations de Maxwell

C'est en 1873 que Maxwell publia sous une forme achevée les quatre équations couplées qui permettent d'interpréter pratiquement tous les phénomènes rencontrés en électromagnétisme

Lorsqu'on considère des charges variables dans le temps et dans l'espace, on constate qu'elles créent un champ électrique et une induction magnétique. Ces deux grandeurs sont vectorielles et varient dans l'espace et le temps. On les notera respectivement :

$$\vec{E}(\vec{r}, t) \quad \text{et} \quad \vec{B}(\vec{r}, t)$$

Les charges, variables dans le temps ou dans l'espace, qui ont donné naissance au champ électromagnétique sont appelées les sources. Elles peuvent apparaître sous la forme d'une densité volumique de charges notée $\rho(\vec{r}, t)$ ou d'une densité de courant notée $\vec{j}(\vec{r}, t)$.

■ Équations de Maxwell locales

Les équations de Maxwell expriment le comportement du champ électromagnétique en relation avec les sources qui lui ont donné naissance. Ces équations différentielles contiennent toute l'information permettant de résoudre les problèmes d'électromagnétisme.

Dans le vide elles s'écrivent sous la forme :

$$\text{rot}(\vec{E}) = -\frac{\partial \vec{B}}{\partial t} \quad [2.1]$$

$$\text{div}(\vec{B}) = 0 \quad [2.2]$$

$$\text{rot}(\vec{B}) = \mu_0 \left(\vec{j} + \epsilon_0 \frac{\partial \vec{E}}{\partial t} \right) \quad [2.3]$$

$$\text{div}(\vec{E}) = \frac{\rho}{\epsilon_0} \quad [2.4]$$

Ce système d'équations couplées lie les dérivées spatiales et temporelles du champ électrique et de l'induction magnétique à leurs sources. Toutes les grandeurs varient avec l'espace et le temps. Une des difficultés de l'électromagnétisme vient de cette représentation de grandeurs vectorielles, variables dans un espace à quatre dimensions (trois dimensions d'espace et une pour le temps). Heureusement des outils puissants de calculs vectoriels existent qui permettent la résolution des problèmes d'électromagnétisme, comme nous allons le voir dans la suite.

Les deux premières équations ne font intervenir que le champ électromagnétique. Elles sont appelées équations intrinsèques. Les deux suivantes contiennent sa relation aux sources.

Dans ces équations, il apparaît deux constantes caractéristiques du vide :

- La permittivité du vide est notée ϵ_0 . C'est une grandeur constante qui caractérise électriquement le vide. Sa valeur, dans les unités du système international (noté par la suite : U.S.I.) est :

$$\epsilon_0 = \frac{1}{36\pi \cdot 10^9} \text{ U.S.I.}$$

- La perméabilité du vide est notée μ_0 . Elle caractérise le vide d'un point de vue magnétique. Sa valeur est :

$$\mu_0 = 4\pi \cdot 10^{-7} \text{ U.S.I.}$$

Ces deux constantes sont liées à la vitesse de la lumière ($c = 3 \cdot 10^8 \text{ ms}^{-1}$) :

$$c^2 = \frac{1}{\epsilon_0 \mu_0}$$

Les équations [2.1] à [2.4] sont locales car valables en chaque point de l'espace.

Dans l'équation [2.3], le terme $\vec{j}$ est le courant de conduction et le terme qui lui est homogène

$\epsilon_0 \frac{\partial \vec{E}}{\partial t}$ est appelé courant de déplacement.

□ Conservation de la charge

La conservation de la charge est un grand principe de physique qui est contenu dans les équations de Maxwell.

Prenons la divergence de l'équation [2.3]. Sachant que la divergence d'un rotationnel est toujours nulle et que μ_0 et ϵ_0 sont des constantes, on obtient :

$$\mu_0 \operatorname{div}(\vec{j}) + \mu_0 \frac{\partial}{\partial t}(\epsilon_0 \operatorname{div}(\vec{E})) = 0$$

$$\text{soit} \quad \operatorname{div}(\vec{j}) + \frac{\partial}{\partial t}(\rho) = 0 \quad [2.5]$$

Cette équation est appelée équation locale de conservation de la charge. Il n'est pas évident, sous cette forme, de comprendre cette notion de conservation. Pour bien interpréter cette équation, il suffit de la transformer en sa forme intégrale. Considérons un volume (V) limité par une surface (S) comme sur la figure 2.1.

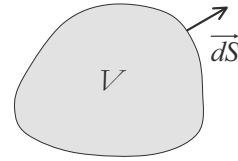


Figure 2.1 – Volume d'intégration pour le théorème de la divergence.

L'équation locale est valable en chaque point du volume. On en déduit donc :

$$\iiint_V (\operatorname{div}(\vec{j}) + \frac{\partial}{\partial t}(\rho)) dv = 0$$

Puisque les coordonnées d'espace et de temps sont indépendantes et considérant le volume (V) fixe dans le temps, on en déduit :

$$\iiint_V (\operatorname{div}(\vec{j})) dv = -\frac{d}{dt} \iiint_V \rho dv$$

Appliquant le théorème de la divergence, appelé aussi théorème d'Ostrogradski et notant Q la charge totale contenue dans le volume (V), on obtient :

$$\iint_S \vec{j} \cdot \vec{dS} = -\frac{dQ}{dt}$$

L'interprétation de cette équation intégrale implique que le flux sortant du vecteur densité de courant est égal à la diminution de la charge totale contenue dans (V) par unité de temps. Il n'y a donc pas d'accumulation de charge puisque toutes les charges qui disparaissent (ou apparaissent) au cours du temps sont compensées par un flux de courant sortant (ou respectivement entrant). On comprend mieux sur cette forme intégrale le sens des termes « conservation de la charge ».

Le flux de la densité de courant sortant à travers une surface est défini comme l'intensité du courant sortant du volume, notée I_S :

$$I_s = \iint_S \vec{j} \cdot \vec{dS}$$

Nous obtenons alors l'équation bien connue en électricité :

$$I_s = -\frac{dQ}{dt}$$

■ Formes intégrales des équations de Maxwell

L'interprétation des équations de Maxwell sous leurs formes intégrales conduit à des résultats importants. Nous allons voir, sur chacune d'elles, ce que cela implique.

□ Conservation du flux d'induction magnétique

Considérons un volume V , limité par une surface fermée S , orientée vers l'extérieur (figure 2.2). Cette surface fermée s'appuie sur les lignes de champ magnétique définissant le tube de champ et est fermée aux deux extrémités par deux surfaces orientées S_1 et S_2 .

En intégrant l'équation [2.2] sur le volume V , nous obtenons :

$$\iiint_V \text{div}(\vec{B}) dv = 0$$

Le théorème de la divergence permet de montrer que le flux magnétique à travers la surface fermée est nul. L'intégrale est la somme de trois intégrales :

$$\iint_{S_1} \vec{B} \cdot d\vec{S} + \iint_{S_2} \vec{B} \cdot d\vec{S} + \iint_{S_{lat}} \vec{B} \cdot d\vec{S} = 0$$

La dernière intégrale est nulle car elle s'effectue sur une surface dont les vecteurs élémentaires sont perpendiculaires au champ magnétique. Orientons les deux surfaces S_1 et S_2 dans le même sens. Pour cela posons $d\vec{S}_2 = -d\vec{S}_1$. Il vient :

$$\iint_{S_1} \vec{B} \cdot d\vec{S} = \iint_{S_2} \vec{B} \cdot d\vec{S}$$

Cette relation de conservation du flux magnétique permet d'affirmer que, dans les zones où les lignes de champ se resserrent, l'induction magnétique est plus forte.

□ Équation de Maxwell – Gauss

Considérons l'équation [2.4] qui fait apparaître la divergence du champ électrique. Le théorème de la divergence permet de passer de l'intégration en volume à l'intégration sur la surface qui entoure complètement le volume V (figure 2.1). La relation intégrale qui en découle est :

$$\iiint_V \text{div}(\vec{E}) dV = \iiint_V \frac{\rho}{\epsilon_0} dV$$

Soit encore, en appliquant le théorème d'Ostrogradski :

$$\oiint_S \vec{E} \cdot d\vec{S} = \frac{Q_{\text{int}}}{\epsilon_0}$$

Où le terme Q_{int} représente les charges comprises à l'intérieur du volume V .

Cette expression traduit le fait que le flux du champ électrique à travers une surface fermée est égal au quotient des charges intérieures par la permittivité du vide. Le théorème de Gauss permet dans bien des cas de calculer le champ électrique.

□ Équation de Maxwell – Faraday

Considérons l'équation [2.1]. Cette équation est intégrée sur la surface S définie sur la figure 2.3. La surface S s'appuie sur la courbe fermée C , orientée. La surface est orientée conformément à la courbe C . C'est-à-dire qu'un vecteur élémentaire de la surface doit voir la courbe tourner dans le sens direct. La surface S n'est pas fermée.

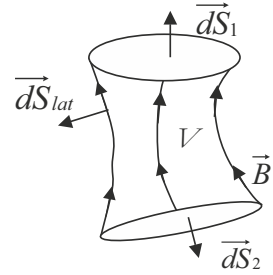


Figure 2.2 – Tube de champ magnétique.

L'intégration sur la surface donne :

$$\iint_S \text{rot}(\vec{E}) \cdot \vec{dS} = \iint_S -\frac{\partial \vec{B}}{\partial t} \cdot \vec{dS}$$

La surface S est maintenue constante dans le temps. Cela permet de transformer le membre de droite. Le flux de l'induction magnétique est défini par :

$$\Phi = \iint_S \vec{B} \cdot \vec{dS}$$

Le membre de gauche est transformé grâce au théorème de Stokes :

$$\oint_C \vec{E} \cdot \vec{dl} = -\frac{d\Phi}{dt}$$

Ce qui signifie que la circulation du champ électrique est égale à l'opposé de la variation du flux magnétique Φ par rapport au temps. Cette loi est à la base de l'interprétation des phénomènes d'induction. Elle constitue une forme de la loi générale de modulation, bien connue en physique, dans la mesure où le système réagit par la création de la force électromotrice e qui s'oppose à la cause qui lui a donné naissance, avec :

$$e = -\frac{d\Phi}{dt}$$

□ Équation de Maxwell – Ampère

L'équation [2.3] va être intégrée selon le même principe. Selon la figure 2.3, on déduit l'expression intégrée suivante :

$$\iint_S \text{rot}(\vec{B}) \cdot \vec{dS} = \mu_0 \iint_S \left(\vec{j} + \epsilon_0 \frac{\partial \vec{E}}{\partial t} \right) \cdot \vec{dS}$$

Soit encore, grâce au théorème de Stokes :

$$\oint_C \vec{B} \cdot \vec{dl} = \mu_0 \iint_S \vec{j} \cdot \vec{dS} + \mu_0 \epsilon_0 \iint_S \frac{\partial \vec{E}}{\partial t} \cdot \vec{dS}$$

On retrouve, dans le cas statique, le théorème d'Ampère, très utilisé en magnétostatique qui exprime la circulation de l'induction magnétique le long d'une courbe fermée comme le produit de la perméabilité du vide par l'intensité du courant électrique créant l'induction.

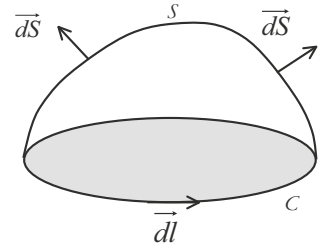


Figure 2.3 – Circulation de l'induction magnétique.

2.1.2 Équations de Maxwell dans la matière

Les équations de Maxwell dans la matière prennent en compte la nature des constituants matériels, formés d'atomes, de molécules ou d'ions. Nous simplifierons ici l'étude de ces équations qui peut être très complexe. Notre propos est simplement de modéliser des composants rayonnants. On se placera dans le cas de modèles de matériaux parfaits. Pour cela on définit un vecteur déplacement électrique $\vec{D}(\vec{r}, t)$ qui tient compte des propriétés électriques du matériau. C'est un vecteur macroscopique, c'est-à-dire que les propriétés électriques microscopiques sont moyennées dans un volume très petit, autour d'un point. Le vecteur déplacement est lié au champ électrique par :

$$\vec{D} = \epsilon \vec{E} \quad [2.6]$$

La permittivité du matériau est ϵ . La permittivité du matériau est liée à sa permittivité relative ϵ_r et à la permittivité du vide par la relation $\epsilon = \epsilon_0 \epsilon_r$.

De la même façon, pour les propriétés magnétiques, on définit le champ magnétique $\vec{H}(\vec{r}, t)$ qui est lié à la perméabilité magnétique et à l'induction magnétique :

$$\vec{B} = \mu \vec{H} \quad [2.7]$$

La perméabilité magnétique du matériau est liée à sa perméabilité relative μ_r et à celle du vide :

$$\mu = \mu_0 \mu_r$$

Moyennant ces définitions, les équations de Maxwell s'écrivent sous la forme :

$$\text{rot}(\vec{E}) = -\frac{\partial \vec{B}}{\partial t} \quad [2.8]$$

$$\text{div}(\vec{B}) = 0 \quad [2.9]$$

$$\text{rot}(\vec{H}) = \vec{j} + \frac{\partial \vec{D}}{\partial t} \quad [2.10]$$

$$\text{div}(\vec{D}) = \rho \quad [2.11]$$

■ Caractéristiques des matériaux

Les relations [2.6] et [2.7] sont appelées **relations constitutives** car elles lient la réponse du matériau à une excitation à laquelle il est soumis soit sous forme d'un champ électrique, soit sous forme d'un champ magnétique.

La permittivité ou la perméabilité des matériaux peuvent dépendre de l'espace et du temps (ou de la fréquence, dans le domaine fréquentiel). Lorsqu'elles ne dépendent pas de l'espace, le matériau est **homogène**.

Lorsque le matériau présente des pertes diélectriques, la permittivité dépend du temps, et la réponse $\vec{D}$ au champ électrique présente un retard par rapport à celui-ci. Dans le domaine de Fourier, cela se traduit par une permittivité complexe. Les mêmes remarques s'appliquent à la perméabilité.

Si le matériau est **isotrope**, ses propriétés sont les mêmes dans toutes les directions et la permittivité (ou la perméabilité) est une grandeur scalaire. Si le matériau est **anisotrope**, la permittivité (ou la perméabilité) est une grandeur matricielle.

Un matériau est **linéaire** si sa permittivité est indépendante du champ électrique. Le vecteur déplacement est proportionnel au champ électrique. Dans le cas contraire, le matériau est dit **non linéaire**. Il en va de même pour les grandeurs magnétiques.

Un matériau est dit **parfait** s'il est linéaire, homogène, isotrope et sans pertes. C'est une approximation qui permet, dans un premier temps, une conception rapide des dispositifs.

Dans la suite, la perméabilité relative des matériaux sera prise égale à l'unité. Peu d'antennes comportent, en effet, des matériaux magnétiques comme des ferrites...

■ Forme harmonique des équations de Maxwell

Nous utiliserons souvent les équations de Maxwell sous leur forme harmonique, en considérant que les grandeurs varient sinusoïdalement. L'utilisation de la notation complexe permet de poser une variation temporelle de chaque grandeur sous la forme du terme : $e^{j\omega t}$. On notera alors les grandeurs sous la forme :

$$\vec{E}(\vec{r}, t) = \vec{E}(\vec{r})e^{j\omega t}$$

Les équations de Maxwell [2.8] et [2.10] contenant des dérivées temporelles se transforment alors en :

$$\text{rot}(\vec{E}) = -j\omega \vec{B} \quad [2.12]$$

$$\text{rot}(\vec{H}) = \vec{j} + j\omega \vec{D} \quad [2.13]$$

Ces deux équations utilisent la convention de signe positif dans l'exponentielle. Avec la convention de signe opposée, les signes provenant de la dérivation temporelle auraient été opposés. Dans la suite de cet ouvrage, nous adopterons la convention des équations [2.12] et [2.13]. Lorsque les champs se trouvent dans le vide, les mêmes notations sont utilisées.

2.1.3 Conditions de passage à la traversée de deux milieux

Les équations de Maxwell entraînent certaines contraintes sur le champ électromagnétique lorsqu'il existe deux milieux différents. La figure 2.4 représente la décomposition des champs dans un milieu, sur laquelle nous allons nous appuyer.

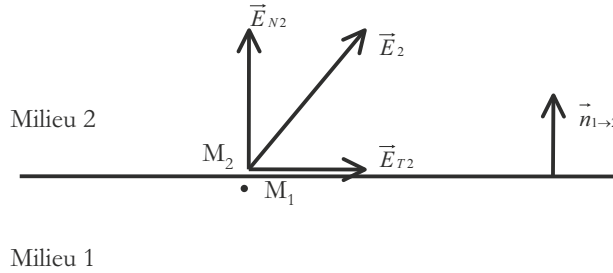


Figure 2.4 – Définition des projections des champs.

Afin d'établir les relations de passage à la traversée de deux milieux, nous allons considérer deux points M_1 et M_2 très proches de l'interface, appartenant respectivement aux milieux 1 et 2.

Le champ en chacun de ces points est projeté selon sa composante tangentielle $\vec{E}_T$ et sa composante normale $\vec{E}_N$.

Chacune des équations de Maxwell impose une condition de passage. Les démonstrations sont faites sur des dimensions infinitésimales. Cela permet de considérer l'interface comme plane et les équations obtenues, comme locales.

■ Condition sur le champ électrique

Considérons l'équation [2.8] et intégrons-la sur la surface définie par la courbe (C) de la figure 2.5

$$\iint \text{rot}(\vec{E}) \cdot \vec{dS} = - \iint \frac{\partial \vec{B}}{\partial t} \cdot \vec{dS}$$

Grâce au théorème de Stokes et en plaçant la dérivée temporelle à l'extérieur de l'intégrale, nous obtenons :

$$\oint_C \vec{E} \cdot \vec{dl} = - \frac{d}{dt} \iint \vec{B} \cdot \vec{dS}$$

En exprimant cette relation sous forme différentielle et tenant compte des ordres des infiniment petits, on obtient :

$$\vec{E}_1 \cdot \vec{dl}_1 + \vec{E}_2 \cdot \vec{dl}_2 = 0$$

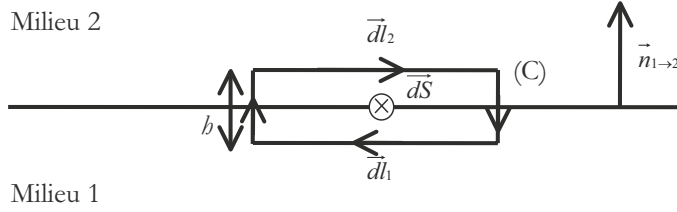


Figure 2.5 – Circulation du champ.

Soit, compte tenu de la géométrie de la figure 2.5 :

$$\vec{E}_{T1} = \vec{E}_{T2} \quad [2.14]$$

Les équations de Maxwell entraînent la continuité de la composante tangentielle du champ électrique à la traversée de deux milieux.

■ Condition sur le champ magnétique

Considérons l'équation [2.10] et intégrons-la sur la surface définie par la figure 2.5, nous obtenons :

$$\iint \vec{\text{rot}}(\vec{H}) \cdot \vec{dS} = \iint \vec{j} \cdot \vec{dS} + \iint \frac{\partial \vec{D}}{\partial t} \cdot \vec{dS}$$

Le théorème de Stokes conduit à :

$$\oint_C \vec{H} \cdot \vec{dl} = \iint \vec{j} \cdot \vec{dS} + \frac{d}{dt} \iint \vec{D} \cdot \vec{dS}$$

Cette équation exprimée de façon différentielle à la surface fait intervenir le courant de surface :

$$\iint \vec{j} \cdot \vec{dS} = \vec{j}_S \cdot (\vec{dl} \wedge \vec{n}_{1 \rightarrow 2})$$

La relation complète s'écrit :

$$\vec{H}_{T2} - \vec{H}_{T1} = \vec{j}_S \wedge \vec{n}_{1 \rightarrow 2} \quad [2.15]$$

Les composantes tangentielles du champ magnétique sont continues à la traversée de deux milieux, s'il n'existe pas de courant électrique surfacique. S'il existe un courant électrique surfacique (cas où l'un des matériaux est un métal parfait), les composantes tangentielles du champ magnétique présentent une discontinuité.

■ Condition sur l'induction magnétique

Intégrons l'équation [2.9] sur le volume défini sur la figure 2.6.

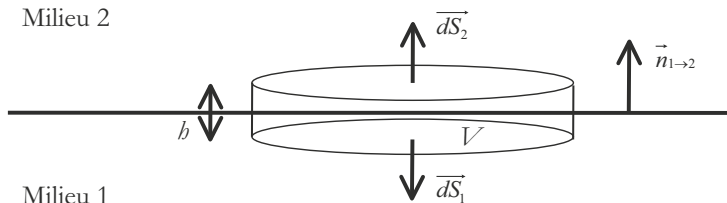


Figure 2.6 – Volume d'intégration à la traversée de deux milieux.

L'intégrale s'écrit sous la forme :

$$\iiint \operatorname{div}(\vec{B}) dv = 0$$

Grâce au théorème de la divergence :

$$\oiint_S \vec{B} \cdot d\vec{S} = 0$$

L'équation s'exprime de façon différentielle :

$$\vec{B}_1 \cdot d\vec{S}_1 + \vec{B}_2 \cdot d\vec{S}_2 = 0$$

Les surfaces étant opposées, on déduit :

$$\vec{B}_{N1} = \vec{B}_{N2} \quad [2.16]$$

Les composantes normales de l'induction magnétique sont continues à la traversée de deux milieux.

■ Condition sur le déplacement électrique

L'équation [2.11] est traitée de la même façon. Son intégration en volume conduit à :

$$\iiint \operatorname{div}(\vec{D}) dv = \iiint \rho dv$$

Le théorème de la divergence entraîne :

$$\oiint_S \vec{D} \cdot d\vec{S} = \iiint \rho dv$$

Cette équation exprimée de façon différentielle fait intervenir la densité surfacique de charges :

$$\iiint \rho dv = \sigma_S d\vec{S} \cdot \vec{n}_{1 \rightarrow 2}$$

On déduit donc :

$$\vec{D}_{N2} - \vec{D}_{N1} = \sigma_S \vec{n}_{1 \rightarrow 2} \quad [2.17]$$

Lorsqu'il n'existe pas de densité de charges à l'interface entre les deux matériaux, la composante normale du vecteur déplacement est continue.

2.1.4 Conditions sur une surface métallique

Plus la fréquence augmente, plus les ondes électromagnétiques pénètrent difficilement dans les métaux qui sont de bons conducteurs, de conductivité finie σ . Elles s'atténuent sur une distance appelée épaisseur de peau δ :

$$\delta = \sqrt{\frac{2}{\omega \mu \sigma}}$$

L'épaisseur de peau diminue lorsque la fréquence augmente. On considère alors que le courant existant à la surface du métal est réduit à une nappe surfacique, comme celui qui existe dans le modèle du conducteur parfait pour lequel on admet que :

$$\sigma \rightarrow \infty$$

Le modèle du conducteur parfait est très utilisé pour la conception d'antenne. Il permet de simplifier les calculs.

Dans un métal parfait, le champ électromagnétique est nul et on vérifie donc les conditions aux limites suivantes. De [2.14], on déduit :

$$\vec{E}_{T1} = \vec{E}_{T2} = 0$$

Soit encore, en notant $\vec{E}$, le champ électrique à la surface du métal :

$$\vec{E} \wedge \vec{n} = 0$$

De [2.15], on déduit le champ magnétique à la surface du métal :

$$\vec{H}_T = \vec{j}_S \wedge \vec{n}_{1 \rightarrow 2}$$

Le champ magnétique normal est nul. En effet d'après [2.16]

$$\vec{B}_{N1} = \vec{B}_{N2} = 0$$

Prenant la perméabilité du métal comme étant égale à celle du vide, on montre que la composante normale du champ magnétique est nulle, soit :

$$\vec{H} \cdot \vec{n} = 0$$

2.1.5 Propagation des ondes dans des milieux infinis

Nous supposons que les sources sont rejetées à l'infini et n'apparaissent donc pas dans les équations. Les hypothèses prises ici supposent que les ondes se propagent dans un milieu infini linéaire, homogène et isotrope. Les équations de Maxwell permettent de déduire l'équation de propagation du champ électromagnétique :

$$\Delta \vec{E} = \epsilon \mu \frac{\partial^2 \vec{E}}{\partial t^2} \quad \text{et} \quad \Delta \vec{H} = \epsilon \mu \frac{\partial^2 \vec{H}}{\partial t^2}$$

Ces équations appartiennent à la famille des équations de propagation qui s'expriment de façon générale par un premier membre couplant les dérivées spatiales aux dérivées temporelles par l'intermédiaire de la célérité c de l'onde. L'équation homogène de propagation se met sous la forme générale :

$$\Delta \vec{E} - \frac{1}{c^2} \frac{\partial^2 \vec{E}}{\partial t^2} = 0$$

D'où la vitesse de propagation de l'onde dans le milieu infini :

$$c = \frac{1}{\sqrt{\epsilon \mu}}$$

Cherchant des solutions sous forme de champs variant dans le temps sous forme sinusoïdale, on impose une variation sous la forme complexe :

$$e^{\pm j\omega t}$$

Les solutions répondent à l'équation d'Helmholtz :

$$\Delta \vec{E}(\vec{r}) + k^2 \vec{E}(\vec{r}) = 0 \quad \text{avec} \quad k = \frac{\omega}{c}$$

La forme élémentaire de la solution de l'équation d'Helmholtz est :

$$\vec{E}_0 e^{\pm j \vec{k} \cdot \vec{r}}$$

$\vec{E}_0$ est un vecteur constant. $\vec{k}$ est le vecteur d'onde, de norme k .

Rappelons que la convention adoptée en $e^{+j\omega t}$ sera celle adoptée dans tout cet ouvrage. La forme générale de la solution élémentaire se met sous la forme :

$$\vec{E}_0 e^{j(\omega t \pm \vec{k} \cdot \vec{r})}$$

Le signe – dans la phase indique une propagation dans le sens du vecteur $\vec{k}$. Le signe opposé correspond à une propagation dans le sens inverse.

Cette solution représente une onde plane car tous les points d'un plan d'onde, perpendiculaire au vecteur de propagation, ont le même état vibratoire.

Il est fondamental de bien connaître les propriétés de l'onde plane, car elles constituent une base, au sens mathématique du terme, du développement de nombreuses ondes électromagnétiques.

Signalons les propriétés des champs associés à ces ondes planes :

- Le champ électromagnétique est contenu dans le plan d'onde. Il n'a donc pas de composante perpendiculaire au plan d'onde.
- Le champ électrique, le champ magnétique et le vecteur de propagation forment un trièdre direct :

$$\vec{H} = \frac{\vec{k} \wedge \vec{E}}{\omega \mu}$$

- Les modules des champs électrique et magnétique sont liés par une relation définissant l'impédance d'onde Z :

$$Z = \sqrt{\frac{\mu}{\epsilon}}$$

2.2 Potentiels électromagnétiques

Au paragraphe précédent, les équations de Maxwell ont été développées dans le cas où les sources étaient à l'infini. L'antenne représentant une source du champ électromagnétique, les équations doivent être développées avec des sources à distance finie. Le milieu de propagation choisi pour cette étude est le vide.

Dans de nombreux domaines de la physique, reposant sur la résolution d'équations différentielles, il est d'usage de passer par des fonctions auxiliaires appelées potentiels. Dans le cas de la résolution des équations de Maxwell, deux potentiels apparaissent :

- le potentiel vecteur, appelé aussi potentiel magnétique
- le potentiel scalaire, appelé aussi potentiel électrique.

La dérivation vectorielle de ces potentiels conduit aux champs électriques et magnétiques.

2.2.1 Potentiel vecteur

L'équation de Maxwell [2.2] imposant à la divergence de l'induction magnétique d'être nulle, conduit à la définition du potentiel vecteur $\vec{A}$:

$$\vec{B} = \text{rot } \vec{A} \quad [2.18]$$

On montre, en effet, que la divergence d'un rotationnel est nulle quel que soit le vecteur sur lequel on applique cette double dérivation vectorielle. Cette définition du potentiel vectoriel permet bien de retrouver l'induction magnétique à partir de la dérivation du potentiel. $\vec{A}$ est défini à une constante vectorielle près relativement aux variations spatiales. $\vec{A}$ dépend de l'espace et du temps.

2.2.2 Potentiel scalaire

En introduisant la définition du potentiel vecteur [2.18] dans l'équation de Maxwell [2.1] relative au rotationnel du champ électrique, on obtient :

$$\text{rot}(\vec{E} + \frac{\partial \vec{A}}{\partial t}) = 0$$

Or le rotationnel du gradient d'un champ scalaire est nul quel que soit ce champ. On introduit donc le potentiel scalaire φ comme répondant à la relation :

$$\vec{E} + \frac{\partial \vec{A}}{\partial t} = -\vec{\text{grad}}\varphi$$

Ce potentiel est défini à une constante scalaire près.

On remarque le signe moins devant le gradient du potentiel scalaire. Ce choix implique une décroissance du potentiel le long des lignes de champ électrique.

Le terme $\frac{\partial \vec{A}}{\partial t}$ est homogène à un champ électrique. Il est appelé champ de Neumann ou champ électromoteur et rend compte des phénomènes d'induction. Il exprime le couplage entre le champ électrique et le champ magnétique. Il est très important dans les phénomènes de propagation et ce d'autant plus que la fréquence est élevée.

Le champ électrique se déduit donc des deux potentiels par :

$$\vec{E} = -\frac{\partial \vec{A}}{\partial t} - \vec{\text{grad}}\varphi \quad [2.19]$$

2.2.3 Équation de propagation du potentiel vecteur

De même que les champs électriques et magnétiques se propagent dans le vide à la vitesse de la lumière, nous allons montrer que les potentiels vecteur et scalaire se propagent tous deux aussi à cette vitesse.

Dans toute la suite, on utilisera la propriété mathématique sur les doubles dérivations partielles qui permet d'inverser l'ordre des dérivations par rapport à des variables indépendantes. Ce sera le cas des variables d'espace et de temps.

Exprimons dans [2.3] les champs en fonction des potentiels par [2.18] et [2.19] :

$$\vec{\text{rot}}(\text{rot } \vec{A}) = \mu_0 \vec{j} + \epsilon_0 \mu_0 \frac{\partial}{\partial t}(-\frac{\partial \vec{A}}{\partial t} - \vec{\text{grad}}\varphi)$$

On rappelle la propriété mathématique suivante :

$$\Delta \vec{A} = \vec{\text{grad}}(\text{div } \vec{A}) - \vec{\text{rot}}(\text{rot } \vec{A})$$

L'équation précédente s'écrit alors, en remplaçant le double rotationnel :

$$\vec{\text{grad}}(\text{div } \vec{A}) - \Delta \vec{A} = \mu_0 \vec{j} - \epsilon_0 \mu_0 \frac{\partial^2 \vec{A}}{\partial t^2} - \epsilon_0 \mu_0 \vec{\text{grad}}(\frac{\partial \varphi}{\partial t})$$

Soit encore :

$$\Delta \vec{A} - \epsilon_0 \mu_0 \frac{\partial^2 \vec{A}}{\partial t^2} = -\mu_0 \vec{j} + \epsilon_0 \mu_0 \vec{\text{grad}}(\frac{\partial \varphi}{\partial t}) + \vec{\text{grad}}(\text{div } \vec{A})$$

On reconnaît une équation de propagation dans le membre de gauche. Le membre de droite est complexe et contient, entre autres, la source qui impose la solution. Le membre de droite se simplifie grâce à l'équation de jauge de Lorentz.

■ Jauge de Lorentz

La jauge de Lorentz contraint les potentiels selon la relation suivante :

$$\epsilon_0 \mu_0 \frac{\partial \varphi}{\partial t} + \text{div } \vec{A} = 0 \quad [2.20]$$

Elle traduit une loi de conservation de nature relativiste.

Cette équation permet de simplifier le membre de droite de l'équation de propagation précédente sous la forme :

$$\Delta \vec{A} - \epsilon_0 \mu_0 \frac{\partial^2 \vec{A}}{\partial t^2} = -\mu_0 \vec{j} \quad [2.21]$$

2.2.4 Équation de propagation du potentiel scalaire

L'équation de Maxwell [2.4] s'exprime en fonction des potentiels, d'après [2.19] :

$$\text{div} \left(\frac{\partial \vec{A}}{\partial t} + \text{grad } \varphi \right) = -\frac{\rho}{\epsilon_0}$$

L'équation de jauge dérivée par rapport au temps se met sous la forme :

$$\text{div} \left(\frac{\partial \vec{A}}{\partial t} \right) = -\epsilon_0 \mu_0 \frac{\partial^2 \varphi}{\partial t^2}$$

En utilisant la propriété du laplacien, valable pour tout champ scalaire :

$$\Delta \varphi = \text{div} (\vec{\text{grad}} \varphi)$$

On obtient l'équation de propagation avec source pour le potentiel scalaire :

$$\Delta \varphi - \epsilon_0 \mu_0 \frac{\partial^2 \varphi}{\partial t^2} = -\frac{\rho}{\epsilon_0} \quad [2.22]$$

Les potentiels répondent à des équations de propagation analogues du type :

$$\Delta f - \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = S$$

où c représente la célérité de l'onde et S le terme source. Tous les termes de ces équations dépendent de l'espace et du temps. La vitesse de propagation est donc telle que :

$$\epsilon_0 \mu_0 c^2 = 1, \text{ soit } c = 3 \times 10^8 \text{ ms}^{-1}$$

Nous allons voir quelles sont les solutions de ces équations en tenant compte des sources.

2.2.5 Potentiels retardés

Les solutions des équations [2.21] et [2.22] se mettent respectivement sous la forme :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} \iiint_{\text{sources}} \frac{\vec{j}(\vec{r}', t')}{||\vec{R}||} d\vec{r}' \quad [2.23]$$

$$\varphi(\vec{r}, t) = \frac{1}{4\pi\epsilon_0} \iiint_{\text{sources}} \frac{\rho(\vec{r}', t')}{||\vec{R}||} d\vec{r}' \quad [2.24]$$

On constate que ces équations ont la même forme pour le potentiel scalaire et pour le potentiel vecteur. R représente la distance du point source élémentaire au point de calcul du potentiel

(figure 2.7), r' représente la distance du point source à l'origine, t' est donné par :

$$t' = t - \frac{R}{c}$$

$\frac{R}{c}$ est le temps que met le signal pour aller du point source au point d'observation M. L'état de la source est ainsi considéré à un instant t' , antérieur au temps d'observation t , prenant en compte le temps de propagation entre la source et le point d'observation, qui ne peut pas être infiniment petit car la vitesse de la lumière n'est pas infinie.

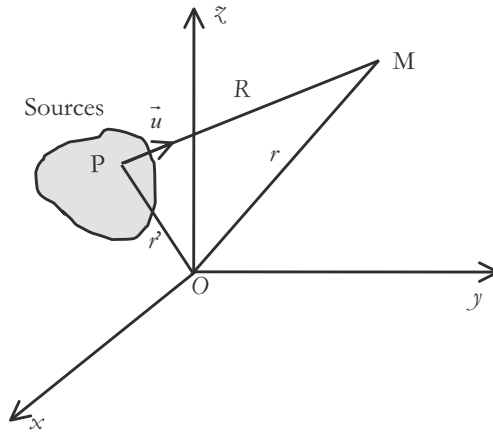


Figure 2.7 – Représentation géométrique des sources.

Le nom de potentiels retardés est donné à cette forme de solution en raison du retard temporel à considérer, dû à la propagation du signal.

Dans la suite, nous allons démontrer la formule des potentiels retardés.

Nous nous appuyerons sur l'identité de Green qui lie deux fonctions scalaires Φ et Ψ , définies sur V et dérivables deux fois. Cette identité est une forme intégrale à trois dimensions.

■ Identité de Green

Considérons un volume V fini, entouré par une surface fermée S , orientée vers l'extérieur (figure I.1)

Le développement de la divergence du produit de Φ par le gradient de Ψ s'écrit :

$$\text{div}(\Phi \overrightarrow{\text{grad}} \Psi) = \Phi \Delta \Psi + \overrightarrow{\text{grad}} \Psi \cdot \overrightarrow{\text{grad}} \Phi$$

La forme symétrique s'écrit en échangeant le rôle des fonctions et par différence :

$$\text{div}(\Phi \overrightarrow{\text{grad}} \Psi) - \text{div}(\Psi \overrightarrow{\text{grad}} \Phi) = \Phi \Delta \Psi - \Psi \Delta \Phi$$

Cette relation, valable en chaque point, est intégrée sur le volume V . On utilise le théorème de la divergence qui permet d'utiliser l'équation suivante :

$$\iiint_V \text{div}(\Phi \overrightarrow{\text{grad}} \Psi - \Psi \overrightarrow{\text{grad}} \Phi) dv = \oint_S (\Phi \overrightarrow{\text{grad}} \Psi - \Psi \overrightarrow{\text{grad}} \Phi) \cdot d\vec{s}$$

En remarquant que :

$$\overrightarrow{\text{grad}}\Phi \cdot \overrightarrow{n} = \frac{\partial \Phi}{\partial \overrightarrow{n}}$$

$\overrightarrow{n}$ étant la normale à la surface, on obtient alors l'identité de Green :

$$\iiint_V (\Phi \Delta \Psi - \Psi \Delta \Phi) dv = \oint_S (\Phi \frac{\partial \Psi}{\partial \overrightarrow{n}} - \Psi \frac{\partial \Phi}{\partial \overrightarrow{n}}) ds \quad [2.25]$$

Cette identité permet de remplacer un calcul en volume par la connaissance de conditions surfaciques sur la fonction et sur sa dérivée normale. Ce qui revient à dire que les phénomènes sur une surface fermée et leur variation normalement à la surface, traduisent les variations à l'intérieur du volume.

■ Calcul des potentiels retardés

Pour ce calcul, nous allons rappeler des propriétés de la transformée de Fourier. $g(\overrightarrow{r}, t)$ et sa transformée de Fourier $g_\omega(\overrightarrow{r})$ sont liées par :

$$g(\overrightarrow{r}, t) = \int_{-\infty}^{\infty} g_\omega(\overrightarrow{r}) e^{j\omega t} d\omega \quad \text{et} \quad g_\omega(\overrightarrow{r}) = \frac{1}{2\pi} \int_{-\infty}^{\infty} g(\overrightarrow{r}, t) e^{-j\omega t} dt$$

De même, la transformée de Fourier de $\varphi(\overrightarrow{r}, t)$ et son inverse sont liées par :

$$\varphi(\overrightarrow{r}, t) = \int_{-\infty}^{\infty} \varphi_\omega(\overrightarrow{r}) e^{j\omega t} d\omega \quad [2.26]$$

$$\text{et} \quad \varphi_\omega(\overrightarrow{r}) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \varphi(\overrightarrow{r}, t) e^{-j\omega t} dt \quad [2.27]$$

Nous allons résoudre [2.22], soit en explicitant les variables espace et temps :

$$\Delta \varphi(\overrightarrow{r}, t) - \frac{1}{c^2} \frac{\partial^2 \varphi(\overrightarrow{r}, t)}{\partial t^2} = -\frac{\rho(\overrightarrow{r}, t)}{\epsilon_0} = -g(\overrightarrow{r}, t)$$

Afin de s'affranchir de la variable temps, on effectue une transformée de Fourier de cette équation :

$$\Delta \varphi_\omega(\overrightarrow{r}) + \frac{\omega^2}{c^2} \varphi_\omega(\overrightarrow{r}) = -g_\omega(\overrightarrow{r}) \quad [2.28]$$

Dans le domaine de Fourier, la variable temps n'apparaît plus. Cependant, les fonctions sont paramétrées par la pulsation ω de l'onde et la résolution doit s'effectuer pour chaque valeur de cette pulsation pour reconstituer le signal temporel. Il reste donc à résoudre cette équation spatialement.

Pour cela, on associe une équation différentielle ayant la même forme que l'équation à résoudre mais avec un second membre constitué d'une distribution de Dirac :

$$\Delta G(\overrightarrow{r}, \overrightarrow{r}') + \frac{\omega^2}{c^2} G(\overrightarrow{r}, \overrightarrow{r}') = -\delta(\overrightarrow{r} - \overrightarrow{r}') \quad [2.29]$$

Les solutions de l'équation [2.29] sont de la forme :

$$G(\overrightarrow{r}, \overrightarrow{r}') = \frac{e^{\pm jk|\overrightarrow{r} - \overrightarrow{r}'|}}{4\pi|\overrightarrow{r} - \overrightarrow{r}'|}, \quad [2.30]$$

avec $k = \frac{\omega}{c}$

Cette fonction s'appelle la fonction de **Green du vide**. Elle représente la réponse en un point placé en $\vec{r}$, à une impulsion imposée en un point source, placé en $\vec{r}'$ pour un système dans le vide, comme c'est le cas d'une source simple de rayonnement.

Connaissant la fonction de Green, il suffit de considérer chaque point source comme un Dirac, affecté d'un poids correspondant à l'amplitude du point source considéré. Nous allons expliciter cette opération dans la suite.

Calculons l'intégrale I définie par :

$$I = \iiint_V [\varphi_\omega(\vec{r}') \cdot \Delta G(\vec{r}, \vec{r}') - G(\vec{r}, \vec{r}') \cdot \Delta \varphi_\omega(\vec{r}')] dv'$$

L'intégrale porte sur le volume des sources. En utilisant les équations [2.28] et [2.29], les laplaciens peuvent être remplacés. On obtient alors :

$$I = \iiint_V [-\varphi_\omega(\vec{r}') \cdot \delta(\vec{r} - \vec{r}') + G(\vec{r}, \vec{r}') \cdot g_\omega(\vec{r}')] dv'$$

Or la distribution de Dirac a la propriété suivante, quelle que soit la fonction f :

$$\iiint_{V_\infty} f(\vec{r}') \delta(\vec{r} - \vec{r}') dv' = f(\vec{r})$$

Il suffit d'étendre la définition de l'intégrale I à l'infini. Ce qui est possible car on ne considère ici que le cas de sources à distance finie. On déduit :

$$I_\infty = -\varphi_\omega(\vec{r}) + \iiint_{V_\infty} G(\vec{r}, \vec{r}') \cdot g_\omega(\vec{r}') dv'$$

Donc :

$$\varphi_\omega(\vec{r}) = -I_\infty + \iiint_{V_\infty} G(\vec{r}, \vec{r}') \cdot g_\omega(\vec{r}') dv'$$

Nous allons montrer *a posteriori* que I_∞ est nulle. Alors on déduit, en remplaçant la fonction de Green par sa valeur :

$$\varphi_\omega(\vec{r}) = \iiint_{V_\infty} \frac{e^{\pm jk|\vec{r} - \vec{r}'|}}{4\pi|\vec{r} - \vec{r}'|} \cdot g_\omega(\vec{r}') dv' \quad [2.31]$$

L'intégrale peut alors uniquement porter sur le volume des sources puisque la fonction $g_\omega(\vec{r})$ est nulle en dehors.

Montrons que I_∞ s'annule à l'infini. Pour cela utilisons l'identité de Green :

$$I_\infty = \iint_{S_\infty} (\varphi_\omega \frac{\partial G(\vec{r}, \vec{r}')}{\partial \vec{n}} - G(\vec{r}, \vec{r}') \frac{\partial \varphi_\omega}{\partial \vec{n}}) ds$$

Cherchons un majorant de cette expression :

$$|I_\infty| \leq \iint_{S_\infty} \left(\left| \varphi_\omega \frac{\partial G(\vec{r}, \vec{r}')}{\partial \vec{n}} \right| + \left| G(\vec{r}, \vec{r}') \frac{\partial \varphi_\omega}{\partial \vec{n}} \right| \right) ds$$

Les fonctions φ et G varient en $\frac{1}{r}$. Donc leur dérivée varie en $\frac{1}{r^2}$. Chaque terme de la parenthèse a un équivalent en $\frac{1}{r^3}$ à l'infini. La surface variant en r^2 , l'ensemble est équivalent à un terme en $\frac{1}{r}$ sur la surface de l'infini. L'intégrale étant toujours inférieure à un terme tendant vers zéro à l'infini, tend elle-même vers zéro.

Revenons à l'expression [2.26] de la transformée de Fourier de la fonction φ . Appliquons la transformée inverse :

$$\varphi(\vec{r}, t) = \int_{-\infty}^{\infty} \iiint_{V_{\infty}} \frac{e^{\pm jk|\vec{r} - \vec{r}'|}}{4\pi|\vec{r} - \vec{r}'|} \cdot g_{\omega}(\vec{r}') dv' \cdot e^{j\omega t} d\omega$$

Soit encore :

$$\varphi(\vec{r}, t) = \int_{-\infty}^{\infty} \iiint_{V_{\infty}} \frac{e^{\pm jk|\vec{r} - \vec{r}'|} \cdot e^{j\omega t}}{4\pi|\vec{r} - \vec{r}'|} \cdot g_{\omega}(\vec{r}') dv' \cdot d\omega$$

En inversant les intégrales temporelle et spatiale, ce qui est possible car les variables sont indépendantes de même que les bornes, on obtient :

$$\varphi(\vec{r}, t) = \iiint_{V_{\infty}} \int_{-\infty}^{\infty} \frac{e^{\pm jk|\vec{r} - \vec{r}'|} \cdot e^{j\omega t}}{4\pi|\vec{r} - \vec{r}'|} \cdot g_{\omega}(\vec{r}') d\omega dv'$$

Introduisons le temps t' défini par :

$$t' = t \mp \frac{1}{c} |\vec{r} - \vec{r}'|$$

On obtient :

$$\varphi(\vec{r}, t) = \iiint_{V_{\infty}} \int_{-\infty}^{\infty} \frac{1}{4\pi|\vec{r} - \vec{r}'|} \cdot g_{\omega}(\vec{r}') \cdot e^{j\omega t'} d\omega dv'$$

Soit en utilisant la propriété de la transformée de Fourier :

$$\varphi(\vec{r}, t) = \iiint_{V_{\infty}} \frac{g(\vec{r}', t')}{4\pi|\vec{r} - \vec{r}'|} dv' \quad [2.32]$$

Cette démonstration, tout à fait générale, fait apparaître un signe $\mp$ dans la définition du temps t' . Lorsqu'on choisit le signe moins, on a bien affaire à un potentiel retardé. Avec la convention de signe sur l'exponentielle temporelle prise au départ, un signe $-$ devant k correspond à une onde se propageant dans le sens du vecteur de propagation. Elle est appelée onde sortante. C'est le cas des antennes considérées à l'émission.

Pour la convention de signe inverse sur le temps, on parle de potentiel avancé. L'impulsion passant en M à l'instant t sera reçue par l'antenne de réception au temps $t' = t + \frac{1}{c} |\vec{r} - \vec{r}'|$. Cette convention s'applique aux antennes en réception. C'est la convention d'onde entrante.

La démonstration précédente porte sur des grandeurs scalaires. Pour l'étendre au potentiel vecteur, il suffit de considérer chacune des composantes du potentiel vecteur. Le résultat conduit donc à l'équation [2.23].

2.3 Vecteur de Poynting

Dans ce paragraphe, sont rappelées les notions relatives à la puissance rayonnée, qui sont fondamentales pour l'étude des antennes. Le vecteur caractéristique de la puissance rayonnée est le vecteur de Poynting, défini, soit dans le domaine temporel, soit dans le domaine fréquentiel.

2.3.1 Vecteur de Poynting dans le domaine temporel

Les grandeurs électromagnétiques sont variables avec l'espace et le temps. Dans cette partie, les notations ne font apparaître que la variation temporelle afin de ne pas alourdir les notations. Le

champ électromagnétique en temporel est noté :

$$\vec{e}(t) \vec{h}(t)$$

Le but de ce paragraphe est de montrer que la puissance rayonnée $p(t)$ par le champ électromagnétique à travers une surface S est égale au flux du vecteur de Poynting :

$$p(t) = \iint_S \left(\vec{e}(t) \wedge \vec{h}(t) \right) \cdot \vec{ds} \quad [2.33]$$

Le vecteur de Poynting est défini par :

$$\vec{\Pi}(t) = \vec{e}(t) \wedge \vec{h}(t) \quad [2.34]$$

Ses variations temporelles sont liées à celles du champ électromagnétique.

La démonstration permettant d'interpréter les propriétés du vecteur de Poynting est détaillée ci-après.

Calculons la divergence du produit vectoriel du champ électrique par le champ magnétique. Le développement de la dérivation du produit vectoriel conduit à :

$$\text{div} \left(\vec{e}(t) \wedge \vec{h}(t) \right) = \vec{h}(t) \cdot (\vec{\text{rot}} \vec{e}(t)) - \vec{e}(t) \cdot (\vec{\text{rot}} \vec{h}(t))$$

En utilisant les équations de Maxwell dans le vide, on obtient :

$$\text{div} \left(\vec{e}(t) \wedge \vec{h}(t) \right) = \vec{h}(t) \cdot \left(-\mu_0 \frac{\partial \vec{h}(t)}{\partial t} \right) - \vec{e}(t) \cdot \left(\vec{j}(t) + \epsilon_0 \frac{\partial \vec{e}(t)}{\partial t} \right) \quad [2.35]$$

La densité d'énergie électrique est définie par :

$$u_e(t) = \frac{1}{2} \epsilon_0 (\vec{e}(t))^2 \quad [2.36]$$

De même pour la densité d'énergie magnétique :

$$u_m(t) = \frac{1}{2} \mu_0 (\vec{h}(t))^2 \quad [2.37]$$

Dans l'expression [2.35], ce sont donc les dérivées par rapport au temps de la densité d'énergie qui interviennent, selon :

$$\text{div} \left(\vec{e}(t) \wedge \vec{h}(t) \right) = -\frac{\partial u_m}{\partial t} - \frac{\partial u_e}{\partial t} - \vec{e}(t) \cdot \vec{j}(t)$$

En posant la densité d'énergie électromagnétique :

$$u(t) = u_e(t) + u_m(t) \quad [2.38]$$

On obtient l'expression :

$$-\frac{\partial u}{\partial t} = \text{div} \left(\vec{e}(t) \wedge \vec{h}(t) \right) + \vec{e}(t) \cdot \vec{j}(t)$$

Cette équation locale est intégrée dans un volume V , constant dans le temps et limité par la surface S selon la figure I.1.

On obtient après intégration :

$$-\iiint_V \frac{\partial u}{\partial t} dv = \iiint_V \text{div} \left(\vec{e}(t) \wedge \vec{h}(t) \right) dv + \iiint_V \vec{e}(t) \cdot \vec{j}(t) dv$$

$U(t)$ représentant l'énergie contenue dans le volume V à l'instant t , l'équation intégrale suivante, obtenue après intégration selon le théorème d'Ostrogradski, permet de donner une interprétation physique au vecteur de Poynting :

$$- \iint_S (\vec{e}(t) \wedge \vec{h}(t)) \cdot \vec{ds} = \frac{dU(t)}{dt} + \iiint_V \vec{e}(t) \cdot \vec{j}(t) dv$$

Le terme du membre de gauche représente le flux entrant du vecteur de Poynting à travers S , en raison du signe négatif. Le premier terme du membre de droite est la variation d'énergie du volume V au cours du temps, donc la puissance électromagnétique disponible dans V . Le second terme représente la puissance fournie aux courants se trouvant dans V . Cette dernière peut être liée à l'effet Joule ou bien à l'accélération de particules chargées en mouvement. L'équation traduit l'équilibre énergétique du volume V : la puissance rayonnée entrant dans le volume V est égale à la somme de l'accroissement d'énergie de V par unité de temps et de la puissance fournie aux particules qui s'y trouvent. Le flux du vecteur de Poynting est ainsi interprété comme la puissance rayonnée à travers une surface selon [2.33].

Cette équation écrite dans un autre ordre conduit à une interprétation utile dans le cadre de l'étude des antennes :

$$-\frac{dU(t)}{dt} = \iint_S (\vec{e}(t) \wedge \vec{h}(t)) \cdot \vec{ds} + \iiint_V \vec{e}(t) \cdot \vec{j}(t) dv$$

En effet, la perte de puissance électromagnétique dans V est égale à la puissance rayonnée sortant du volume V , à travers S , ajoutée à la puissance cédée aux particules qui s'y trouvent.

Le vecteur de Poynting répond à l'équation locale :

$$\text{div } \vec{\Pi}(t) + \frac{\partial u}{\partial t} = -\vec{e}(t) \cdot \vec{j}(t) \quad [2.39]$$

qui est l'équation locale de conservation de l'énergie, d'après l'interprétation qui vient d'être faite.

2.3.2 Vecteur de Poynting dans le domaine fréquentiel

Dans le domaine fréquentiel, les champs sont considérés à une fréquence donnée. Ils varient sinusoïdalement. On utilise la représentation complexe :

$$\vec{e}(t) = \text{Rél}(\vec{E} e^{j\omega t}) \quad \text{et} \quad \vec{h}(t) = \text{Rél}(\vec{H} e^{j\omega t})$$

$\vec{E}$ et $\vec{H}$ sont les amplitudes complexes portant la phase et variant dans l'espace.

La puissance instantanée définie précédemment oscille au cours du temps. On s'intéresse donc plutôt à sa valeur moyenne sur une période. Ainsi, la puissance rayonnée à travers une surface S , moyennée dans le temps est donnée par :

$$\bar{P} = \frac{1}{2} \text{Rél} \iint_S (\vec{E}(t) \wedge \vec{H}^*(t)) \cdot \vec{ds} \quad [2.40]$$

On définit alors le vecteur de Poynting complexe :

$$\vec{\Pi} = \frac{1}{2} \vec{E}(t) \wedge \vec{H}^*(t) \quad [2.41]$$

La partie réelle de son flux à travers une surface unité est la densité surfacique de puissance S_r :

$$S_r = \frac{1}{2} \text{Rél} (\vec{E}(t) \wedge \vec{H}^*(t)) \cdot \vec{n} \quad [2.42]$$

Cette grandeur est la base de l'étude du rayonnement des antennes à grande distance.

2.4 Théorèmes importants de l'électromagnétisme

Grâce aux outils de calcul vectoriel, il est possible de démontrer, à partir des équations de Maxwell, différents théorèmes qui apparaissent comme indispensables lors de l'étude des antennes. Le but ici n'est pas de donner tous les théorèmes d'électromagnétisme, qui s'énoncent d'ailleurs chacun sous différentes formes, mais d'en démontrer quelques-uns des plus utilisés qui permettent la compréhension des phénomènes liés au fonctionnement des antennes.

2.4.1 Généralisation des équations de Maxwell

Les équations de Maxwell sont des équations différentielles liant le champ électromagnétique aux sources, apparaissant sous la forme de densités de courant et de densités de charges et vérifiant la relation de conservation de la charge [2.5].

Dans ces équations, il n'apparaît que des sources électriques qui sont les seules ayant une existence réelle. Ces sources sont constituées de particules chargées électriquement. On ne connaît pas de source matérielle magnétique équivalente.

Cependant, un artifice mathématique consiste à introduire des sources magnétiques dans les équations de Maxwell qui deviennent :

$$\vec{\text{rot}}(\vec{E}) = -\frac{\partial \vec{B}}{\partial t} - \vec{M} \quad [2.43]$$

$$\text{div}(\vec{B}) = -\tau \quad [2.44]$$

$$\vec{\text{rot}}(\vec{H}) = \vec{j} + \frac{\partial \vec{D}}{\partial t} \quad [2.45]$$

$$\text{div}(\vec{D}) = \rho \quad [2.46]$$

Le terme $\vec{M}$ représente dans ce cas les courants magnétiques, et τ les charges magnétiques. Ce sont des sources au sens mathématique du terme. Dans la réalité, ces sources n'existent pas matériellement. Cependant, dans certaines circonstances, des portions de l'espace dans lesquelles il existe un champ électrique, peuvent être assimilées à une zone de courants magnétiques. Nous détaillerons ce point plus loin.

Nous voyons que l'introduction de ces termes dans les équations intrinsèques (qui ne le sont plus du fait de l'introduction des sources magnétiques) rend symétriques deux à deux les équations.

2.4.2 Théorème d'unicité

Ce théorème permet d'expliciter les conditions sur les champs qui permettent de s'assurer qu'une solution en champ électromagnétique est unique.

Considérons des sources électriques et magnétiques $(\vec{J}, \vec{M})$ qui créent dans un volume V un champ électromagnétique $(\vec{E}_1, \vec{H}_1)$. Si la solution n'est pas unique, ces mêmes sources devraient créer un autre champ électromagnétique : $(\vec{E}_2, \vec{H}_2)$. Nous ferons la démonstration dans l'espace fréquentiel, sachant qu'on repasse dans le domaine temporel par une transformée de Fourier.

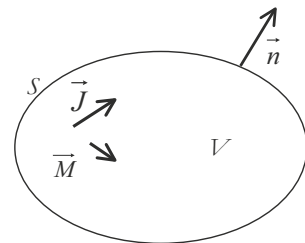


Figure 2.8 – Sources du champ électromagnétique dans le volume V .

Les milieux sont considérés comme linéaires. Ces deux solutions répondent alors aux équations de Maxwell :

$$\vec{\text{rot}} \vec{E}_1 = -j\omega\mu\vec{H}_1 - \vec{M} \quad [2.47]$$

$$\vec{\text{rot}} \vec{H}_1 = j\omega\epsilon\vec{E}_1 + \vec{J} \quad [2.48]$$

$$\vec{\text{rot}} \vec{E}_2 = -j\omega\mu\vec{H}_2 - \vec{M} \quad [2.49]$$

$$\vec{\text{rot}} \vec{H}_2 = j\omega\epsilon\vec{E}_2 + \vec{J} \quad [2.50]$$

La différence entre les équations [2.49] et [2.47] donne :

$$\vec{\text{rot}} (\vec{E}_2 - \vec{E}_1) = -j\omega\mu (\vec{H}_2 - \vec{H}_1)$$

De même, on obtient :

$$\vec{\text{rot}} (\vec{H}_2 - \vec{H}_1) = j\omega\epsilon (\vec{E}_2 - \vec{E}_1)$$

Posons :

$$\delta\vec{E} = \vec{E}_2 - \vec{E}_1 \quad \text{et} \quad \delta\vec{H} = \vec{H}_2 - \vec{H}_1$$

Le théorème de la divergence sur le volume V , limité par la surface fermée S , nous permet d'écrire :

$$\oint_S (\delta\vec{E} \wedge \delta\vec{H}^*) \vec{dS} = \iiint_V \text{div} (\delta\vec{E} \wedge \delta\vec{H}^*) dv$$

En développant :

$$\oint_S (\delta\vec{E} \wedge \delta\vec{H}^*) \vec{dS} = \iiint_V (\delta\vec{H}^* \vec{\text{rot}} \delta\vec{E} - \delta\vec{E} \vec{\text{rot}} \delta\vec{H}^*) dv$$

En utilisant les propriétés établies pour les différences des champs :

$$\oint_S (\delta\vec{E} \wedge \delta\vec{H}^*) \vec{dS} = -j\omega \iiint_V \left(\mu \|\delta\vec{H}\|^2 - \epsilon^* \|\delta\vec{E}\|^2 \right) dv$$

Rappelons les définitions de la permittivité et de la perméabilité effectives complexes :

$$\epsilon = \epsilon' + j\epsilon''$$

$$\mu = \mu' + j\mu''$$

Donc la partie réelle du premier membre de cette intégrale est calculable :

$$\text{Réel} \oint_S (\delta\vec{E} \wedge \delta\vec{H}^*) \vec{dS} = \omega \iiint_V \left(\mu'' \|\delta\vec{H}\|^2 + \epsilon'' \|\delta\vec{E}\|^2 \right) dv$$

Pour des matériaux dissipatifs, on montre que :

$$\epsilon'' > 0 \quad \text{et} \quad \mu'' > 0$$

Donc la partie réelle de l'intégrale est forcément positive ou nulle. Elle est nulle si et seulement si les deux solutions en champ électromagnétique sont confondues. Cette condition assure l'unicité de la solution.

Donc, si on impose :

$$\iint_S (\delta \vec{E} \wedge \delta \vec{H}^*) \vec{dS} = 0$$

les champs $\vec{E}_1$ et $\vec{E}_2$ doivent être identiques, de même que les champs magnétiques associés. En faisant apparaître la normale à la surface S , cette expression s'écrit de plusieurs façons :

$$\oint_S (\delta \vec{E} \wedge \delta \vec{H}^*) \vec{dS} = \oint_S (\vec{n} \wedge \delta \vec{E}) \delta \vec{H}^* dS = \oint_S (\delta \vec{H}^* \wedge \vec{n}) \delta \vec{E} dS = 0$$

Ce sont les conditions pour assurer l'unicité de la solution.

Plusieurs possibilités assurent donc l'unicité de la solution dans un volume V :

- Imposer le champ électrique tangentiel sur la surface S , indépendamment de sa composante normale et du champ magnétique. Alors : $\vec{n} \wedge \vec{E} \text{ imposé sur } S \Rightarrow \vec{n} \wedge \delta \vec{E} = 0 \text{ sur } S$
- Imposer le champ magnétique tangentiel sur la surface S , indépendamment de sa composante normale et du champ électrique. Alors : $\vec{n} \wedge \vec{H} \text{ imposé sur } S \Rightarrow \vec{n} \wedge \delta \vec{H} = 0 \text{ sur } S$
- Imposer soit le champ électrique tangentiel sur une portion de la surface S et le champ magnétique tangentiel sur la portion complémentaire.

Cette démonstration repose sur les propriétés des matériaux à pertes. On pourra considérer les matériaux sans pertes comme des cas limites de ceux-ci. Remarquons que naturellement tous les matériaux présentent des pertes, même minimales.

2.4.3 Théorème de réciprocité de Lorentz

Ce théorème considère les réponses d'un système linéaire à deux excitations différentes et démontre une réciprocité entre les réponses.

Considérons un milieu linéaire et isotrope, de permittivité ϵ et de perméabilité μ .

On considère deux états indépendants :

- L'un, défini par une distribution volumique de courant $\vec{J}_1$ créant dans le volume (V) un champ électromagnétique : $(\vec{E}_1, \vec{H}_1)$.
- L'autre, défini par une autre distribution de courant $\vec{J}_2$ créant dans le volume (V) un champ électromagnétique : $(\vec{E}_2, \vec{H}_2)$.

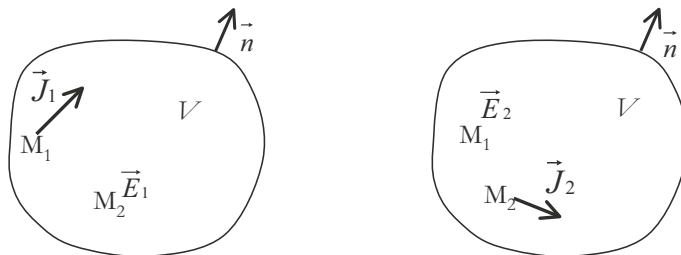


Figure 2.9 – Les deux états d'excitation du système contenu dans le volume V .

Les équations de Maxwell s'écrivent en régime harmonique, dans le cas de matériaux linéaires, pour le premier état :

$$\vec{\text{rot}} \vec{E}_1 = -j\omega\mu\vec{H}_1 \quad \text{et} \quad \vec{\text{rot}} \vec{H}_1 = j\omega\epsilon\vec{E}_1 + \vec{J}_1$$

Pour le second état :

$$\vec{\text{rot}} \vec{E}_2 = -j\omega\mu\vec{H}_2 \quad \text{et} \quad \vec{\text{rot}} \vec{H}_2 = j\omega\epsilon\vec{E}_2 + \vec{J}_2$$

Considérons la combinaison vectorielle des champs électriques et magnétiques des deux états selon l'expression :

$$(\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1)$$

Sa dérivation conduit à :

$$\text{div} (\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1) = \vec{H}_2 \vec{\text{rot}} \vec{E}_1 - \vec{E}_1 \vec{\text{rot}} \vec{H}_2 - \vec{H}_1 \vec{\text{rot}} \vec{E}_2 + \vec{E}_2 \vec{\text{rot}} \vec{H}_1$$

Soit en remplaçant les rotationnels grâce aux équations de Maxwell :

$$\text{div} (\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1) = -\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1$$

L'intégration de cette équation locale conduit à :

$$\iiint_V \text{div} (\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1) dv = \iiint_V (-\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1) dv$$

Soit en utilisant le théorème de la divergence :

$$\oint_S (\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1) \cdot \vec{ds} = \iiint_V (-\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1) dv \quad [2.51]$$

Cette expression constitue le théorème de réciprocité de Lorentz.

Ce théorème prend une forme particulière lorsque le premier membre s'annule. Ce qui est vrai dans les cas suivants qui couvrent de nombreuses situations :

1. Le volume (V) représente tout l'espace. Alors la surface (S) d'intégration est repoussée à l'infini. Si les sources restent à distance finie, les champs créés à l'infini sont nuls et l'intégrale de surface est nulle.
2. La surface (S) entoure intégralement toutes les sources. Dans ce cas, puisque les sources sont contenues dans le volume V , on peut étendre le volume d'intégration à l'infini, le terme source étant nul dans l'extension du volume :

$$\iiint_V (-\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1) dv = \iiint_{V_\infty} (-\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1) dv$$

L'intégration sur la surface de l'infini est donc nulle pour les mêmes raisons que précédemment.

3. (S) est parfaitement conductrice. Dans ce cas, le champ électrique est perpendiculaire à la surface. Donc le produit vectoriel du champ électrique par le champ magnétique est perpendiculaire au vecteur surface et l'intégrale de surface est identiquement nulle.
4. La surface (S) présente une impédance de surface définie par :

$$\vec{E}_T = -Z \vec{n} \wedge \vec{H}$$

L'intégrande de l'intégrale de surface de [2.51] se met alors sous la forme :

$$[(\vec{n} \wedge \vec{H}_1) \wedge \vec{H}_2] \cdot \vec{ds} - [(\vec{n} \wedge \vec{H}_2) \wedge \vec{H}_1] \cdot \vec{ds}$$

Le développement des doubles produits vectoriels de cette expression conduit à une expression nulle pour le premier membre.

Dans tous les cas cités précédemment, pour lesquels le premier membre de l'expression [2.51] est nulle, on déduit que :

$$\iiint_V \vec{E}_1 \vec{J}_2 dv = \iiint_V \vec{E}_2 \vec{J}_1 dv \quad [2.52]$$

qui s'énonce comme suit. L'effet de la source 1 ($\vec{E}_1$) pris à l'endroit de la source 2, en M_2 et multiplié par celle-ci est égal à l'effet de la source 2 ($\vec{E}_2$), pris à l'endroit de la source 1, en M_1 , multiplié par la valeur de la source 1. Il y a donc réciprocité entre les deux situations d'excitation. La réciprocité apparaît encore plus nettement si on attribue aux deux sources la valeur unité. La réciprocité permet d'effectuer certains calculs de champs qui ne seraient pas possibles sinon. Si les deux sources ont la même l'amplitude, les projections des champs sont égales.

■ Généralisation du théorème de réciprocité de Lorentz

Dans les cas où il est nécessaire de généraliser les équations de Maxwell, en introduisant les courants magnétiques $\vec{M}$, le théorème se généralise.

Considérons les deux états suivants :

- Les sources électriques et magnétiques ($\vec{J}_1, \vec{M}_1$) créent les champs électriques et magnétiques : ($\vec{E}_1, \vec{H}_1$) pour l'état 1.
- Les sources électriques et magnétiques ($\vec{J}_2, \vec{M}_2$) créent les champs électriques et magnétiques : ($\vec{E}_2, \vec{H}_2$) pour l'état 2.

Les équations de Maxwell sont modifiées en prenant en compte des courants magnétiques :

$$\text{rot } \vec{E}_1 = -j\omega\mu\vec{H}_1 - \vec{M}_1 \quad \text{et} \quad \text{rot } \vec{E}_2 = -j\omega\mu\vec{H}_2 - \vec{M}_2$$

La démonstration précédente peut être reprise sur le même principe et l'on obtient le théorème de Lorentz :

$$\iint_S (\vec{E}_1 \wedge \vec{H}_2 - \vec{E}_2 \wedge \vec{H}_1) \vec{ds} = \iiint_V (-\vec{E}_1 \vec{J}_2 + \vec{E}_2 \vec{J}_1 + \vec{H}_1 \vec{M}_2 - \vec{H}_2 \vec{M}_1) dv \quad [2.53]$$

Dans les quatre cas cités pour lesquels le premier membre est nul, nous obtenons :

$$\iiint_V (\vec{E}_1 \vec{J}_2 - \vec{H}_1 \vec{M}_2) dv = \iiint_V (\vec{E}_2 \vec{J}_1 - \vec{H}_2 \vec{M}_1) dv \quad [2.54]$$

L'interprétation est la même que précédemment en prenant en compte les courants magnétiques. Ce théorème permet de vérifier certaines propriétés des objets diffractants. En particulier, la réciprocité entraîne le fait que les propriétés du rayonnement d'une antenne sont les mêmes en émission et en réception.

2.4.4 Théorie des images

L'utilisation des charges images est courante en électromagnétisme. Le principe est connu depuis longtemps en électrostatique. Il s'énonce simplement en trouvant une équivalence au problème d'une charge ponctuelle en présence d'un plan infini parfaitement conducteur. Les deux problèmes représentés sur la figure 2.10 sont équivalents dans le demi-plan supérieur.

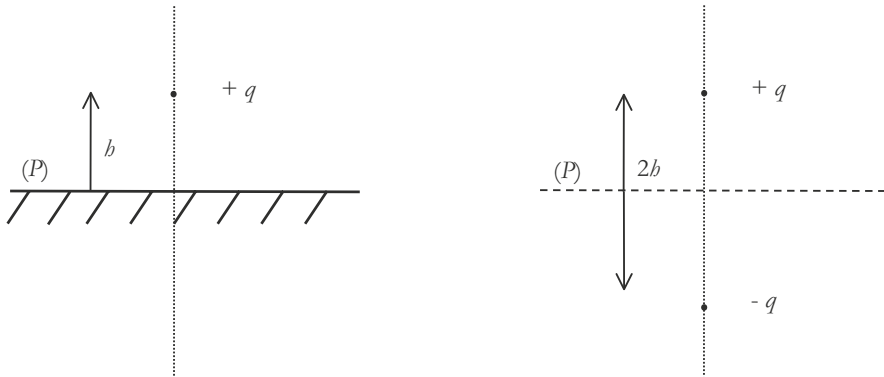


Figure 2.10 – Principe d'équivalence des charges.

Lorsque deux charges égales en valeur absolue, mais de signe opposé se trouvent placées à égale distance (b) d'un plan fictif infini (P), celui-ci contient tous les points de l'espace portés à un potentiel nul. Le problème équivalent est donc constitué d'une charge égale à l'une des deux charges précédentes, placée en regard d'un plan métallique (P), portée à un potentiel nul, situé à la même position que le plan (P) médiateur des deux charges.

Les charges $-q$ et $+q$ sont dites images l'une de l'autre.

Nous constatons, sur cet exemple de base, que le problème équivalent ne donne un résultat que dans un demi-espace. L'équivalence n'est pas complète. Il faut traiter un autre problème équivalent pour l'autre demi-espace.

En électromagnétisme, les sources sont des charges variables.

■ Dipôle parallèle au plan de masse

Un dipôle est placé parallèlement au plan de masse, à une distance b (figure 2.11). Sachant qu'il est équivalent à deux charges opposées, placées à l'extrémité du dipôle, nous allons préciser le problème équivalent par rapport à ces charges.

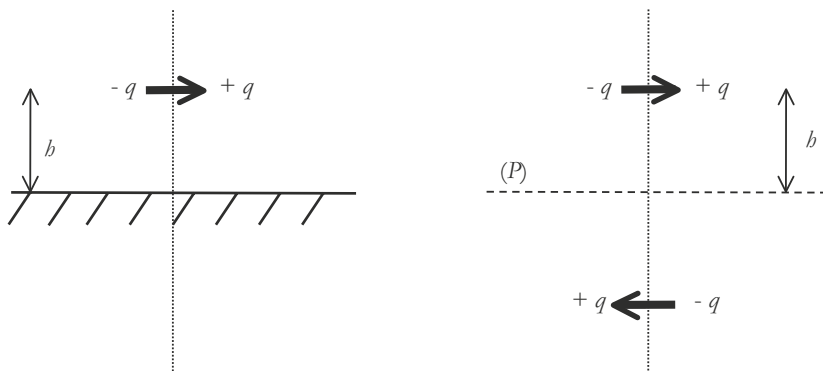


Figure 2.11 – Dipôle horizontal en présence d'un plan de masse (à gauche) et le système équivalent (à droite).

La charge $-q$ du dipôle a une image $+q$ placée symétriquement par rapport au plan (P). De même pour la charge de signe opposé du dipôle. Le dipôle image apparaît donc comme parallèle au dipôle objet et de signe opposé.

■ Dipôle perpendiculaire au plan de masse

Un dipôle est placé perpendiculairement au plan de masse, à une distance h (figure 2.12).

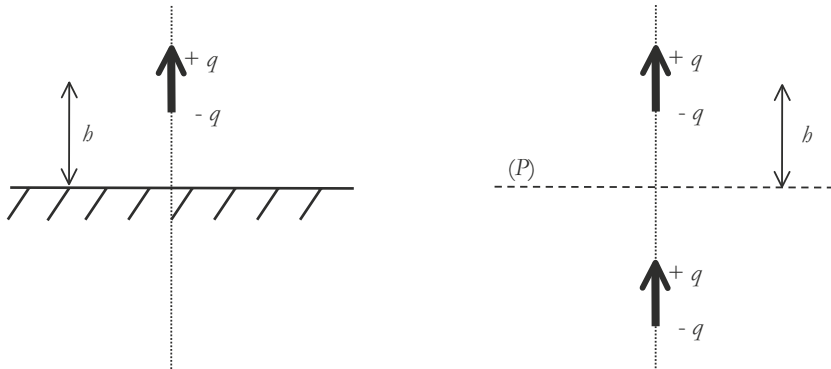


Figure 2.12 – Dipôle vertical en présence d'un plan de masse (à gauche) et le système équivalent (à droite).

En utilisant la même méthode que pour le dipôle horizontal, on constate que le dipôle image est parallèle au dipôle objet. Dans ce cas, les deux dipôles ont le même sens.

3 • CALCUL DU RAYONNEMENT D'ANTENNES DE RÉFÉRENCE

Dans ce chapitre, nous montrerons comment calculer le rayonnement d'antennes de base, qui constituent une référence quant aux méthodes utilisées.

3.1 Rayonnement créé par des courants

Nous avons vu au début de cet ouvrage que les antennes pouvaient être classées en deux catégories :

- celles qui émettent un rayonnement engendré par des courants variables,
- celles dont le rayonnement dépend du champ électromagnétique créé sur des ouvertures.

Nous présentons dans ce chapitre les calculs du champ électromagnétique créé par la première catégorie d'antennes.

3.1.1 Méthode de calcul

Les sources de rayonnement, situées dans une région de l'espace, à distance finie de l'observateur sont placées au point P. Elles sont repérées par le vecteur $\vec{r}'$. L'observateur est situé en M, repéré par le vecteur $\vec{r}$ (figure 3.1). Le vecteur $\vec{R} = \vec{r} - \vec{r}'$ apparaît dans les calculs car il permet de repérer le point d'observation par rapport aux sources.

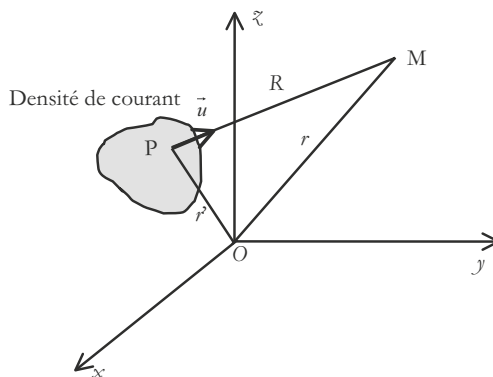


Figure 3.1 – Position des courants sources par rapport au point d'observation.

Le potentiel vecteur au point M est donné par l'expression [2.23] :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} \iiint_{sources} \frac{\vec{j}(\vec{r}', t')}{\|\vec{R}\|} dv'$$

$\vec{j}(\vec{r}', t')$ représente la source du champ électromagnétique. Dans notre cas c'est une densité volumique de courant électrique.

En régime harmonique, les sources s'expriment sous la forme :

$$\vec{j}(\vec{r}', t') = \vec{j}(\vec{r}')e^{j\omega t'}$$

Avec $t' = t - \frac{R}{c}$ et $R = ||\vec{R}||$. D'où :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} e^{j\omega t} \iiint_{sources} \vec{j}(\vec{r}') \frac{e^{-jkR}}{R} dv'$$

Posons :

$$\vec{A}(M) = \frac{\mu_0}{4\pi} \iiint_{sources} \psi(R) \vec{j}(M') dv'$$

Les sources de courant sont notées : $\vec{j}(M')$.

Avec les notations :

$$\psi(R) = \frac{e^{-jkR}}{R}, \quad \text{et} \quad k = \frac{\omega}{c}$$

Le champ électrique, exprimé à partir du potentiel vecteur $\vec{A}$ et du potentiel scalaire φ se met sous la forme :

$$\vec{E} = -j\omega \vec{A} - \vec{\text{grad}}\varphi$$

La jauge de Lorentz [2.20] lie le potentiel scalaire au potentiel vecteur. Elle s'exprime dans le domaine de Fourier par :

$$\epsilon_0 \mu_0 j\omega \varphi + \text{div} \vec{A} = 0$$

On déduit donc l'expression du champ électrique :

$$\vec{E} = \frac{1}{j\omega\epsilon_0\mu_0} (\vec{\text{grad}}(\text{div} \vec{A}) + k^2 \vec{A}) \quad [3.1]$$

Pour le champ magnétique :

$$\vec{H} = \frac{1}{\mu_0} \vec{\text{rot}} \vec{A}$$

Ces expressions sont calculées au point d'observation M et les dérivées apparaissant dans ces expressions sont effectuées par rapport au point M.

Le calcul complet du champ électrique doit tenir compte de l'intégration sur les sources. Cela donne :

$$\vec{E}(M) = \frac{1}{4\pi j\omega\epsilon_0} \iiint_{sources} (\vec{\text{grad}}(\text{div}) + k^2) [\psi(R) \vec{j}(M')] dv'$$

Dans cette expression, la source de courant est considérée en M'. La fonction ψ dépend de la distance entre les sources et le point d'observation.

Remarquons que :

$$\vec{\text{grad}}\psi(R) = \psi'(R) \vec{u}$$

$\vec{u}$ est le vecteur unitaire du rayon partant de la source vers le point d'observation

Calculons la divergence de $\psi(R) \vec{j}(M')$:

$$\text{div} [\psi(R) \vec{j}(M')] = \vec{j}(M') \cdot \vec{\text{grad}}\psi(R)$$

Avec la remarque précédente :

$$\text{div} \left[\psi(R) \vec{j}(M') \right] = \psi'(R) \vec{j}(M') \cdot \vec{u}$$

Pour dériver cette expression, utilisons le fait que :

$$\overrightarrow{\text{grad}} \left[\vec{j}(M') \cdot \vec{u} \right] = \frac{\vec{j}(M')}{R} - \frac{(\vec{j}(M') \cdot \vec{u}) \vec{u}}{R}$$

Il vient :

$$\overrightarrow{\text{grad}}(\text{div} [\psi(R) \vec{j}(M')]) = \frac{\psi'(R)}{R} \vec{j}(M') + \left(\psi''(R) - \frac{\psi'(R)}{R} \right) (\vec{j}(M') \cdot \vec{u}) \vec{u}$$

Les dérivations successives de la fonction ψ conduisent aux expressions suivantes :

$$\psi'(R) = - \left(jk + \frac{1}{R} \right) \psi(R)$$

$$\psi''(R) = \left(-k^2 + \frac{2jk}{R} + \frac{2}{R^2} \right) \psi(R)$$

L'expression complète du champ électrique en M est donc :

$$\begin{aligned} \vec{E}(M) = \frac{k^2}{4\pi j\omega\epsilon_0} \iiint_{\text{sources}} & \left[\left(1 + \frac{1}{jkR} - \frac{1}{k^2 R^2} \right) \vec{j}(M') \right. \\ & \left. - \left(1 + \frac{3}{jkR} - \frac{3}{k^2 R^2} \right) (\vec{j}(M') \cdot \vec{u}) \vec{u} \right] \psi(R) dv' \end{aligned} \quad [3.2]$$

Le calcul du champ magnétique est plus simple. Il utilise le fait que :

$$\overrightarrow{\text{rot}} \left[\psi(R) \vec{j}(M') \right] = \psi'(R) \vec{u} \wedge \vec{j}(M')$$

L'expression du champ magnétique en M est donc :

$$\vec{H}(M) = \frac{jk}{4\pi} \iiint_{\text{sources}} \left(1 + \frac{1}{jkR} \right) (\vec{j}(M') \wedge \vec{u}) \psi(R) dv' \quad [3.3]$$

3.1.2 Champ créé à grande distance

Les expressions précédentes se simplifient à grande distance.

Considérons $MM' \gg \lambda$, c'est-à-dire $kR \gg 1$ dans les expressions des champs électrique et magnétique, démontrées dans le paragraphe précédent [3.2] et [3.3], on néglige les termes en $\frac{1}{kR}$ et leurs puissances, devant l'unité. Les champs s'expriment ainsi sous la forme :

$$\begin{aligned} \vec{E}(M) &= \frac{k^2}{4\pi j\omega\epsilon_0} \iiint_{\text{sources}} \left[\vec{j}(M') - (\vec{j}(M') \cdot \vec{u}) \vec{u} \right] \psi(R) dv' \\ &= \frac{-k^2}{4\pi j\omega\epsilon_0} \iiint_{\text{sources}} \left[\vec{j}(M') \wedge \vec{u} \right] \psi(R) dv' \\ \vec{H}(M) &= \frac{jk}{4\pi} \iiint_{\text{sources}} (\vec{j}(M') \wedge \vec{u}) \psi(R) dv' \end{aligned}$$

On constate alors que :

$$\vec{H}(M) = \frac{\vec{u} \wedge \vec{E}(M)}{Z} \quad \text{avec} \quad Z = \sqrt{\frac{\mu_0}{\epsilon_0}} \approx 377 \, \Omega$$

Le champ électromagnétique a localement la même structure qu'une onde plane, puisque le champ électrique est perpendiculaire au champ magnétique. Il ne faut cependant pas oublier que les champs varient selon l'inverse de la distance aux sources. L'onde est dite localement plane, en restreignant cette propriété dans l'espace au voisinage immédiat du point.

Supposons que le volume contenant les sources est de dimensions petite par rapport à R et proche de l'origine. Ainsi, selon la figure 3.2, considérant que M est très loin des sources, le vecteur $\vec{R}$ est pratiquement parallèle à $\vec{r}$.

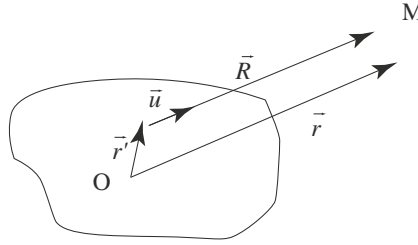


Figure 3.2 – Position de M à grande distance.

Partant du fait que :

$$\vec{R} = \vec{r} - \vec{r}'$$

et que $\vec{R}$ et $\vec{r}$ sont pratiquement parallèles, on peut alors écrire :

$$R \approx r - \vec{u} \cdot \vec{r}'$$

La fonction $\psi(R)$ peut être développée selon :

$$\psi(R) \approx e^{j\vec{k} \cdot \vec{r}'} \frac{e^{-jkr}}{r - \vec{u} \cdot \vec{r}'}$$

$\vec{r}$ étant très supérieur à $\vec{r}'$, on obtient :

$$\psi(R) \approx \psi(r) e^{j\vec{k} \cdot \vec{r}'}$$

Les champs s'expriment ainsi sous la forme :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \iiint_{sources} \left[\vec{j}(M') \wedge \vec{u} \right] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} dv' \quad [3.4]$$

$$\vec{H}(M) = \frac{jk}{4\pi} \psi(r) \iiint_{sources} \left(\vec{j}(M') \wedge \vec{u} \right) e^{j\vec{k} \cdot \vec{r}'} dv' \quad [3.5]$$

Ces formules du champ électromagnétique expriment le fait que le champ résultant en un point M provient de l'intégration, en volume, d'une expression vectorielle portant sur les sources de courant. Ces formules sont transposables sans difficulté au cas de courant surfacique ou linéique. Seuls les éléments d'intégration et l'ordre d'intégration changent dans ces cas. Nous allons, dans la suite, montrer comment utiliser ces formules générales dans le cas de sources de courant linéiques.

■ Application à une répartition linéique de courant

Prenons l'exemple d'un fil rectiligne parcouru par une intensité I dont la valeur dépend du point source considéré M' , variant sinusoïdalement dans le temps avec une pulsation ω . Ce fil est porté par l'axe z et de longueur finie (figure 3.3).

En tenant compte de la relation entre la densité de courant $\vec{j}$ et l'intensité I pour un courant filiforme, le terme $\vec{j} dv'$ est équivalent à $I \vec{dl}'$, vecteur élémentaire porté par le fil. Le champ électrique en M est alors calculé selon l'expression :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \int_{fil} [I(M') \vec{u}_z \wedge \vec{u}] \wedge \vec{u} e^{i\vec{k} \cdot \vec{r}'} \vec{dl}'$$

Soit encore :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \sin \theta \int_{z_1}^{z_2} I(z') e^{jkz' \cos \theta} dz' \vec{u}_\theta \quad [3.6]$$

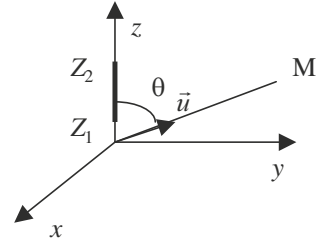


Figure 3.3 – Source de courant filiforme.

3.1.3 Rayonnement d'un dipôle électrique

Lorsque l'élément de courant est très petit devant la longueur d'onde, on dit que c'est un dipôle électrique, appelé aussi doublet électrique. Ce nom vient du fait qu'aux deux extrémités du fil se développent des charges opposées. Le principe de conservation de la charge s'exprime par :

$$I = \frac{dq}{dt}$$

Cela montre que le courant n'étant pas nul à l'extrémité du fil, des charges ponctuelles doivent apparaître pour satisfaire à ce principe.

La longueur l du dipôle étant très petite, on pourra considérer que, dans l'intégration, l'intensité est constante, car on vérifie toujours $l \ll \lambda$.

Pour la même raison, kz' est toujours très petit devant l'unité, d'où :

$$e^{jkz' \cos \theta} \approx 1$$

Dans la mesure où la longueur l du fil reste très petite par rapport aux distances d'observation, on assimilera R à r . L'intégration ne porte que sur l'élément de longueur du fil. Il est possible de déterminer le champ électromagnétique que ce soit à proximité du dipôle ou avec les approximations de champ lointain.

■ Rayonnement d'un dipôle à distance quelconque

Tenant compte du fait que :

$$\vec{u}_z = \vec{u} \cos \theta - \vec{u}_\theta \sin \theta$$

et avec les remarques précédentes concernant l'intégration, on obtient les coordonnées sphériques du champ rayonné, en utilisant [3.2] :

$$E_r = \frac{jkZ}{2\pi} (Il) \left(\frac{1}{jkR} - \frac{1}{k^2 R^2} \right) \frac{e^{-jkr}}{r} \cos \theta \quad [3.7]$$

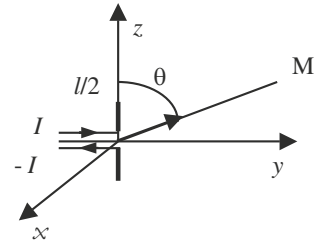


Figure 3.4 – Dipôle électrique.

$$E_{\theta} = \frac{jkZ}{4\pi} (Il) \left(1 + \frac{1}{jkR} - \frac{1}{k^2 R^2} \right) \frac{e^{-jkr}}{r} \sin \theta \quad [3.8]$$

$$H_{\varphi} = \frac{jk}{4\pi} (Il) \left(1 + \frac{1}{jkR} \right) \frac{e^{-jkr}}{r} \sin \theta \quad [3.9]$$

On remarque que le champ magnétique n'a qu'une composante qui « tourne » autour de l'élément de courant, conformément à ce qu'on sait de la création d'un champ magnétique par un courant. Cela vient du fait que l'induction magnétique est à divergence nulle. Le champ électromagnétique possède une symétrie radiale et s'annule dans la direction du fil ($\theta = 0$ ou $\theta = \pi$).

■ Rayonnement d'un dipôle en champ lointain

Lorsque la distance d'observation est grande devant la longueur d'onde, les termes en $\frac{1}{kR}$ sont négligeables et le champ électromagnétique prend la forme :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \sin \theta (Il) \vec{u}_{\theta} \quad [3.10]$$

$$\vec{H}(M) = \frac{jk}{4\pi} \psi(r) \sin \theta (Il) \vec{u}_{\varphi} \quad [3.11]$$

On retrouve bien la structure d'une onde localement plane :

$$\vec{E} = Z(\vec{H} \wedge \vec{u}) \quad \text{avec} \quad Z = 377 \, \Omega$$

Compte tenu de la définition du vecteur de Poynting complexe :

$$\vec{\Pi} = \vec{E} \wedge \vec{H}^*$$

Il s'exprime, à grande distance par :

$$\vec{\Pi} = \frac{k^2}{16\pi^2} Z |\psi(r)|^2 \sin^2 \theta (Il)^2 \vec{u}_r \quad [3.12]$$

La densité surfacique de puissance a donc pour valeur :

$$S_r(\theta, \varphi) = \frac{k^2}{32\pi^2} Z |\psi(r)|^2 \sin^2 \theta (Il)^2 \quad [3.13]$$

On constate que la densité de puissance est nulle dans la direction $\theta = 0$ et qu'elle est maximale dans la direction $\theta = \frac{\pi}{2}$. Cette dernière direction correspond à l'axe de rayonnement de l'antenne dipolaire. La transmission du signal est donc maximale lorsque le dipôle d'émission et le dipôle de réception sont parallèles et que leurs perpendiculaires sont communes.

La fonction caractéristique de rayonnement qui représente la répartition angulaire de la densité de puissance S_r [2.42] sur une sphère de rayon r , s'exprime par :

$$F(\theta, \varphi) = \frac{k^2}{32\pi^2} Z \sin^2 \theta (Il)^2 \quad [3.14]$$

Sur la figure 3.5, le diagramme de rayonnement (représentation de la fonction caractéristique de rayonnement) est tracé en trois dimensions. On se rend compte, sur cette figure, de la symétrie axiale de la fonction caractéristique de l'antenne. Le diagramme de rayonnement en deux dimensions de la figure 3.6 permet d'estimer précisément la valeur de la fonction caractéristique.

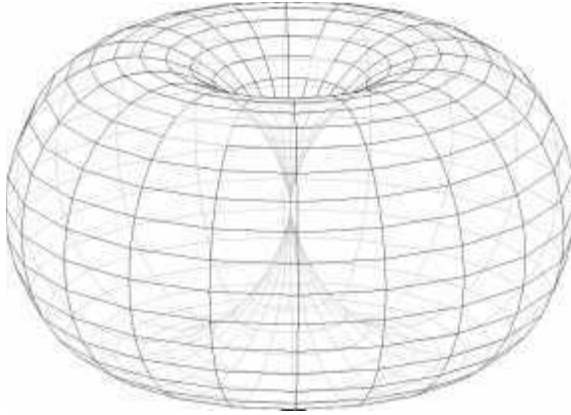


Figure 3.5 – Diagramme de rayonnement d'un dipôle en trois dimensions.

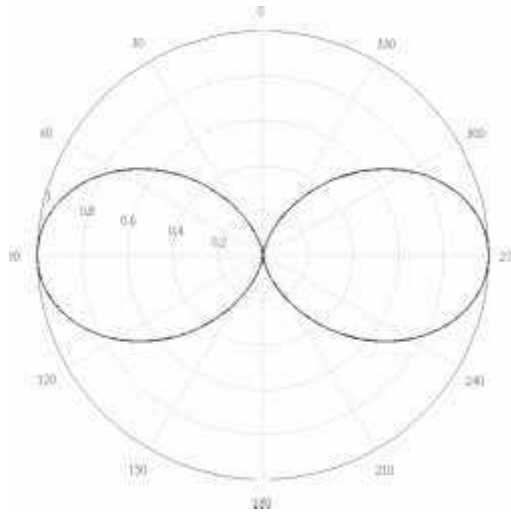


Figure 3.6 – Diagramme de rayonnement normalisé d'un dipôle en deux dimensions en fonction de l'angle θ .

■ Résistance de rayonnement

La puissance totale rayonnée par l'antenne à travers une sphère de rayon r est donnée par :

$$P_r = \frac{k^2}{32\pi^2} Z \frac{(Il)^2}{r^2} \int_0^\pi 2\pi r^2 \sin^2 \theta \sin \theta d\theta$$

Le résultat de l'intégration donne :

$$P_r = \frac{k^2}{12\pi} Z (Il)^2$$

Soit encore, en rapportant la taille du dipôle à la longueur d'onde λ :

$$P_r = \frac{\pi}{3} \frac{l^2}{\lambda^2} Z I^2 \quad [3.15]$$

La puissance est proportionnelle au carré de l'intensité, ce qui nous permet de considérer que l'antenne a une résistance, appelée résistance de rayonnement dont la valeur est donnée par rapport à l'intensité efficace I_{eff} :

$$P_r = R_r I_{eff} = R_r \frac{1}{2} I^2$$

Soit

$$R_r = \frac{2\pi}{3} Z \frac{l^2}{\lambda^2} \approx 800 \frac{l^2}{\lambda^2} \quad [3.16]$$

Un calcul rapide donne un ordre de grandeur de la résistance d'un dipôle. Si le rapport de la taille du dipôle à la longueur d'onde est de 1/10, la résistance de rayonnement est de 8Ω . On constate aussi que, plus la dimension du dipôle est petite, plus la puissance rayonnée est faible.

La directivité de l'antenne est définie, de façon générale, par :

$$D(\theta) = \frac{4\pi F(\theta)}{\iint_{espace} F(\theta) d\Omega}$$

Le calcul de la directivité du dipôle est donné par :

$$D(\theta) = \frac{3}{2} \sin^2 \theta \quad [3.17]$$

La directivité maximale est donc de 1,5. Le doublet ayant un rayonnement radialement uniforme, sa directivité n'est pas très grande.

3.1.4 Rayonnement des antennes filaires à ondes stationnaires

La puissance rayonnée par un dipôle étant très faible, on préfère, pour transmettre une puissance plus importante, tout en conservant un diagramme de rayonnement large, utiliser une antenne filaire. Cette antenne de longueur totale l est alimentée en son centre par un courant alternatif.

■ Expression du champ à grande distance

L'antenne filaire est placée à l'origine, parallèlement à l'axe z selon la figure 3.7.

On rappelle l'expression générale du champ électrique créé, à grande distance, par une répartition de courant filaire variant sinusoïdalement :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \sin \theta \int_{z_1}^{z_2} I(z') e^{jkz' \cos \theta} dz' \vec{u}_\theta \quad [3.18]$$

La dimension du fil étant de l'ordre de grandeur de la longueur d'onde, le calcul de l'intégrale doit prendre en compte la forme du courant en fonction de z .

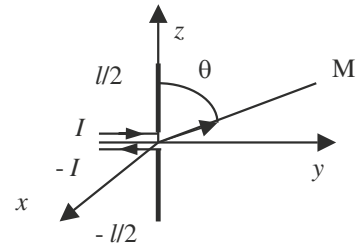


Figure 3.7 – Antenne filaire.

■ Antennes à ondes stationnaires

L'antenne est alimentée par une ligne bifilaire sur laquelle s'installe un régime d'onde stationnaire. Compte tenu de la forme symétrique de l'antenne et de son alimentation au centre du fil, la répartition du courant prend la forme suivante :

$$I(z') = I_0 \sin \left[k \left(\frac{l}{2} - |z'| \right) \right] \quad [3.19]$$

On remarque qu'au centre de l'antenne, la valeur du courant est donnée par :

$$I(0) = I_0 \sin \left(k \frac{l}{2} \right) \quad [3.20]$$

À l'extrémité de l'antenne, l'intensité est nulle.

On est donc conduit à calculer la valeur de l'intégrale de l'expression [3.18] en tenant compte de la fonction de répartition du courant. Expriment la valeur absolue de z' , l'intégrale fait apparaître deux termes :

$$\begin{aligned} \int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' &= I_0 \int_{-l/2}^0 \sin \left[k \left(\frac{l}{2} + z' \right) \right] e^{jkz' \cos \theta} dz' \\ &\quad + I_0 \int_0^{l/2} \sin \left[k \left(\frac{l}{2} - z' \right) \right] e^{jkz' \cos \theta} dz' \end{aligned}$$

Compte tenu de la symétrie des termes contenus dans les deux intégrales :

$$\int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' = 2I_0 \int_0^{l/2} \sin \left[k \left(\frac{l}{2} - z' \right) \right] \cos(kz' \cos \theta) dz'$$

Exprimant le produit des fonctions trigonométriques en une somme :

$$\begin{aligned} \int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' &= I_0 \int_0^{l/2} \sin k \left[\frac{l}{2} - z'(1 - \cos \theta) \right] dz' \\ &\quad + I_0 \int_0^{l/2} \sin k \left[\frac{l}{2} - z'(1 + \cos \theta) \right] dz' \end{aligned}$$

Il est alors aisé de trouver des primitives :

$$\int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' = I_0 \left[\frac{-\cos k \left[\frac{l}{2} - z'(1 - \cos \theta) \right]}{-k(1 - \cos \theta)} + \frac{-\cos k \left[\frac{l}{2} - z'(1 + \cos \theta) \right]}{-k(1 + \cos \theta)} \right]_{0}^{l/2}$$

L'intégration des deux termes conduit donc à :

$$\int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' = I_0 \left[\cos \left(k \frac{l}{2} \cos \theta \right) - \cos k \frac{l}{2} \right] \left[\frac{1}{k(1 - \cos \theta)} + \frac{1}{k(1 + \cos \theta)} \right]$$

Soit encore à :

$$\int_{-l/2}^{l/2} I(z') e^{jkz' \cos \theta} dz' = \frac{2I_0}{k} \frac{\cos \left(k \frac{l}{2} \cos \theta \right) - \cos k \frac{l}{2}}{\sin^2 \theta} \quad [3.21]$$

Le champ électromagnétique s'exprime donc sous la forme :

$$\vec{E}(M) = \frac{jZ}{2\pi} I_0 \frac{e^{-jkr}}{r} \frac{\cos \left(k \frac{l}{2} \cos \theta \right) - \cos k \frac{l}{2}}{\sin \theta} \vec{u}_\theta \quad [3.22]$$

□ Diagrammes de rayonnement des antennes filaires

Le vecteur de Poynting complexe a pour expression :

$$\vec{E} \wedge \vec{H}^* = \frac{I_0^2}{4\pi^2 r^2} \left[\frac{\cos\left(k\frac{l}{2} \cos \theta\right) - \cos k\frac{l}{2}}{\sin \theta} \right]^2$$

On rappelle que :

$$\frac{kl}{2} = \frac{\pi l}{\lambda}$$

En posant $l = n\lambda$, la fonction caractéristique de rayonnement s'exprime par :

$$F(\theta) = \left(\frac{\cos(n\pi \cos \theta) - \cos n\pi}{\sin \theta} \right)^2$$

La forme du diagramme de rayonnement dépend donc de la taille de l'antenne rapportée à la longueur d'onde. Ses propriétés sont les suivantes :

- Dans la direction $\theta = 0$, le diagramme de rayonnement est nul.
- Si $l \leq \lambda$, il n'existe qu'un seul lobe présentant un maximum pour $\theta = \pi/2$.
- Si $l > \lambda$, il existe plusieurs lobes. Le diagramme de rayonnement présente des zéros pour les angles tels que :

$$\cos \theta_0 = \pm 1 + 2m \frac{\lambda}{l}$$

m étant un entier. Ces zéros se trouvent à $\theta_0 = \pi/2$, si $l = 2m\lambda$.

Quelques diagrammes de rayonnement sont présentés dans la suite pour plusieurs longueurs d'antennes.

□ Résistance de rayonnement d'une antenne demi-onde

Le calcul de la résistance de rayonnement des antennes filaires s'effectue de la même façon que dans le cas des antennes dipolaires. Le résultat s'exprime sous la forme :

$$R_r = \frac{Z}{4\pi} C_i(2\pi) = 73 \, \Omega \quad \text{avec} \quad C_i(x) = - \int_x^\infty \frac{\cos u}{u} du$$

La résistance de rayonnement de l'antenne demi-onde est environ dix fois plus grande que celle du dipôle ($8 \, \Omega$ pour une longueur de l'ordre du dixième de longueur d'onde).

□ Directivité d'une antenne demi-onde

On rappelle la définition de la directivité :

$$D(\theta) = \frac{4\pi F_n(\theta)}{\iint_{\text{espace}} F_n(\theta) d\Omega}$$

Le calcul de la directivité de l'antenne demi-onde conduit à :

$$D(\theta) = 1,64 \left(\frac{\cos(\pi/2 \cos \theta)}{\sin \theta} \right)^2 \quad [3.23]$$

Donc la directivité au maximum de cette antenne est à peine supérieure à celle du dipôle (1,5).

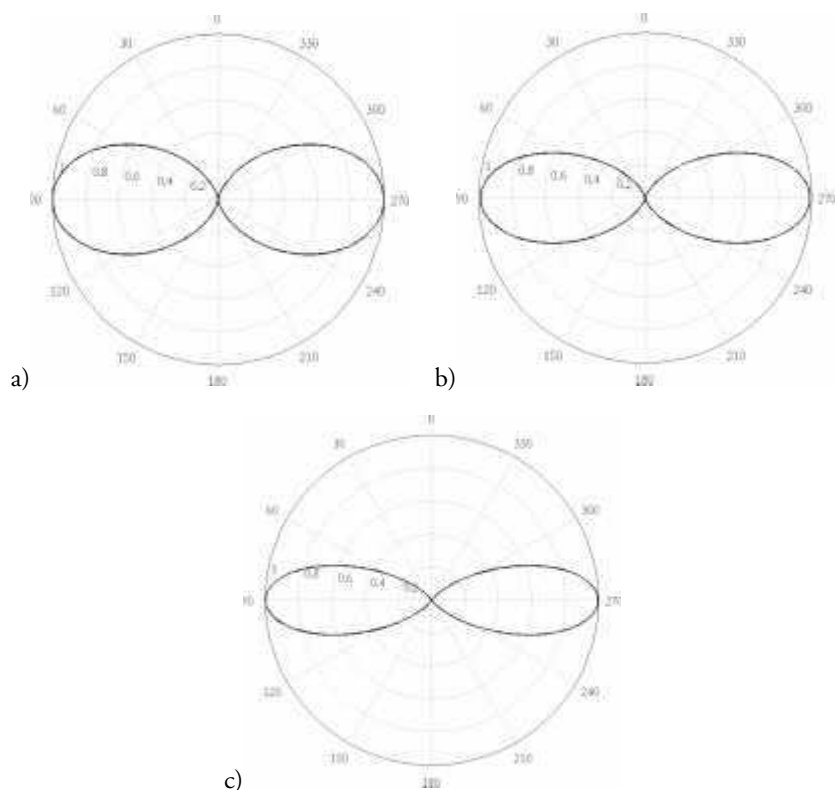


Figure 3.8 – a) Diagramme de rayonnement normalisé d'une antenne demi-onde ($n = 1/2$) en fonction de θ .

b) Diagramme de rayonnement normalisé d'une antenne trois quarts de longueur d'onde ($n = 3/4$) en fonction de θ .

c) Diagramme de rayonnement normalisé d'une antenne onde ($n = 1$) en fonction de θ .

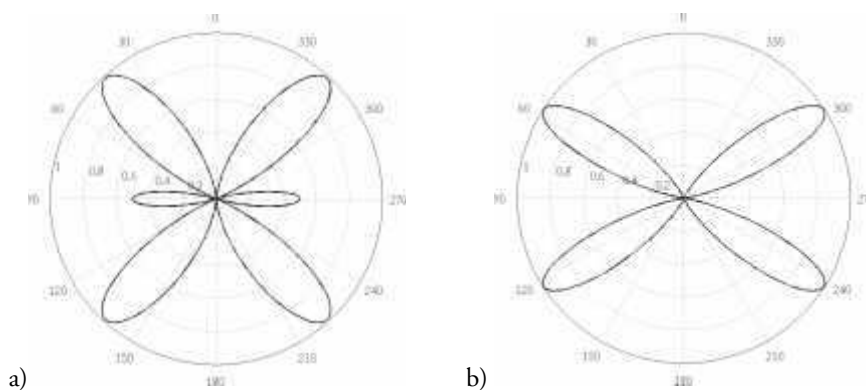


Figure 3.9 – a) Diagramme de rayonnement normalisé d'une antenne d'une longueur d'onde et demi ($n = 3/2$) en fonction de θ .

b) (3.24) Diagramme de rayonnement normalisé d'une antenne deux longueurs d'onde ($n = 2$) en fonction de θ .

3.1.5 Rayonnement des antennes filaires de type boucle

Nous allons appliquer dans ce paragraphe les méthodes de calcul vues précédemment au cas de l'antenne boucle. La boucle, de rayon a , est parcourue par un courant d'intensité I (figure 3.10).

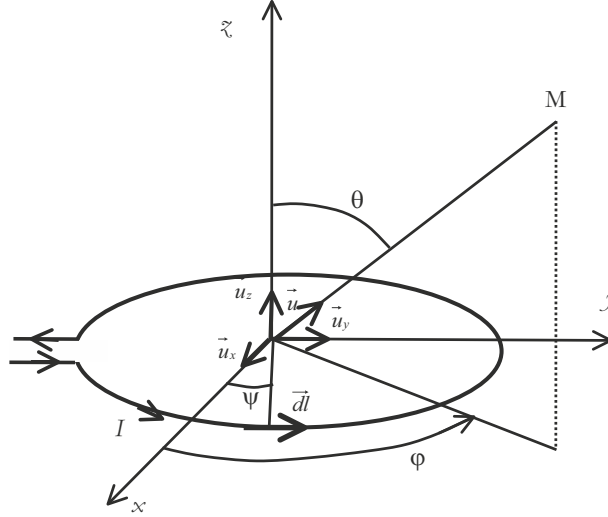


Figure 3.10 – Géométrie d'une antenne boucle.

Chaque élément de courant porté par $\vec{dl}$, crée un champ électrique élémentaire dont on somme les contributions.

■ Cas de la boucle de dimension petite par rapport à la longueur d'onde

Dans ce cas, l'intensité est pratiquement constante le long de la boucle. Seule sera considérée sa variation temporelle :

$$I(t) = I e^{j\omega t}$$

Le champ électrique lointain se déduit de [3.4], appliquée aux courants filiformes :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) I \int_{\text{boucle}} \left[\vec{dl} \wedge \vec{u} \right] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} d\psi$$

L'élément de courant a pour expression :

$$\vec{dl} = a \left[-\sin \psi \vec{u}_x + \cos \psi \vec{u}_y \right] d\psi$$

Soit encore :

$$\vec{dl} \wedge \vec{u} = a \left[(\sin \varphi \sin \psi + \cos \varphi \cos \psi) \vec{u}_\theta + \cos \theta (\cos \varphi \sin \psi - \sin \varphi \cos \psi) \vec{u}_\varphi \right] d\psi$$

D'où :

$$\vec{dl} \wedge \vec{u} = a \left[\cos(\psi - \varphi) \vec{u}_\theta + \cos \theta \sin(\psi - \varphi) \vec{u}_\varphi \right] d\psi$$

Et :

$$\left(\vec{dl} \wedge \vec{u} \right) \wedge \vec{u} = a \left[-\cos(\psi - \varphi) \vec{u}_\varphi + \cos \theta \sin(\psi - \varphi) \vec{u}_\theta \right] d\psi$$

Par ailleurs, le terme de déphasage se met sous la forme :

$$\vec{k} \cdot \vec{r}' = ka \sin \theta (\cos \varphi \cos \psi + \sin \varphi \sin \psi)$$

Soit encore :

$$\vec{k} \cdot \vec{r}' = ka \sin \theta \cos(\psi - \varphi)$$

Le champ électrique s'exprime donc par :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) I a \int_{\text{boucle}} [-\cos(\psi - \varphi) \vec{u}_\varphi + \cos \theta \sin(\psi - \varphi) \vec{u}_\theta] e^{jka \sin \theta \cos(\psi - \varphi)} d\psi$$

L'intégrale doit porter sur un tour complet de la boucle, c'est-à-dire que la variation de Ψ doit être de 2π . Après un changement de variable, prenant la direction φ comme origine, l'intégrale devient :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) I a \int_0^{2\pi} [-\cos \psi \vec{u}_\varphi + \cos \theta \sin \psi \vec{u}_\theta] e^{jka \sin \theta \cos \psi} d\psi$$

La composante en θ est nulle car :

$$\int_0^{2\pi} [\cos \theta \sin \psi] e^{jka \sin \theta \cos \psi} d\psi = - \int_0^{2\pi} [\cos \theta] e^{jka \sin \theta \cos \psi} d(\cos \psi)$$

Et la primitive donne une fonction de période 2π .

La composante en φ donne, après développement de l'exponentielle complexe :

$$\vec{E}(M) = \frac{k}{2\pi} Z \psi(r) I a \int_0^\pi \cos \psi \sin(ka \sin \theta \cos \psi) d\psi \vec{u}_\varphi \quad [3.24]$$

Dans le cas d'une petite boucle, appelée aussi dipôle magnétique, l'approximation suivante est valable :

$$\sin(ka \sin \theta \cos \psi) \approx ka \sin \theta \cos \psi$$

$$\vec{E}(M) = \frac{k}{2\pi} Z \psi(r) I a \int_0^\pi (\cos \psi)^2 ka \sin \theta d\psi \vec{u}_\varphi$$

Après intégration :

$$\vec{E}(M) = \frac{k^2}{4} Z \psi(r) I a^2 \sin \theta \vec{u}_\varphi$$

La densité de puissance rayonnée est donc :

$$S_r = \frac{(ka)^4 I^2}{32r^2} Z \sin^2 \theta$$

On constate que le diagramme du rayonnement d'une petite boucle (figure 3.11) est analogue à celui d'un dipôle électrique. Cependant la polarisation des champs est différente. Alors que, pour un dipôle électrique, la polarisation du champ électrique est selon $\vec{u}_\theta$, c'est-à-dire parallèle au dipôle, la polarisation du champ électrique créé par la boucle est selon $\vec{u}_\varphi$, c'est-à-dire parallèle au plan de la boucle.

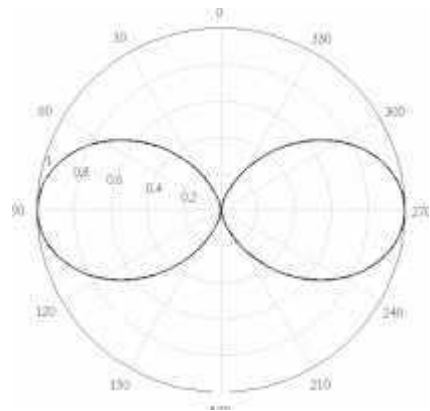


Figure 3.11 – Diagramme de rayonnement d'une boucle de courant très petite en fonction de θ .

La résistance de rayonnement de la boucle a pour valeur :

$$R_r = 320\pi^6 \left(\frac{a}{\lambda}\right)^4$$

Pour $\frac{a}{\lambda} = \frac{1}{10}$, on trouve $R_r \approx 32 \Omega$.

La résistance de rayonnement de la boucle est du même ordre que celle d'un dipôle de dimension comparable.

La valeur de la résistance de rayonnement peut être multipliée par un facteur n^2 en bobinant n spires de façon très serrée afin de limiter leur encombrement.

Lorsque le courant I peut être considéré comme constant, mais que la boucle a une taille assez grande qui ne permet pas d'effectuer le développement limité présenté précédemment, l'intégrale de l'expression [3.24] conduit à une fonction de Bessel :

$$\vec{E}(M) = \frac{kZ}{2} \psi(r) I a J_1(ka \sin \theta) \vec{u}_\varphi$$

Le diagramme de rayonnement dans ce cas est représenté sur la figure 3.12, pour une valeur $a = 0,25 \lambda$.

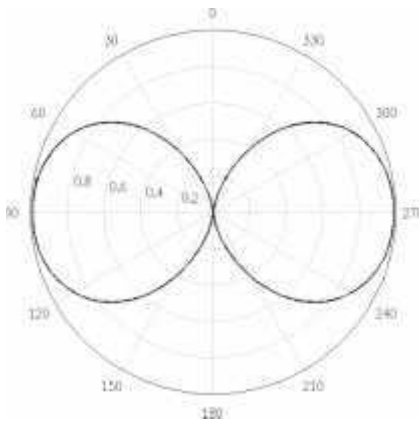


Figure 3.12 – Diagramme de rayonnement normalisé d'une boucle magnétique de rayon $a = 0,25 \lambda$, en deux dimensions, en fonction de θ .

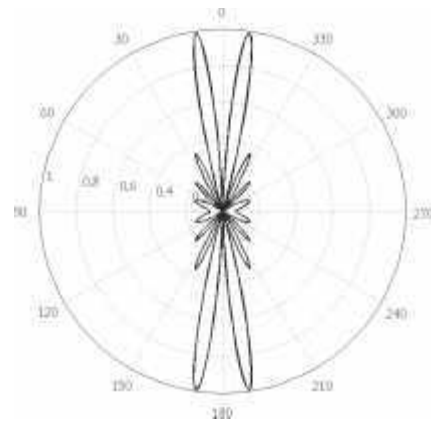


Figure 3.13 – Diagramme de rayonnement normalisé d'une boucle magnétique de rayon $a = 2 \lambda$, en deux dimensions, en fonction de θ .

En augmentant la taille de la boucle tout en maintenant le courant constant, on voit apparaître plusieurs lobes secondaires (figure 3.13). Il faut remarquer que l'alimentation doit être conçue, dans ce cas, de façon à assurer une intensité constante dans la boucle.

Lorsque la boucle est plus grande, le courant ne peut plus être considéré comme constant sur la boucle. On a affaire à une antenne à onde progressive dont l'étude est présentée au cours du paragraphe 8.3.

3.2 Diffraction par une ouverture et principe d'équivalence

Nous allons préciser, dans ce chapitre, certains calculs concernant la diffraction par une ouverture. Ces notions qui sont la base des méthodes de l'optique physique sont valables dès que les dimensions de l'ouverture diffractante sont supérieures à quelques longueurs d'ondes. La méthode présentée n'intègre pas les courants de bords sur le pourtour de l'ouverture, qui doivent exister en toute rigueur, mais n'ont d'influence sur le résultat que dans des directions très inclinées, là où le champ a des valeurs minimales. On retrouve dans cette méthode le principe d'Huyghens qui remplace l'excitation globale de l'ouverture par une infinité de sources secondaires, de déphasage constant par rapport à l'onde incidente.

Considérons une onde plane éclairant une surface plane opaque dans laquelle une ouverture a été découpée. Pour des raisons de simplicité, nous considérerons que l'onde incidente est normale à la surface et se trouve dans le demi-espace $z < 0$ (figure 3.14). Elle vient illuminer l'ouverture plane, notée (S). Le phénomène de diffraction entraîne une modification de la structure de l'onde dans le demi-espace $z > 0$, que nous allons étudier en faisant référence aux coordonnées sphériques.

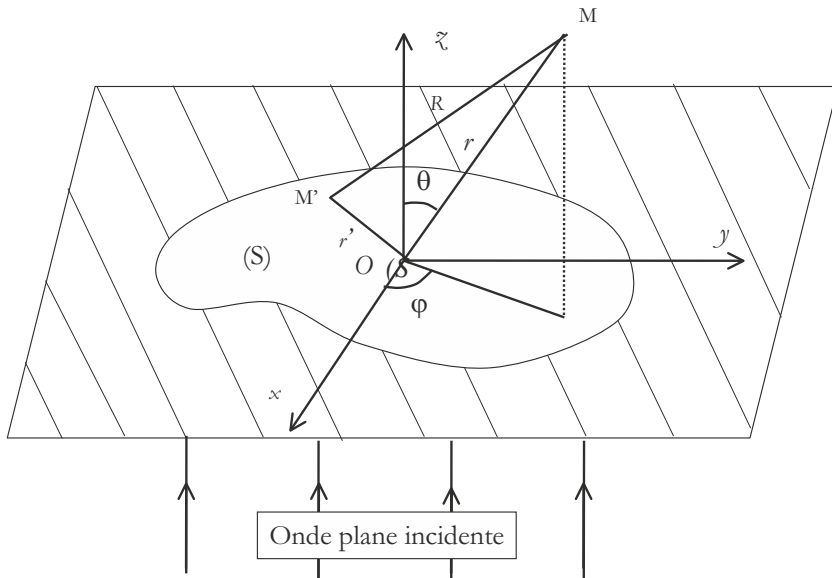


Figure 3.14 – Rayonnement d'une ouverture plane.

Le champ diffracté est considéré comme la somme des champs rayonnés par des sources secondaires placées sur la surface (S) (principe d'Huyghens). On néglige dans ce calcul le champ dû aux courants équivalents situés sur le contour.

Nous allons tout d'abord présenter des principes généraux d'étude que nous appliquons, à titre d'exemple, au cas de la diffraction par une ouverture plane.

3.2.1 Principe d'équivalence

Le but du principe d'équivalence est de remplacer le problème posé par un problème équivalent, pour lequel il est possible de trouver la solution. En fait, nous verrons que l'équivalence n'est valable que dans le demi-espace supérieur.

Considérons le champ électromagnétique de l'onde incidente polarisé rectilignement, le champ électrique $\vec{E}_i$ étant parallèle à l'axe Ox et le champ magnétique $\vec{H}_i$, parallèle à l'axe Oy (figure 3.15)

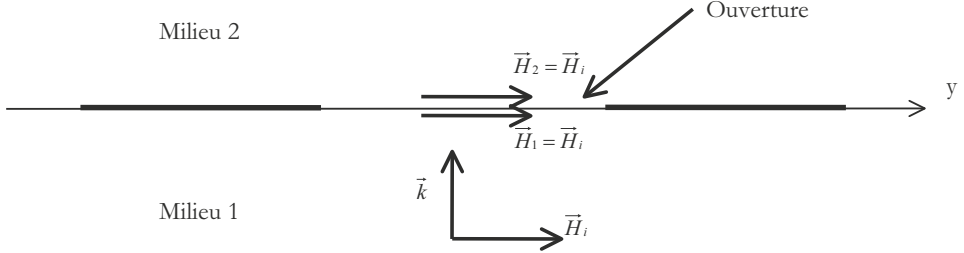


Figure 3.15 – Problème original de la diffraction.

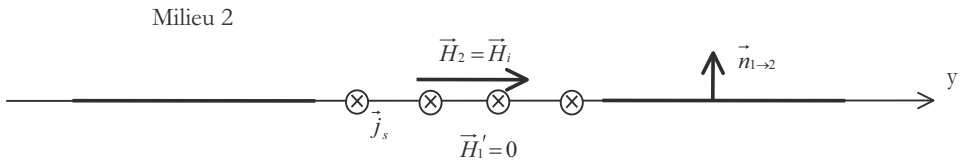
Imaginons une surface fictive passant par le plan de l'ouverture. Elle est matérialisée par une partie du plan Oxy . Le champ magnétique est continu à la traversée de cette surface puisqu'il n'existe pas de densité surfacique de courant. On vérifie donc l'égalité des champs magnétiques de part et d'autre du plan :

$$\vec{H}_1 = \vec{H}_2 = \vec{H}_i$$

Afin de calculer le champ dans le demi-espace $z > 0$, ce problème est remplacé par le problème équivalent suivant :

- dans le demi-espace inférieur $z < 0$, on place un milieu dans lequel le champ magnétique est nul ;
- le demi-espace supérieur reste identique.

En particulier, si nous voulons obtenir la même solution que celle du problème initial dans le demi-espace supérieur, il faut que le champ magnétique tangentiel soit le même que celui existant dans le problème initial, selon le théorème d'unicité. Ce champ est nul sur la partie située en dessous de l'écran et il est égal au champ incident sur la partie située au-dessus de l'ouverture (figure 3.16) dans le milieu 2. Remarquons qu'avec l'incidence normale choisie pour cette démonstration, le champ tangentiel se confond avec le champ total.



Champ magnétique nul

Figure 3.16 – Problème équivalent pour le champ magnétique.

Dans ce cas, il est nécessaire de compenser la discontinuité du champ magnétique sur l'ouverture par un courant électrique de surface. Selon les conditions de passage à la traversée de deux milieux, résultant des équations de Maxwell, introduisant la densité surfacique de courant électrique $\vec{j}_s$, la condition suivante doit être vérifiée :

$$\vec{H}_2 - \vec{H}_1 = \vec{j}_s \wedge \vec{n}_{1 \rightarrow 2}$$

Le champ magnétique dans le demi-espace inférieur est nul. On vérifie donc :

$$\vec{H}_2 = \vec{j}_s \wedge \vec{n}_{1 \rightarrow 2} \quad \text{soit encore} \quad \vec{j}_s \wedge \vec{n}_{1 \rightarrow 2} = \vec{H}_i$$

Ce qui donne la définition du courant électrique surfacique en fonction du champ magnétique incident :

$$\vec{j}_s = \vec{n}_{1 \rightarrow 2} \wedge \vec{H}_i$$

Grâce à la connaissance de ces courants électriques surfaciques équivalents, on connaît une partie des sources du champ électromagnétique.

Un raisonnement analogue conduit à des résultats semblables sur le champ électrique. La difficulté, dans ce cas est de justifier de la discontinuité du champ électrique dans le problème équivalent. Cela est rendu possible par l'introduction de courants magnétiques de surface $\vec{M}_s$ à la traversée de la surface fictive. La figure 3.17 donne le schéma de ce problème équivalent.

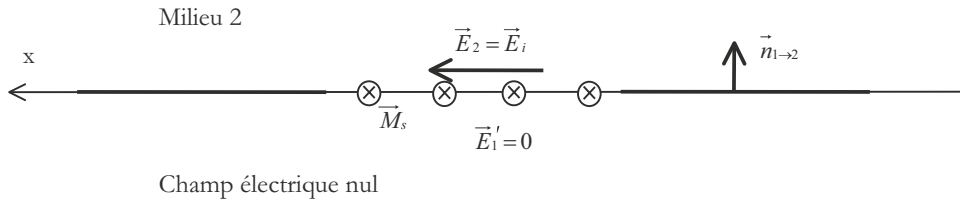


Figure 3.17 – Problème équivalent pour le champ électrique.

Selon la condition de passage sur le champ électrique tangentiel, on vérifie :

$$\vec{E}_2 - \vec{E}_1' = -\vec{M}_s \wedge \vec{n}_{1 \rightarrow 2}$$

D'où la valeur des courants surfaciques équivalents :

$$\vec{M}_s = -\vec{n}_{1 \rightarrow 2} \wedge \vec{E}_i$$

Le problème équivalent est complètement défini, puisque la condition d'unicité est bien vérifiée, grâce à l'identité entre les composantes tangentielles du champ électromagnétique entre les deux problèmes. Cette équivalence n'est que partielle, puisqu'elle ne se vérifie que dans le demi-espace supérieur.

On a donc remplacé le problème initial par le problème consistant en le calcul du champ créé par les sources de courants surfaciques électriques et magnétiques $\vec{j}_s$ et $\vec{M}_s$, dans le demi espace supérieur.

De façon plus générale, le principe d'équivalence repose sur l'idée qu'un problème complexe est remplacé par un problème qu'on sait résoudre dans une partie de l'espace. Il y a de ce fait une perte d'information, mais la possibilité de résoudre les problèmes par morceaux.

Enfin, précisons que le principe d'équivalence revêt des formes différentes, selon les problèmes posés qui peuvent relever, soit de l'étude de la diffraction, soit de calculs portant sur les antennes. Dans le cas traité en exemple, le champ incident est supposé connu. Le principe est explicité par rapport au champ incident.

3.2.2 Dualité des grandeurs électriques et magnétiques

Les solutions des équations de Maxwell, dans le cas de sources de courant électrique et de courant magnétique, utilisent le principe de dualité entre les grandeurs électriques et magnétiques.

□ Solution avec sources électriques

Le champ électromagnétique $(\vec{E}_1, \vec{H}_1)$ créé par les sources électriques a été calculé au paragraphe précédent. Rappelons qu'il est solution des équations de Maxwell avec sources électriques qu'on écrira sous la forme :

$$\text{rot}(\vec{E}_1) = -\frac{\partial \vec{B}_1}{\partial t} \quad [3.25]$$

$$\text{div}(\vec{B}_1) = 0 \quad [3.26]$$

$$\text{rot}(\vec{H}_1) = \vec{j} + \frac{\partial \vec{D}_1}{\partial t} \quad [3.27]$$

$$\text{div}(\vec{D}_1) = \rho \quad [3.28]$$

Le champ électromagnétique est dérivé du potentiel vecteur, exprimé en fonction des sources de courant en volume sous la forme :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} \iiint_{\text{sources}} \frac{\vec{j}(\vec{r}', t')}{\|\vec{R}\|} d\vec{r}'$$

Dans le cas du calcul de la diffraction par une ouverture, le terme source correspondant à la densité surfacique de courant $\vec{j}_s$ donne la valeur du potentiel :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} \iint_{\text{ouverture}} \frac{\vec{j}_s(\vec{r}', t')}{\|\vec{R}\|} dx' dy'$$

Pour une source en régime harmonique définie par :

$$\vec{j}_s(\vec{r}', t') = \vec{j}_s(\vec{r}') e^{j\omega t'}$$

le potentiel prend la forme :

$$\vec{A}(\vec{r}, t) = \frac{\mu_0}{4\pi} \iint_{\text{ouverture}} \frac{\vec{j}_s(\vec{r}') e^{-jk\|\vec{R}\|}}{\|\vec{R}\|} dx' dy'$$

Selon la méthode exposée au paragraphe 3.1.2, les solutions en champ électrique et en champ magnétique sont données en champ lointain par :

$$\vec{E}_1(M) = \frac{jk}{4\pi} Z \psi(r) \iint_{\text{ouverture}} \left[\vec{j}_s(\vec{r}') \wedge \vec{u} \right] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} dx' dy'$$

$$\vec{H}_1(M) = \frac{jk}{4\pi} \psi(r) \iint_{\text{ouverture}} \left(\vec{j}_s(\vec{r}') \wedge \vec{u} \right) e^{j\vec{k} \cdot \vec{r}'} dx' dy'$$

On rappelle la notation :

$$R = \|\vec{R}\|$$

□ **Solution avec sources magnétiques**

Le principe d'équivalence nous a conduit à introduire des sources de courant magnétique. Si ce sont les seules sources, les équations de Maxwell s'écrivent :

$$\text{rot}(\vec{E}_2) = -\frac{\partial \vec{B}_2}{\partial t} - \vec{M} \quad [3.29]$$

$$\text{div}(\vec{B}_2) = -\tau \quad [3.30]$$

$$\text{rot}(\vec{H}_2) = \frac{\partial \vec{D}_2}{\partial t} \quad [3.31]$$

$$\text{div}(\vec{D}_2) = 0 \quad [3.32]$$

On constate que ces équations sont similaires aux équations de Maxwell avec sources électriques [3.25] à [3.28]. Dans les équations [3.29] à [3.32], le fait de faire disparaître les sources électriques et de placer des sources magnétiques, fait jouer un rôle aux grandeurs électriques, analogue à celui que jouent les grandeurs magnétiques dans les équations du premier groupe, au signe près. C'est le principe de la dualité entre grandeurs électriques et magnétiques.

Cette remarque va nous permettre de ne pas refaire tous les calculs pour les champs.

En effet les transformations :

$$\vec{H}_2 = \frac{1}{jZ} \vec{E}_1 \quad \text{et} \quad \vec{E}_2 = jZ \vec{H}_1$$

conduisent à une analogie complète entre les deux systèmes d'équations, tout en respectant l'homogénéité des grandeurs, à condition de poser :

$$\vec{M} = -jZ \vec{j}$$

Pour des courants surfaciques, cette expression s'écrit :

$$\vec{M}_s = -jZ \vec{j}_s$$

Les solutions sont donc données par :

$$\begin{aligned} \vec{E}_2(M) &= jZ \vec{H}_1(M) = -\frac{jk}{4\pi} \Psi(r) \iint_{\text{ouverture}} \left(\vec{M}_s(\vec{r}') \wedge \vec{u} \right) e^{j\vec{k} \cdot \vec{r}'} dx' dy' \\ \text{et} \quad \vec{H}_2(M) &= \frac{1}{jZ} \vec{E}_1(M) = -\frac{k}{4\pi} \Psi(r) \iint_{\text{ouverture}} \left[\frac{\vec{M}_s(\vec{r}')}{jZ} \wedge \vec{u} \right] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} dx' dy' \end{aligned}$$

□ **Solution globale**

La solution totale résulte du théorème de superposition : la superposition des deux états, l'un avec source électrique, l'autre avec source magnétique, donne la solution globale. C'est, en fait, la linéarité des équations de Maxwell qui est traduite dans ce théorème. Le champ électrique lointain s'exprime alors par :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \Psi(r) \iint_{\text{ouverture}} \left[\left(\vec{j}_s(\vec{r}') \wedge \vec{u} \right) \wedge \vec{u} - \frac{\vec{M}_s(\vec{r}')}{Z} \wedge \vec{u} \right] e^{j\vec{k} \cdot \vec{r}'} dx' dy' \quad [3.33]$$

et pour le champ magnétique :

$$\vec{H}(M) = \frac{jk}{4\pi}\psi(r) \iint_{\text{ouverture}} \left[\vec{j}_S(\vec{r}') \wedge \vec{u} + \left(\frac{\vec{M}_S(\vec{r}')}{Z} \wedge \vec{u} \right) \wedge \vec{u} \right] e^{j\vec{k} \cdot \vec{r}'} dx' dy' \quad [3.34]$$

□ Application au rayonnement d'une ouverture plane

Soit $(\vec{E}_i, \vec{H}_i)$, le champ électromagnétique de l'onde incidente. Les courants surfaciques s'expriment, en fonction des vecteurs unitaires $\vec{u}_x$ et $\vec{u}_y$, par :

$$\vec{j}_S = -H_i \vec{u}_x \quad \text{et} \quad \vec{M}_S = -ZH_i \vec{u}_y$$

La valeur du champ électrique est donc :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) \iint_{\text{ouverture}} [(-H_i \vec{u}_x \wedge \vec{u}) + H_i \vec{u}_y] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} dx' dy'$$

Soit en introduisant les coordonnées sphériques du point d'observation M :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) H_i \iint_{\text{ouverture}} [(\sin \varphi \vec{u}_\theta + \cos \theta \cos \varphi \vec{u}_\varphi) - \vec{u}_y] \wedge \vec{u} e^{j\vec{k} \cdot \vec{r}'} dx' dy'$$

Le résultat s'exprime finalement comme :

$$\vec{E}(M) = \frac{jk}{2\pi} Z \psi(r) H_i \left(\frac{1 + \cos \theta}{2} \right) (\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) \iint_{\text{ouverture}} e^{j\vec{k} \cdot \vec{r}'} dx' dy' \quad [3.35]$$

Dans cette expression, la partie vectorielle sort de l'intégration car elle ne dépend que du point d'observation et donne la polarisation de l'onde. L'intégrale dépend de la forme de la surface.

Le terme $\frac{(1 + \cos \theta)}{2}$ est appelé le terme d'obliquité.

3.2.3 Diffraction par une ouverture rectangulaire

Dans le cas d'une ouverture rectangulaire de dimensions a selon Ox et b selon Oy (figure 3.18), l'intégrale prend la forme :

$$I = \iint_{\text{ouverture}} e^{j\vec{k} \cdot \vec{r}'} dx' dy' = \int_{-a}^a e^{-jkx' \sin \theta \cos \varphi} dx' \int_{-b}^b e^{-jky' \sin \theta \sin \varphi} dy'$$

Le résultat de l'intégration donne :

$$I = 4ab \frac{\sin(ka \sin \theta \cos \varphi)}{ka \sin \theta \cos \varphi} \frac{\sin(kb \sin \theta \sin \varphi)}{kb \sin \theta \sin \varphi}$$

Cette variation du champ électrique est caractéristique de la diffraction par une ouverture rectangulaire.

Le champ total s'exprime sous la forme :

$$\vec{E}(M) = \frac{jk}{4\pi} Z \psi(r) H_i (4ab) (1 + \cos \theta) (\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) \frac{\sin(ka \sin \theta \cos \varphi)}{ka \sin \theta \cos \varphi} \frac{\sin(kb \sin \theta \sin \varphi)}{kb \sin \theta \sin \varphi}$$

Posons :

$$E = jk H_i Z \frac{e^{-jkr}}{2\pi r} (4ab)$$

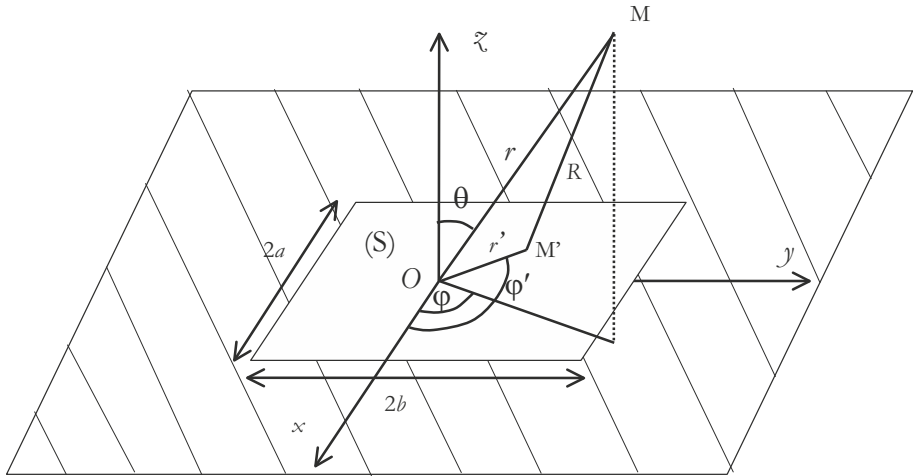


Figure 3.18 – Diffraction par une ouverture rectangulaire.

L'expression du champ est alors :

$$\vec{E}(M) = E \frac{1 + \cos \theta}{2} (\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) \frac{\sin(ka \sin \theta \cos \varphi)}{ka \sin \theta \cos \varphi} \frac{\sin(kb \sin \theta \sin \varphi)}{kb \sin \theta \sin \varphi} \quad [3.36]$$

La puissance rayonnée est proportionnelle au module au carré du champ électrique. Le module du champ est représenté dans l'espace sur la figure 3.19 pour des dimensions de l'ouverture $2a = 2\lambda$ et $2b = 3\lambda$. On remarque que ce diagramme n'a pas de symétrie axiale.

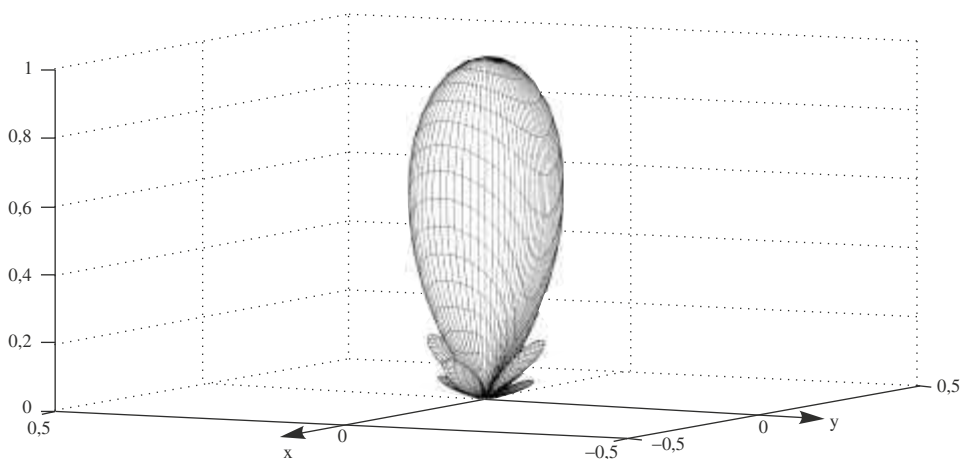


Figure 3.19 – Module du champ électrique diffracté par une ouverture rectangulaire de dimensions $2a = 2\lambda$ et $2b = 3\lambda$, illuminée par une onde plane.

La répartition de la puissance dans le plan (Ox, Oz) , plan E, est représentée en fonction du sinus de l'angle entre l'axe Oz et le rayon r sur la figure 3.20, alors que la répartition dans le plan H est donnée sur la figure 3.21. On constate sur ces figures l'influence de la taille de l'ouverture.

Le champ électrique dans le plan E ($\varphi = 0$) prend la forme :

$$\vec{E}(M) = E \frac{1 + \cos \theta}{2} \frac{\sin(ka \sin \theta)}{ka \sin \theta} \vec{u}_\theta$$

Le premier zéro de cette fonction est obtenu pour :

$$\sin \theta = \frac{\lambda}{2a}$$

C'est bien ce qu'on observe sur la figure 3.19 où $a = \lambda$.

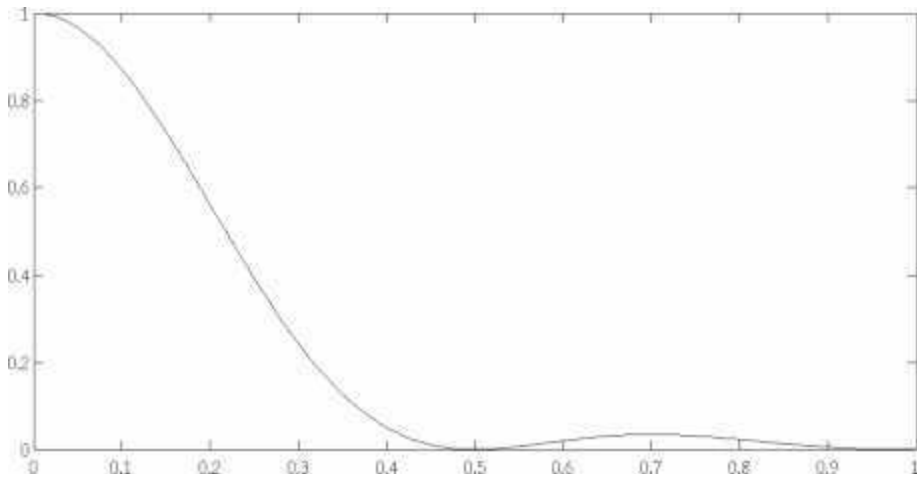


Figure 3.20 – Puissance normalisée rayonnée par une ouverture rectangulaire, dans le plan E en fonction de la variation du sinus θ ($a = \lambda$).

Le champ électrique dans le plan H ($\varphi = \pi/2$) prend la forme :

$$\vec{E}(M) = -E \frac{1 + \cos \theta}{2} \frac{\sin(kb \sin \theta)}{kb \sin \theta} \vec{u}_\varphi$$

Le premier zéro de cette fonction apparaît pour :

$$\sin \theta = \frac{\lambda}{2b}$$

Sur la figure 3.21, le zéro apparaît pour $\sin \theta = \frac{1}{3}$

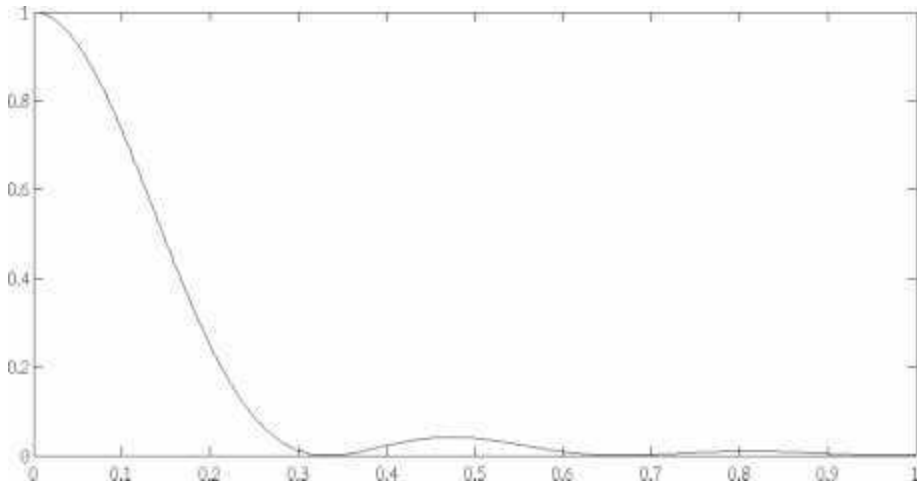


Figure 3.21 – Puissance normalisée rayonnée par une ouverture rectangulaire, dans le plan H en fonction de la variation du sinus θ ($2b = 3\lambda$).

3.2.4 Diffraction par une ouverture circulaire

Considérons l'ouverture circulaire diffractant dans le demi-espace $z > 0$ (figure 3.22).

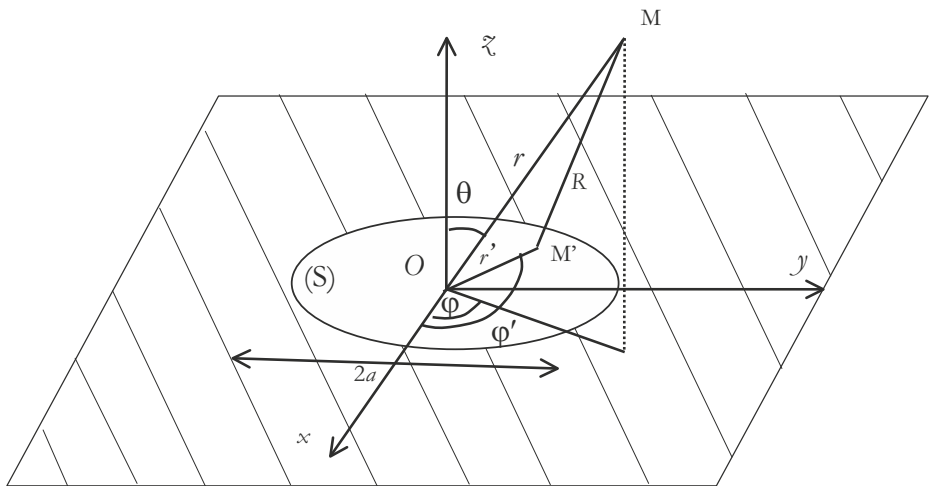


Figure 3.22 – Diffraction par une ouverture circulaire.

Reprenons l'expression [3.35], donnant l'expression du champ à grande distance. En raison de la forme de l'ouverture, un point P' de l'ouverture est repéré par r' et son azimut φ' . L'intégrale sur l'ouverture devient :

$$I = \int_0^a \int_0^{2\pi} e^{jkr' \sin \theta \cos(\varphi - \varphi')} r' dr' d\varphi'$$

On rappelle la propriété de la fonction de Bessel d'ordre 0 :

$$\int_0^{2\pi} e^{iu \cos \theta} d\theta = 2\pi J_0(u)$$

On en déduit :

$$I = 2\pi \int_0^a J_0(kr' \sin \theta) r' dr'$$

Les fonctions de Bessel d'ordre 0 et d'ordre 1 sont liées par la relation :

$$\int_0^d r J_0(r) dr = d J_1(d)$$

L'intégrale I se calcule :

$$I = 2\pi a^2 \frac{J_1(ka \sin \theta)}{ka \sin \theta}$$

L'expression globale du champ électrique se met sous la forme :

$$\vec{E}(M) = \frac{jk}{2} Z \psi(r) H_i(1 + \cos \theta) (\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) a^2 \frac{J_1(ka \sin \theta)}{ka \sin \theta}$$

La norme du champ électrique à grande distance présente une symétrie axiale en raison de la forme de l'ouverture (figure 3.23)

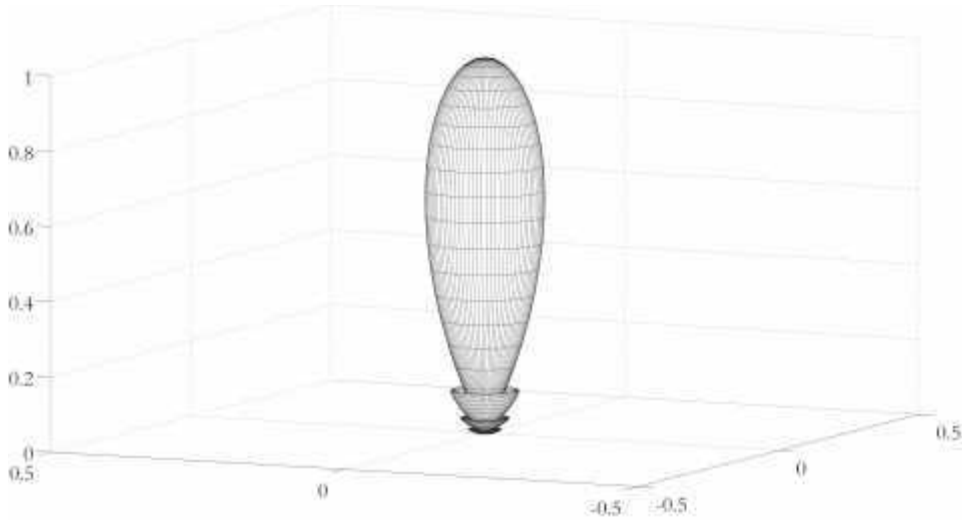


Figure 3.23 – Module du champ électrique normalisé rayonné par une ouverture circulaire à grande distance en trois dimensions, pour une ouverture de rayon 2λ .

La répartition de puissance dépend uniquement de l'angle θ (figure 3.24).

Le premier zéro de cette répartition de champ est obtenu pour le premier zéro de la fonction de Bessel, qui correspond à une valeur de la variable égale, environ, à 3,83, soit :

$$\sin \theta = 0,61 \frac{\lambda}{a}$$

C'est un résultat bien connu en optique qui caractérise la tache d'Airy.

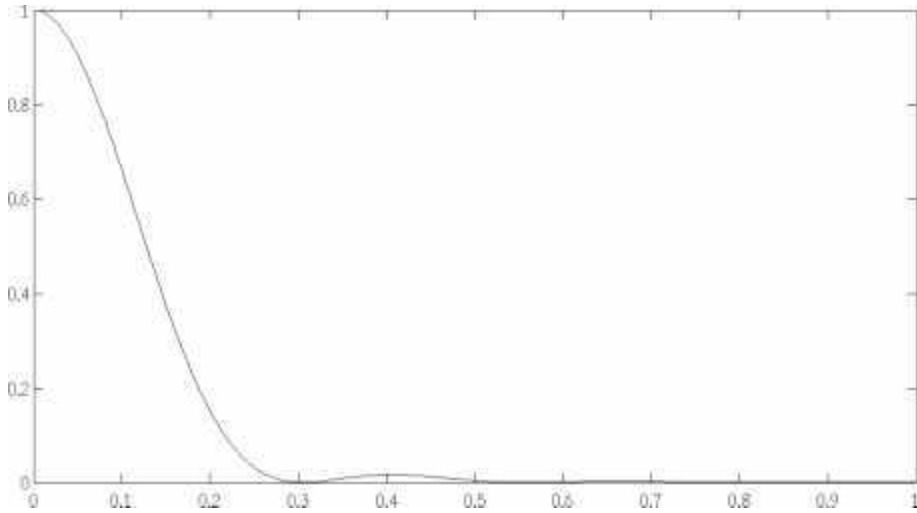


Figure 3.24 – Puissance normalisée rayonnée par une ouverture circulaire en 2D en fonction de la variation du sinus θ ($a = 2\lambda$).

3.3 Rayonnement des ouvertures planes

Le rayonnement de certaines antennes est calculé à partir de la connaissance du champ électrique sur une région donnée, appelée ouverture de l'antenne. La méthode de calcul du champ rayonné, la plus générale, repose sur l'utilisation de la transformée de Fourier. Au paragraphe 3.2, nous avons vu une méthode de calcul du champ diffracté par une ouverture plane, illuminée par une onde plane. Ce paragraphe a développé une méthode approximative qui donne de bons résultats si l'ouverture a une taille supérieure à quelques longueurs d'onde.

La méthode qui va être présentée est appelée la formulation vectorielle, car elle ne fait appel à aucune approximation sur le champ diffracté et considère les propriétés du champ d'un point de vue vectoriel.

3.3.1 Calcul général du champ rayonné

Le problème posé consiste à calculer le champ rayonné pour tous les points M du demi-espace $z > 0$, connaissant le champ sur l'ouverture (S) de l'antenne (figure 3.25). Par exemple, le champ créé par une antenne cornet est calculable à partir de la connaissance du champ sur le plan situé juste à l'embouchure du cornet.

On utilisera les coordonnées sphériques pour exprimer le champ rayonné. La surface rayonnante est supposée plane.

■ Rappels des propriétés de la transformée de Fourier

On rappelle ici les définitions des transformées de Fourier et certaines de leurs propriétés utilisées dans les calculs du champ.

□ Définition des transformées de Fourier à une dimension

La transformée de Fourier ($T.F.$) d'une fonction de carré sommable est définie par :

$$T.F. [u(x)] = \int_{-\infty}^{+\infty} u(x) e^{jk_x x} dx$$

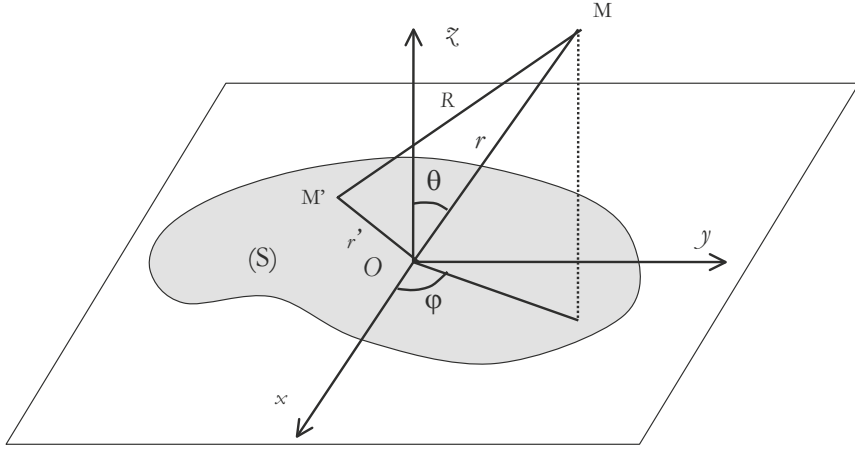


Figure 3.25 – Représentation de l'ouverture rayonnante dans l'espace.

On notera :

$$u(k_x) = \int_{-\infty}^{+\infty} u(x) e^{jk_x x} dx$$

La transformée inverse permet d'écrire :

$$u(x) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} u(k_x) e^{-jk_x x} dk_x$$

□ Définition des transformées de Fourier à deux dimensions

Les propriétés de la transformée de Fourier se généralisent à deux dimensions pour une fonction à deux variables, x et y :

$$T.F. [u(x, y)] = \iint u(x, y) e^{j(k_x x + k_y y)} dx dy$$

De même qu'à une dimension, on notera :

$$u(k_x, k_y) = \iint u(x, y) e^{j(k_x x + k_y y)} dx dy \quad [3.37]$$

Par transformée inverse on obtient :

$$u(x, y) = \frac{1}{4\pi^2} \iint u(k_x, k_y) e^{-j(k_x x + k_y y)} dk_x dk_y \quad [3.38]$$

Les bornes d'intégration vont de $-\infty$ à $+\infty$ pour chacune des variables.

La transformée de Fourier de la dérivée première d'une fonction conduit à :

$$T.F. \left[\frac{\partial u(x, y)}{\partial x} \right] = -jk_x T.F. [u(x, y)]$$

Pour la dérivée seconde, on obtient :

$$T.F. \left[\frac{\partial^2 u(x, y)}{\partial x^2} \right] = (-jk_x)^2 T.F. [u(x, y)]$$

Ces propriétés font de la transformée de Fourier un outil très puissant qui permet de s'affranchir d'une ou plusieurs dérivations qui s'expriment alors par l'intermédiaire des variables k_x et k_y , dans l'espace de Fourier.

■ Expression générale du champ électrique

On rappelle l'équation d'Helmholtz dans l'espace à trois dimensions qui traduit la propagation des ondes :

$$\Delta \vec{E}(x, y, z) + k^2 \vec{E}(x, y, z) = 0 \quad \text{avec} \quad k = \frac{\omega}{c}$$

L'espace étant vide de charge, l'une des équations de Maxwell permet d'écrire :

$$\text{div } \vec{E}(x, y, z) = 0$$

Afin de calculer le champ électrique, on lui applique une transformée de Fourier à deux dimensions. Posons :

$$\vec{E}(k_x, k_y, z) = T.F. \left[\vec{E}(x, y, z) \right]$$

La transformée de Fourier de l'équation d'Helmholtz conduit à :

$$T.F. \left[\Delta \vec{E}(x, y, z) + k^2 \vec{E}(x, y, z) \right] = 0$$

En utilisant les propriétés des dérivées, l'équation d'Helmholtz s'écrit dans le domaine de Fourier :

$$\left(-k_x^2 - k_y^2 + \frac{\partial^2}{\partial z^2} \right) \vec{E}(k_x, k_y, z) + k^2 \vec{E}(k_x, k_y, z) = 0$$

Posons :

$$k_z^2 = k^2 - k_x^2 - k_y^2 \quad [3.39]$$

On est alors conduit à :

$$\frac{\partial^2}{\partial z^2} \vec{E}(k_x, k_y, z) + k_z^2 \vec{E}(k_x, k_y, z) = 0$$

Les solutions de cette équation différentielle s'expriment en fonction d'exponentielles complexes sous la forme :

$$\vec{E}(k_x, k_y, z) = \vec{f}(k_x, k_y) e^{\pm j k_z z} \quad [3.40]$$

Le signe moins de l'exponentielle correspond à une onde qui se propage dans le sens des z positifs, c'est-à-dire qui part de l'ouverture. On parle alors d'onde sortante. Dans le cas contraire, le signe plus correspond à des ondes se dirigeant vers l'ouverture. On parle d'ondes entrantes. Dans la suite, on se limitera au cas d'ondes sortantes. L'autre cas se traite en utilisant la même méthode.

L'équation de la divergence se traite de la même façon. On en prend la transformée de Fourier à deux dimensions. Cela conduit à :

$$k_x E_x(k_x, k_y, z) + k_y E_y(k_x, k_y, z) + j \frac{\partial}{\partial z} E_z(k_x, k_y, z) = 0$$

En tenant compte de l'expression [3.40] du champ transformé, on obtient :

$$k_x f_x + k_y f_y + k_z f_z = 0$$

Les composantes de la fonction $\vec{f}(k_x, k_y)$ sont notées (f_x, f_y, f_z) .

Soit

$$\vec{f} \cdot \vec{k} = 0 \quad [3.41]$$

La divergence nulle du champ impose aux deux vecteurs d'être perpendiculaires.

L'équation [3.41] impose une relation entre les composantes (f_x, f_y, f_z) qui ne sont donc pas indépendantes.

Pour obtenir le champ dans l'espace réel, on applique la transformée inverse :

$$\vec{E}(x, y, z) = \frac{1}{4\pi^2} \iint \vec{E}(k_x, k_y, z) e^{-j(k_x x + k_y y)} dk_x dk_y$$

Utilisant la solution [3.40] imposée par l'équation d'Helmholtz :

$$\vec{E}(x, y, z) = \frac{1}{4\pi^2} \iint \vec{f}(k_x, k_y) e^{-j(k_x x + k_y y + k_z z)} dk_x dk_y$$

Soit encore :

$$\vec{E}(x, y, z) = \frac{1}{4\pi^2} \iint \vec{f}(k_x, k_y) e^{-j\vec{k} \cdot \vec{r}} dk_x dk_y \quad [3.42]$$

Cette expression montre que la connaissance de $\vec{f}(k_x, k_y)$ permet de déterminer le champ en tout point de l'espace $z > 0$.

La fonction :

$$\vec{f}(k_x, k_y) e^{-j\vec{k} \cdot \vec{r}}$$

représente le champ d'une onde plane de vecteur d'onde $\vec{k}$. L'équation [3.42] exprime alors la décomposition du champ sous la forme d'une infinité continue d'ondes planes. L'amplitude spectrale de chacune des ondes est donnée par

$$\vec{f}(k_x, k_y)$$

C'est ce qu'on appelle le spectre continu d'ondes planes.

■ Propriétés du spectre d'ondes planes

Si $k_x^2 + k_y^2 > k^2$, alors $k_z^2 < 0$ selon [3.39]. Dans ce cas, les ondes sont évanescentes. Elles ne se propagent pas et restent localisées au voisinage de l'ouverture. Les termes correspondant à ce cas interviennent dans le calcul du champ et se retrouvent dans le champ très proche de l'ouverture. Si $k_x^2 + k_y^2 < k^2$, alors $k_z^2 > 0$. Dans ce cas, les ondes se propagent. Ce sont les ondes qu'on retrouve loin de l'ouverture.

Ces deux cas sont illustrés dans la figure 3.26.

La démonstration suivante permet le calcul du champ électrique à partir de la connaissance du champ tangentiel sur l'ouverture, noté :

$$\vec{E}_o(x, y)$$

Ce champ s'exprime, en utilisant la décomposition de Fourier, par :

$$\vec{E}_o(x, y) = \vec{E}_T(x, y, 0) = \frac{1}{4\pi^2} \iint \vec{f}_T(k_x, k_y) e^{-j\vec{k} \cdot \vec{r}} dk_x dk_y$$

$\vec{f}_T(k_x, k_y)$ est une fonction vectorielle à deux dimensions, dont les composantes sont selon les axes Ox et Oy . On l'obtient par transformée de Fourier inverse :

$$\vec{f}_T(k_x, k_y) = \iint_S \vec{E}_o(x, y) e^{j(k_x x + k_y y)} dx dy$$

C'est donc la répartition du champ tangentiel dans l'ouverture qui détermine la fonction $\vec{f}_T(k_x, k_y)$ représentant le spectre du champ tangentiel dans l'ouverture, selon le théorème d'unicité démontré en 2.4.2.

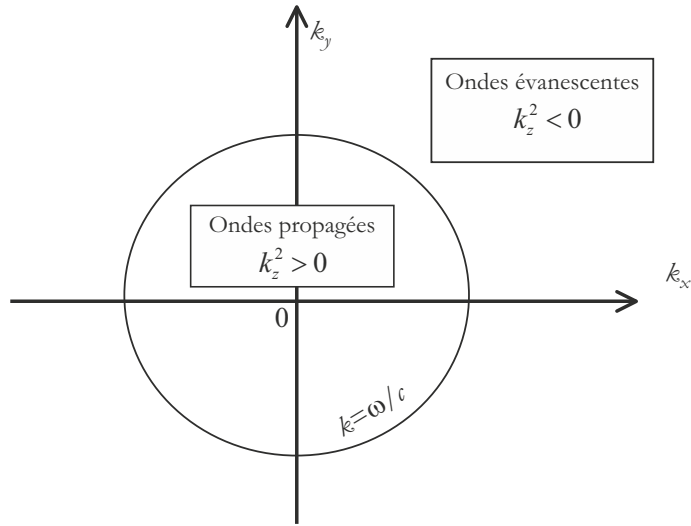


Figure 3.26 – Zones définissant les ondes propagées et évanescentes dans le diagramme des vecteurs d'onde.

On doit vérifier la relation :

$$f_z = \frac{-\vec{k}_T \cdot \vec{f}_T}{k_z} = \frac{-k_x f_x - k_y f_y}{k_z}$$

Donc la fonction $\vec{f}(k_x, k_y)$ est complètement déterminée par la connaissance de ses composantes tangentielle et de sa composante longitudinale. On déduit ainsi le champ pour tout point du demi-espace $z > 0$ par l'équation [3.42]. Finalement, c'est bien la répartition tangentielle du champ dans l'ouverture qui détermine la répartition spectrale et donc le champ en tout point, pour $z > 0$.

Résumons la méthode. Le champ tangentiel dans l'ouverture est décomposé selon ses composantes spectrales dans l'espace à deux dimensions des vecteurs de propagation transverses. À partir de ce spectre, le champ est calculable en tous les points de l'espace, par « propagation des composantes spectrales ». Seul le champ tangentiel sur l'ouverture est utile pour le calcul du champ dans le demi-espace $z > 0$. C'est lui qui porte toute l'information.

3.3.2 Calcul du champ à grande distance

D'après ce que nous venons de démontrer, le champ est calculé selon l'équation :

$$\vec{E}(\vec{r}) = \frac{1}{4\pi^2} \iint \vec{f}(k_x, k_y) e^{-j\vec{k} \cdot \vec{r}} dk_x dk_y$$

Dans cette intégrale, les variables d'intégration sont k_x et k_y . Le vecteur $\vec{r}$ est un paramètre d'intégration. À grande distance l'argument de l'exponentielle est grand et l'exponentielle présente des variations très rapides.

Lorsque la phase de l'exponentielle est stationnaire, c'est-à-dire qu'elle présente un extremum, l'intégrande oscille plus lentement. C'est donc au voisinage de ce point que les valeurs de l'intégrande contribuent à la valeur de l'intégrale. Ailleurs qu'en ce point, les parties positives et négatives de l'intégrande se compensent. Sur la figure 3.27 la partie réelle de l'exponentielle intervenant dans le calcul du champ électrique est représentée en fonction de k_x , pour une longueur

d'onde de 1 m et en un point situé à 100 m de la source. La phase présente un maximum pour $k_x = 0$ dans ce cas, c'est-à-dire pour $\theta = 0$.

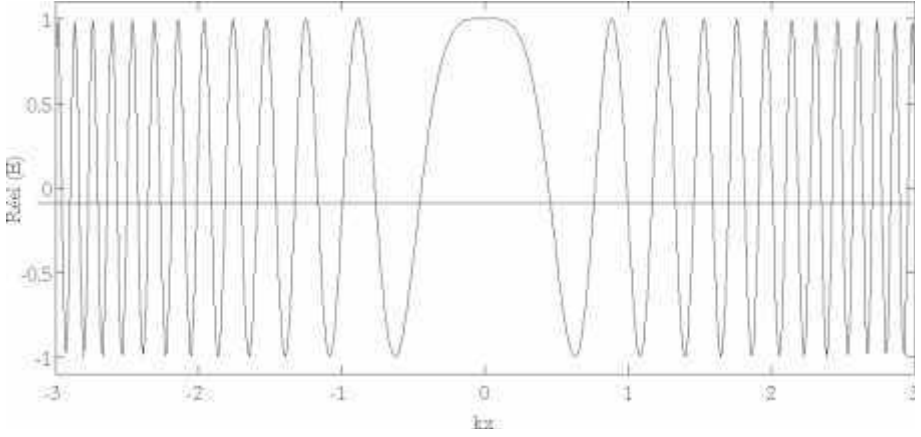


Figure 3.27 – Partie réelle normalisée de l'exponentielle complexe à 100 m de la source en fonction de k_x ($\lambda = 1$ m).

Le point de phase stationnaire est obtenu pour :

$$\frac{\partial(\vec{k} \cdot \vec{r})}{\partial k_x} = 0 \quad \text{et} \quad \frac{\partial(\vec{k} \cdot \vec{r})}{\partial k_y} = 0$$

En ce point, on posera :

$$k_x = k_1 \quad \text{et} \quad k_y = k_2$$

La fonction $\vec{f}_T(k_x, k_y)$ varie lentement par rapport à l'exponentielle. On évaluera sa valeur en k_1 et k_2 .

Étant donné que :

$$x = r \sin \theta \cos \phi \quad y = r \sin \theta \sin \phi \quad z = r \cos \theta$$

Le calcul de $\vec{k} \cdot \vec{r}$ conduit à :

$$\vec{k} \cdot \vec{r} = r(k_x \sin \theta \cos \phi + k_y \sin \theta \sin \phi + \sqrt{k^2 - k_x^2 - k_y^2} \cos \theta)$$

Le point de phase stationnaire est donc obtenu pour :

$$k_x = k_1 = k \sin \theta \cos \phi \quad \text{et} \quad k_y = k_2 = k \sin \theta \sin \phi$$

Le développement de Taylor au voisinage du point de phase stationnaire permet le calcul au deuxième ordre :

$$\vec{k} \cdot \vec{r} = k \cdot r + \frac{1}{2} \frac{\partial^2(\vec{k} \cdot \vec{r})}{\partial^2 k_x} (k_x - k_1)^2 + \frac{1}{2} \frac{\partial^2(\vec{k} \cdot \vec{r})}{\partial^2 k_y} (k_y - k_2)^2 + \frac{\partial^2(\vec{k} \cdot \vec{r})}{\partial k_x \partial k_y} (k_x - k_1)(k_y - k_2)$$

Posons :

$$u = k_x - k_1 \quad \text{et} \quad v = k_y - k_2$$

Soit :

$$\vec{k} \cdot \vec{r} = k \cdot r - (Au^2 + Bv^2 + Cuv)$$

Puisque la fonction $\vec{f}(k_x, k_y)$ varie lentement, l'expression du champ électrique est donnée par l'équation :

$$\vec{E}(\vec{r}) = \frac{e^{-j\vec{k} \cdot \vec{r}}}{4\pi^2} \vec{f}(k \sin \theta \cos \phi, k \sin \theta \sin \phi) \iint e^{j(Au^2 + Bv^2 + Cuv)} du dv$$

Pour calculer l'intégrale, on utilise la relation :

$$\int_{-\infty}^{+\infty} e^{j\gamma(x-x_0)^2} dx = \sqrt{\frac{\pi}{\gamma}} e^{j\pi/4}$$

Ce qui conduit à :

$$\iint e^{j(Au^2 + Bv^2 + Cuv)} du dv = 2\pi j \frac{k}{r} \cos \theta$$

Finalement le champ s'exprime par :

$$\vec{E}(\vec{r}) = \frac{jk}{2\pi r} \cos \theta e^{-jk.r} \vec{f}(k \sin \theta \cos \phi, k \sin \theta \sin \phi)$$

Dans la suite, le champ électrique sera exprimé en fonction des composantes tangentielles du spectre d'ondes. Pour cela on rappelle :

$$\vec{f}_T = f_x \vec{u}_x + f_y \vec{u}_y \quad \text{et} \quad f_z = -\frac{k f_x \sin \theta \cos \phi + k f_y \sin \theta \sin \phi}{k \cos \theta}$$

Sachant que :

$$\vec{u}_x = \cos \phi \sin \theta \vec{u}_r + \cos \phi \cos \theta \vec{u}_\theta - \sin \phi \vec{u}_\phi$$

$$\vec{u}_y = \sin \phi \sin \theta \vec{u}_r + \sin \phi \cos \theta \vec{u}_\theta + \cos \phi \vec{u}_\phi$$

$$\vec{u}_z = \cos \theta \vec{u}_r - \sin \theta \vec{u}_\theta$$

L'expression du champ électrique se met alors sous la forme :

$$\vec{E}(\vec{r}) = \frac{jke^{-jk.r}}{2\pi r} [(f_x \cos \phi + f_y \sin \phi) \vec{u}_\theta + (f_y \cos \phi - f_x \sin \phi) \cos \theta \vec{u}_\phi] \quad [3.43]$$

Cette équation montre que le champ électrique est transversal et qu'il a deux composantes selon $\vec{u}_\theta$ et $\vec{u}_\phi$. C'est ainsi qu'on détermine la polarisation du champ. Le champ à grande distance a bien une structure d'onde localement plane.

3.3.3 Calcul du champ rayonné à grande distance par une ouverture rectangulaire

Le champ rayonné par une ouverture rectangulaire est calculé ici par la méthode vectorielle, présentée dans le paragraphe précédent. Le champ est supposé uniforme dans l'ouverture de dimensions $2a$ sur $2b$ et polarisé rectilignement selon x (figure 3.28).

Le champ dans l'ouverture est noté :

$$\vec{E}_o = E_0 \vec{u}_x$$

La composante tangentielle du spectre d'onde s'écrit :

$$\vec{f}_T = E_0 \vec{u}_x \int_{-a}^{+a} \int_{-b}^{+b} e^{j(k_x x + k_y y)} dx dy$$

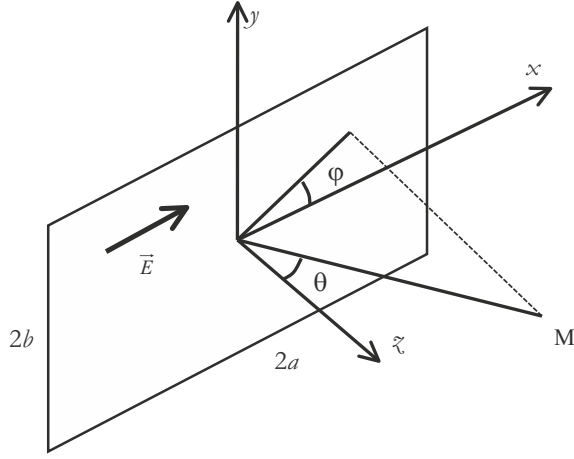


Figure 3.28 – Ouverture plane rectangulaire.

Le spectre d'onde contenu dans le champ tangentiel devient :

$$\vec{f}_T = 4abE_0 \vec{u}_x \frac{\sin(k_x a)}{k_x a} \frac{\sin(k_y b)}{k_y b}$$

$$\text{Avec } k_x = k_1 = k \sin \theta \cos \phi \quad k_y = k_2 = k \sin \theta \sin \phi$$

D'où l'expression du champ à grande distance :

$$\vec{E}(\vec{r}) = \frac{jke^{-jk.r}}{2\pi r} 4abE_0 \frac{\sin(ka \sin \theta \cos \phi)}{ka \sin \theta \cos \phi} \frac{\sin(kb \sin \theta \sin \phi)}{kb \sin \theta \sin \phi} [\cos \phi \vec{u}_\theta - \sin \phi \cos \theta \vec{u}_\phi] \quad [3.44]$$

Cette expression est très générale. On retrouve un résultat comparable à [3.36] de la formulation scalaire lorsque θ est proche de zéro. En effet, si θ est petit :

$$\frac{1 + \cos \theta}{2} \approx 1$$

Le facteur d'obliquité apparaît dans la formulation scalaire. Rappelons que celle-ci n'est valable que si les dimensions de l'ouverture sont supérieures à la longueur d'onde. Or dans ce cas, le premier zéro de la fonction de diffraction apparaît pour des angles petits et l'approximation est donc bien vérifiée.

3.4 Différentes zones de rayonnement

Le calcul du champ diffracté par une ouverture a été présenté au paragraphe 3.2, en utilisant une méthode très générale, valable quelle que soit la taille de l'ouverture diffractante. C'est la formulation vectorielle. Le résultat est valable en tout point de l'espace. C'est une méthode lourde qui peut être remplacée, sous certaines hypothèses, par une formulation scalaire consistant à traiter indépendamment toutes les composantes du champ électromagnétique.

Lorsque la taille de l'ouverture est plus grande que quelques longueurs d'ondes, on montre qu'on peut utiliser cette formulation scalaire simplifiée. Nous allons en exposer les grandes lignes qui nous serviront à définir différentes zones de rayonnement pour le champ diffracté. Cette théorie est basée sur l'intégrale de Fresnel-Kirchhoff. Elle est utilisée pour déterminer l'intensité diffractée par une ouverture en optique.

Avec les notations de la figure 3.29, cette intégrale permet de calculer les composantes du champ par l'expression :

$$E(r, \theta, \varphi) = \frac{1}{4\pi} \iint_{\text{ouverture}} E_o(x', y') \frac{e^{-jkR}}{R} \left[\left(\frac{1}{R} + jk \right) \cos \theta + jk \cos \theta' \right] dx' dy'$$

θ est l'angle formé entre la normale à la surface et le rayon allant du point source au point d'observation. θ' est l'angle formé par les rayons incidents et la normale.

3.4.1 Zone de Rayleigh

La zone de champ très proche, appelée aussi zone de Rayleigh, est celle dans laquelle le champ électromagnétique est très complexe. Dans cette région, il existe des ondes évanescentes qui restent localisées au voisinage de l'ouverture. La théorie scalaire ne s'applique pas dans ce cas. La méthode vectorielle doit être utilisée sans approximation. Le champ électrique, dans la partie centrale de cette zone, présente une structure organisée sous forme de plans d'onde car les bords de l'ouverture sont vus avec un angle proche de 90° . Dans cette zone, le champ ne présente pas de divergence. Le faisceau est pratiquement cylindrique. On admet que cette zone s'étend jusqu'à une distance de l'ordre de la longueur d'onde.

3.4.2 Zone de Fresnel

La zone de Fresnel est située au-delà de la zone de champ très proche et à une distance au moins égale à une longueur d'onde de l'ouverture. Lorsque $R \gg \frac{\lambda}{2\pi}$, le terme en $1/R$ disparaît. L'amplitude du champ est donc purement imaginaire et non plus complexe comme elle l'était dans l'expression dans la zone de Rayleigh. Cela traduit la disparition des ondes évanescentes. L'expression du champ s'écrit :

$$E(r, \theta, \varphi) = \frac{1}{4\pi} \iint_{\text{ouverture}} E_o(x', y') \frac{e^{-jkR}}{R} jk [\cos \theta + \cos \theta'] dx' dy'$$

La théorie scalaire donne des résultats acceptables dans cette zone qui constitue une transition entre le champ très proche et le champ lointain. Le champ y présente de grandes variations.

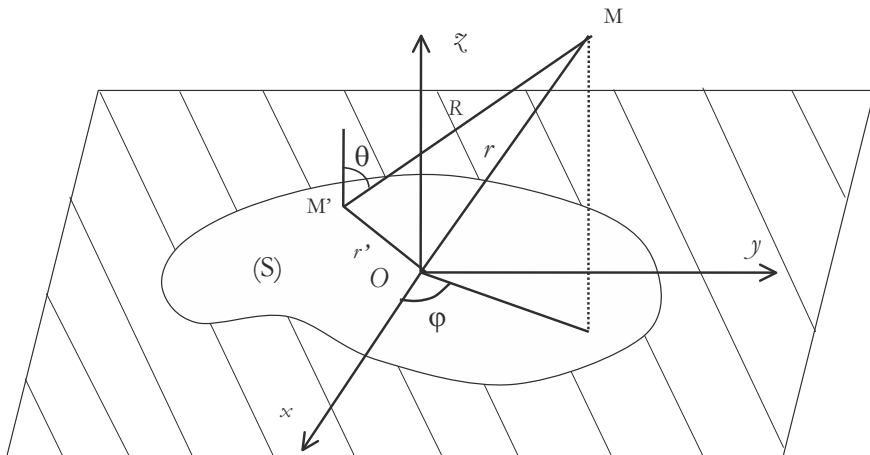


Figure 3.29 – Position du point d'observation par rapport à l'ouverture.

Selon la figure 3.29, pour l'incidence normale, le champ électrique est donné par :

$$E(r, \theta, \varphi) = j \frac{1 + \cos \theta}{2\lambda} \frac{e^{-jkr}}{r} I(r, \theta, \varphi)$$

Avec :

$$I(r, \theta, \varphi) = \iint_{\text{ouverture}} E_a(x', y') e^{\left[\frac{-jk}{2r}(x'^2 + y'^2)\right]} e^{[jk \sin \theta (x' \cos \varphi + y' \sin \varphi)]} dx' dy' \quad [3.45]$$

Le coefficient d'obliquité est défini par :

$$\frac{1 + \cos \theta}{2}$$

Lorsqu'on trace la puissance reçue sur l'axe z pour une ouverture carrée de dimension d (égale à huit longueurs d'onde pour la figure 3.30), on obtient un champ très oscillant juste derrière l'ouverture.

Sur la figure 3.30, la puissance est représentée pour une distance comprise entre une longueur d'onde et trente longueurs d'onde.

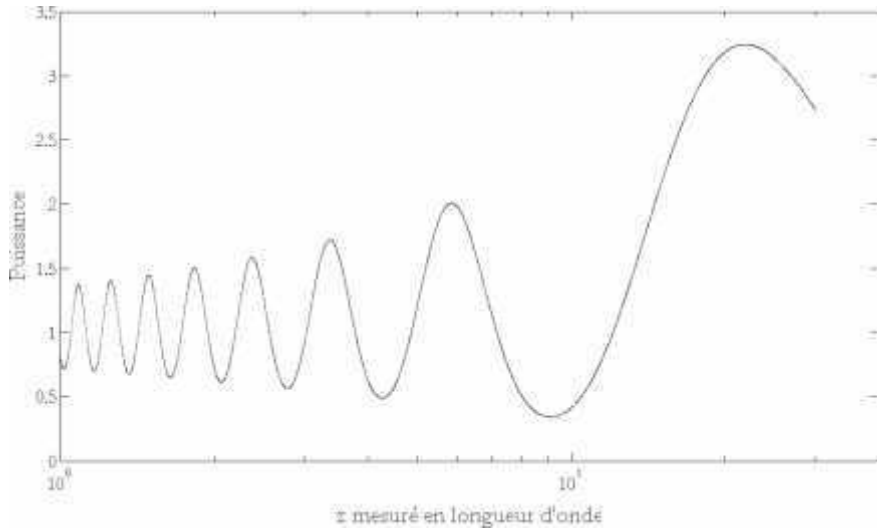


Figure 3.30 – Puissance électromagnétique sur l'axe z perpendiculaire à l'ouverture rectangulaire ($d = 8\lambda$), pour des valeurs de z comprises entre une et trente longueurs d'onde.

En traçant la puissance pour des distances plus grandes, les oscillations disparaissent. À partir d'une certaine distance, la puissance décroît ensuite selon l'inverse du carré de la distance (figure 3.31), caractéristique du champ lointain.

3.4.3 Zone de Fraunhofer

Lorsqu'on se trouve suffisamment loin, la puissance décroît selon l'inverse du carré de la distance. Cette variation est obtenue lorsque la distance r entre le point d'observation et l'ouverture, supposée carrée, de dimension d , est telle que :

$$r \geq \frac{2d^2}{\lambda}$$

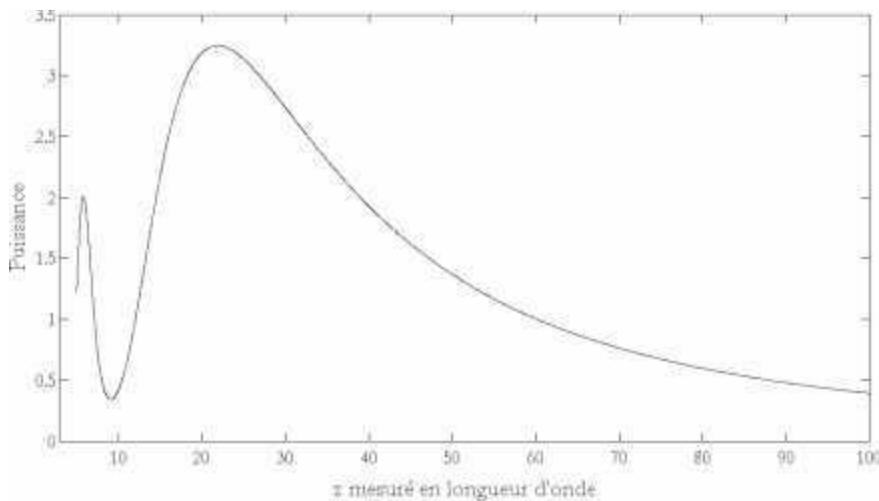


Figure 3.31 – Puissance électromagnétique sur l'axe z perpendiculaire à l'ouverture rectangulaire ($d = 8 \lambda$), pour des valeurs de z allant de 5 à 100 longueurs d'onde.

Cette condition définit la zone de champ lointain d'une antenne. C'est à partir de cette valeur que les antennes sont utilisées dans la majorité des cas.

Nous allons montrer la valeur de cette limite pour le champ lointain. Dans l'expression [3.45], l'exponentielle est la cause des oscillations à faible distance. Ces oscillations sont dues à une phase variant rapidement en fonction de la distance r . Lorsque la distance augmente, la phase portée par cette expression tend vers zéro. Cette stabilisation est obtenue pour :

$$\frac{k}{2r}(x'^2 + y'^2) \leq \frac{\pi}{4}$$

Dans cette zone, de champ lointain, on parle de diffraction de Fraunhofer. L'expression du champ est alors donnée par :

$$E(r, \theta, \varphi) = j \frac{e^{-jkr}}{\lambda r} I(r, \theta, \varphi)$$

Avec :

$$I(r, \theta, \varphi) = \iint_{\text{ouverture}} E_a(x', y') e^{[jk \sin \theta (x' \cos \varphi + y' \sin \varphi)]} dx' dy'$$

Le terme d'obliquité n'apparaît plus car les observations ayant lieu à grande distance, on est dans le cas où :

$$\frac{1 + \cos \theta}{2} \approx 1$$

On retrouve l'expression du champ démontrée avec la théorie vectorielle à grande distance.

Bibliographie

BALANIS C.A. – *Advanced Engineering Electromagnetics*, New Jersey, A. John Wiley & Sons, Inc., 1989.

BALANIS C.A. – *Antenna Theory Analysis and Design*, New Jersey, A. John Wiley & Sons, Inc., 2005.

COMBES P.F. – *Micro-ondes 2. Circuits passifs, propagation, antennes*, Paris, Dunod, 1997.

ELIOTT R.S. – *Antenna Theory and design*, IEEE Press, New Jersey, A. John Wiley & Sons, Inc., 2003.

STUTZMANN W. L., GARY A.T. – *Antenna Theory and Design*, New Jersey, A. John Wiley & Sons, Inc., 1998.

ULABY F. T., MOORE R. K., FUNG A.K. – *Microwave Remote Sensing, Active and Passive, Vol. I*, Norwood, Artech House, 1981.

4 • CARACTÉRISTIQUES DES ANTENNES

Dans ce chapitre, seront définies les caractéristiques des antennes utiles pour le dimensionnement des systèmes d'émission réception. Ces caractéristiques sont essentiellement liées à la forme du rayonnement dans l'espace.

Le fonctionnement normal d'une antenne est d'émettre ou de recevoir le rayonnement à grande distance. On ne s'intéressera donc dans ce chapitre qu'aux propriétés du rayonnement électromagnétique en champ lointain.

Les antennes présentées sont réciproques, c'est pourquoi, il ne sera pas spécifié si les antennes fonctionnent en émission ou en réception.

4.1 Fonction caractéristique de rayonnement

La figure 4.1 rappelle la situation de l'étude.

L'antenne est située en O . Le rayonnement est observé au point M , situé en champ lointain de l'antenne. Dans ces conditions, le champ électromagnétique est assimilable, localement, à celui d'une onde plane. Les vecteurs $(\vec{E}, \vec{H}, \vec{k})$ sont directement perpendiculaires entre eux. Les champs électriques et magnétiques sont contenus dans un plan perpendiculaire au vecteur de propagation. Le champ électrique n'ayant pas de composante selon le vecteur radial s'exprime par :

$$\vec{E} = E_\theta \vec{u}_\theta + E_\phi \vec{u}_\phi$$

Le champ magnétique s'en déduit par :

$$\vec{H} = \frac{1}{Z} (\vec{u} \wedge \vec{E})$$

Avec l'impédance du vide :

$$Z = \sqrt{\frac{\mu_0}{\epsilon_0}} \approx 377 \, \Omega$$

La densité surfacique de puissance est définie en fonction du champ électrique d'après la formule [2.42]

$$S_r = \frac{1}{2Z} (|E_\theta|^2 + |E_\phi|^2)$$

C'est la grandeur fondamentale sur laquelle repose l'étude du rayonnement électromagnétique d'une antenne.

Le champ électromagnétique issu de l'antenne contient un terme variant avec la distance, provenant du fait que le rayonnement en M est la somme d'ondes sphériques émises par les sources

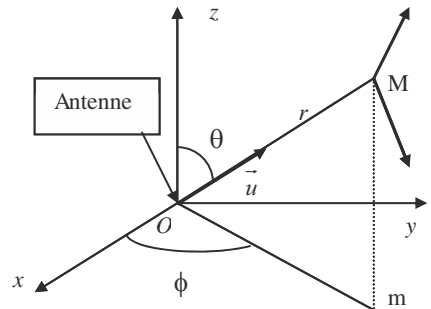


Figure 4.1 – Situation de l'antenne par rapport au point M d'observation.

élémentaires de l'antenne. Les composantes du champ électrique se mettent donc sous la forme :

$$E_\theta = \frac{e^{-jkr}}{r} f_1(\theta, \phi) \quad \text{et} \quad E_\phi = \frac{e^{-jkr}}{r} f_2(\theta, \phi)$$

Selon ces expressions, la densité surfacique de puissance s'exprime par :

$$S_r = \frac{1}{2Z} \frac{1}{r^2} \left(|f_1(\theta, \phi)|^2 + |f_2(\theta, \phi)|^2 \right) \quad [4.1]$$

Cette expression présente un découplage entre la variation angulaire et la variation radiale. Cette dernière, représentée par une variation en inverse du carré de la distance, est identique pour tous les types d'antennes et donne l'atténuation liée à l'expansion d'une onde sphérique. La variation angulaire est spécifique d'un type d'antenne. Dans la suite, on s'intéressera essentiellement à cette variation angulaire.

Ainsi, la définition de la fonction caractéristique de rayonnement est donnée par :

$$F(\theta, \phi) = \frac{1}{2Z} \left(|f_1(\theta, \phi)|^2 + |f_2(\theta, \phi)|^2 \right) \quad [4.2]$$

Cette définition ne fait intervenir que les variables angulaires.

La fonction caractéristique de rayonnement normalisée est définie par :

$$F_n(\theta, \phi) = \frac{F(\theta, \phi)}{F_{\max}(\theta, \phi)}$$

Soit en fonction de la densité surfacique de puissance :

$$F_n(\theta, \phi) = \frac{S_r(\theta, \phi)}{S_{r\max}(\theta, \phi)}$$

La fonction normalisée s'exprime soit de façon linéaire, soit de façon logarithmique. Cette valeur est souvent grande et résulte du rapport de deux puissances. L'échelle logarithmique en décibel (dB) permet de mieux apprécier les variations à des échelles différentes.

4.2 Diagramme de rayonnement

La représentation graphique de la fonction caractéristique de l'antenne porte le nom de diagramme de rayonnement. La direction du maximum de rayonnement est appelée l'**axe de rayonnement** de l'antenne.

La représentation de cette fonction donne les caractéristiques du rayonnement dans l'espace. Classiquement, on a pris l'habitude de représenter le diagramme de rayonnement dans deux plans perpendiculaires qui sont : **le plan E et le plan H**. Le plan E est défini comme le plan contenant l'axe de l'antenne et le champ électrique. Le plan H est défini comme le plan contenant l'axe de l'antenne et le champ magnétique.

Certaines représentations en trois dimensions ont l'avantage de montrer toutes les directions de rayonnement dans l'espace (figure 4.2) mais permettent difficilement une appréciation quantitative. La figure est en coordonnées logarithmiques. Ceci permet de mieux voir les détails pour les faibles valeurs, dans les lobes latéraux.

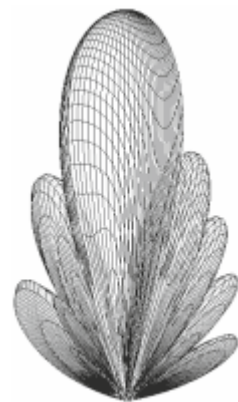


Figure 4.2 – Exemple de diagramme de rayonnement 3D en valeurs logarithmiques.

Le diagramme de rayonnement, généralement en coordonnées logarithmiques, est présenté soit en coordonnées rectangulaires, soit en coordonnées polaires, dans les deux plans perpendiculaires (E et H). Le diagramme de rayonnement de la figure 4.2 est présenté ci-dessous sous ces différentes formes. Les coordonnées rectangulaires sont utilisées sur la figure 4.3 pour présenter le diagramme de rayonnement normalisé dans le plan E et, sur la figure 4.4, dans le plan H. Sur la figure 4.5, la représentation en coordonnées polaires est utilisée pour représenter ces diagrammes

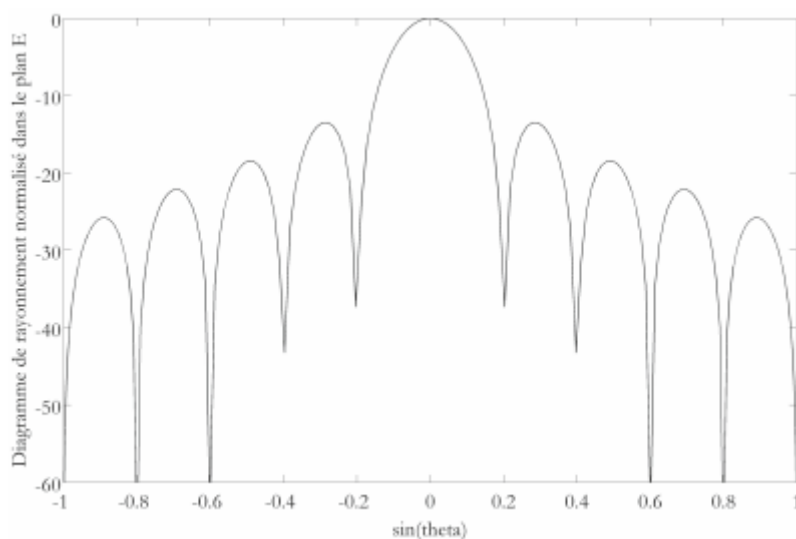


Figure 4.3 – Diagramme de rayonnement normalisé dans le plan E, en dB.

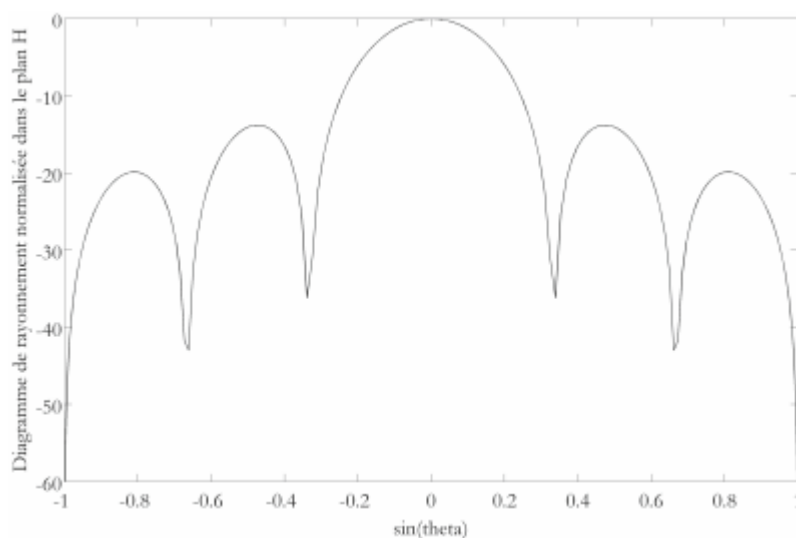


Figure 4.4 – Diagramme de rayonnement normalisé dans le plan H, en dB.

dans les plans E et H, respectivement. La valeur du maximum est de 0 dB, obtenu dans la direction du maximum qui correspond à de l'axe de l'antenne.

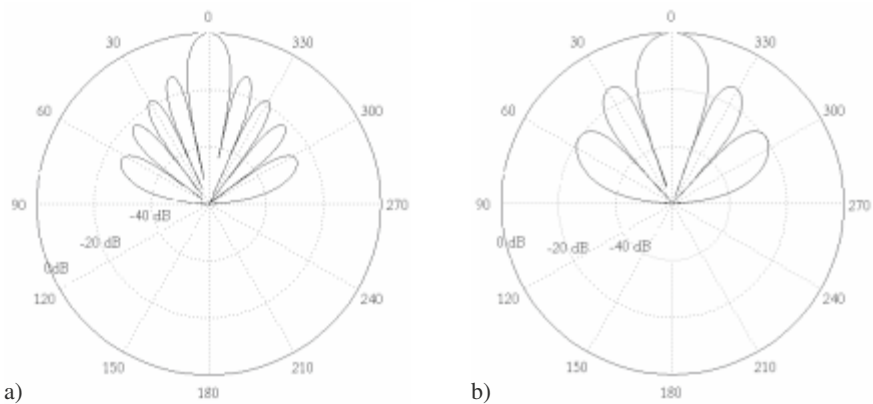


Figure 4.5 – a) Diagramme de rayonnement normalisé dans le plan E en coordonnées polaires, en fonction de θ , en dB.
b) Diagramme de rayonnement normalisé dans le plan H en coordonnées polaires, en fonction de θ , en dB.

Le **lobe principal** est défini entre les deux minima de chaque côté du maximum. Des maxima secondaires apparaissent de chaque côté. Ils constituent les **lobes secondaires**.

■ Ouverture de l'antenne

L'ouverture angulaire à mi-hauteur ou ouverture à 3 dB est définie par l'écart angulaire existant entre les deux directions situées de chaque côté de l'axe, pour lesquelles la puissance est divisée par deux.

L'ouverture angulaire n'a de sens que dans un plan. Pour évaluer l'ouverture d'une antenne, il faut donner l'ouverture dans deux plans perpendiculaires. Donc on parle d'ouverture dans le plan E, ou dans le plan H. C'est alors le produit de ces angles qui donnerait une idée de l'ouverture en trois dimensions. On utilise plutôt la notion d'angle solide qui est une généralisation de la notion d'angle valable en deux dimensions.

On rappelle la définition de l'angle solide $d\Omega$ sous lequel est vue une surface $d\vec{s}$ à partir d'un point O (figure 4.6) :

$$d\Omega = \frac{\vec{u} \cdot d\vec{s}}{r^2}$$

où r représentent la distance entre le point O et la surface élémentaire. $\vec{u}$ est le vecteur élémentaire partant de O vers le point d'observation.

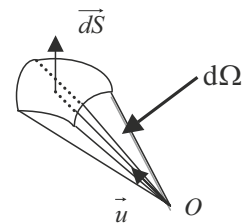


Figure 4.6 – Définition de l'angle solide.

■ Angle solide de l'antenne

L'angle solide de l'antenne donne une valeur qui correspond à l'angle moyen dans lequel l'antenne émet. Il est donné par la relation suivante :

$$\Omega_p = \iint_{4\pi} F_n(\theta, \phi) d\Omega \quad [4.3]$$

Considérons la fonction créneau F_0 , telle que :

$$\begin{cases} F_0 = 1 & \text{si } \Omega < \Omega_p \\ F_0 = 0 & \text{si } \Omega > \Omega_p \end{cases}$$

On vérifie :

$$\Omega_p = \iint_{4\pi} F_n(\theta, \phi) d\Omega = \iint_{4\pi} F_0 d\Omega$$

L'angle solide de l'antenne apparaît donc comme celui d'une antenne imaginaire de diagramme de rayonnement unité à l'intérieur de l'angle solide Ω_p et nulle à l'extérieur.

L'angle solide de l'antenne donne une idée de la directivité d'une antenne. Une antenne très directive, vise dans une direction très précise. Son ouverture est faible et son angle solide est petit. Une antenne peu directive présente au contraire un angle solide grand.

Une autre indication utile pour évaluer la forme du rayonnement est l'angle solide du lobe principal. L'intégration en angle solide se limite alors au lobe principal :

$$\Omega_M = \iint_{\text{lobe principal}} F_n(\theta, \phi) d\Omega$$

On définit de même l'angle solide dans les lobes secondaires :

$$\Omega_m = \iint_{\text{lobe secondaire}} F_n(\theta, \phi) d\Omega$$

Avec :

$$\Omega_p = \Omega_M + \Omega_m$$

On définit l'**efficacité dans le lobe principal** par :

$$\eta_M = \frac{\Omega_M}{\Omega_p} \quad [4.4]$$

Une antenne isotrope est une antenne qui rayonnerait la même puissance dans toutes les directions. Cette antenne ne peut pas exister à moins d'être ponctuelle. Elle aurait une fonction caractéristique de rayonnement égale à 1 dans toutes les directions. Son angle solide serait donc de :

$$\Omega_p = \Omega_M = 4\pi$$

■ Intensité

Si dS représente la surface élémentaire perpendiculaire au vecteur de propagation, la puissance élémentaire s'exprime par :

$$dP = S_r dS$$

La surface élémentaire est liée à l'angle solide correspondant. Cela permet d'exprimer la puissance élémentaire par :

$$dP = S_r r^2 d\Omega$$

L'intensité I dans une direction donnée est définie par la puissance rayonnée par unité d'angle solide :

$$dP = I(\theta, \phi) d\Omega = F(\theta, \phi) d\Omega$$

4.3 Directivité

4.3.1 Définition

On parle d'une antenne plus ou moins directive. Afin de quantifier cette propriété la notion de directivité a été introduite.

La directivité dans une direction est le rapport entre la valeur de la fonction caractéristique de rayonnement dans cette direction à sa valeur moyenne dans tout l'espace :

$$D(\theta, \phi) = \frac{F_n(\theta, \phi)}{\frac{1}{4\pi} \iint F_n(\theta, \phi) d\Omega} \quad [4.5]$$

Une valeur faible pour la moyenne de la fonction caractéristique de rayonnement, correspond à une antenne directive : la puissance n'est envoyée que dans un cône d'angle solide petit. Dans ce cas, d'après la définition [4.5], la valeur moyenne étant au dénominateur entraîne une valeur élevée de la directivité.

La directivité s'exprime aussi en fonction de la densité surfacique de puissance :

$$D(\theta, \phi) = \frac{S_r(\theta, \phi)}{\frac{1}{4\pi} \iint S_r(\theta, \phi) d\Omega} \quad [4.6]$$

La directivité dans une direction permet de comparer la densité de puissance rayonnée dans cette direction à la puissance moyenne rayonnée dans l'espace.

La directivité étant égale à un rapport de puissances, peut être exprimée soit sur une échelle linéaire, soit sur une échelle logarithmique. Dans ce dernier cas, elle s'exprime en décibel (dB), comme dix fois son logarithme en base dix.

Remarquons que la directivité peut être inférieure ou supérieure à 1 sur une échelle linéaire, ou bien positive ou négative sur une échelle logarithmique.

Le maximum de directivité est obtenu dans la direction de l'axe de l'antenne, lorsque la fonction caractéristique normalisée est égale à 1. Soit :

$$D_0 = \frac{1}{\frac{1}{4\pi} \iint F_n(\theta, \phi) d\Omega} \quad [4.7]$$

D_0 est appelé la directivité de l'antenne sans préciser de direction. La directivité exprime la proportion de puissance ramenée au voisinage de l'axe par rapport à la puissance totale rayonnée.

Dans cette formule, la directivité au maximum de la fonction caractéristique est égale à l'inverse de la valeur moyenne de celle-ci. En introduisant l'angle solide de l'antenne défini précédemment [4.3] :

$$D_0 = \frac{4\pi}{\Omega_p}$$

En utilisant cette valeur dans l'expression générale de la directivité, il vient :

$$D(\theta, \phi) = D_0 F_n(\theta, \phi)$$

La directivité est une mesure relative de la puissance rayonnée dans une direction par rapport à la puissance totale rayonnée. Dans un paragraphe suivant, la puissance rayonnée sera comparée à la puissance d'alimentation de l'antenne.

4.3.2 Surface effective

Considérons une antenne rayonnant dans la direction z dont le champ dans l'ouverture est polarisé selon x . Selon l'expression [3.43], le champ à grande distance se met sous la forme :

$$\vec{E}(\vec{r}) = \frac{jke^{-jk.r}}{2\pi r} [\cos \phi \vec{u}_\theta - \sin \phi \cos \theta \vec{u}_\phi] f_x(\theta, \phi)$$

Avec :

$$f_x(\theta, \phi) = \iint_{\text{ouverture}} E(x, y) e^{j(k_x x + k_y y)} dx dy \quad [4.8]$$

La densité surfacique de puissance à distance r est :

$$S_r = \frac{1}{2Z} EE^*$$

Soit, en considérant θ petit :

$$S_r(\theta, \phi) \approx \frac{k^2}{2Z} \frac{1}{4\pi^2 r^2} |f_x(\theta, \phi)|^2$$

Ou encore :

$$S_r(\theta, \phi) \approx \frac{1}{2Z} \frac{1}{\lambda^2 r^2} |f_x(\theta, \phi)|^2 \quad [4.9]$$

La puissance totale rayonnée P_r est l'intégrale de la densité de puissance sur la sphère de rayon r :

$$P_r = \iint_{4\pi} S_r(\theta, \phi) ds \quad [4.10]$$

Or la directivité, qui représente la densité de puissance dans une direction sur sa moyenne dans tout l'espace, s'exprime par :

$$D(\theta, \phi) = \frac{S_r(\theta, \phi)}{\frac{1}{4\pi} \iint S_r(\theta, \phi) d\Omega}$$

En utilisant l'expression de la puissance totale :

$$D(\theta, \phi) = \frac{S_r(\theta, \phi)}{P_r} 4\pi r^2$$

En utilisant l'expression [4.9] de la densité de puissance :

$$D(\theta, \phi) = \frac{1}{2Z} \frac{|f_x(\theta, \phi)|^2}{\lambda^2} \frac{4\pi}{P_r}$$

Par ailleurs, la puissance totale rayonnée sur la sphère de rayon r est aussi égale à la puissance contenue dans l'ouverture puisque l'ouverture est la source de puissance :

$$P_r = \frac{1}{2Z} \iint_{\text{ouverture}} |E(x, y)|^2 ds \quad [4.11]$$

La directivité s'exprime alors par :

$$D(\theta, \phi) = \frac{|f_x(\theta, \phi)|^2}{\lambda^2} \frac{4\pi}{\iint_{\text{ouverture}} |E(x, y)|^2 ds}$$

Avec [4.8], il vient :

$$D(\theta, \phi) = \frac{4\pi}{\lambda^2} \frac{\left| \iint_{\text{ouverture}} E(x, y) e^{j(k_x x + k_y y)} dx dy \right|^2}{\iint_{\text{ouverture}} |E(x, y)|^2 dx dy}$$

Si S est la surface géométrique de l'ouverture, l'inégalité de Schwartz s'écrit :

$$S \iint |E(x, y)|^2 dx dy \geq \left| \iint E(x, y) e^{j(k_x x + k_y y)} dx dy \right|^2 \quad [4.12]$$

Pour se convaincre de la validité de cette relation, il suffit d'utiliser la représentation de Fresnel des nombres complexes, qui tient compte de leur phase. Dans le membre de droite, les termes sont sommés en phase, alors que dans le membre de gauche, ce sont les normes qui sont sommées.

On en déduit, dans la direction perpendiculaire à la surface de l'ouverture :

$$D(0,0) \leq \frac{4\pi S}{\lambda^2} \quad [4.13]$$

D'après [4.12], l'égalité est vérifiée si le champ est constant sur la surface.

Dans le cas contraire, la surface effective de rayonnement S_{eff} est introduite selon la relation suivante :

$$D(0,0) = \frac{4\pi S_{\text{eff}}}{\lambda^2} \quad [4.14]$$

Ce qui impose :

$$S_{\text{eff}} \leq S$$

Dans le cas où l'égalité de l'expression [4.12] est vérifiée, la **surface effective** S_{eff} est égale à la surface géométrique et :

$$D(0,0) = \frac{4\pi S}{\lambda^2}$$

Et on définit l'**efficacité de rayonnement** η de l'antenne :

$$\eta = \frac{S_{\text{eff}}}{S} \quad [4.15]$$

La valeur de l'efficacité maximale est égale à 1. Elle mesure la capacité d'une antenne à présenter la directivité maximale dans l'axe. L'efficacité est calculable à partir de la connaissance de la répartition des champs dans l'ouverture :

$$\eta = \frac{1}{S} \frac{\left| \iint_{\text{ouverture}} E(x, y) e^{j(k_x x + k_y y)} dx dy \right|^2}{\iint_{\text{ouverture}} |E(x, y)|^2 dx dy}$$

Le tableau 4.1 donne différentes valeurs de l'efficacité, de la largeur de l'ouverture et des niveaux de lobes secondaires en fonction de la répartition de l'amplitude du champ électrique, appelée aussi la loi d'illumination.

On constate, sur ce tableau, que l'efficacité maximale est obtenue pour une répartition constante du champ électrique. Une variation du champ importante nuit à une bonne efficacité. L'efficacité la plus faible est obtenue pour la variation en cosinus carré.

Tableau 4.1 – Différentes caractéristiques d'antennes en fonction de la répartition du champ sur l'ouverture.

Représentation du champ en fonction de x	$E(x)$	η	Largeur à mi-hauteur en degrés	Niveau des lobes secondaires en dB
	Constant	1	$51 \lambda/a$	-13,3
	$1 - \frac{2 x }{a}$	0,75	$73 \lambda/a$	-26,5
	$\cos \frac{\pi x}{a}$	0,81	$68 \lambda/a$	-23,1
	$\cos^2 \frac{\pi x}{a}$	0,67	$82 \lambda/a$	-31,7
	$\frac{1}{2} \left(1 + 2 \cos^2 \frac{\pi x}{a} \right)$	0,89	$61 \lambda/a$	-26

Les lobes secondaires ont tendance à remonter lorsque la discontinuité du champ aux bords est forte. Ainsi, la remontée du lobe secondaire pour la répartition constante est la plus importante parmi toutes les distributions mentionnées. Des champs dont la distribution s'annule lentement aux bords, comme c'est le cas de la répartition en cosinus carré ou celle en cosinus surélevé, donnent des remontées de lobes secondaires très faibles.

Le cône d'ouverture de l'antenne est lié à la répartition du champ. Pour une répartition sur l'ouverture présentant un champ fort central, les points au centre participent davantage à la construction du diagramme de rayonnement. La surface équivalente de rayonnement est donc plus petite et l'angle d'ouverture plus grand. L'ouverture en cosinus donne un champ plus réparti sur l'ouverture que pour le cas du cosinus carré pour lequel le champ est plus concentré. En conséquence, l'ouverture dans ce dernier cas est plus grande et donc la directivité plus faible.

Toutes ces propriétés sont liées au phénomène de diffraction dont la compréhension est fondamentale pour la conception d'antenne. Les dimensions et formes d'une antenne sont choisies en général au terme d'une optimisation entre les différentes caractéristiques visées, comme le montre le tableau 4.1 pour quelques exemples.

4.4 Gain d'une antenne

Soit P_t la puissance d'alimentation d'une antenne. Cette puissance est transformée en une puissance rayonnée P_0 . Dans le sens de l'émission, la puissance rayonnée est inférieure à la puissance d'alimentation. L'antenne est un transformateur imparfait. Il y a des pertes lors de la transformation d'énergie, comme dans tout système. L'efficacité de l'antenne est définie par :

$$\eta_l = \frac{P_0}{P_t} \quad [4.16]$$

Elle permet de mesurer le taux de transformation. C'est un rendement au sens thermodynamique du terme :

$$\eta_l \leq 1$$

À la réception, la transformation a lieu en sens inverse. La puissance P_r reçue sur le récepteur est inférieure à la puissance P_0 rayonnée arrivant sur l'antenne.

Le gain dans une direction est défini par le rapport de la densité de puissance rayonnée dans une direction à la densité de puissance S_{ri} qui serait rayonnée par une antenne isotrope sans pertes, les deux antennes étant alimentées par la même puissance et placées à la même position.

$$G(\theta, \phi) = \frac{S_r(\theta, \phi)}{S_{ri}} \quad [4.17]$$

Le gain s'exprime en décibel (dB). On utilise quelquefois la notation dBi pour préciser la référence au rayonnement isotrope. La puissance totale rayonnée par l'antenne est donnée par :

$$P_0 = \iint_{4\pi} S_r(\theta, \phi) r^2 d\Omega \quad [4.18]$$

P_0 est liée à la puissance d'alimentation par l'efficacité de l'antenne [4.16]. La densité de puissance isotrope S_{ri} se déduit de la puissance d'alimentation par :

$$P_t = P_{0i} = 4\pi r^2 S_{ri} \quad [4.19]$$

Donc, en utilisant [4.16] et [4.19] :

$$S_{ri} = \frac{1}{4\pi\eta_l} \iint S_r(\theta, \phi) d\Omega$$

En reportant cette valeur dans l'expression du gain :

$$G(\theta, \phi) = 4\pi\eta_l \frac{S_r(\theta, \phi)}{\iint S_r(\theta, \phi) d\Omega} \quad [4.20]$$

Soit encore :

$$G(\theta, \phi) = \eta_l D(\theta, \phi) \quad [4.21]$$

Le gain est proportionnel à la directivité. Il porte la même information sur les directions de rayonnement. Le gain tient compte du rendement de transformation entre la puissance d'alimentation et la puissance rayonnée.

Au maximum de directivité, on a simplement :

$$G_0 = \eta_l D_0$$

On déduit de cette démonstration la densité de puissance rayonnée dans une direction, comme étant la densité de puissance émise par l'antenne isotrope multipliée par le gain dans la direction considérée qui tient compte de la répartition angulaire imposée par la géométrie de l'antenne et

de l'efficacité :

$$S_r(\theta, \phi) = P_t \frac{G(\theta, \phi)}{4\pi r^2} \quad [4.22]$$

La **puissance isotrope rayonnée équivalente (PIRE)** est la puissance de l'antenne isotrope qui rayonnerait la même puissance que l'antenne réelle sur son axe.

■ Aire d'absorption

La notion de surface effective définie par rapport à la directivité est transposée pour le gain. On appelle souvent cette surface l'aire d'absorption A . Elle est définie par :

$$G(0,0) = \frac{4\pi A}{\lambda^2}$$

□ Application numérique

Considérons une antenne d'émission demi-onde, de directivité 1,64, dont la puissance rayonnée est de 2 W, d'efficacité 0,9. On se propose de calculer la densité de puissance incidente sur une surface située à 100 m de l'antenne, sur l'axe de l'antenne, puis la valeur du champ électrique au niveau de cette surface.

La densité surfacique de puissance au maximum de directivité est donnée par :

$$S_r = \frac{P}{4\pi r^2} D_0 = 2,6 \cdot 10^{-5} \text{ Wm}^{-2}$$

Cela correspond à une puissance d'alimentation de $2/0,9 = 2,22$ W.

La norme du champ électrique correspondant à cette densité de puissance s'obtient en utilisant [2.42]. Soit :

$$|\vec{E}| = \sqrt{2ZS_r} = 0,14 \text{ Vm}^{-1}$$

4.5 Facteur d'antenne

Le facteur d'antenne F_A est utilisé en mesure. Il donne le rapport entre le champ électrique efficace rayonné et la tension mesurée juste derrière l'antenne.

Les calculs de rayonnement utilisent l'amplitude du champ électrique (champ crête). Dans ce paragraphe, nous utilisons le champ efficace :

$$E_{\text{eff}} = \frac{E}{\sqrt{2}}$$

Considérons le rayonnement provenant sur une antenne de réception. Il est caractérisé par la densité surfacique de puissance S_r .

D'après les définitions précédentes :

$$S_r = \frac{E_{\text{eff}}^2}{Z}$$

La puissance reçue P_r par l'antenne est fonction de la densité surfacique de puissance et de l'aire effective de réception A_r :

$$P_r = S_r A_r$$

Or l'aire effective de réception dépend du gain de l'antenne de réception, G_r , de la longueur d'onde λ :

$$A_r = \frac{\lambda^2}{4\pi} G_r \quad \text{donc} \quad P_r = \frac{E_{\text{eff}}^2}{Z} \frac{\lambda^2}{4\pi} G_r$$

Par ailleurs, la puissance reçue directement derrière l'antenne vérifie aussi la relation entre la tension et l'impédance caractéristique Z_0 de la ligne :

$$P_r = \frac{V^2}{Z_0} = \frac{E_{\text{eff}}^2}{Z} \frac{\lambda^2}{4\pi} G_r$$

En introduisant le facteur d'antenne F_A , comme le rapport entre le champ électrique efficace reçu et la tension induite aux bornes de l'antenne :

$$F_A = \frac{E_{\text{eff}}}{V} \quad [4.23]$$

On obtient :

$$F_A = \frac{1}{\lambda} \sqrt{\frac{Z}{Z_0} \frac{4\pi}{G_r}} \quad [4.24]$$

Ce facteur, qui indique le taux de transformation entre le champ reçu et la tension transformée, fait intervenir le rapport de l'impédance de l'air à l'impédance de la ligne, le gain de l'antenne et la longueur d'onde. Il sert à comparer des antennes à une antenne étalon, en mesurant le facteur d'antenne en champ constant.

4.6 Polarisation d'une onde

La polarisation d'une onde est une donnée fondamentale pour l'étude des antennes. En effet selon la constitution de l'antenne, elle ne recevra qu'une certaine forme de polarisation.

Donc si la polarisation de l'antenne de réception n'est pas accordée sur la polarisation de l'antenne d'émission, la puissance reçue ne sera pas maximale. Nous rappelons ici les caractéristiques de l'onde électromagnétique pour sa polarisation.

4.6.1 Polarisation rectiligne

La polarisation rectiligne est la plus simple à étudier. C'est celle d'un champ électromagnétique dont l'orientation reste la même au cours de la propagation (figure 4.7). Le champ électrique est parallèle au vecteur unitaire $\vec{u}_x$ et le champ magnétique est parallèle au vecteur unitaire $\vec{u}_y$. Le vecteur de propagation $\vec{k}$ est alors parallèle au vecteur unitaire $\vec{u}_z$.

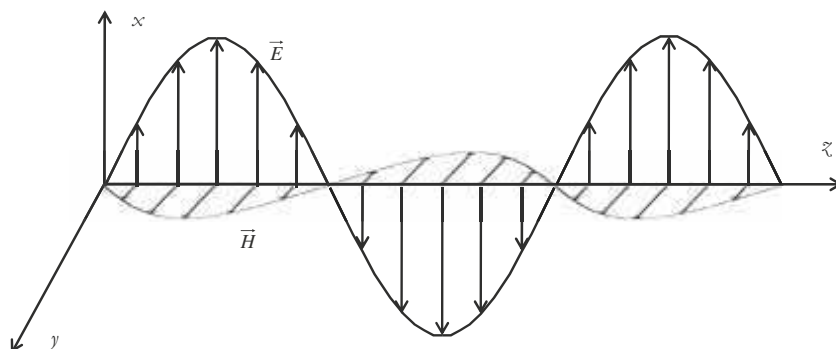


Figure 4.7 – Champ électromagnétique d'une onde plane polarisée rectilignement.

L'exemple choisi est celui d'une onde plane, par souci de simplicité. Une onde sphérique peut aussi présenter une polarisation rectiligne.

■ Polarisations H et V

Ces termes sont les abréviations de polarisations horizontale et verticale. Pour définir l'orientation de la polarisation, on choisit de repérer l'onde par rapport à la surface vers laquelle elle se propage (figure 4.8). Cette surface est assimilée au sol. On parle donc de polarisation horizontale si le champ électrique est parallèle au sol. Si le champ électrique est contenu dans un plan perpendiculaire au sol, on parle de polarisation verticale. Le champ magnétique est alors parallèle au sol.

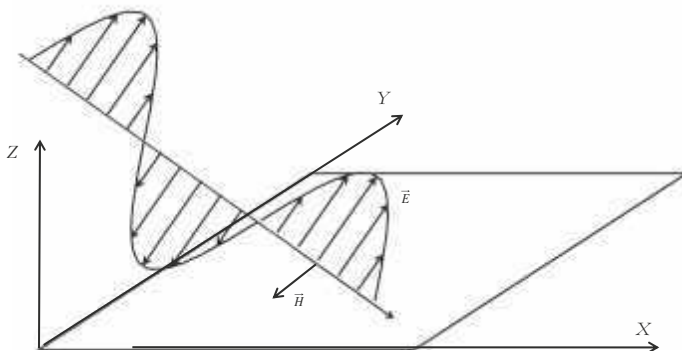


Figure 4.8 – Onde plane à polarisation verticale.

Toute onde polarisée rectilignement est la combinaison d'une onde à polarisation verticale et d'une onde à polarisation horizontale en phase l'une par rapport à l'autre.

■ Polarisations TE et TM

Une autre façon de définir la polarisation rectiligne est de repérer les champs par rapport au plan d'incidence, défini par la direction d'incidence et la normale au plan vers lequel se dirige l'onde. Lorsque le champ électrique est perpendiculaire au plan d'incidence, la polarisation est dite transverse électrique (TE). Si au contraire, elle est dans ce plan, c'est le champ magnétique qui est perpendiculaire au plan d'incidence. La polarisation est dite transverse magnétique (TM).

4.6.2 Polarisation elliptique

La polarisation d'une onde est décrite de façon générale par le lieu de l'extrémité du champ électrique lors de sa propagation. Pour décrire l'état de polarisation d'une onde, on regarde l'onde arriver. Lorsque les pertes sont nulles, l'extrémité du champ électrique décrit alors une courbe fermée qui est, dans le cas le plus général, une ellipse. La polarisation rectiligne est le cas particulier d'une ellipse d'excentricité nulle.

Considérons un champ électrique quelconque repéré par rapport au référentiel sphérique de la figure 4.9.

L'expression du champ électrique en notation complexe dans ce référentiel est :

$$\vec{E} = (E_\theta \vec{u}_\theta + E_\varphi \vec{u}_\varphi) \cdot e^{j\omega t} \quad [4.25]$$

Les composantes sont complexes et s'expriment sous la forme :

$$\begin{aligned} E_\theta &= E_\theta^r + jE_\theta^i \\ \text{et } E_\varphi &= E_\varphi^r + jE_\varphi^i \end{aligned}$$

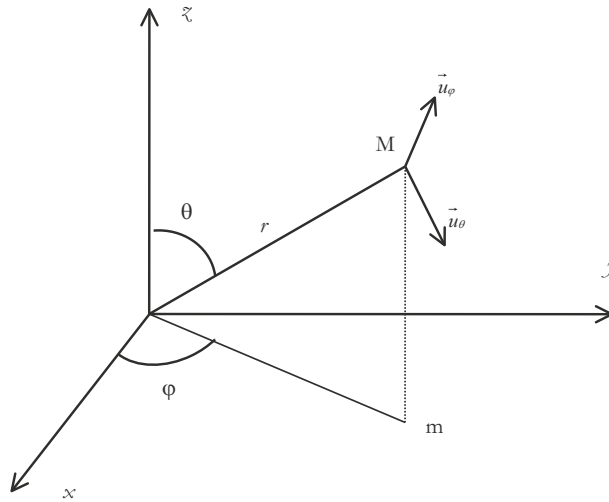


Figure 4.9 – Référentiel du champ électrique.

On rappelle que le champ réel, noté $\vec{e}(r, \theta, \varphi, t)$, est obtenu par :

$$\vec{e}(r, \theta, \varphi, t) = \text{R  el} \left[\vec{E} \right]$$

Soit
$$\vec{e}(r, \theta, \varphi, t) = \text{R  el} \left[((E_\theta^r + jE_\theta^i) \vec{u}_\theta + (E_\varphi^r + jE_\varphi^i) \vec{u}_\varphi \cdot e^{j\omega t}) \right]$$

D'o   l'expression du champ r  el :

$$\vec{e}(r, \theta, \varphi, t) = (E_\theta^r \cos \omega t - E_\theta^i \sin \omega t) \vec{u}_\theta + (E_\varphi^r \cos \omega t - E_\varphi^i \sin \omega t) \vec{u}_\varphi$$

Posons :

$$A = \sqrt{(E_\theta^r)^2 + (E_\theta^i)^2} \quad \text{et} \quad B = \sqrt{(E_\varphi^r)^2 + (E_\varphi^i)^2}$$

$$\alpha = \arctan \frac{E_\theta^i}{E_\theta^r} \quad \text{et} \quad \beta = \arctan \frac{E_\varphi^i}{E_\varphi^r}$$

L'expression du champ   lectrique se met sous la forme :

$$\vec{e} = A \cos(\omega t + \alpha) \vec{u}_\theta + B \cos(\omega t + \beta) \vec{u}_\varphi$$

En choisissant l'origine des temps telle que $\alpha = 0$, l'expression du champ s'  crit :

$$\vec{e} = A \cos \omega t \vec{u}_\theta + B \cos(\omega t + \beta) \vec{u}_\varphi \quad [4.26]$$

Suivons l'extr  mit   du champ au cours du temps. Pour cela, le temps prend les valeurs successives $\left(0, \frac{T}{4}, \frac{T}{2}, \frac{3T}{4}\right)$, o   T repr  sente la p  riode de l'onde. La figure 4.10 est trac  e en prenant $A = 1$, $B = 3$ et $\beta = \pi/3$. Le sens de polarisation, dans ce cas, est le sens r  trograde.

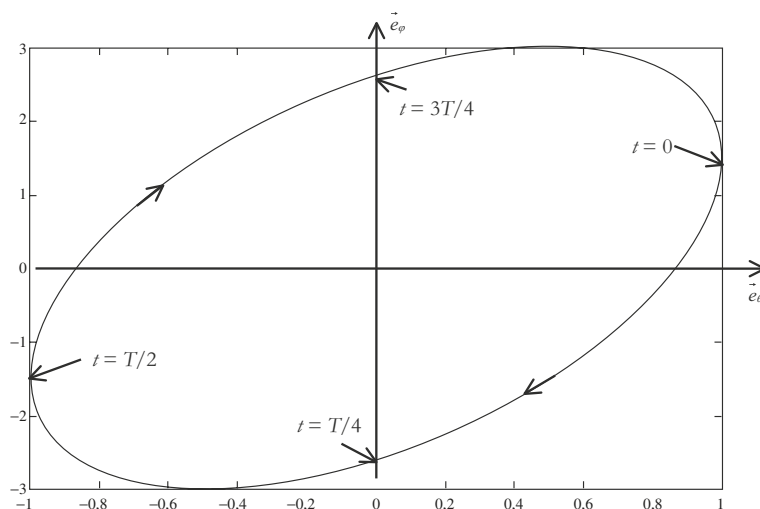


Figure 4.10 – Exemple d'ellipse de polarisation.

■ Sens de polarisation

Lorsque l'observateur voit arriver l'onde vers lui, si le champ électrique tourne dans le sens des aiguilles d'une montre, l'onde est dite polarisée **gauche ou rétrograde**. Dans l'autre sens on définit la polarisation **droite ou directe**.

Le sens de rotation de l'ellipse est fonction de β :

- Si $\beta > 0$ la polarisation est gauche
- Si $\beta < 0$ la polarisation est droite

La norme du champ électrique est donnée par :

$$\|\vec{e}\| = e = \sqrt{A^2 \cos^2 \omega t + B^2 \cos^2(\omega t + \beta)}$$

Les extrema de cette fonction sont obtenus pour les valeurs du temps telles que la dérivée est nulle. Le maximum correspond au demi-grand axe de l'ellipse et le minimum au demi-petit axe. Ces points correspondent aux valeurs données par la relation :

$$\tan 2\omega t = \frac{-B^2 \sin 2\beta}{A^2 + B^2 \cos 2\beta}$$

La figure 4.11 représente une ellipse de polarisation dans sa forme générale. Elle donne la définition d'un certain nombre d'angles qui permettent de définir la polarisation. Nous allons voir qu'il est équivalent de définir la polarisation soit par rapport au couple (ϵ, τ) soit par rapport au couple (γ, β) .

Sur cette figure, le grand axe est représenté par l'axe OX et le petit axe par OY . Le vecteur d'onde est supposé porté par l'axe Z et dans le même sens.

L'angle τ représente le décalage angulaire (tilt en anglais) de l'ellipse par rapport au repère de base. C'est l'angle défini entre l'axe $\vec{u}_\theta$ et le grand axe.

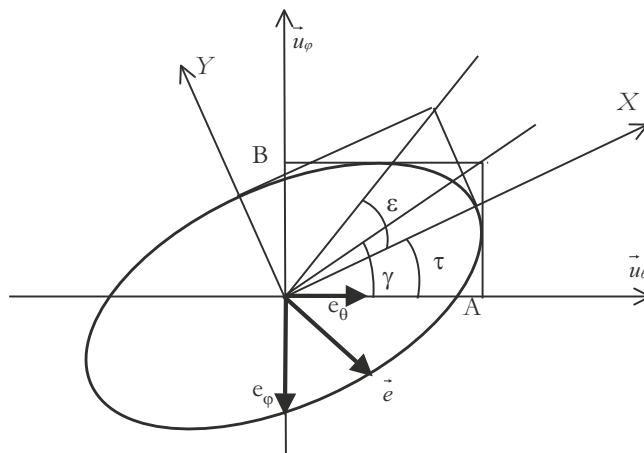


Figure 4.11 – Ellipse de polarisation.

■ Rapport axial

Le rapport du demi-grand axe au demi-petit axe de l'ellipse est appelé rapport axial. On le compte positivement pour les ondes gauches et négativement pour les ondes droites. Il vérifie la relation suivante :

$$1 \leq |AR| \leq \infty$$

Lorsque $AR = 1$, la polarisation est circulaire. Lorsque AR est infini, la polarisation est rectiligne. AR est souvent mesuré en dB :

$$AR_{dB} = 20 \log_{10}(|AR_{linéaire}|)$$

On admet qu'une antenne fournissant une polarisation telle que son rapport axial soit inférieur à 3 dB est une antenne à polarisation circulaire.

L'angle ϵ est donc défini comme :

$$\epsilon = \cot^{-1}(-AR) \quad -45^\circ \leq \epsilon \leq 45^\circ \quad [4.27]$$

■ Paramètres de l'ellipse

Reprenons [4.26]. Notons le champ dans le repère lié à l'onde sous la forme :

$$\vec{e} = e_\theta \vec{u}_\theta + e_\phi \vec{u}_\phi$$

L'équation de l'ellipse se met sous la forme :

$$\left(\frac{e_\theta}{A}\right)^2 + \left(\frac{e_\phi}{B}\right)^2 - 2\frac{e_\theta e_\phi}{AB} \cos \beta - \sin^2 \beta = 0 \quad [4.28]$$

Pour déterminer l'équation de l'ellipse par rapport à ses axes (X et Y), effectuons le changement de repère en fonction de l'angle fixant la position de l'ellipse :

$$e_\theta = X \cos \tau - Y \sin \tau$$

$$e_\phi = X \sin \tau + Y \cos \tau$$

Dans le nouveau repère l'équation de l'ellipse s'écrit :

$$\frac{X^2}{a^2} + \frac{Y^2}{b^2} - 1 = 0 \quad [4.29]$$

En annulant le terme croisé en XY , provenant de la transformation de l'équation [4.28], on obtient :

$$\tan(2\tau) = \frac{2AB \cos \beta}{A^2 - B^2} \quad [4.30]$$

Par ailleurs, on vérifie :

$$\gamma = \tan^{-1} \frac{B}{A} \quad [4.31]$$

Utilisant cette expression, [4.30] se transforme alors en :

$$\tan(2\tau) = \cos \beta \tan(2\gamma)$$

Des considérations géométriques sur l'ellipse permettent de montrer que :

$$\cos(2\gamma) = \cos(2\varepsilon) \cos(2\tau) \quad [4.32]$$

Il est donc équivalent d'utiliser le couple (ε, τ) ou le couple (γ, β) qui se déduisent l'un de l'autre par les relations précédentes.

On vérifie aussi les relations suivantes :

$$\tan(2\varepsilon) = \tan \beta \sin(2\tau) \quad [4.33]$$

$$\sin(2\varepsilon) = \sin(2\gamma) \sin \beta \quad [4.34]$$

■ La sphère de Poincaré

Considérons la sphère de Poincaré, de rayon unité, sur laquelle les angles précédemment définis sont reportés (figure 4.12) de façon à ce que la longitude soit égale à 2τ et la latitude à 2ε .

Compte tenu des relations entre les angles, le point P est repéré soit par le couple (ε, τ) soit par le couple (γ, β) . En particulier, on vérifie bien la relation [4.32].

À chaque valeur des angles est associé un état de polarisation sur la sphère de Poincaré représentée sur la figure 4.13.

Sur l'équateur $\varepsilon = 0$. Cela entraîne $\beta = 0$. Donc sur cette ligne, la polarisation est rectiligne. En fonction de sa longitude sur l'équateur, l'orientation de la droite de polarisation tourne.

Lorsque $\varepsilon = 45^\circ$, $\beta = 90^\circ$ (pôle nord) et la polarisation est circulaire gauche. Au pôle sud, elle est circulaire droite.

Sur la figure, la polarisation elliptique est aussi représentée. L'orientation de l'ellipse varie en fonction de la longitude.

La sphère de Poincaré permet de calculer le facteur d'adaptation (F) entre l'antenne d'émission et l'antenne de réception quant à la polarisation. Considérons l'état de polarisation de l'onde d'émission représentée par un point E et celui de l'antenne de réception par un point R . La puissance moyenne reçue doit tenir compte des polarisations relatives. Soit ER l'angle au centre interceptant l'arc ER sur la sphère. Le facteur F s'exprime par :

$$F = \cos^2 \frac{(ER)}{2} \quad [4.35]$$

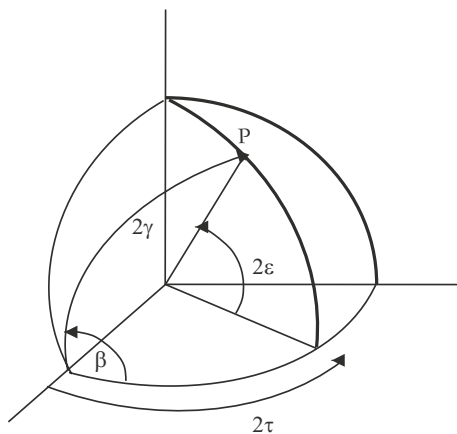


Figure 4.12 – Représentation des angles de polarisation sur une sphère.

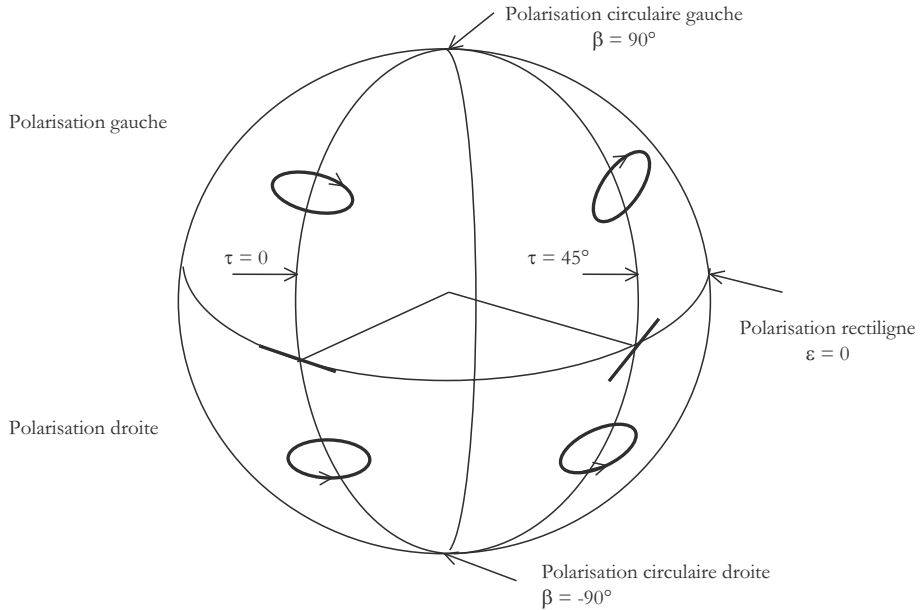


Figure 4.13 – Différents états de polarisation sur la sphère de Poincaré.

Ainsi deux antennes de même état de polarisation conduisent à une réception maximum. Au contraire, pour deux antennes dont les états de polarisation sont diamétralement opposés sur la sphère de Poincaré, la puissance reçue est nulle. C'est le cas par exemple de deux antennes dont les polarisations sont rectilignes et à 90° . Leurs points représentatifs se situent sur l'équateur, à 180° .

■ Paramètres de Stokes

Afin de caractériser la polarisation, on introduit les paramètres de Stokes définis ci-dessous :

$$\begin{aligned} s_0 &= A^2 + B^2 \\ s_1 &= A^2 - B^2 \\ s_2 &= 2AB \cos \beta = 2\text{Réel}(e_\theta e_\phi^*) \\ s_3 &= 2AB \sin \beta = 2\text{Im}(e_\theta e_\phi^*) \end{aligned}$$

Ces coefficients ne sont pas indépendants. Pour un rayonnement cohérent, ils sont liés par :

$$s_0^2 = s_1^2 + s_2^2 + s_3^2$$

Calculons le rapport entre s_2 et s_1 :

$$\frac{s_2}{s_1} = \frac{2AB \cos \beta}{A^2 - B^2}$$

En utilisant [4.30] :

$$\frac{s_2}{s_1} = \tan(2\tau)$$

Calculons le rapport entre s_3 et s_0 :

$$\frac{s_3}{s_0} = \frac{2AB \sin \beta}{A^2 + B^2}$$

En utilisant la définition de l'angle γ :

$$\frac{s_3}{s_0} = \frac{2 \tan(\gamma)}{1 + \tan^2(\gamma)} \sin \beta$$

Grâce à [4.34], on obtient :

$$\frac{s_3}{s_0} = \sin(2\varepsilon)$$

Calculons le rapport entre s_1 et s_0 :

$$\frac{s_1}{s_0} = \frac{A^2 - B^2}{A^2 + B^2}$$

En utilisant la définition de l'angle γ :

$$\frac{s_1}{s_0} = \frac{1 - \tan^2 \gamma}{1 + \tan^2 \gamma} = \cos(2\gamma)$$

D'où, d'après [4.32] :

$$\frac{s_1}{s_0} = \cos(2\tau) \cos(2\varepsilon)$$

De la même façon, on trouve :

$$\frac{s_2}{s_0} = \sin(2\tau) \cos(2\varepsilon)$$

Les paramètres de Stokes s_1 , s_2 , s_3 , normalisés à s_0 sont donc les axes de la sphère de Poincaré.

Le vecteur de Stokes définissant l'état de polarisation du rayonnement est constitué par ces paramètres :

$$\begin{pmatrix} s_0 \\ s_1 \\ s_2 \\ s_3 \end{pmatrix} = \begin{pmatrix} I \\ I \cos(2\tau) \cos(2\varepsilon) \\ I \sin(2\tau) \cos(2\varepsilon) \\ I \sin(2\varepsilon) \end{pmatrix}$$

Une polarisation rectiligne horizontale est caractérisée par $(1100)^t$.

Une polarisation rectiligne verticale est caractérisée par $(1, -1, 0, 0)^t$.

Une polarisation rectiligne à 45° est caractérisée par $(1, 0, 1, 0)^t$.

Une polarisation circulaire gauche est caractérisée par $(1, 0, 0, 1)^t$.

■ Co-polarisation et polarisation croisée

On définit la co-polarisation comme la polarisation attendue en rapport avec l'orientation du champ électrique rayonné par l'antenne.

Les différentes formes des antennes et leur environnement créent des perturbations qui génèrent un rayonnement dans une direction de polarisation orthogonale à celle attendue. Cette polarisation est appelée polarisation croisée.

On mesure le rapport des puissances de ces deux polarisations. C'est une indication du degré de pureté du rayonnement en termes de polarisation.

4.7 Impédance d'une antenne

Les éléments essentiels d'un émetteur sont un générateur d'ondes connecté à une ligne de propagation, elle-même reliée à l'antenne. Chacune de ces trois parties présente une impédance propre complexe. Pour des antennes réciproques, l'impédance de l'antenne est la même en émission ou en réception.

4.7.1 Impédance d'entrée

On appelle impédance d'entrée de l'antenne l'impédance vue à l'entrée de ce composant. Elle est représentée par :

$$Z_A = R_A + jX_A \quad [4.36]$$

L'impédance de l'antenne est influencée par les objets environnants, en particulier par des objets ou des plans métalliques proches ou par d'autres antennes. Dans ce dernier cas, on parle d'impédances mutuelles entre éléments rayonnants.

Nous ne traiterons dans ce paragraphe que de l'impédance propre de l'antenne, c'est-à-dire celle de l'antenne placée seule et rayonnant dans l'espace vide infini.

La résistance d'entrée R_A représente un terme de dissipation. Il est lié, d'une part à la puissance rayonnée et d'autre part, à la puissance perdue par effet Joule. Cette dernière est en général petite par rapport à la puissance rayonnée pour assurer le fonctionnement optimal de l'antenne. Cependant les pertes par effet Joule peuvent représenter des valeurs non négligeables en fonction de la géométrie de l'antenne. Les pertes dans le plan de masse sont aussi à prendre en compte.

La réactance X_A est liée à la puissance réactive stockée au voisinage de l'antenne.

Rappelons la définition de la puissance dissipée dans une antenne :

$$P_e = \frac{1}{2} R_A I_A^2$$

La résistance d'antenne est constituée des deux termes correspondant d'une part au rayonnement R_r et d'autre part à l'effet Joule R_j :

$$R_A = R_r + R_j$$

La puissance rayonnée est donnée par :

$$P_r = \frac{1}{2} R_r I_A^2$$

On rappelle, à titre d'exemple, la valeur de la résistance de rayonnement d'un dipôle élémentaire [3.16] :

$$R_r \approx 800 \frac{l^2}{\lambda^2}$$

Cette résistance, très faible pour une longueur l petite du dipôle par rapport à la longueur d'onde λ , entraîne une puissance de rayonnement faible. Les pertes Joule sont alors à évaluer précisément. L'efficacité d'une antenne est définie par le rapport de la puissance rayonnée à la puissance à l'entrée de l'antenne :

$$\eta = \frac{P_r}{P_e}$$

C'est un rendement. Il s'exprime par rapport aux différentes résistances par :

$$\eta = \frac{R_r}{R_r + R_j} \quad [4.37]$$

Évaluons la résistance liée à la puissance perdue par effet Joule pour un dipôle. Aux fréquences de travail, le courant est considéré comme surfacique en raison de l'effet de peau. La résistance de surface a pour expression :

$$R_S = \sqrt{\frac{\omega\mu}{2\sigma}}$$

Avec ω la pulsation de l'onde, μ la perméabilité et σ la conductivité du conducteur. La résistance du dipôle s'exprime en fonction de sa longueur l et du rayon a du fil :

$$R_J = \frac{l}{2\pi a} R_S$$

Pour un dipôle de longueur d'un dixième de longueur d'onde à 1 GHz, et de rayon égal à un millimètre, la résistance est égale à :

$$R_J \approx 0,03 \, \Omega$$

Pour ce dipôle, la résistance de rayonnement est d'environ $8 \, \Omega$. Les pertes par effet Joule sont négligeables. Remarquons que la résistance de rayonnement croît comme le carré de la longueur, alors que la résistance ohmique croît comme la longueur du dipôle. Pour des dipôles petits, l'efficacité est donc moindre.

4.7.2 Adaptation

L'impédance d'entrée de l'antenne est utilisée pour insérer cet élément de façon optimale dans la chaîne de l'émetteur (ou du récepteur). Si l'impédance caractéristique de la ligne de propagation est Z_0 et l'impédance d'entrée de l'antenne Z_A , le signal se réfléchit à l'entrée de l'antenne avec un coefficient Γ dont l'expression est :

$$\Gamma = \frac{Z_A - Z_0}{Z_A + Z_0}$$

Dans le cas où le coefficient de réflexion est non nul, un système d'ondes stationnaires apparaît et la puissance émise par le générateur n'est pas transmise de façon optimale à l'antenne.

C'est donc ce cas de réflexion minimale à l'entrée de l'antenne qu'on visera. Il correspond à un paramètre de transmission de la matrice de répartition (S_{21}) proche de 1 et à un coefficient de réflexion (S_{11}) proche de 0 (en valeurs linéaires). On admet qu'une bonne adaptation est obtenue lorsque le coefficient de réflexion est inférieur à -10 dB. Cela correspond à un rapport d'ondes stationnaires (VSWR, *Voltage Standing Wave Ratio*), compris entre 1 et 1,2. Le rapport d'ondes stationnaires est défini comme le rapport de la tension maximale à la tension minimale sur une ligne. Pour améliorer l'adaptation d'une antenne, tous les moyens associés aux techniques hyperfréquences sont utilisables. En particulier, pour des antennes planaires, la gravure de lignes laisse une grande marge pour la conception, et les tronçons de lignes d'adaptation sont couramment introduits pour permettre l'adaptation entre la ligne et l'antenne. Des tronçons de lignes en parallèle peuvent aussi être ajoutés pour modifier la réactance.

Dans le cas des antennes filaires, des composants localisés sont ajoutés à l'entrée de l'antenne.

L'impédance interne du générateur intervient dans le fonctionnement d'un émetteur. Le schéma électrique simplifié de ce système peut se résumer, d'une part, à la partie génératrice ayant une impédance Z_G et d'autre part à l'antenne associée à sa ligne d'accès ayant une impédance Z .

La puissance maximale transmise à l'antenne est obtenue lorsque l'impédance du générateur est égale à l'impédance

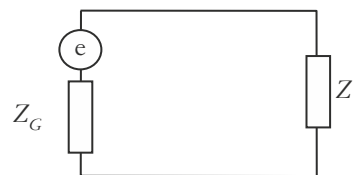


Figure 4.14 – Schéma électrique d'une antenne en présence du générateur.

conjuguée de l'antenne :

$$Z_G = Z_*$$

Dans la majorité des cas, l'impédance du générateur vaut $50 \, \Omega$. C'est une impédance standard en hyperfréquences.

4.8 Alimentation

Comme nous venons de le voir, l'adaptation est un point fondamental dans la conception des antennes. Nous allons présenter quelques exemples montrant l'influence de la forme du dispositif d'alimentation.

Lorsque la taille de l'antenne est très petite par rapport à la longueur d'onde, l'impédance d'entrée est fixée, comme nous l'avons vu dans le cas du dipôle. Lorsque la taille de l'antenne est plus grande, l'impédance d'entrée dépend du point d'excitation. Les antennes planaires qui ont une dimension de l'ordre de la demi-longueur d'onde, afin de permettre au phénomène de résonance de s'établir, présentent des exemples variés d'alimentations. Il est possible de choisir différentes formes d'alimentation et différents points d'excitation.

■ Excitation par une ligne coaxiale

Un patch rectangulaire (figure 4.15) est formé d'un rectangle de métal très fin déposé sur un diélectrique de faible permittivité. De l'autre côté du diélectrique, une couche métallique très fine tient lieu de plan de masse. L'alimentation, pour le cas présenté sur cette figure, est assurée par un coaxial. Pour cela une cavité est usinée dans le diélectrique et le plan de masse, de façon à faire passer l'âme du coaxial qui est soudée sur le plan métallique rayonnant (patch). Le dispositif permettant la traversée s'appelle un *via hole*.

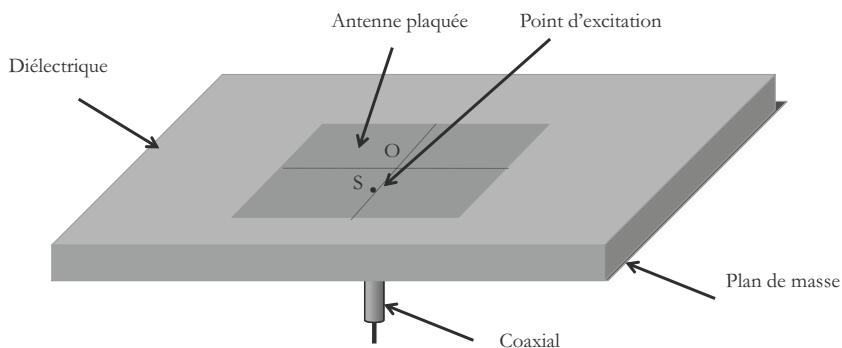


Figure 4.15 – Alimentation d'une antenne plaquée rectangulaire par coaxial.

Le point S d'excitation se trouve sur une des médiatrices du rectangle. On montre que l'impédance d'entrée augmente pour un point d'excitation s'écartant du centre de l'antenne. Ceci permet d'avoir une certaine latitude pour le choix du point d'excitation afin d'obtenir la meilleure adaptation.

Si le patch est presque carré et excité sur sa diagonale (figure 4.16), il est possible d'obtenir une excitation circulaire. Le montage de la ligne coaxiale est du même type que dans l'exemple précédent.

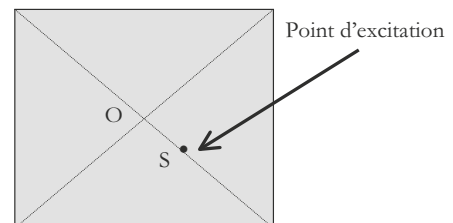


Figure 4.16 – Excitation d'un patch presque carré sur sa diagonale pour l'obtention d'une polarisation circulaire.

Une excitation circulaire est aussi obtenue à partir d'un patch circulaire alimenté par un coaxial (figure 4.17).

Dans les différents cas présentés, l'adaptation de l'antenne est très sensible à la position du point d'excitation. En simulation, on constate que le coefficient de réflexion varie très rapidement en fonction de la géométrie. C'est un point à étudier finement. De plus il est rare que, lors de la réalisation, il n'y ait pas une légère différence entre la position souhaitée de l'excitation et sa position réelle, en raison de la taille du point de soudure et de la réalisation du *via hole*.

Ce type d'alimentation est mécaniquement délicat à réaliser. Pour ces différentes raisons, on choisit, lorsque c'est possible, une excitation par ligne micro ruban.

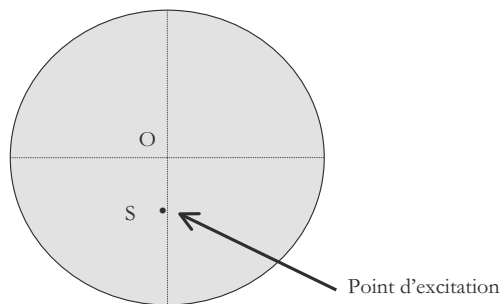


Figure 4.17 – Excitation d'un patch circulaire pour l'obtention d'une polarisation circulaire.

■ Alimentation par ligne micro ruban

L'alimentation la plus simple consiste à utiliser une ligne micro ruban sur le même plan que le patch rayonnant (figure 4.18).



Figure 4.18 – Alimentation par ligne micro ruban.

Cette disposition présente un inconvénient si la ligne rayonne. C'est le cas en très haute fréquence. Le rayonnement de la ligne perturbe alors celui de l'antenne qui ne présente pas la même pureté de polarisation. Cependant pour les cas usuels, cette technique très utilisée, présente le grand avantage de la simplicité de réalisation.

Pour améliorer l'adaptation entre la ligne micro ruban et l'antenne, il est courant de réaliser des encoches (figure 4.19) dont la taille est à calculer afin d'obtenir une meilleure adaptation.

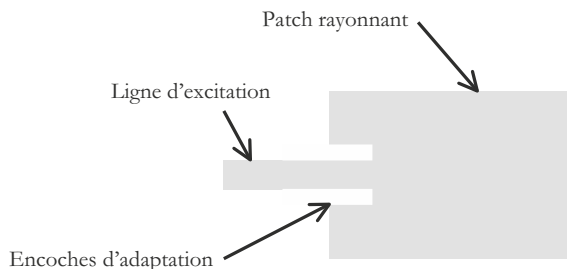


Figure 4.19 – Encoches d'adaptation.

La polarisation circulaire est obtenue avec ce type d'alimentation en excitant un patch presque carré selon sa diagonale (figure 4.20)

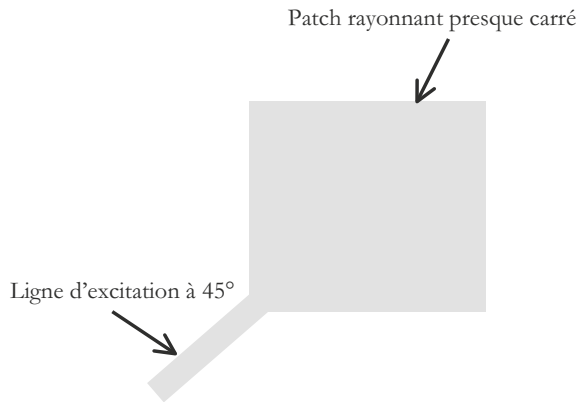


Figure 4.20 – Alimentation pour un rayonnement à polarisation circulaire.

Une autre façon d'obtenir la polarisation circulaire est de réaliser deux lignes d'accès à 90°, débouchant symétriquement sur un patch carré, dont l'alimentation présente un déphasage de 90° (figure 4.21).

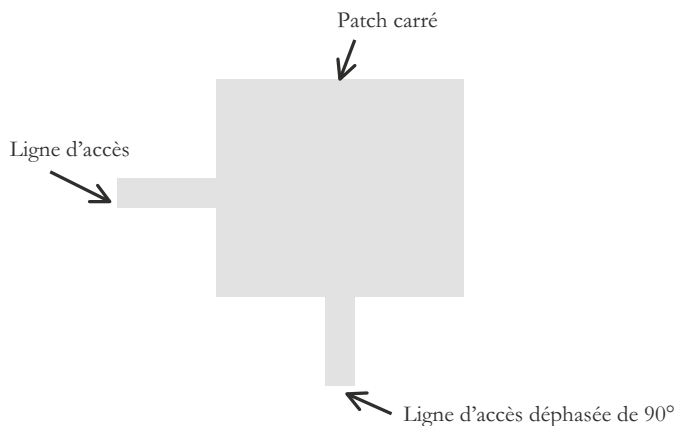


Figure 4.21 – Alimentation d'un patch carré par deux lignes déphasées de 90° pour l'obtention d'un rayonnement circulaire.

■ Alimentation par fente

L'alimentation par une ligne micro ruban sur le même plan que l'élément rayonnant est, dans certains cas, un désavantage par rapport à la qualité du rayonnement. La solution à ce problème consiste à placer la ligne d'excitation sur un plan inférieur (figure 4.22). La ligne écrantée par le plan de masse ne rayonne pratiquement pas.

La ligne est gravée sur la face arrière d'un diélectrique possédant un plan de masse sur la face avant, dans lequel une ouverture est pratiquée. Le plan de masse est recouvert par un deuxième diélectrique au-dessus duquel est gravé l'élément rayonnant métallique. Le diélectrique se trouvant au-dessus est généralement de faible permittivité afin de favoriser le rayonnement. Le diélectrique

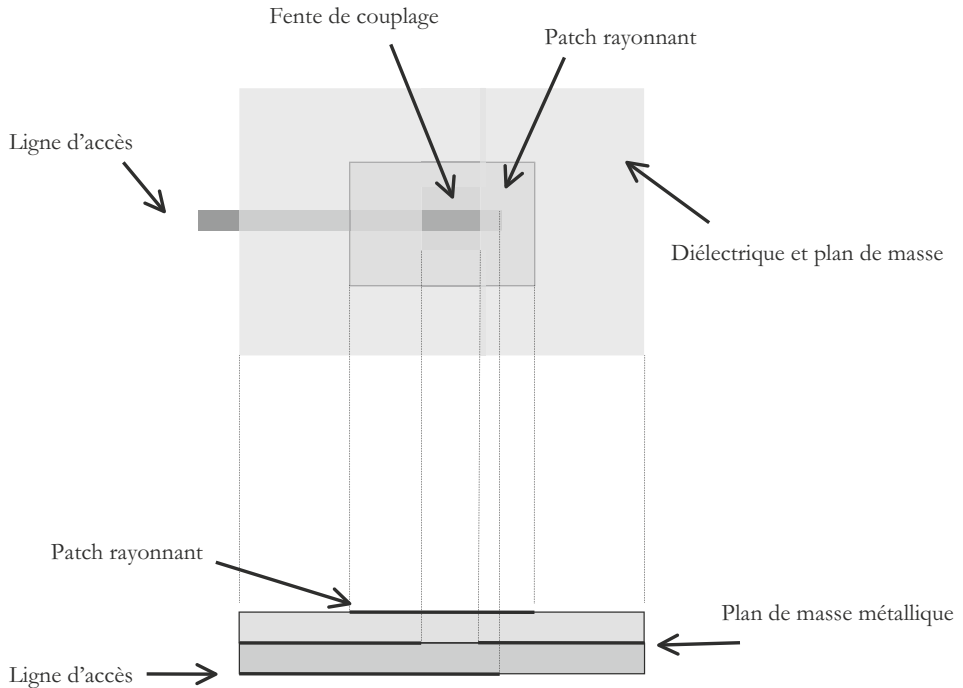


Figure 4.22 – Alimentation par fente.

supportant la ligne est de permittivité plus élevée afin de jouer le rôle d'écran de la ligne d'excitation et de concentrer le champ électrique. Le couplage existant est un couplage magnétique à travers l'ouverture. Cette ouverture introduit un élément inductif qui est compensé par l'effet capacitif plus ou moins prononcé de la ligne micro ruban. Relativement à ce dernier paramètre, il suffit d'ajuster la longueur de la ligne afin d'obtenir l'adaptation adéquate.

Ces différentes descriptions des formes d'adaptation entre l'élément rayonnant et la ligne d'accès montrent la grande variété des solutions qui s'offrent lors de la conception.

Cette description pourrait être encore longue, car les dispositifs d'adaptation dépendent à la fois des lignes d'accès et de l'élément rayonnant. Il existe par exemple des techniques différentes pour les accès coplanaires.

La conception d'une antenne ne peut être faite indépendamment de sa ligne d'adaptation. Dans les simulations numériques sur lesquelles s'appuie la conception, la partie de la ligne d'excitation est modélisée en même temps que le dispositif rayonnant.

4.9 Largeur de bande

La largeur de bande, appelée aussi bande passante, d'une antenne définit le domaine de fréquences dans lequel le rayonnement de l'antenne présente les caractéristiques requises. Il s'agit la plupart du temps de la puissance transmise par l'antenne, mais on peut définir d'autres caractéristiques exigées pour le fonctionnement d'une antenne telle que la polarisation. Il se peut par exemple qu'une polarisation circulaire soit recherchée et obtenue seulement dans une bande de fréquence. Dans la suite, nous allons définir la bande de fréquence relative à la puissance de rayonnement. La valeur des limites sur les critères de fonctionnement de l'antenne définit un domaine de fréquences situé entre une valeur minimale f_1 et une valeur maximale f_2 . La bande de fréquence

Δf est définie par la différence entre ces deux fréquences :

$$\Delta f = f_2 - f_1$$

La largeur relative de bande B_r est un pourcentage exprimant le rapport de la bande à la fréquence centrale f_0 .

$$B_r = \frac{f_2 - f_1}{f_0}$$

La largeur de bande est aussi définie par le rapport entre les deux fréquences extrêmes :

$$B_f = \frac{f_2}{f_1}$$

Dans le cas d'antennes de faible largeur de bande, la première définition est plus utilisée, alors que la seconde est plutôt utilisée pour les antennes larges bandes.

Pour connaître la largeur de bande d'une antenne relativement au rayonnement, on trace le paramètre S_{11} de réflexion en fonction de la fréquence. On admet généralement que si ce paramètre est inférieur à -10 dB, la puissance de rayonnement est suffisante. Il suffit alors de repérer sur la courbe les valeurs de la fréquence correspondant à cette valeur.

Les antennes résonantes ont généralement des largeurs de bandes faibles. Une antenne de type dipôle d'une longueur égale à la demi-longueur d'onde a une largeur de l'ordre de 10 %.

Le paramètre de réflexion à l'entrée d'une antenne planaire simple est donné sur la figure 4.23. Définissant la bande de fréquence pour un coefficient de réflexion inférieur à 10 dB, les fréquences limites de la bande sont :

$$f_1 = 7,27 \text{ GHz} \quad \text{et} \quad f_2 = 7,51 \text{ GHz}$$

La bande relative est donc un peu supérieure à 3 %.

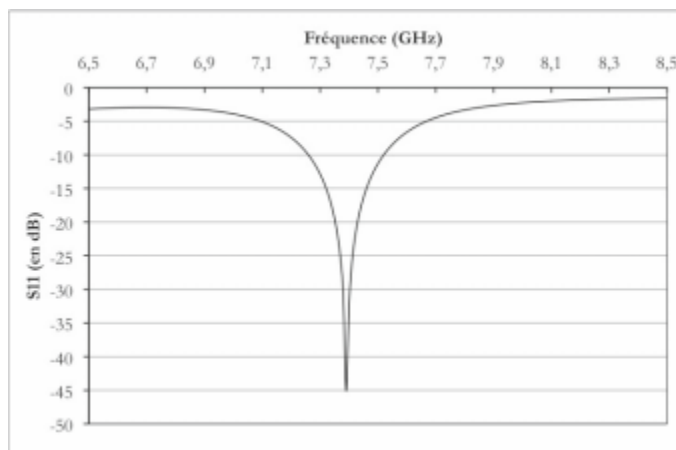


Figure 4.23 – Coefficient de réflexion à l'entrée d'une antenne patch en fonction de la fréquence.

Dans le cas d'une antenne à résonance, la fréquence centrale est déterminée par les dimensions de l'antenne et les matériaux la constituant. À la résonance, l'impédance d'entrée de l'antenne est réelle. Si la fréquence s'écarte légèrement de cette fréquence centrale, la partie réelle varie et la partie imaginaire devient différente de zéro. L'adaptation de l'antenne, en général conçue pour

la fréquence centrale, n'est plus parfaite de part et d'autre de celle-ci. La désadaptation entraîne alors une limite de fonctionnement en fréquence.

Les antennes à ondes progressives ont des grandes largeurs de bande. On donne la dénomination d'antenne large bande lorsque le rapport B_f est supérieur à 2, ce qui correspond à une octave.

Signalons enfin que les phénomènes de couplage élargissent la bande passante. Pour créer une largeur de bande plus grande il est possible d'associer deux résonateurs de fréquences de résonance proches.

4.10 Bilan de liaison

L'utilisation d'antennes se justifie pour transmettre une onde d'un point à un autre. Cette onde, généralement appelée la porteuse, est le support de l'information transportée. Il est donc indispensable de connaître les caractéristiques de sa transmission.

Dans ce chapitre, nous étudions le cas simple de deux antennes, en espace libre, en regard l'une de l'autre, correctement orientées, c'est-à-dire que leur axe de rayonnement est commun et enfin on suppose qu'elles sont bien orientées par rapport à leur polarisation. On suppose qu'elles sont suffisamment espacées pour considérer qu'elles sont en champ lointain l'une de l'autre.

Ce n'est souvent qu'une approximation car le rayonnement d'une antenne peut donner lieu à des réflexions multiples sur des objets se trouvant sur son trajet. Ces cas relèvent des études de propagation qui seront évoqués ultérieurement.

Cependant, le cas du rayonnement en espace libre est très important car il traduit les propriétés propres aux antennes. C'est pourquoi lorsqu'on s'intéressera aux mesures d'antennes (chapitre 10), les diagrammes de rayonnement seront donnés en espace libre. Deux cas seront envisagés :

- celui de l'espace libre effectif, difficile à obtenir car nécessitant un espace important,
- celui de l'espace libre équivalent, obtenu dans une chambre anéchoïde.

La démonstration qui suit permet de connaître la valeur de la puissance reçue par une antenne en fonction de la puissance d'émission, en tenant compte de la distance des antennes et de leur gain. Les antennes sont supposées en champ lointain l'une de l'autre dans tout ce paragraphe.

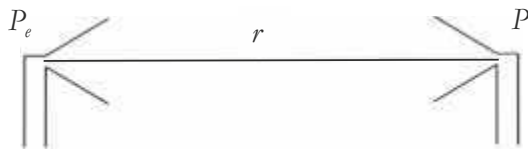


Figure 4.24

Considérons une antenne d'émission dont la puissance d'émission est P_e et le gain G_e , la densité surfacique de puissance au voisinage de l'antenne de réception est donnée par l'expression [4.22]. Soit E_0 la norme du champ électrique au voisinage de l'antenne de réception, on obtient :

$$\frac{E_0^2}{2Z} = \frac{P_e}{4\pi r^2} G_e$$

La puissance reçue par l'antenne P_r est proportionnelle à la densité surfacique de puissance et à la surface d'absorption de l'antenne de réception S_{ar} :

$$P_r = \frac{E_0^2}{2Z} S_{ar}$$

Selon la définition du paragraphe 4.4, la surface d'absorption est liée au gain par :

$$S_{ar} = \frac{\lambda^2}{4\pi} G_r$$

On en déduit la relation entre la puissance reçue et la puissance émise :

$$\frac{P_r}{P_e} = \frac{\lambda^2}{(4\pi r)^2} G_e G_r$$

Les gains dans cette formule sont des gains linéaires.

Prenons l'exemple d'une puissance d'émission de 2 W, émise par une antenne ayant un gain de valeur 2 et un récepteur, situé à 100 m, ayant un gain de valeur 10. La fréquence de fonctionnement vaut 1 GHz. Cela correspond à une longueur d'onde de 0,3 m. La puissance reçue dans ce cas est :

$$P_r = 2,3 \mu\text{W}$$

Si l'on introduit la surface d'absorption S_{ae} de l'antenne d'émission qui est liée au gain d'émission selon :

$$S_{ae} = \frac{\lambda^2}{4\pi} G_e$$

le rapport entre la puissance reçue et la puissance émise prend la forme :

$$\frac{P_r}{P_e} = \frac{S_{ae} S_{ar}}{r^2 \lambda^2}$$

■ Formule de Friis

Cette dernière relation exprimée en échelle logarithmique s'appelle la formule de Friis. Elle s'exprime sous la forme suivante, lorsque les puissances sont exprimées en dBm, les gains en dB, la distance en km, la fréquence en MHz :

$$P_r = P_e + G_e + G_r - 20 \log r - 20 \log f - 32,44$$

4.11 Température de bruit d'une antenne

Une antenne reçoit tous les rayonnements provenant d'objets émetteurs qui l'entourent. Elle est plus sensible au rayonnement provenant dans un angle solide proche de son axe et égal à l'angle solide de l'antenne. En dehors de l'angle solide de l'antenne, le rayonnement peut être reçu, dans une moindre mesure, dans la direction des lobes secondaires.

En général, les antennes fonctionnent deux à deux : l'une émet, l'autre reçoit. Un signal utile est transmis entre les deux antennes. Lorsqu'une antenne reçoit de la puissance en provenance d'autres sources qui viennent perturber le fonctionnement de la transmission, cette puissance constitue du bruit. L'ensemble du bruit pondéré par la fonction caractéristique de rayonnement constitue la puissance de bruit B captée par l'antenne. Elle s'exprime pour une bande de fréquence donnée Δf par :

$$B = k T_B \Delta f$$

$k = 1,3810^{-23} \text{ J K}^{-1}$ est la constante de Boltzmann.

La température T_B est appelée température de bruit de l'antenne. La notion de température d'antenne est reprise en détail dans le paragraphe 6.5.5.

Un satellite reçoit par exemple du bruit en provenance de la Terre, qui émet un rayonnement comme tout corps ayant une température non nulle, mais aussi en provenance de l'atmosphère et du ciel. Ce dernier correspond au rayonnement extragalactique de l'ordre de 3 K, bien connu des astrophysiciens.

Bibliographie

BALANIS C.A. – *Advanced Engineering Electromagnetics*, New Jersey, A. John Wiley & Sons, Inc., 1989.

BALANIS C.A. – *Antenna Theory Analysis and Design*, New Jersey, A. John Wiley & Sons, Inc., 2005.

COMBES P.F. – *Micro-ondes 2. Circuits passifs, propagation, antennes*, Paris, Dunod, 1997.

ELIOTT R.S. – *Antenna Theory and design*, IEEE Press, New Jersey, A. John Wiley & Sons, Inc., 2003.

STUTZMANN W. L., GARY A.T. – *Antenna Theory and Design*, New Jersey, A. John Wiley & Sons, Inc., 1998.

ULABY F. T., MOORE R. K., FUNG A.K. – *Microwave Remote Sensing, Active and Passive, Vol. I*, Norwood, Artech House, 1981.

5 • RÉSEAUX D'ANTENNES

5.1 Les réseaux d'antennes

Un réseau d'antennes est par définition l'association régulière d'antennes identiques pour créer un rayonnement de forme particulière. La puissance rayonnée est donc plus grande puisqu'on multiplie le nombre d'éléments rayonnants. Le rayonnement résulte de l'addition en phase des champs provenant de chaque élément. Les combinaisons possibles sont donc nombreuses et entraînent une grande souplesse dans la conception de réseaux.

Les applications des réseaux d'antennes sont nombreuses et utilisent tout type d'éléments : cornets, antennes filaires, antennes plaquées, etc.

Le réseau occupant un espace plus important que l'antenne élémentaire, son diagramme de rayonnement est plus étroit puisque sa directivité augmente avec sa surface, selon l'équation [4.14]. On parvient facilement à augmenter le gain de l'antenne élémentaire de 10 à 15 dB. Le réseau est donc globalement plus puissant et plus directif que l'antenne élémentaire.

Un autre avantage du réseau d'antennes tient au fait que le choix d'un déphasage régulier entre les éléments fixe une orientation du faisceau, dans l'espace dans certaines limites d'angles.

L'organisation spatiale des antennes d'une part, et le mode d'alimentation de chacune des antennes d'autre part, confèrent au réseau des propriétés de rayonnement bien définies. Ces propriétés sont modifiables dans certains cas, grâce essentiellement à la possibilité d'agir sur la phase et l'amplitude de l'alimentation de chaque antenne. On obtient alors un réseau d'antennes reconfigurable.

Lors de la conception de réseau, le couplage entre les antennes élémentaires est un point délicat car ce couplage modifie légèrement les caractéristiques de rayonnement et d'adaptation. En particulier la bande passante du réseau est un peu plus large que celle de l'antenne élémentaire du fait des couplages.

La figure 5.1 donne la configuration pour l'étude générale d'un réseau. Les coordonnées sphériques sont utilisées dans les mêmes conditions que dans les chapitres précédents.

La difficulté de modélisation d'un réseau est due à sa taille, qui peut être importante lorsque le nombre d'éléments est grand.

Comme nous allons le voir plus loin dans différentes applications, le réseau permet de reconstituer une loi d'illumination sur l'ouverture et d'orienter le faisceau dans une direction particulière.

Idéalement une conception de réseau commence par la synthèse, suivie d'une modélisation du facteur de réseau. Une simulation globale est ensuite nécessaire. Elle peut être associée à des mesures.

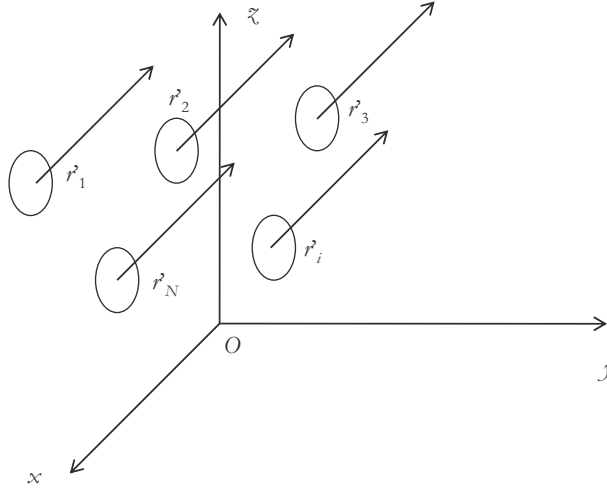


Figure 5.1 – Exemple de la géométrie d'un réseau d'antennes.

5.1.1 Facteur de réseau

Considérons un réseau à N éléments identiques, ayant la même orientation. L'antenne affectée de l'indice i est repérée par le vecteur $\vec{r}_i$. On note R_i la distance entre cette antenne et le point d'observation M, placé à grande distance.

Le centre de phase d'une antenne est son point de référence pour la phase. On le choisit, généralement, de façon à vérifier une répartition équilibrée des phases par rapport à ce point. L'antenne élémentaire a un centre de phase. Le réseau a aussi son centre de phase.

Considérons une antenne de référence placée à l'origine O , indépendamment du fait qu'une antenne réelle y soit placée. O est choisi comme centre de phase pour cette antenne et pour le réseau. Le champ électrique créé par cette antenne est donné par :

$$\vec{E}_0(r, \theta, \phi) = \vec{f}(\theta, \phi) \frac{e^{-jkr}}{4\pi r}$$

Toutes les antennes sont physiquement identiques. Elles sont alimentées avec une amplitude c_i et une phase α_i . Le champ total rayonné est la somme vectorielle des champs créés par chacune des antennes :

$$\vec{E}(r, \theta, \phi) = \sum_{i=1}^N c_i e^{j\alpha_i} \vec{f}(\theta, \phi) \frac{e^{-jkR_i}}{4\pi R_i} \quad [5.1]$$

Le point d'observation M étant situé à grande distance, la relation entre les grandeurs géométriques devient :

$$R_i \approx r - \vec{u} \cdot \vec{r}_i$$

En introduisant cette relation dans l'expression du champ total :

$$\vec{E}(r, \theta, \phi) = \sum_{i=1}^N c_i e^{j\alpha_i} \vec{f}(\theta, \phi) \frac{e^{-jk(r - \vec{u} \cdot \vec{r}_i)}}{4\pi r_i}$$

Comme nous l'avons remarqué au chapitre 4, la fonction inverse varie beaucoup moins vite que la fonction exponentielle. Ceci nous permet d'écrire :

$$\vec{E}(r, \theta, \phi) = \frac{\vec{f}(\theta, \phi)}{4\pi r} e^{-jkr} \sum_{i=1}^N c_i e^{j\alpha_i + jk\vec{u} \cdot \vec{r}_i} \quad [5.2]$$

Rappelons que la densité surfacique de puissance dans une direction est proportionnelle à la norme du champ électrique au carré, selon l'expression [2.42]. Le facteur de réseau F_R est alors défini par :

$$F_R(\theta, \phi) = \left| \sum_{i=1}^N c_i e^{j\alpha_i + jk\vec{u} \cdot \vec{r}_i} \right|^2 \quad [5.3]$$

Ce facteur de réseau est généralement normalisé en divisant cette expression par la valeur du maximum de F_R .

La densité surfacique de courant s'exprime donc sous la forme :

$$S_r = \frac{1}{2Z} \frac{|\vec{f}(\theta, \phi)|^2}{(4\pi r)^2} \left| \sum_{i=1}^N c_i e^{j\alpha_i + jk\vec{u} \cdot \vec{r}_i} \right|^2 \quad [5.4]$$

En reprenant la notation [4.2], définissant la fonction caractéristique $F(\theta, \phi)$ de l'antenne élémentaire, la densité de puissance est égale à :

$$S_r = \frac{1}{2Z} \frac{F(\theta, \phi)}{r^2} F_R(\theta, \phi)$$

La fonction caractéristique de l'antenne constituée par le réseau est donc :

$$F_{\text{réseau}}(\theta, \phi) = F(\theta, \phi) F_R(\theta, \phi)$$

Cette fonction est le résultat de la fonction caractéristique de l'antenne, modulée par le facteur de réseau. Le diagramme de rayonnement du réseau est obtenu par la représentation de la fonction caractéristique normalisée :

$$F_n(\theta, \phi) = \frac{F(\theta, \phi) F_R(\theta, \phi)}{F_{\text{max}}(\theta, \phi) F_{R\text{max}}(\theta, \phi)} \quad [5.5]$$

Le calcul précédent suppose que toutes les antennes sont rigoureusement identiques, sans couplage entre elles. En fait les antennes présentent deux à deux un couplage, caractérisé par une impédance mutuelle. Celle-ci est plus grande pour des antennes proches que pour des antennes éloignées et dépend de la position des antennes élémentaires dans le réseau.

La complexité de conception d'un réseau augmente rapidement avec son nombre d'éléments. Certains réseaux peuvent, en effet, comporter plus d'une centaine d'éléments.

5.1.2 Exemples de calcul de réseaux

■ Réseau uniforme à une dimension

Le principe de calcul du diagramme de rayonnement d'un réseau est appliqué dans ce paragraphe à un réseau linéaire, à $N + 1$ éléments. On dispose les éléments régulièrement espacés de la distance d , le long de l'axe Ox , en plaçant le premier élément en O . Pour fixer la structure on supposera que l'antenne élémentaire est un dipôle, ce qui n'enlève rien à la généralité de la démonstration, puisqu'on se focalise, dans cette partie, sur le facteur de réseau.

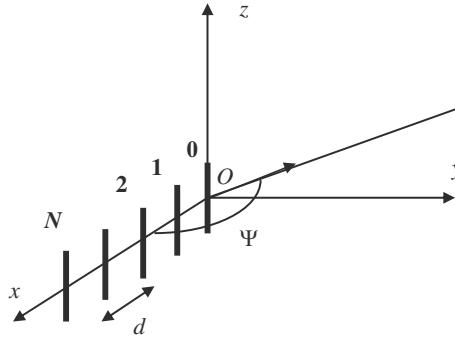


Figure 5.2 – Réseau linéaire de $N + 1$ éléments rayonnants.

Dans cet exemple, chaque dipôle est alimenté avec la même amplitude, avec une phase variant linéairement selon x :

$$\alpha_n = n\alpha d$$

Le diagramme de réseau est donné par :

$$F_R(\theta, \phi) = I_0 \left| \sum_{n=0}^N e^{j(n\alpha d + jknd \cos \psi)} \right|^2$$

$$\text{avec } \cos \psi = \sin \theta \cos \varphi$$

Il apparaît dans le calcul du facteur de réseau la suite géométrique donnée par la formule suivante :

$$\sum_{n=0}^N w^n = \frac{1 - w^{N+1}}{1 - w}$$

Soit, en appliquant ce calcul au facteur de réseau :

$$F_R(\theta, \phi) = I_0 \left| \frac{1 - e^{j(N+1) \cdot (\alpha + k \cos \psi)d}}{1 - e^{j(\alpha + k \cos \psi)d}} \right|^2$$

Donc :

$$F_R(\theta, \phi) = I_0 \left| \frac{\sin [(N+1) (\alpha + k \cos \psi)d/2]}{\sin [(\alpha + k \cos \psi)d/2]} \right|^2$$

Posons :

$$u = kd \cos \psi \quad \text{et} \quad u_0 = \alpha d$$

On obtient :

$$F_R(\theta, \phi) = I_0 \left| \frac{\sin [(N+1) (u + u_0)/2]}{\sin [(u + u_0)/2]} \right|^2$$

Le facteur d'antenne est maximum pour :

$$u = -u_0 \quad \text{et} \quad (u + u_0)/2 = m\pi$$

avec m entier. À ces valeurs correspondent des maxima principaux du diagramme de rayonnement. Pour ces maxima, le facteur de réseau est égal à :

$$F_R = I_0 (N + 1)^2$$

Cette valeur fait apparaître la puissance résultant du fonctionnement des $N + 1$ éléments du réseau, en phase. Entre les maxima, il existe N minima, obtenus pour :

$$\frac{(N + 1)(u + u_0)}{2} = l\pi$$

Entre ceux-ci la fonction passe par des maxima secondaires, dont le nombre est $N - 1$, qui correspondent aux lobes secondaires du réseau. La fonction représentant le facteur de réseau est tracée sur la figure 5.3 pour $N = 5$, en fonction de u , avec $u_0 = 0$, sans restriction sur le domaine variation de cette variable.

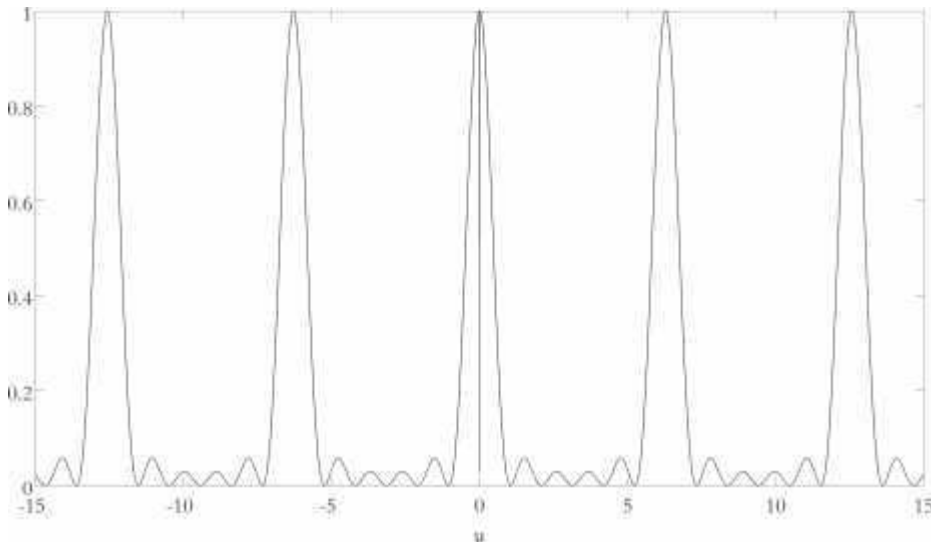


Figure 5.3 – Fonction réseau normalisée pour $N = 5$ et $u_0 = 0$.

Si u_0 est différent de zéro, la courbe de la figure 5.3 est traduite.

Plus le nombre d'éléments du réseau est grand, plus les faisceaux correspondant aux maxima principaux sont fins.

Cependant, le cosinus ayant un domaine de variation compris entre -1 et $+1$ pour des angles réels, la variable u présente un domaine physique de définition, tel que :

$$-kd \leq u \leq kd$$

Cette zone s'appelle la zone de visibilité. En dehors de cette zone, la fonction n'a pas de sens puisqu'elle ne correspond à aucune direction réelle de rayonnement. Sur la figure 5.3, il faudrait donc réduire l'espace de représentation de la fonction à la zone de visibilité. Comme la constante de propagation est liée à la longueur d'onde, cette condition s'écrit sous la forme :

$$-2\pi \frac{d}{\lambda} \leq u \leq 2\pi \frac{d}{\lambda}$$

Ainsi l'espacement des dipôles fixe le nombre de lobes de réseau réels. Le choix de la distance entre éléments est en général tel qu'il n'existe qu'un lobe principal. La variation du paramètre u_0 , c'est-à-dire la variation linéique de phase entre les éléments, fixe la position du lobe principal dans cette zone de visibilité. Ce paramètre permet d'orienter le faisceau dans une direction particulière. Pour ne voir apparaître qu'un seul lobe principal, lorsque $u_0 = 0$, il faut vérifier la condition :

$$d < \lambda$$

Lorsque u_0 est différent de 0, cette condition n'est pas assez restrictive. Pour être sûr de ne voir qu'un lobe principal, il faut la remplacer par :

$$d \leq \frac{\lambda}{2}$$

Le diagramme de rayonnement de la figure 5.4 est celui d'un réseau de six dipôles élémentaires ($N = 5$), espacés de $d = \lambda/2$. La moitié du diagramme est représentée afin de matérialiser les lobes secondaires. Il correspond à des angles θ compris entre 0 et 90°.

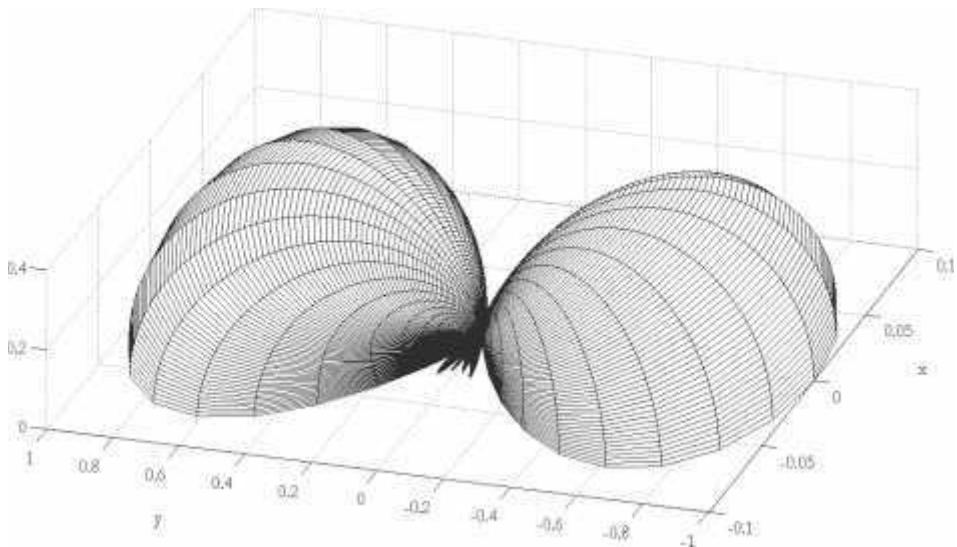


Figure 5.4 – Demi-diagramme de rayonnement normalisé d'un réseau linéaire de six dipôles élémentaires espacés de $d = \lambda/2$, avec $0 \leq \theta \leq 90^\circ$.

Le diagramme de rayonnement a perdu la symétrie axiale du dipôle. Dans la direction perpendiculaire à l'alignement des dipôles (Oy), le rayonnement est maximal. Dans la direction x , les dipôles étant espacés de $\lambda/2$, ils rayonnent en opposition de phase. Cela explique le zéro de rayonnement dans cette direction. On devine les lobes secondaires entre les deux lobes principaux.

Si la distance entre les dipôles est supérieure à la longueur d'onde λ , il apparaît plusieurs lobes principaux. Cette propriété peut être utilisée pour créer une antenne multifaisceaux. La figure 5.5 représente le diagramme de rayonnement des six dipôles de la figure 5.3 avec un espacement de trois demi-longueurs d'onde, toujours en prenant $u_0 = 0$.

Ce diagramme n'étant pas très précis du fait de la perspective, il est représenté dans le plan $\theta = 90^\circ$, soit le plan (xOy) , sur la figure 5.6a. Les lobes secondaires y sont très peu visibles. Un agrandissement les fait apparaître sur la figure 5.6b.

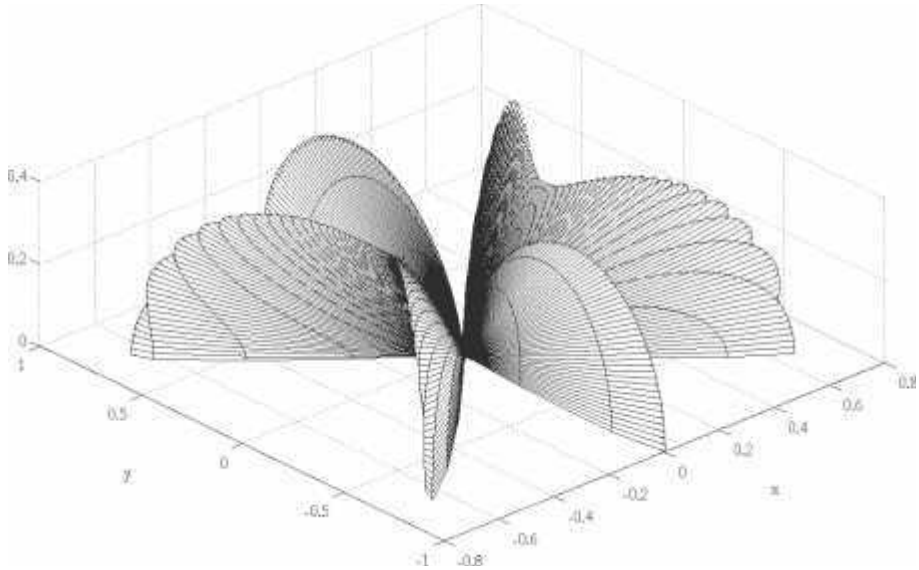


Figure 5.5 – Diagramme de rayonnement d'un réseau linéaire de six dipôles élémentaires espacés de $d = 3\lambda/2$, ($u_0 = 0$).

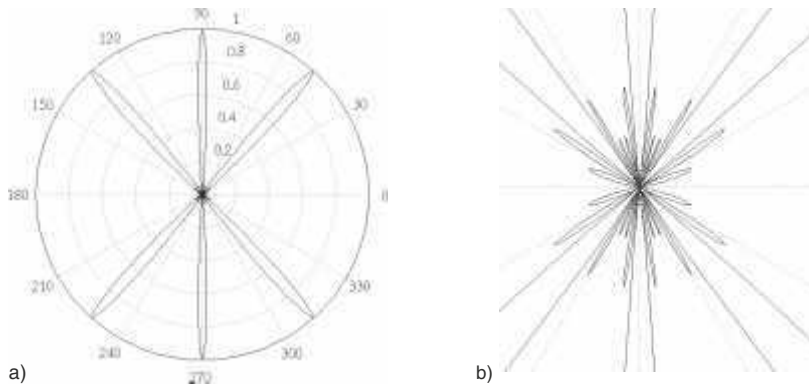


Figure 5.6 – a) Diagramme de rayonnement d'un réseau linéaire de six dipôles élémentaires espacés de $d = 3\lambda/2$, ($u_0 = 0$), dans le plan xOy . b) Agrandissement de la partie centrale du diagramme de la figure a)

■ Réseau à deux éléments

Dans ce paragraphe, la théorie des réseaux est appliquée au cas le plus simple, constitué de deux éléments (figure 5.7). On prendra le cas d'antennes demi-onde afin de comparer les propriétés de ce réseau à celles de l'antenne élémentaire.

Les antennes élémentaires sont alimentées en phase, avec la même amplitude. Le facteur de réseau est alors :

$$F_R(\psi) = |1 + e^{jkd \cos \psi}|^2$$

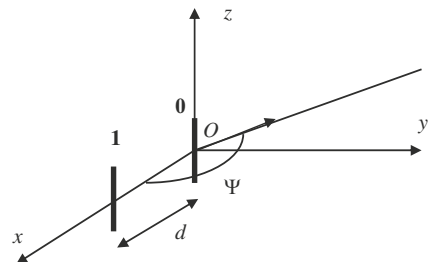


Figure 5.7 – Réseau de deux dipôles rayonnants

Soit :

$$F_R(\psi) = 4 \cos^2 \left(\frac{\pi d}{\lambda} \cos \psi \right)$$

La figure 5.8 représente le facteur de réseau en fonction de :

$$u = kd \cos \psi$$

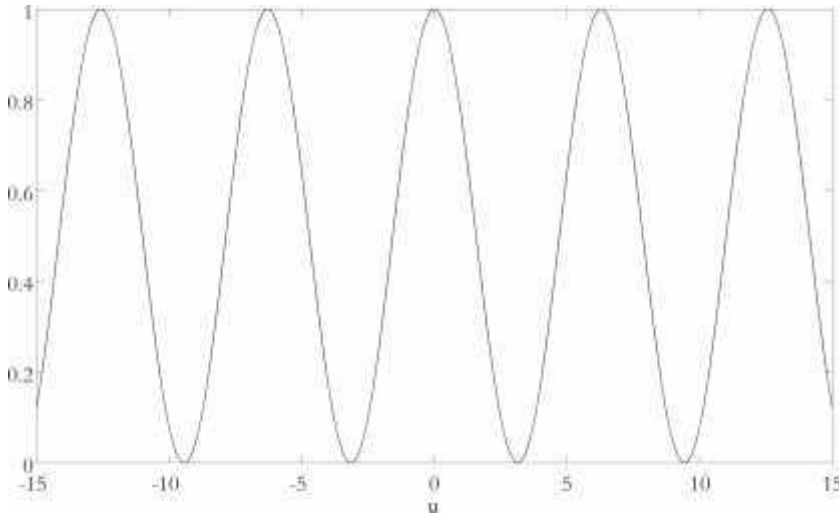


Figure 5.8 – Fonction réseau normalisée pour deux éléments

On retrouve une alternance de maxima de même valeur et de minima. C'est la fonction caractéristique des interférences à deux ondes qu'on retrouve dans l'expérience classique, en optique, des trous d'Young. Les maxima sont situés à :

$$\cos \psi = n \frac{\lambda}{d}$$

La région visible étant déterminée par la condition : $|\cos \psi| \leq 1$, on va diminuer le nombre de maxima dans cette région et le réduire à 1, en prenant :

$$\frac{\lambda}{d} > 1$$

□ Cas particulier de deux antennes demi-onde

Dans ce paragraphe, nous allons montrer comment le facteur de réseau se combine avec le diagramme de rayonnement de l'antenne élémentaire pour obtenir le diagramme de rayonnement global du réseau. Soient deux antennes demi-onde, espacées de $d = \lambda/2$.

On rappelle que la fonction caractéristique de l'antenne demi-onde est donnée par :

$$F(\theta) = \left(\frac{\cos(\pi/2 \cos \theta)}{\sin \theta} \right)^2$$

Cette fonction présente une symétrie axiale en raison de la constitution de l'antenne. Lorsqu'on associe deux antennes de ce type, en décalant leurs axes selon x , comme il est indiqué sur la figure 5.7 la symétrie est rompue, comme on peut le constater sur le diagramme de rayonnement global (figure 5.9).

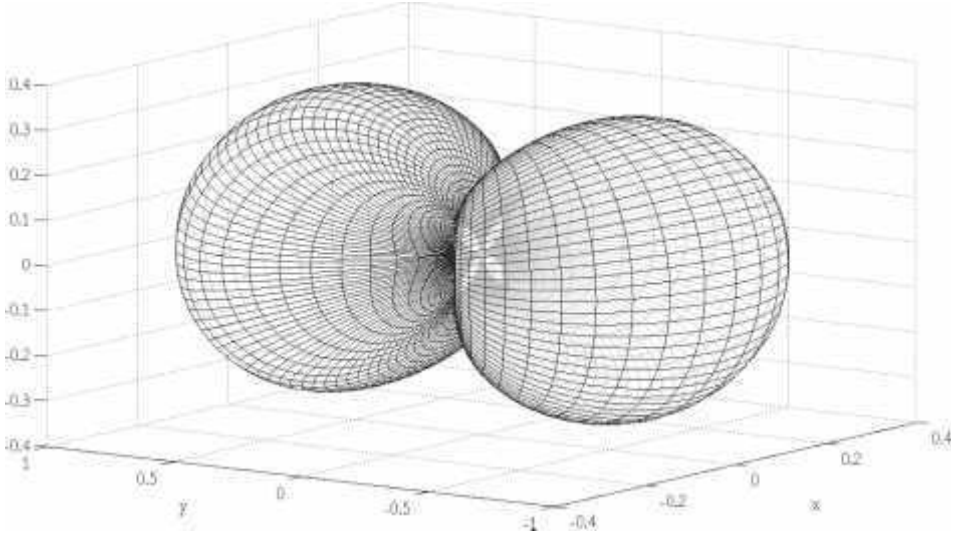


Figure 5.9 – Diagramme de rayonnement normalisé de deux antennes demi-onde placées selon l'axe x distantes de $\lambda/2$.

Le facteur de réseau s'exprime sous sa forme non normalisée :

$$F_R(\theta, \phi) = 4 \cos^2\left(\frac{\pi}{2} \sin \theta \cos \phi\right)$$

Afin d'estimer quantitativement le facteur de réseau dans l'espace, sa représentation est faite dans deux plans perpendiculaires.

Dans le plan $(\vec{Oy}, \vec{Oz})$, défini par : $\phi = \frac{\pi}{2}$, soit $\cos \phi = 0$. C'est un plan E, parallèle au dipôle. Dans ce plan, le facteur de réseau normalisé est constant, $F_R = 1$ (figure 5.10). Le diagramme de rayonnement total aura donc la même forme que celui de l'antenne élémentaire dans ce plan. La fonction caractéristique de rayonnement des deux antennes dans ce plan est donnée par :

$$F_n(\theta, \frac{\pi}{2}) = \left(\frac{\cos(\pi/2 \cos \theta)}{\sin \theta} \right)^2$$

La figure 5.10 représente la coupe du diagramme de la figure 5.9 dans le plan $\phi = \pi/2$.

Dans le plan $(\vec{Ox}, \vec{Oy})$, défini par : $\theta = \pi/2$ soit $\sin \theta = 1$, la fonction caractéristique de l'antenne élémentaire est constante (figure 5.11). Compte tenu de l'expression du facteur de réseau, la fonction caractéristique des deux antennes est donnée par :

$$F_n(\frac{\pi}{2}, \phi) = \cos^2\left(\frac{\pi}{2} \cos \phi\right)$$

Cette fonction s'annule pour $\cos \phi = 1$. Il y a donc un zéro de rayonnement dans la direction $\phi = 0$ qui correspond à l'axe $\vec{Ox}$, comme le montrent les figures 5.9 et 5.11. Ce zéro de rayonnement vient du fait que les deux antennes sont placées sur cet axe et séparées par $d = \lambda/2$. Dans la direction $\vec{Ox}$, les antennes rayonnent en opposition de phase. Dans la direction perpendiculaire, elles rayonnent en phase.

La figure 5.11 représente la coupe du diagramme 5.9 dans le plan $\theta = \pi/2$.

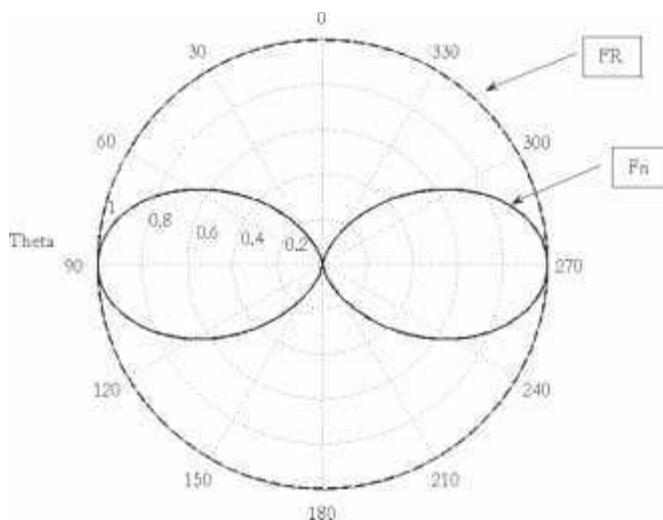


Figure 5.10 – Diagramme de rayonnement dans le plan $\phi = \pi/2$, de deux antennes demi onde, séparées de $\lambda/2$ selon x , en fonction de θ .

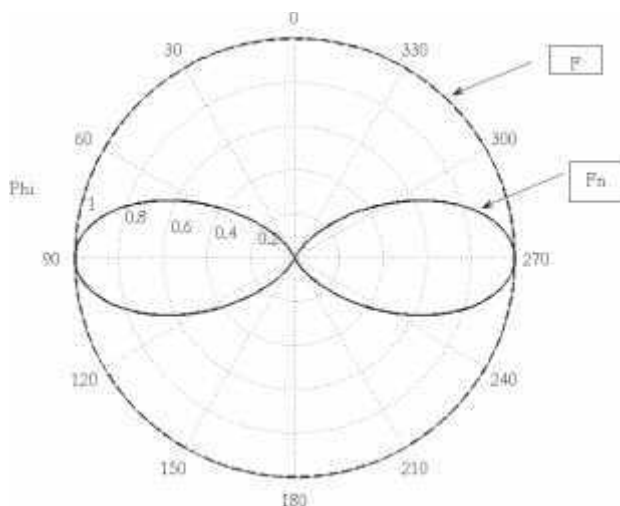


Figure 5.11 – Diagramme de rayonnement dans le plan $\theta = \pi/2$, de deux antennes demi onde, séparées de $\lambda/2$.

On constate que le diagramme de rayonnement de deux antennes est plus fin que celui d'une seule antenne. Le diagramme de la figure 5.10 reproduit le diagramme d'une seule antenne. Cette remarque va dans le sens d'une augmentation de la directivité si le nombre d'antennes du réseau augmente.

5.1.3 Antennes à balayage électronique

Les antennes à balayage permettent de diriger le rayonnement dans différentes directions de l'espace. Elles sont utilisées, par exemple, dans les radars permettant de détecter une cible dans une région de l'espace. L'antenne du radar est, en général, portée par un support vertical tournant, l'antenne étant horizontale. L'axe de l'antenne tourne de façon azimutale. C'est le principe du

balayage mécanique. L'antenne du radar peut ainsi capter le signal retourné par une cible. Par mesure du temps d'aller-retour du signal, on connaît la position de la cible.

Il existe une autre façon d'obtenir un balayage dans l'espace. Celle-ci consiste à introduire des déphasages électroniques variables dans les directions Ox et Oy , comme cela sera expliqué dans la suite.

Le réseau est maintenant construit de façon à introduire un déphasage sur le champ électrique dans l'ouverture, proportionnel à la distance dans la direction Ox et dans la direction Oy . Le déphasage linéique dans la direction Ox est noté α et dans la direction Oy , β . Le champ s'écrit ainsi dans l'ouverture :

$$\vec{E}_O = \vec{E}_0 e^{j\alpha x} e^{j\beta y} \vec{u}_x$$

Pour simplifier, on se limite au cas d'une seule polarisation selon Ox .

La méthode (3.3.2) passe par le calcul du spectre d'onde dans l'ouverture :

$$\vec{f}_T = E_0 \vec{u}_x \int_{-a}^{+a} \int_{-b}^{+b} e^{j(k_x - \alpha)x} e^{j(k_y - \beta)y} dx dy$$

On pose :

$$\alpha a = u_0 \quad \text{et} \quad \beta b = v_0$$

On en déduit le champ électrique à grande distance, par la même méthode qu'au paragraphe 3.3.3. L'expression du champ obtenue est la généralisation de [3.44], en tenant compte des termes de déphasage :

$$\vec{E}(\vec{r}) = \frac{jke^{-jk.r}}{2\pi r} 4abE_0 \frac{\sin(k_x a - u_0)}{k_x a - u_0} \frac{\sin(k_y b - v_0)}{k_y b - v_0} [\cos \phi \vec{u}_\theta - \sin \phi \cos \theta \vec{u}_\phi]$$

Le maximum du champ électrique est obtenu pour :

$$ka \sin \theta \cos \phi = u_0 \quad \text{et} \quad kb \sin \theta \sin \phi = v_0$$

La direction du maximum est donnée par :

$$\tan \phi = \frac{\beta}{\alpha} \quad \text{et} \quad \sin^2 \theta = \frac{\alpha^2 + \beta^2}{k^2}$$

Ces deux expressions déterminent la direction du maximum de puissance.

L'antenne présentée dans le paragraphe 3.2.3, ayant pour dimensions de l'ouverture $2a = 2\lambda$ et $2b = 3\lambda$ est reprise ici en imposant un déphasage tel que :

$$u_0 = \frac{2\pi}{3} \quad \text{et} \quad v_0 = \pi$$

Les figures 5.13 et 5.14 montrent respectivement les diagrammes de rayonnement dans un plan proche du plan (xOz) et dans un plan proche du plan (yOz). Le maximum de rayonnement est obtenu dans la direction telle que :

$$ka \sin \theta \cos \phi = \frac{2\pi}{3} \quad \text{et} \quad kb \sin \theta \sin \phi = \pi$$

Avec les valeurs de a et b , les angles repérant le maximum de rayonnement vérifient les relations :

$$\sin \theta \cos \phi = \frac{1}{3} \quad \text{et} \quad \sin \theta \sin \phi = \frac{1}{3}$$

Le maximum de rayonnement est obtenu dans la direction définie par :

$$\phi = \frac{\pi}{4} \quad \text{et} \quad \theta = 28^\circ$$

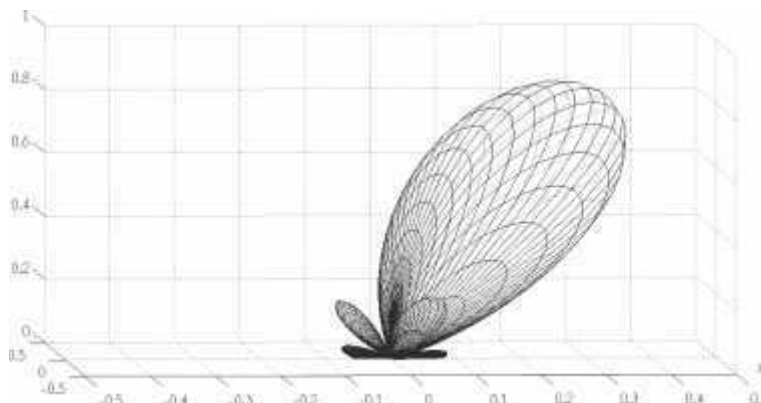


Figure 5.12 – Diagramme de rayonnement de l'antenne définie par une ouverture rectangulaire de dimensions $2a = 2\lambda$ et $2b = 3\lambda$, présentant un déphasage, dans un plan proche du plan (xOz).

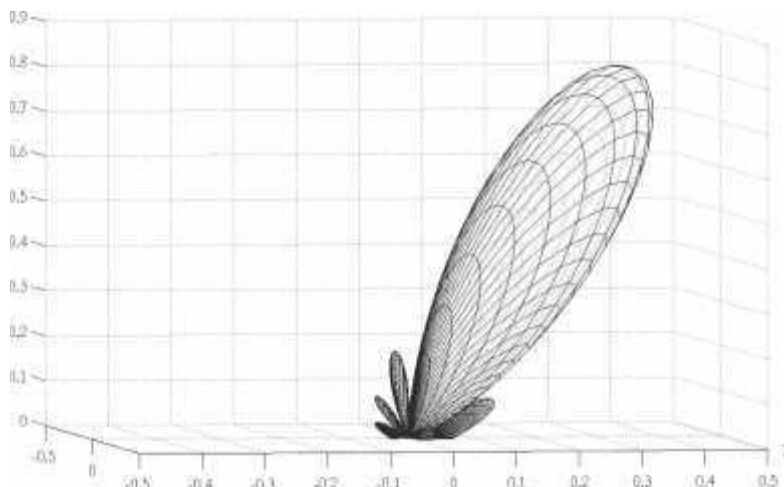


Figure 5.13 – Diagramme de rayonnement de l'antenne définie par une ouverture rectangulaire de dimensions $2a = 2\lambda$ et $2b = 3\lambda$, présentant un déphasage, dans un plan proche du plan (yOz).

La position du faisceau est imposée dans la direction voulue en contrôlant les variations linéaires de phase dans les directions Ox et Oy . Ce principe permet de donner aux antennes une certaine reconfigurabilité. Il est possible de choisir une orientation particulière lorsque le système d'émission ou de réception requiert l'utilisation de directions variables avec le temps, par exemple, pour supporter un trafic de communications important à une période donnée.

Lorsque la polarisation dans l'ouverture n'est pas rectiligne, le principe de calcul est le même, en partant de l'expression plus générale du champ donnée par [3.43].

En faisant varier la phase avec le temps, on obtient une antenne à balayage électronique. Ce type de balayage est utilisé dans les systèmes radar.

Les réseaux sont utilisés pour toutes les applications dans lesquelles on recherche de la puissance, ou une grande directivité, ou de la reconfigurabilité, ou de la formation de faisceaux. On en retrouvera les principes, par exemple dans les antennes multifaisceaux et dans les antennes de stations de bases.

5.2 Antennes multifaisceaux

Les antennes multifaisceaux ont connu un développement considérable avec leur utilisation dans les radars et dans les systèmes satellites. Elles sont aussi utilisées dans les systèmes de communications mobiles. Ces systèmes comportent de nombreuses antennes identiques qui sont reliées à des fonctions électroniques permettant d'imposer une amplitude et une phase à chaque élément. Un système multifaisceaux peut contenir plus d'une centaine d'antennes. On comprend la complexité rencontrée lors de leur conception.

Les éléments les plus courants sont des cornets, des patchs, des éléments à ondes de fuites utilisant des diélectriques ou autres...

La plupart du temps, les antennes sont associées à des réflecteurs ou à des lentilles

Les fonctionnalités des antennes multifaisceaux permettent :

- de réaliser un faisceau d'un contour particulier afin, par exemple, d'obtenir une zone de couverture déterminée.
- de réaliser une répartition de puissance spécifique à l'intérieur du faisceau afin, par exemple, de diminuer les lobes secondaires ou bien de créer une direction dans laquelle la puissance rayonnée est nulle.
- d'orienter, éventuellement temporairement, certains faisceaux dans des directions particulières afin de répondre à des demandes de trafic important.

Dans les deux premiers cas, on parle alors de formation de faisceau. Dans le troisième cas d'antennes reconfigurables.

Les antennes multifaisceaux décrites dans ce paragraphe fonctionnent de la même façon en émission et en réception.

L'augmentation du gain est considérable en raison de la grande taille de l'ouverture équivalente du réseau.

La figure 5.14 donne un exemple de système d'antenne à multifaisceaux.

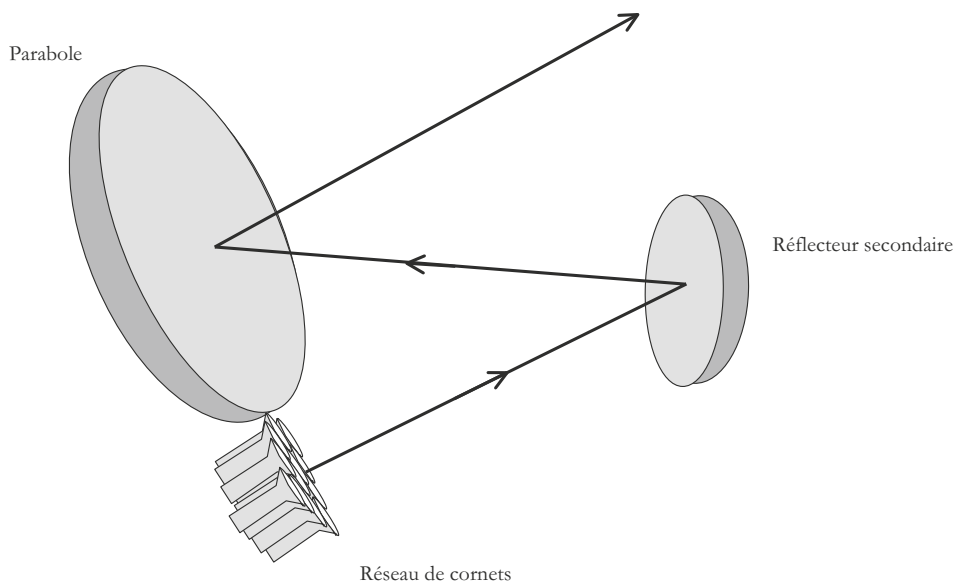


Figure 5.14 – Antenne multifaisceaux utilisant des cornets.

Les différents cornets émettent une onde avec une amplitude et une phase déterminées. Ces ondes sont réfléchies une première fois sur le réflecteur secondaire puis sur la parabole qui renvoie les ondes dans la direction voulue. À la sortie de la parabole, la répartition du champ sur la surface équivalente détermine le diagramme de rayonnement.

5.2.1 Principe de fonctionnement

À chaque élément rayonnant correspond un faisceau élémentaire de sortie. L'ensemble de ces faisceaux reconstitue une illumination particulière dans un cône donné. Celui-ci définit le champ de l'antenne. Rappelons rapidement les propriétés d'une parabole : un faisceau parallèle provenant de l'infini se focalise dans le plan focal de la parabole. Si le faisceau est parallèle à l'axe de la parabole, l'ensemble des rayons se focalise au foyer F . Si le faisceau fait un angle avec l'axe, la focalisation a lieu en un point du plan focal F' décalé par rapport au foyer (figure 5.15)

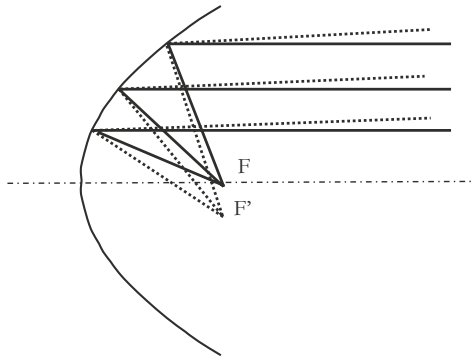


Figure 5.15 – Réflexion sur une parabole.

Chaque antenne est placée de façon à ce que le rayonnement partant de la parabole et issu de cette antenne soit un faisceau parallèle, faisant un angle correspondant au décalage de l'antenne. La forme du réflecteur secondaire est calculée pour que les images des antennes primaires se forment dans le plan focal de la parabole. Finalement, à chaque source correspond un pinceau élémentaire.

Un autre système permet d'obtenir le même résultat. Il utilise une lentille dans le plan focal de laquelle se trouve l'ensemble des antennes (figure 5.16)

À chaque antenne correspond un faisceau parallèle, donc un angle de rayonnement. Le rôle de la lentille est d'introduire un déphasage entre chaque rayon, qui impose l'orientation du plan d'onde en sortie de la lentille. La forme de la lentille est calculée en fonction du matériau la constituant et des déphasages requis.

Ce déphasage peut être réalisé, dans d'autres systèmes, par un retard induit par la différence de chemin parcouru pour chaque antenne. Les antennes, dans ce cas, ne sont pas placées sur un même plan et leur décalage géométrique introduit naturellement le déphasage.

Le déphasage peut aussi être imposé par un dispositif électronique associé.

Les systèmes à réflecteurs sont préférés à ceux qui comportent une lentille lorsque l'ouverture de l'antenne est supérieure à environ cent longueurs d'onde.

Les différents éléments d'une antenne multifaisceaux sont :

- Les antennes et les dispositifs de mise en commun des faisceaux (réflecteurs, lentilles)
- Le réseau de formation de faisceau. Ce réseau est constitué de circuits passifs d'autant plus complexes que le nombre d'antennes est grand. Lorsqu'on a affaire à des guides, les circuits sont lourds et encombrants. Pour les antennes planaires, les circuits sont réalisés dans une

technologie planaire, moins encombrante, mais de conception complexe. Le principe des matrices de Butler est souvent utilisé pour réaliser ce réseau. Les réseaux de formations de faisceaux constituent un point délicat de la conception d'antennes multifaisceaux, en raison de leur complexité. Les formes des lignes présentent des angles, des discontinuités et finalement une longueur non négligeable. Ces points ont pour conséquences une augmentation des pertes dans cette partie.

- Les circuits de contrôle. Ils contiennent les dispositifs permettant de contrôler la phase et l'amplitude. On trouve des combinaisons de déphaseurs et de diviseurs de puissance, en émission, et de combineurs en réception.

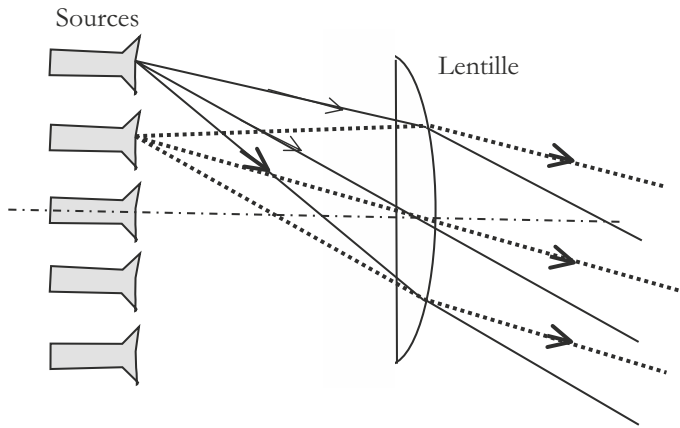


Figure 5.16 – Principe de la lentille.

5.2.2 Formation de faisceau

Il s'agit de déterminer la loi de répartition de l'amplitude du champ électrique dans le faisceau. Le champ électrique de chaque faisceau issu d'une antenne élémentaire est noté :

$$\vec{E}_i(r, \theta, \phi) = c_i e^{j\alpha_i} \vec{f}(\theta, \phi) \frac{e^{-jkR_i}}{4\pi R_i}$$

Puisque chaque antenne élémentaire est identique et, en négligeant l'influence mutuelle entre les antennes, on peut calculer le champ total sous la forme :

$$\vec{E}(r, \theta, \phi) = \frac{\vec{f}(\theta, \phi)}{4\pi r} e^{-jkr} \sum_{i=1}^N c_i e^{j\alpha_i + jk\vec{u} \cdot \vec{r}_i}$$

Prenons l'exemple d'antennes circulaires dont le diagramme élémentaire sans déphasage a été calculé en 3.2.4. :

$$\vec{E}_i(M) = \frac{jk}{2} Z\psi(r) H_i(1 + \cos \theta)(\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) a^2 \frac{J_1(ka \sin \theta)}{ka \sin \theta}$$

On introduit le terme de déphasage u_i lié à la distance d_i entre les antennes sous la forme :

$$u_i = \alpha d_i$$

Chaque antenne est déphasée de u_i . L'expression du champ, dans une direction est :

$$\vec{E}_i(M) = \frac{jk}{2} Z\psi(r) H_i(1 + \cos \theta)(\cos \varphi \vec{u}_\theta - \sin \varphi \vec{u}_\varphi) a^2 \frac{J_1(u - u_i)}{u - u_i}$$

Dans cette direction, l'amplitude résultant de la somme de toutes les antennes élémentaires fait intervenir le terme suivant :

$$\sum_i c_i \frac{J_1(u - u_i)}{u - u_i}$$

C'est ce terme qui détermine les propriétés de la loi d'illumination de l'antenne. En général, on se fixe une figure de mérite et les termes c_i et u_i sont calculés par optimisation pour donner la forme la plus proche de celle-ci.

À l'aide de ces dispositifs, il est possible de réaliser les fonctions recherchées de formation ou d'orientation de faisceau.

5.3 Synthèse de réseaux

Nous avons vu, dans les paragraphes précédents, que les coefficients d'amplitude et de phase jouaient un rôle prépondérant pour la répartition de puissance rayonnée. La synthèse de réseau consiste à trouver les bons coefficients qui permettent d'approcher au mieux un gabarit donné.

La synthèse de réseaux se fixe différents objectifs concernant la qualité du rayonnement d'une antenne formée de multiples éléments. Les critères peuvent porter sur la largeur du faisceau, sur son orientation, sur les remontées de lobes secondaires sur leur orientation. Il est aussi possible d'imposer aux antennes de présenter plusieurs faisceaux, comme c'est le cas des antennes satellites qui doivent couvrir plusieurs zones. Dans cet exemple, les orientations des faisceaux sont imposées, mais aussi leur largeur et la répartition de la puissance dans le faisceau ; l'orientation des lobes secondaires est aussi contrainte, afin d'éviter les phénomènes d'interférences par l'intermédiaire des lobes secondaires.

Certaines applications utilisant des réseaux d'antennes, en télécommunication ou pour des systèmes radars, nécessitent d'éliminer les zéros de puissance dans une zone donnée. On utilise alors des réseaux « cosécantés » dont la fonction de répartition dans le plan vertical s'exprime sous la forme de la fonction cosécante, ne présentant pas de zéro entre 0 et 90° (paragraphe 6.4.2). Cette répartition est obtenue en faisant varier les amplitudes d'alimentation des antennes élémentaires. La synthèse de réseaux repose sur des méthodes d'optimisation. Les choix technologiques sont variés car les paramètres ajustables lors de la conception du réseau sont multiples.

La synthèse peut reposer d'une part sur un ajustement des paramètres d'amplitude. Tous les éléments sont alors alimentés en phase. Ce cas est bien adapté aux réseaux rayonnants dans l'axe. D'autre part la méthode peut consister à déphaser les éléments en maintenant l'amplitude constante. Ce cas est adapté aux réseaux visant dans des directions différentes de l'axe pour des applications de suivi. Le mélange des deux méthodes peut aussi être adopté mais au prix d'une certaine complexité.

Nous ne donnerons ici que quelques idées relatives à la synthèse de réseau, le sujet étant trop vaste pour être traité dans cet ouvrage.

Deux méthodes seront décrites ici. La première permet d'avoir une première approximation de la répartition de l'alimentation en amplitude des éléments, grâce à la transformée de Fourier. La seconde présente les bases des méthodes d'optimisation.

5.3.1 Synthèse par la transformée discrète de Fourier

Reprenons l'étude des réseaux du paragraphe 5.1. Nous allons appliquer la synthèse à un réseau linéaire. La généralisation à deux dimensions repose sur les mêmes principes.

Rappelons la forme et les notations utilisées pour calculer le facteur d'un réseau linéaire :

$$F_R(\theta, \phi) = \left| \sum_{n=0}^N c_n e^{jn\alpha d + jknd \cos \psi} \right|^2$$

Nous simplifions l'étude ici en ne considérant qu'une variation d'amplitude entre les différents éléments. Le facteur de réseau se réduit donc à :

$$F_R(\theta, \phi) = \left| \sum_{n=0}^N c_n e^{jknd \cos \psi} \right|^2$$

Nous allons nous intéresser à la fonction :

$$f(\theta, \phi) = \sum_{n=0}^N c_n e^{jknd \cos \psi}$$

Cette fonction détermine la répartition d'amplitude du champ, donc sa puissance. Elle s'écrit encore, pour un réseau comportant un nombre impair d'éléments :

$$f(\theta, \phi) = \sum_{n=-N}^N c_n e^{j2\pi n \frac{d}{\lambda} \cos \psi} \quad [5.6]$$

Les éléments sont alors placés à $x_n = nd$

Pour un réseau comportant un nombre pair d'éléments, ceux-ci se trouvent en :

$$x_n = \frac{2n-1}{2}d \quad \text{pour} \quad 1 \leq n \leq N$$

$$x_{-n} = -\frac{2n-1}{2}d \quad \text{pour} \quad -N \leq n \leq -1$$

Alors, la fonction s'écrit :

$$f(\theta, \phi) = \sum_{n=1}^N (c_{-n} e^{-j\pi(2n-1) \frac{d}{\lambda} \cos \psi} + c_n e^{j\pi(2n-1) \frac{d}{\lambda} \cos \psi}) \quad [5.7]$$

Le but de la méthode est de reconstruire la fonction caractéristique d'une antenne, appelée fonction objectif, à partir d'un réseau d'antennes. La fonction objectif est déterminée par les spécificités du système. Elle s'exprime en fonction des angles d'observation. Cette fonction objectif se développe dans cette méthode sous la forme d'une série de Fourier :

$$f_O(\theta, \phi) = \sum_{n=-\infty}^{\infty} c'_n e^{j2\pi n \frac{d}{\lambda} \cos \psi}$$

La comparaison de cette forme à l'expression de la fonction réseau de l'antenne réelle [5.6] ou [5.7], permet de proposer une solution constituée des termes de la série de la fonction objectif en nombre fini. En tronquant ainsi la série, le résultat n'est pas parfait. L'erreur est estimée par la différence des deux fonctions :

$$e(\theta, \phi) = |f(\theta, \phi) - f_O(\theta, \phi)|$$

La fonction objectif répond à certaines contraintes. Prenons la contrainte suivante :

$$f_O(\theta, \phi) = 1 \quad \text{pour} \quad |\cos \psi| < l$$

$$f_O(\theta, \phi) = 0 \quad \text{pour} \quad |\cos \psi| > l$$

Les coefficients c'_n sont déterminés par la théorie des séries de Fourier.

Pour les séries impaires, avec les restrictions sur la fonction objectif :

$$c'_n = \frac{d}{\lambda} \int_{-l}^{+l} f_O(\theta, \phi) e^{j2\pi n \frac{d}{\lambda} \cos \psi} d(\cos \psi)$$

Pour les séries paires :

$$c'_n = c'_{-n} = \frac{d}{\lambda} \int_{-l}^{+l} f_O(\theta, \phi) e^{j\pi(2n-1)\frac{d}{\lambda} \cos \psi} d(\cos \psi)$$

Le fait d'utiliser un nombre fini d'éléments entraîne une intégration sur une distance finie, donc une erreur par rapport à la fonction objectif.

Par exemple, pour un réseau linéaire pair dont la fonction objectif est définie par :

$$l = 0,5$$

Les valeurs des coefficients de la série de Fourier sont calculables par :

$$c'_n = 2 \frac{d}{\lambda} l \frac{\sin \left((2n-1) \pi \frac{d}{\lambda} l \right)}{(2n-1) \pi \frac{d}{\lambda} l} \quad \text{pour } 1 \leq n \leq N$$

La figure 5.17 montre la reconstruction de la fonction objectif avec 8 éléments, 20 éléments et 40 éléments. Plus le nombre d'antennes élémentaires est grand, meilleure est la fonction obtenue.

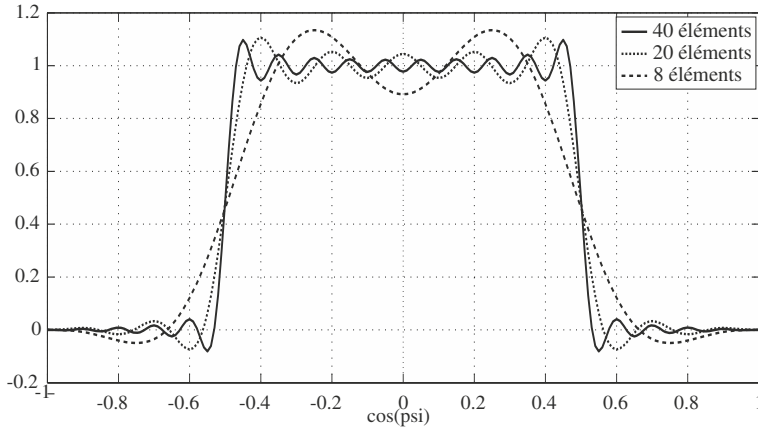


Figure 5.17 – Reconstruction de la fonction objectif d'un réseau à l'aide de 8, 20 ou 40 éléments espacés de $d = \lambda/2$.

La répartition exacte du champ en amplitude est obtenue en multipliant la fonction f par la fonction caractéristique d'un élément.

5.3.2 Synthèse par échantillonnage

Nous reprenons dans ce paragraphe les mêmes notations que précédemment. Cette méthode consiste à échantillonner en direction les fonctions relatives au réseau.

Ainsi pour un réseau défini par les amplitudes d'alimentation et le déphasage de ses éléments, la fonction caractéristique du réseau en amplitude est donnée pour une direction donnée par :

$$f(\theta_i, \phi_i) = \sum_{n=0}^N c_n e^{j\alpha_n d + jknd \cos \psi_i}$$

Pour cette même direction, la fonction objectif est définie par :

$$f_O(\theta_i, \phi_i)$$

Le problème à résoudre est celui de trouver l'erreur minimale sur toutes les directions. Une méthode des moindres carrés peut être utilisée. Elle consiste à chercher le minimum de l'expression :

$$\sum_i |f(\theta_i, \phi_i) - f_O(\theta_i, \phi_i)|$$

Éventuellement, chaque terme de cette somme peut être affecté d'un poids.

L'étude des méthodes de minimisation d'erreur sort du cadre de cet ouvrage. En échantillonnant correctement l'espace, il est donc possible de déterminer des coefficients d'amplitude et de phase à attribuer à chaque élément.

5.4 Alimentation des réseaux

Les types de réseaux étant très variés, nous ne pouvons pas décrire toutes les combinaisons de technologies utilisées. Nous verrons quelques-unes des plus classiques. Ce domaine de conception de réseau est extrêmement complexe et nous nous limitons aux concepts simples.

Lorsque les réseaux contiennent un grand nombre d'éléments, les éléments sont regroupés par sous-réseaux. Cela permet d'optimiser le nombre de composants et de simplifier les fonctions de commande.

5.4.1 Les dispositifs utilisés dans les réseaux

Comme nous l'avons vu précédemment, les différents éléments sont alimentés en phase et en amplitude. La base de l'alimentation de réseau est donc constituée de déphaseurs et de dispositifs permettant de diviser la puissance entre les différents éléments lors de l'émission ou bien de combiner la puissance en réception.

■ Les déphaseurs

Les déphaseurs sont placés en parallèle ou en série (figures 5.18 et 5.19)

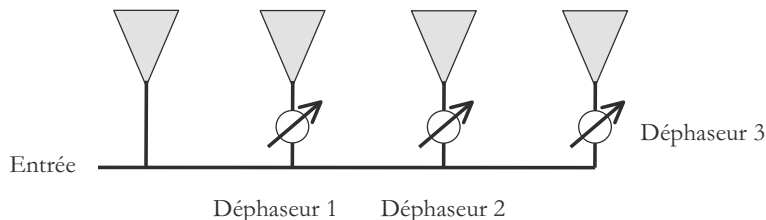


Figure 5.18 – Déphaseurs en parallèle.

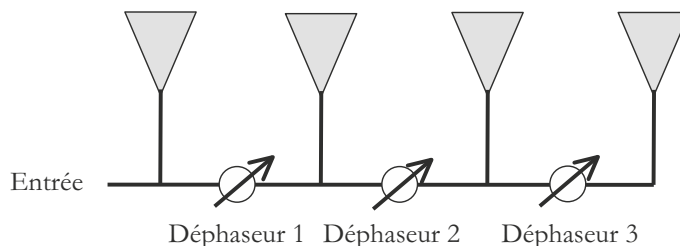


Figure 5.19 – Déphaseurs en série.

On distingue classiquement deux types de déphaseurs :

- Les déphaseurs à lignes
- Les déphaseurs à ferrites

Les premiers sont constitués d'un ensemble de lignes de longueurs différentes. Des commutateurs permettent d'insérer l'impédance de chaque tronçon de ligne selon la phase désirée. Ces dispositifs en technologies micro ruban ou coplanaire peuvent éventuellement être réalisés en technologie intégrée.

Les seconds sont constitués d'un barreau de ferrite autour duquel est enroulé un fil qui sert à fixer la valeur du champ magnétique du barreau. En fonction de l'intensité d'excitation, le champ magnétique est plus ou moins grand. Celui-ci fixe le déphasage de l'onde. Ce dispositif, placé à l'intérieur d'un guide d'onde, permet d'obtenir en sortie un déphasage contrôlé.

■ Les diviseurs ou combineurs

Selon le type d'élément et selon le choix de l'alimentation, les combinaisons peuvent être extrêmement variées. Nous donnons dans la suite un aperçu de quelques techniques.

5.4.2 Technologie en guide

De nombreux réseaux d'antennes sont constitués à base d'antennes alimentées par des guides d'onde. C'est le cas de la plupart des antennes utilisées dans le domaine spatial, pour des raisons de robustesse, de puissance supportée et de CEM.

Nous allons présenter le principe d'un diviseur de puissance en guide. La variation de puissance est obtenue grâce à des déphaseurs variables et à un coupleur. Un Té magique est placé en tête du diviseur pour obtenir la séparation en phase du signal d'entrée (figure 5.20).

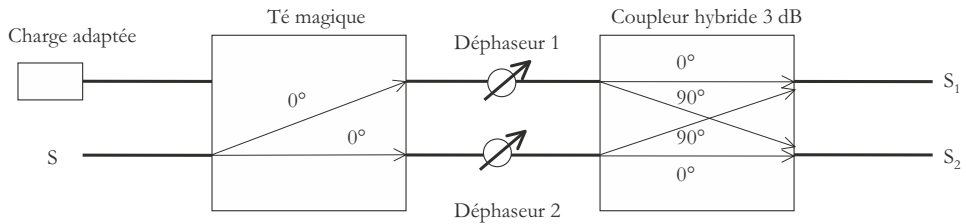


Figure 5.20 – Déphaseur variable.

Les déphaseurs 1 et 2 introduisent respectivement des déphasages α_1 et α_2 .

Grâce aux combinaisons en phase des signaux, le signal d'entrée S est réparti en S_1 et S_2 sur les deux sorties, avec :

$$S_1 = S \sin \left(\frac{\alpha_2 - \alpha_1}{2} + \frac{\pi}{4} \right) e^{j \left(\frac{\alpha_2 + \alpha_1}{2} - \frac{\pi}{4} \right)}$$

$$S_2 = S \cos \left(\frac{\alpha_2 - \alpha_1}{2} + \frac{\pi}{4} \right) e^{j \left(\frac{\alpha_2 + \alpha_1}{2} - \frac{\pi}{4} \right)}$$

L'ajustement des déphasages permet d'obtenir une valeur déterminée de la puissance en sortie qui sert d'alimentation à l'élément d'antenne. Dans ce dispositif, le coupleur hybride est un coupleur classique, associant deux guides ayant une paroi commune qui communique sur une portion commune.

5.4.3 Technologie planaire

Nous allons montrer le principe de la répartition pour un réseau planaire. La synthèse du réseau a permis de déterminer les coefficients de pondération des signaux à affecter à chaque élément. Pour répartir les signaux en phase dans un réseau linéaire, il est possible de proposer le schéma d'alimentation en ligne micro ruban de la figure 5.21.

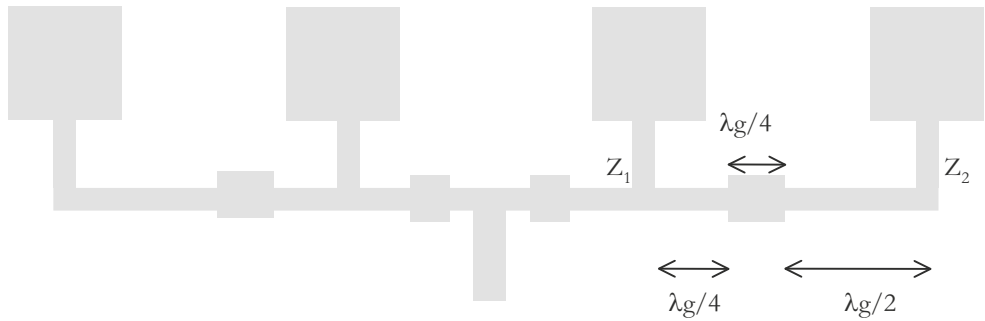


Figure 5.21 – Schéma de répartition de puissance pour un réseau linéaire de quatre éléments en technologie micro ruban.

Le réseau est symétrique. Chacun des éléments a une impédance d'entrée notée Z_1 et Z_2 . Ces impédances sont égales, mais nous les distinguons lors du calcul.

Les antennes sont reliées par des tronçons de ligne de longueurs égales à la demi-longueur d'onde guidée ou au quart de longueur d'onde guidée.

Analysons les impédances dans chacun des plans significatifs de ce schéma agrandi sur la figure 5.22.

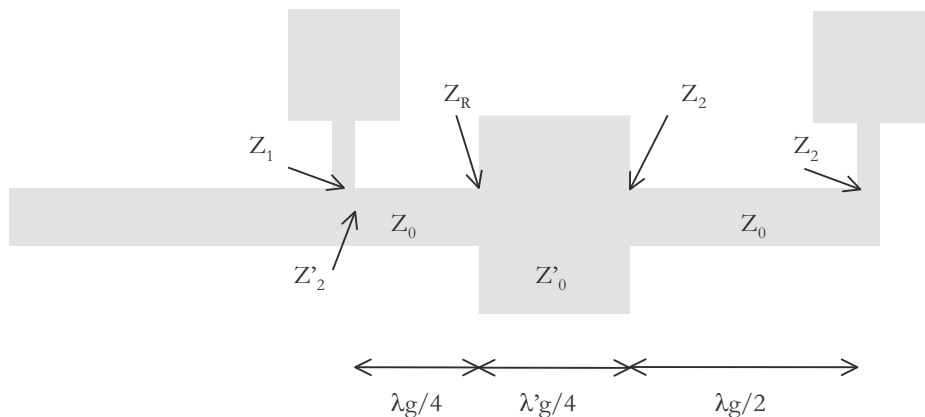


Figure 5.22 – Extrémité de l'alimentation.

Les tronçons de ligne ont des impédances caractéristiques Z_0 et Z'_0 . Partons de l'antenne 2 d'impédance d'entrée Z_2 . À l'entrée du tronçon de ligne élargie son impédance ramenée est la même puisque la longueur de ligne vaut une demi-longueur d'onde guidée. À gauche de l'élargissement, l'impédance ramenée est Z_R :

$$Z_R = \frac{Z_0'^2}{Z_2}$$

De la même façon, l'impédance Z'_2 est égale à :

$$Z'_2 = \frac{Z_0^2}{Z_R} \quad \text{d'où} \quad Z'_2 = \frac{Z_0^2}{Z_0^2} Z_2$$

Le rapport des puissances P_1 et P_2 alimentant chacun des éléments est fixé par la synthèse réalisée au préalable. On déduit donc :

$$x^2 = \frac{P_2}{P_1} = \frac{U^2/Z'_2}{U^2/Z_1} = \frac{Z_0^2}{Z_1^2}$$

Dans cette expression, il est tenu compte du fait que les antennes ont la même impédance d'entrée. Le rapport des amplitudes apparaît comme le rapport des impédances caractéristiques des tronçons de ligne. Il est donc possible de dimensionner la largeur des lignes en fonction des longueurs d'ondes guidées et du rapport d'impédance recherché. Ce travail est le résultat de compromis sur la géométrie des lignes, car généralement l'espacement des antennes est fixé par le gabarit imposé lors de la conception.

La méthode présentée est généralisable à un nombre quelconque d'éléments.

5.4.4 Distribution et contrôle optique

Les avantages des réseaux d'antennes sont nombreux comme nous venons de le voir. Cependant, l'alimentation des éléments, en amplitude et en phase, constitue une étape limitant la réalisation de grands réseaux. C'est pourquoi, une solution consiste à alimenter ces réseaux par voie optique. Les avantages attendus en sont :

- une plus grande souplesse de conception
- un gain en poids
- des pertes plus faibles
- une grande largeur de bande
- la possibilité d'intégration de composants optiques donc de réduction de l'encombrement.

Les composants optiques étant par ailleurs insensibles aux champs électromagnétiques, le système est rendu moins vulnérable aux perturbations.

Le principe est de transférer les informations d'amplitude et de phase portées par un signal optique au signal micro-ondes qui alimente les éléments.

Le principe de fonctionnement est le suivant : une source laser envoie une onde optique dans une fibre, munie d'un modulateur. Ce dernier module l'onde optique à la fréquence micro-onde correspondant à la fréquence de fonctionnement du réseau d'antennes. Le signal se propageant sur la fibre est ensuite réparti sur plusieurs fibres. Un traitement optique est effectué sur chacune de ces fibres afin d'introduire les variations d'amplitude et des déphasages sur le signal. Afin de récupérer le signal micro-onde, chaque fibre est terminée par un photodétecteur, suivi d'un amplificateur qui est placé juste avant l'élément rayonnant.

On distingue deux sortes de modulateurs :

- Les modulateurs directs qui modulent le courant source du laser
- Les modulateurs externes dont le principe repose sur l'effet Pockels : l'indice d'un matériau électro-optique varie en fonction de la tension appliquée, en l'occurrence, le signal micro-onde. Il est ainsi possible de transmettre une modulation micro-onde sur le signal optique. Le modulateur de ce type le plus utilisé est le modulateur de Mach-Zender dont le fonctionnement est basé sur l'interférence de deux signaux se propageant selon deux chemins différents et subissant une modulation différente. Les matériaux électro-optiques les plus utilisés sont le niobate de lithium et l'arséniure de gallium.

Les déphasages introduits sur les fibres peuvent être obtenus par différents moyens, soit par modulation, soit par l'introduction de longueurs de fibre bien définies. La figure 5.23 donne un exemple de fonctionnement de ce type de dispositifs qui nécessite l'introduction de commutateurs optiques.

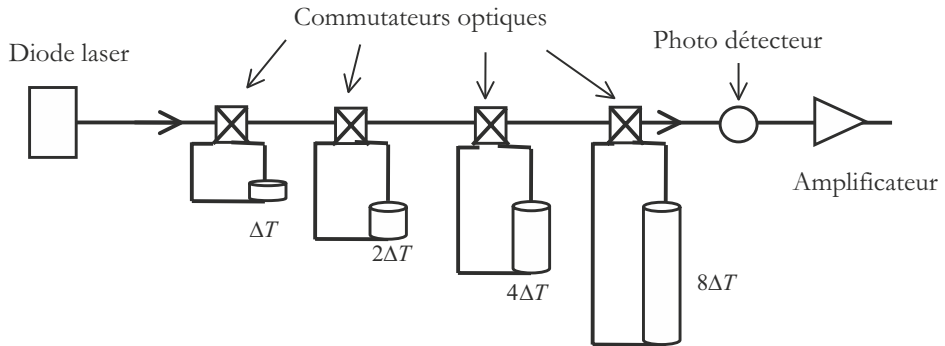


Figure 5.23 – Principe des lignes à retard utilisées pour le déphasage des signaux

Chaque commutateur est relié à une fibre optique d'une longueur égale à un multiple pair de la plus petite longueur. Les retards introduits sont donc des multiples quelconques du plus petit retard ΔT . Avec N commutateurs, il est possible de réaliser 2^N combinaisons. Un organe de commande permet d'actionner les commutateurs pour obtenir un déphasage au plus proche de celui souhaité.

Ces dispositifs sont souvent utilisés pour contrôler des sous-réseaux.

Bien d'autres possibilités existent pour contrôler optiquement les systèmes d'antennes. Nous nous sommes restreints ici à donner un bref aperçu des dispositifs opto-électroniques utilisables.

Bibliographie

BALANIS C.A. – *Antenna Theory Analysis and Design*, New Jersey, A. John Wiley & Sons, Inc., 2005.

BENJAMIN R., SEEDS A.J. – "Optical beam forming techniques for phased array antennas", *IEE Proceeding* – H, vol. 139, n°6, Déc. 1992, pp 526-534.

HANSEN R. C. – *Phased Array Antennas*, New Jersey, A. John Wiley & Sons, Inc., 1998.

RUDGE A. W., MILNE K., OLIVER A. D. , KNIGHT P. – *The Handbook of Antenna Design*, vol. 1, Peter Peregrinus Ltd, London, UK.

STUTZMANN W. L., GARY A.T. – *Antenna Theory and Design*, New Jersey, A. John Wiley & Sons, Inc., 1998.

ULABY F. T., MOORE R. K., FUNG A.K. – *Microwave Remote Sensing, Active and Passive, Vol. I*, Norwood, Artech House, 1981.

Van TREES H. L. – *Optimum Array Processing*, New York, John Wiley & Sons, 2002.

6 • DIFFÉRENTS DOMAINES D'UTILISATION DES ANTENNES

Dans ce chapitre nous présenterons un panorama des différentes applications utilisant des antennes, en les replaçant dans leur contexte. Les systèmes seront décrits brièvement afin d'expliquer les points sensibles à prendre en compte lors de la conception des antennes.

6.1 Gestion du spectre radiofréquence

Le spectre électromagnétique a déjà été décrit globalement (figure I.1). Les antennes n'utilisent qu'une partie de ce spectre, communément appelé spectre radiofréquence (abréviation pour spectre des fréquences radioélectriques) qui s'étend approximativement de 1 kHz à 300 GHz. Ce spectre, très étendu, puisqu'il couvre environ neuf décades, est utilisé dans de nombreuses applications qui vont être développées plus loin. Les applications sont différentes en fonction de la fréquence et sont en général relatives à un domaine de fréquence particulier. C'est pourquoi il est d'usage de désigner les différentes bandes de ce spectre par une abréviation faisant référence à la position de la bande de fréquences. Cette dénomination est reconnue par les spécialistes du domaine et elle est apparue au fur et à mesure du développement des applications et varie au cours du temps.

Il existe actuellement un organisme international l'UIT (Union internationale des télécommunications) qui gère les ressources radiofréquences, en tant qu'institution spécialisée des Nations Unies. Cet organisme a son siège à Genève et se réunit régulièrement pour allouer (ou retirer) certaines bandes de fréquences. Il regroupe 191 états membres et plus de 700 membres associés. Son rôle est de gérer le spectre, d'établir des normes et de suivre le développement des télécommunications dans le monde. Par exemple, l'UIT a récemment élaboré une norme permettant l'harmonisation des systèmes mobiles et des nouveaux dispositifs à débit élevé.

En France, l'ANFR (Agence nationale des fréquences) contrôle l'utilisation des fréquences. Elle coordonne les demandes d'allocation de fréquence et représente la position de la France dans les négociations internationales sur la gestion du spectre.

6.1.1 Le spectre radiofréquence

La figure 6.1 présente l'ensemble du spectre radiofréquence et la dénomination des différentes bandes qui correspond à une division par décades. La répartition du spectre en fonction des applications y est mentionnée de façon approximative, car les largeurs des bandes de fréquences allouées par applications sont variables et, en général, très petites. De plus, les bandes allouées peuvent être différentes d'une région du monde à l'autre.

L'utilisation du spectre radiofréquence est vaste. Les applications en sont la télévision, les communications fixes ou mobiles, l'aéronautique, la marine, les émissions radioamateur, la police, l'armée, la météorologie, la télédétection, la radioastronomie, la radiométrie. La figure 6.1 donne un aperçu de cette répartition en Europe. La bande de chaque application est définie très précisément selon les pays par les documents de l'UIT. Par exemple, on trouvera dans la documentation de

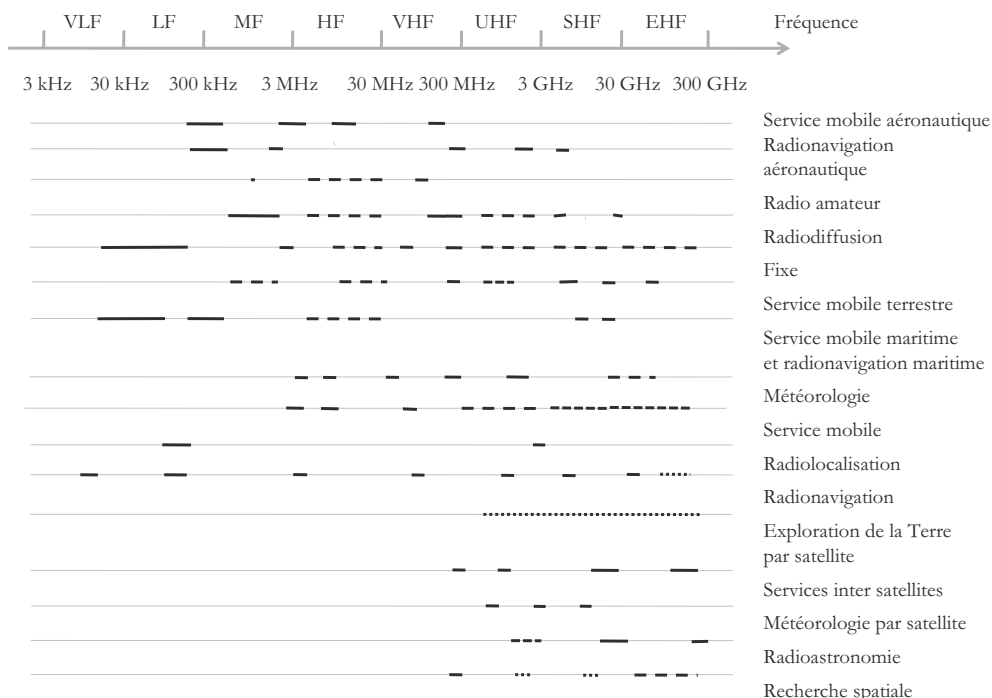


Figure 6.1 – Aperçu de l'encombrement du spectre radioélectrique.

l'ANFR, citée dans la bibliographie, la répartition précise de l'utilisation des bandes de fréquences en pourcentage en fonction des activités pour la France.

Certaines bandes sont utilisables plus librement à condition de respecter des niveaux de puissance limites et de ne pas empiéter sur les bandes adjacentes. Ces bandes appelées ISM (Industriel, scientifique et médical) sont ouvertes à de nombreuses applications et utilisables gratuitement. Elles sont réparties dans tout le spectre. Utilisant les basses fréquences, on trouve des applications domestiques, nécessitant un faible débit d'informations comme les télécommandes, les transmissions de données de capteurs. La bande des 13,56 MHz est celle des cartes à puce sans contact. Les bandes autour de 433 MHz et de 870 MHz sont très utilisées pour l'identification radiofréquence. La bande des 2,4 GHz est actuellement bien connue pour le WIFI et les fours à micro-ondes. Ces derniers ont de nombreuses applications industrielles, liées au chauffage, au séchage et à la polymérisation. Mais on y trouve aussi de nombreuses autres applications de transmissions sans fil (transmetteurs, caméra de vidéo surveillance). D'autres bandes ISM sont réservées pour des fréquences plus élevées, autour de 5,8 GHz, 24 GHz, 61 GHz, 122 GHz, 245 GHz.

■ Applications par bandes de fréquence

La bande VLF (*Very Low Frequency*) est utilisée pour les transmissions avec les sous-marins. Les ondes électromagnétiques pénètrent de plus en plus mal dans l'eau de mer au fur et à mesure que leur fréquence augmente, c'est pourquoi on utilise des ondes basses fréquences pour cette application. Cette bande est aussi utilisée pour des systèmes de radionavigation de type Omega avec réflexion sur l'ionosphère. La bande LF (*Low Frequency*) est utilisée pour certaines applications de communications, de localisations de véhicules (LoranC). Dans le haut de cette bande on trouvait les signaux de certaines balises. La bande MF (*Medium Frequency*) est consacrée pour

une grande partie à la radiodiffusion (bande de modulation d'amplitude). Les ondes longues sont dans cette bande. Dans le bas de la bande, on trouve des communications maritimes, de même que des fréquences réservées aux radioamateurs. La bande HF (*High Frequency*) est très morcelée, constituée de nombreuses sous-bandes aux applications très diverses. On y trouve des communications longues distances, de la radiodiffusion (ondes courtes), des utilisations terrestres (sécurité, armée, police), des bandes de radioamateurs et une bande ISM. Les bandes ISM sont réservées pour des applications industrielles, scientifiques et médicales. Elles sont moins contraintes au niveau des normes de puissance afin de permettre des applications qui peuvent nécessiter des niveaux de puissance élevée. La bande VHF (*Very High Frequency*) contient des applications pour la radiodiffusion (modulation de fréquence radio et télévision), les communications aériennes. La bande UHF contient des sous-bandes réservées à la télévision, aux communications avec les mobiles et les systèmes aéroportés. On trouve dans cette bande de nombreuses applications radar. L'atmosphère étant transparente dans cette bande de fréquences, elle est aussi utilisée pour les services utilisant les transmissions satellites et pour la radioastronomie. La bande SHF (*Super High Frequency*) est aussi utilisée pour des services satellites avec des applications en communications, télédiffusion, et télédétection. De nombreux systèmes de radars et radiomètres se trouvent dans cette bande. Plusieurs bandes sont réservées aux applications ISM. La bande la plus haute, EHF (*Extremely High Frequency*), qui va en principe de 30 à 300 GHz, n'est allouée que jusqu'à 275 GHz. Dans le bas de la bande, on trouve des applications classiques de communications. La recherche spatiale, la télédétection et la radioastronomie sont bien présentes dans cette bande. On trouve aussi une utilisation particulière concernant les services intersatellites qui se répartissent sur toute la bande.

Les modes de propagation diffèrent selon les fréquences. En basses fréquences, jusqu'à la bande MF, la propagation a lieu par onde de sol. En montant en fréquence, en bandes HF et VHF, les ondes se réfléchissent sur l'ionosphère. Ce mode de propagation est utilisé pour couvrir tout le globe, grâce à plusieurs rebonds (système de type Oméga). Au-delà de la bande VHF, les transmissions s'effectuent en trajet direct, incluant éventuellement des phénomènes de diffraction.

6.1.2 Le spectre micro-onde

Les micro-ondes ou hyperfréquences correspondent à une bande comprise entre 300 MHz et 300 GHz. Les longueurs d'ondes associées sont comprises entre 1 m et 1 mm. C'est pourquoi on parle d'ondes métriques, décamétriques, centimétriques et millimétriques. Pour des plus petites longueurs d'ondes, le domaine est appelé sub-millimétrique ou quasi-optique. Il correspond à des technologies intermédiaires entre les micro-ondes et l'optique.

Le spectre micro-onde est partagé en différentes bandes et correspond à des standards de guides d'onde bien identifiés (tableau 6.1).

Tableau 6.1 – Bandes de fréquence micro-ondes

Nom de la bande	Fréquences
Bande L	1 à 2 GHz
Bande S	2 à 4 GHz
Bande C	4 à 8 GHz
Bande X	8 à 12 GHz
Bande Ku	12 à 18 GHz
Bande K	18 à 26 GHz
Bande Ka	26 à 40 GHz
Bande V	40 à 75 GHz
Bande W	75 à 111 GHz

6.2 Les systèmes de communications

Les communications représentent un domaine très vaste dans lequel les systèmes utilisent des antennes de formes variées. Au cours des années, avec les progrès de l'électronique, les fréquences ont eu tendance à augmenter, d'une part, en raison de l'encombrement du spectre et, d'autre part, en raison des demandes de transmission de débits élevés. Au début des télécommunications, il paraissait extraordinaire de transmettre de la voix. On est ensuite passé à la transmission de données. Après le développement des technologies sans fil qui ont permis de transmettre internet, il est maintenant envisagé de transmettre la télévision sur les mobiles.

Les types d'antennes ont beaucoup évolué dans chaque domaine d'applications. L'exemple le plus frappant est celui des antennes de terminaux mobiles qui sont passées, en très peu de temps, d'antennes filaires rectilignes aux antennes hélicoïdales, puis aux antennes patch, ceci en satisfaisant des contraintes correspondant à une diminution de la taille et à une esthétique imposée par le marché grand public.

Les antennes, de plus en plus sophistiquées, sont conçues en fonction des applications visées imposant un gabarit spécifique. Les moyens de simulation sont indispensables et de plus en plus perfectionnés. Ils permettent de répondre à des spécifications très particulières qui seront décrites au fur et à mesure de cet ouvrage.

En premier lieu, avant d'aborder les différentes antennes dans les systèmes de communication, nous allons rappeler quelques bases de propagation.

6.2.1 Notions de propagation

Lorsqu'on considère la transmission des ondes, en milieu infini et homogène, entre l'émetteur et le récepteur, le trajet optimal est la ligne droite représentant le trajet minimal entre les deux antennes. La propagation en vue directe fait apparaître une atténuation selon l'inverse du carré de la distance. Ce cas est plutôt rare dans les applications courantes et généralement d'autres phénomènes viennent perturber cette propagation en espace libre. Il peut se produire :

- des réflexions sur le sol,
- des réflexions sur des objets environnants,
- des diffractions sur des obstacles
- des zones d'ombre si les obstacles sont importants.

On distingue généralement plusieurs types d'environnements selon le nombre d'obstacles rencontrés et les modèles de propagation correspondants sont différents. Citons par exemple :

- L'environnement extérieur, lui-même divisé en zones urbaine, suburbaine, rurale.
- L'environnement intérieur aux bâtiments.

Il n'est pas question ici de détailler les modèles, mais de justifier quelques points sensibles pour l'utilisation d'antennes.

■ Les trajets multiples

Les problèmes de transmission viennent très souvent du fait que, d'autres trajets que le trajet direct, sont empruntés par les ondes, faisant intervenir une ou plusieurs réflexions. Ces ondes interfèrent avec l'onde directe. Donc selon la position du récepteur, le déphasage entre les ondes conduit, soit à un renforcement de la puissance, soit à une diminution. Lorsque la puissance est pratiquement nulle, on parle d'évanouissement ou fading.

Nous allons expliquer comment se produit ce phénomène et quels sont les paramètres mis en jeu. Pour simplifier, nous présentons la démonstration pour un sol parfaitement réfléchissant et plan. Nous verrons ensuite comment modifier les résultats si le sol n'est pas parfaitement réfléchissant.

Considérons une antenne émettrice, constituée d'un dipôle placé en O' , à une hauteur h_1 . L'antenne de réception, aussi constituée d'un dipôle est placée en O'' , à la hauteur h_2 . Le trajet direct $O'O''$ fait avec l'horizontale un angle α_d (figure 6.2). Il existe un rayon, émis par O' qui se réfléchit sur le sol et arrive en O'' . Il est obtenu en traçant l'image de O' par rapport au sol : O_i . La droite joignant cette image à O'' coupe le sol en P.

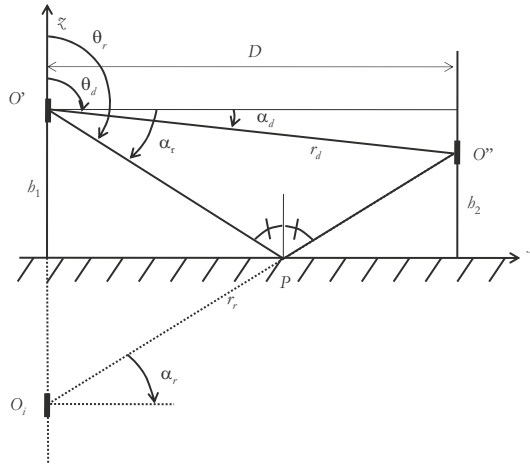


Figure 6.2 – Trajet direct et trajet réfléchi entre une antenne d'émission et une antenne de réception.

L'angle α_r définit la direction du rayon qui se réfléchit en P à l'horizontale.

En O_i se trouve donc le dipôle image de celui placé en O' . Pour trouver le sens du courant dans le dipôle image, il faut revenir à la propriété du dipôle énoncée dans le paragraphe 2.4.4, à savoir que deux charges de signes opposés se développent aux extrémités d'un dipôle (figure 6.3). Il suffit alors de déterminer les images de ces deux charges pour connaître le sens du dipôle image. L'image d'une charge est placée symétriquement par rapport au sol et possède une charge opposée. Par construction, le dipôle image est donc dans le même sens que le dipôle O' .

Le champ créé par un dipôle est donné par l'expression [3.10]. Le champ du dipôle placé en O' a pour valeur :

$$\vec{E}_d = jkZ \frac{Il}{4\pi} \frac{e^{-jkr_d}}{r_d} \sin \theta_d \vec{u}_\theta$$

Ses composantes selon le repère cartésien sont obtenues par projection :

$$E_{dx} = 0$$

$$E_{dy} = -jkZ \frac{Il}{4\pi} \frac{e^{-jkr_d}}{r_d} \cos \alpha_d \sin \alpha_d$$

$$E_{dz} = -jkZ \frac{Il}{4\pi} \frac{e^{-jkr_d}}{r_d} \cos^2 \alpha_d$$

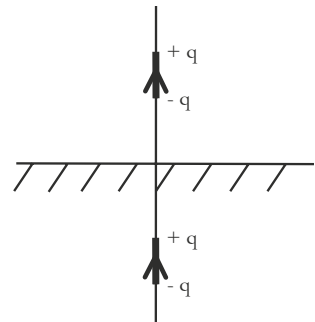


Figure 6.3 – Dipôle image.

De la même façon, le dipôle image créé en O'' un champ $\vec{E}_r$:

$$\vec{E}_r = jkZ \frac{Il}{4\pi} \frac{e^{-jkr_r}}{r_r} \sin \theta_r \vec{u}_\theta$$

Les composantes de ce champ sont :

$$E_{rx} = 0$$

$$E_{ry} = jkZ \frac{Il}{4\pi} \frac{e^{-jkr_r}}{r_r} \cos \alpha_r \sin \alpha_r$$

$$E_{rz} = -jkZ \frac{Il}{4\pi} \frac{e^{-jkr_r}}{r_r} \cos^2 \alpha_r$$

Le champ total $\vec{E}$ en O'' est donc obtenu par la somme des champs créés par le dipôle en O' et par son image en O_i :

$$E_x = 0$$

$$E_y = -jkZ \frac{Il}{4\pi} \left[\frac{e^{-jkr_d}}{r_d} \cos \alpha_d \sin \alpha_d - \frac{e^{-jkr_r}}{r_r} \cos \alpha_r \sin \alpha_r \right]$$

$$E_z = -jkZ \frac{Il}{4\pi} \left[\frac{e^{-jkr_d}}{r_d} \cos^2 \alpha_d + \frac{e^{-jkr_r}}{r_r} \cos^2 \alpha_r \right]$$

En général, les antennes sont placées à grande distance l'une de l'autre et la distance D est bien supérieure aux hauteurs des dipôles. Les angles α_d et α_r sont donc petits. Cette remarque permet des approximations :

$$\cos^2 \alpha_d \approx 1 - \alpha_d^2 \quad \text{et} \quad \sin \alpha_d \approx \alpha_d$$

$$\cos^2 \alpha_r \approx 1 - \alpha_r^2 \quad \text{et} \quad \sin \alpha_r \approx \alpha_r$$

La composante selon Oz est donc prépondérante sur les autres composantes. Ceci justifie le fait de placer le dipôle récepteur verticalement afin de recevoir le maximum de champ.

Par ailleurs, les rayons vecteurs s'expriment, en utilisant ces approximations, par :

$$r_d \approx D \left(1 + \frac{\alpha_d^2}{2} \right) \quad \text{et} \quad r_r \approx D \left(1 + \frac{\alpha_r^2}{2} \right)$$

En ne conservant dans E_z que les développements limités relatifs à la fonction exponentielle qui varie beaucoup plus vite que la fonction inverse, on obtient :

$$E_z \approx -jkZ \frac{Il}{4\pi D} e^{-jkD} \left[e^{-jkD \frac{\alpha_d^2}{2}} + e^{-jkD \frac{\alpha_r^2}{2}} \right]$$

Compte tenu des valeurs très faibles des angles, les approximations suivantes sont valables :

$$\alpha_d \approx \frac{h_1 - h_2}{D}$$

$$\alpha_r \approx \frac{h_1 + h_2}{D}$$

Reportons ces approximations dans l'expression de E_z , il vient :

$$E_z \approx -jkZ \frac{Il}{4\pi D} e^{-jkD} e^{-j\frac{k}{2D}(h_1^2 + h_2^2)} \left[e^{jk\frac{h_1 h_2}{D}} + e^{-jk\frac{h_1 h_2}{D}} \right]$$

Soit encore :

$$E_z \approx -jkZ(Il) \frac{e^{-jkD}}{4\pi D} 2 \cos \left(k \frac{h_1 h_2}{D} \right)$$

En rapportant cette valeur à la composante en z du champ direct, on obtient :

$$\left| \frac{E_z}{E_{dz}} \right| = \left| 2 \cos \left(k \frac{h_1 h_2}{D} \right) \right|$$

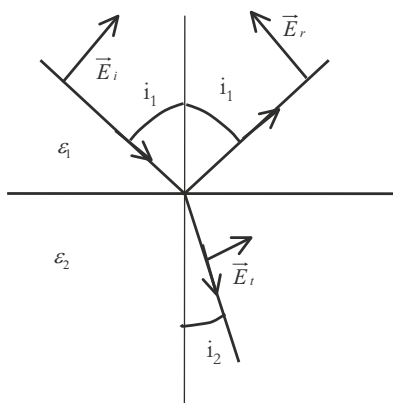
On conçoit donc qu'en fonction des hauteurs des deux antennes et de la distance D , l'onde réfléchie peut entraîner des évanouissements du champ, lorsque le cosinus de l'expression précédente s'annule.

Afin de prendre en compte la nature du sol, il est possible d'introduire dans le calcul le coefficient de réflexion R du sol. Le champ reçu est alors :

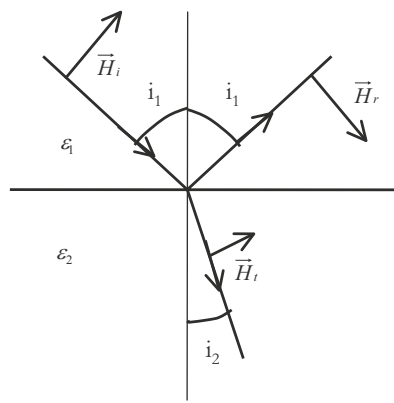
$$E_z \approx -jkZ \frac{Il}{4\pi D} e^{-jkD} e^{-j\frac{k}{2D}(h_1^2 + h_2^2)} \left[e^{jk\frac{h_1 h_2}{D}} + R e^{-jk\frac{h_1 h_2}{D}} \right]$$

Avec les conventions d'orientation de la figure 6.4, le coefficient de réflexion sur un réflecteur parfait est $R = 1$, pour la composante en z du champ électrique.

Les formules de Fresnel donnent les coefficients de réflexion entre deux milieux diélectriques de permittivités différentes (ϵ_1 et ϵ_2), pour les deux polarisations des ondes électromagnétiques, dans la direction spéculaire. Ces formules sont rappelées ci-dessous pour les deux polarisations parallèle et perpendiculaire, définies par la figure 6.4 :



Polarisation parallèle



Polarisation perpendiculaire

Figure 6.4 – Définition des polarisations pour le calcul des coefficients de réflexion (à gauche polarisation parallèle, à droite polarisation perpendiculaire).

La polarisation parallèle correspond au cas où le champ électrique est parallèle au plan d'incidence (défini par la normale et le rayon incident). Cette polarisation est aussi appelée TM (transverse magnétique) car le champ magnétique est perpendiculaire au plan de la figure. Cela correspond aussi à une polarisation verticale.

Le coefficient de réflexion R_{par} est le coefficient de réflexion en amplitude pour le champ électrique défini avec les orientations de la figure 6.4.

$$R_{par} = \frac{\sqrt{\epsilon_2} \cos i_1 - \sqrt{\epsilon_1} \cos i_2}{\sqrt{\epsilon_2} \cos i_1 + \sqrt{\epsilon_1} \cos i_2}$$

La polarisation perpendiculaire correspond au cas où le champ électrique est perpendiculaire au plan d'incidence. C'est aussi une polarisation TE (transverse électrique) ou une polarisation horizontale.

Le coefficient de réflexion R_{perp} est le coefficient de réflexion en amplitude pour le champ électrique défini avec les orientations de la figure 6.4.

$$R_{perp} = \frac{\sqrt{\epsilon_1} \cos i_1 - \sqrt{\epsilon_2} \cos i_2}{\sqrt{\epsilon_1} \cos i_1 + \sqrt{\epsilon_2} \cos i_2}$$

Rappelons la loi de Descartes :

$$\sqrt{\epsilon_1} \sin i_1 = \sqrt{\epsilon_2} \sin i_2$$

Le cas d'une antenne linéaire verticale de type dipolaire correspond à une polarisation parallèle. Lorsque l'incidence est proche de 90°, le coefficient de réflexion, pour des matériaux diélectriques, tend vers -1. Le champ électrique prend alors la forme :

$$E_z \approx -jkZ \frac{Il}{4\pi D} e^{-jkD} e^{-j\frac{k}{D}(h_1^2 + h_2^2)} \left[e^{jk\frac{h_1 h_2}{D}} - e^{-jk\frac{h_1 h_2}{D}} \right]$$

Soit, lorsque la distance entre les antennes est grande par rapport à leur hauteur :

$$\left| \frac{E_z}{E_{dz}} \right| = \left| 2 \sin \left(k \frac{h_1 h_2}{D} \right) \right| \approx 2k \frac{h_1 h_2}{D}$$

Ce coefficient donne une atténuation en puissance proportionnelle à l'inverse du carré de la distance entre les antennes. Étant donné la variation de la puissance rayonnée de l'antenne en l'absence du sol, l'atténuation totale du signal varie comme l'inverse de la puissance quatre de la distance.

On constate, d'après ces calculs de base, que la réflexion d'un signal sur le sol donne des résultats variés selon la nature du sol, la distance entre les antennes et leur hauteur. Nous n'avons pas traité tous les cas possibles, mais simplement donné quelques éléments de base pour le calcul du champ réfléchi. Remarquons cependant deux points :

- Si l'antenne est telle que la polarisation est horizontale (TE), les résultats sont différents, en particulier par le fait que l'image du dipôle horizontal n'a pas le même sens que celui du dipôle induit un signe moins dans les calculs. D'autre part le coefficient de réflexion à prendre en compte est celui qui est relatif à la polarisation horizontale.
- Le sol n'est pas un diélectrique parfait. Il présente une conductivité σ . La permittivité effective est alors définie par :

$$\epsilon_{eff} = \epsilon_r - j \frac{\sigma}{\epsilon_0 \omega}$$

- La permittivité effective dépend de la fréquence. Elle doit être utilisée dans les formules de Fresnel pour déterminer le coefficient de réflexion.

Il n'y a pas que les réflexions sur le sol qui contribuent au champ total reçu. Dans un environnement urbain, de nombreuses réflexions ont lieu sur les bâtiments ou d'autres objets qui peuvent contribuer à ces évanouissements. C'est ce qu'on appelle en propagation le problème des trajets multiples qui ne sont pas faciles à prévoir dans un environnement complexe. On les traite souvent par des méthodes statistiques pour aboutir à différents modèles rendant compte de la complexité

de façon globale, donc approchée. Les évanouissements du signal sont la première conséquence de l'existence de trajets multiples. La seconde conséquence se mesure sur la forme temporelle du signal qui, du fait des différentes réflexions, subit des retards successifs.

Signalons que les principes présentés font l'hypothèse d'une terre plate. Pour plus de précision, lorsque la distance de transmission est grande, il est préférable de tenir compte du modèle de terre sphérique.

■ Ellipsoïde de Fresnel

L'étude précédente permet de comprendre que tout obstacle se trouvant à proximité du trajet en ligne directe perturbe la transmission en réfléchissant ou diffractant l'onde. Une deuxième onde est ainsi créée qui interfère avec celle du trajet direct. Lorsque ces deux ondes ont des différences de marche égales à la demi-longueur d'onde, l'interférence est destructive et le signal s'évanouit. Le lieu des points tels que la différence de marche entre le trajet direct et le trajet ayant subi une réflexion est égale à une valeur donnée est un ellipsoïde ayant pour foyers les deux antennes. Lorsque cette différence est égale à la demi-longueur d'onde, on parle du premier ellipsoïde de Fresnel.

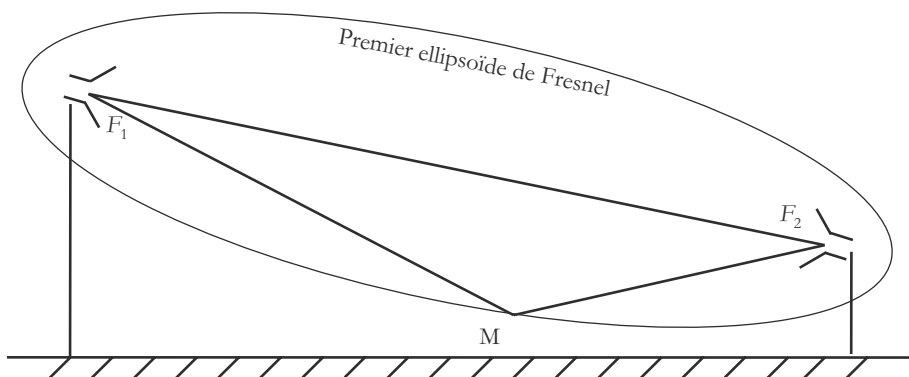


Figure 6.5 – Premier ellipsoïde de Fresnel.

Sur la figure 6.5, le point M est tel que :

$$F_1M + F_2M - F_1F_2 = \frac{\lambda}{2}$$

Tout obstacle entrant dans le premier ellipsoïde de Fresnel abaisse le niveau de puissance reçue. Le choix des sites des antennes de stations hertziennes s'effectue donc en fonction du dégagement du premier ellipsoïde de Fresnel dans l'environnement des antennes.

Les phénomènes de diffraction influent naturellement sur la propagation. Nous les signalons simplement dans le cadre de cet ouvrage

6.2.2 Radioamateurs

L'Union Internationale des Télécommunications définit l'activité des opérateurs des services d'amateur appelés radioamateurs comme « services de radiocommunication ayant pour objet l'instruction individuelle, l'intercommunication et les études techniques, effectuées par des amateurs, c'est-à-dire par des personnes dûment autorisées, s'intéressant à la technique de la radio-électricité à titre uniquement personnel et sans intérêt pécuniaire ».

Dans ce cadre, les radioamateurs ont développé quantité d'applications ; on peut notamment citer :

- Téléphonie.
- Transmission de texte : radio télétype (RTTY, télécriteur), AMTOR.
- Transmission d'image : ATV (TV Amateur), transmission d'image à balayage lent (SSTV), fac-similé (fax).
- Transmissions numériques : réseau de communications numériques (packet-radio, APRS).

Plus précisément, listons quelques activités prises comme exemples. Un des challenges que s'imposent les radioamateurs est la transmission à faible puissance (QRP) pour établir une communication, notamment de type HF pour atteindre d'autres continents. Plus haut en fréquence, on trouve les applications à bases de satellites radioamateurs (OSCAR, *Orbiting Satellites Carrying Amateur Radio*) qui permettent une liaison entre deux points terrestres *via* un relais satellitaire dans les bandes de fréquences VHF et UHF. La réception d'images de la Terre prise depuis un satellite est aussi une activité très prisée. Enfin, on peut mentionner que régulièrement des contacts sont établis entre les navettes américaines (Columbia, Challenger) avec des stations radioamateurs, et cela avec du matériel radioamateur embarqué sur la navette (charge utile SAREX). Tout d'abord, nous listons les bandes autorisées :

Comme on peut le constater, l'attribution de bandes de fréquences occupant tout le spectre hertzien permet d'exploiter les conditions de propagation très différentes et donc toutes les possibilités de transmission (vision directe, réflexion ionosphérique, troposphérique, diffraction...).

Cela conditionne la conception de l'antenne du radioamateur. Les radioamateurs ont beaucoup investigué dans les bandes décimétriques car les antennes sont constituées de fils, donc plus facilement réalisables. De façon générale, les antennes destinées aux émissions radioamateur sont basées sur les techniques utilisées pour la conception des antennes classiques. Toutefois, il faut relever certaines spécificités. En effet, lorsqu'un radioamateur conçoit une antenne décimétrique, il ne s'agit plus de faire une antenne qui soit optimale sur une seule fréquence mais fonctionnelle sur plusieurs bandes. Cela lui permettra de choisir la fréquence optimale en fonction de la liaison envisagée, des conditions météo en altitude... Par exemple, l'utilisation commune des bandes 3,5, 7, 14 et 21 MHz permet d'utiliser un dipôle demi-onde à 3,5 MHz, dipôle qui autorisera un accord onde entière à 7 MHz puis en 2λ à 14 MHz, 3λ à 21 MHz et enfin 4λ à 28 MHz.

Or, les nouvelles bandes récemment attribuées à 10, 18 et 24,9 MHz ne permettent plus cette approche puisque celles-ci n'ont plus la relation de multiples entiers entre elles ni avec les autres bandes. Il faut donc que l'antenne en ondes décimétriques devienne multibandes.

Sur la figure 6.6, on représente la répartition de courant (ondes stationnaires) sur une antenne dont la longueur est comparée à la longueur d'ondes pour différentes fréquences d'utilisation.

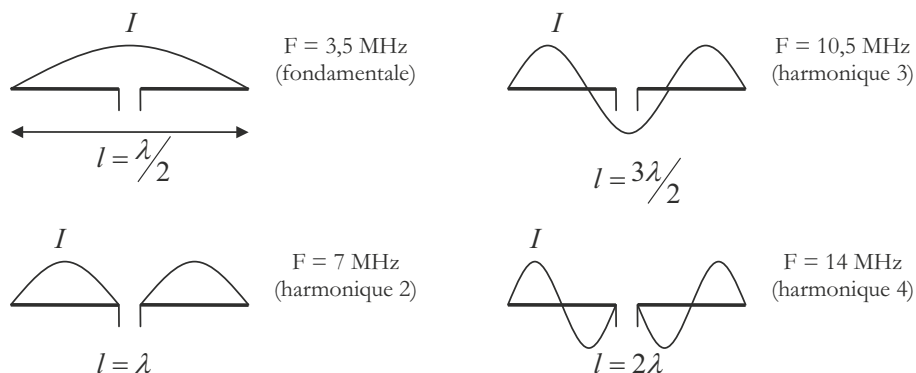


Figure 6.6 – Répartition des ondes de courants sur l'antenne multibandes alimentée au centre.

Tableau 6.2 – Bandes décamétriques

Fréquence inférieure (MHz)	Fréquence supérieure (MHz)	Longueur d'onde moyenne (m)
1,830	1,850	16
3,500	3,800	79 à 86
7,000	7,100	42,5
10,100	10,150	29,5
14,000	14,350	21
18,068	18,168	16,5
21,000	21,450	14
24,890	24,990	12
28,000	29,700	10

Tableau 6.3 – Bandes métriques et décimétriques

Fréquence inférieure (MHz)	Fréquence supérieure (MHz)	Longueur d'onde moyenne (m)
50,2	51,2	6
144	140	2
430	440	0,7
1 240	1 300	0,24
2 300	2 450	0,13

Tableau 6.4 – Bandes centimétriques et millimétriques

Fréquence inférieure (GHz)	Fréquence supérieure (GHz)	Longueur d'onde moyenne (cm)
5,65	5,85	5
10	10,5	3
24	24,25	1,2
47	47,2	0,6
75,5	81	0,375
119,98	120	0,25
142	149	0,2
241	250	0,12

■ Importance de l'emplacement du point d'alimentation

Dans ce fonctionnement multibandes, la position du point d'alimentation prend beaucoup d'importance puisqu'il va conditionner l'impédance d'entrée vue par la ligne connectée à l'antenne.

On considère que le point d'alimentation doit correspondre au point d'intensité maximale (adaptation d'impédance facilitée). Dans le cas d'une alimentation centrée, les fréquences de résonance correspondent aux harmoniques 1, 3, 5, 7... Tandis que, dans le cas de l'alimentation excentrée, ils correspondent aux harmoniques 1, 3, 9, 15..., donc à un nombre réduit de bandes.

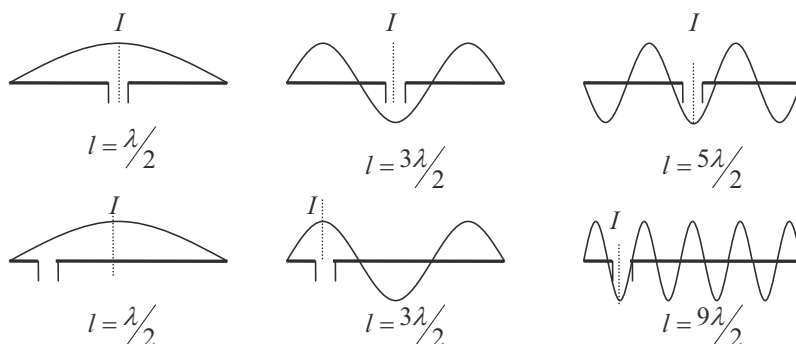


Figure 6.7 – Répartition des ondes de courants sur l'antenne multibandes à alimentation excentrée.

En ce qui concerne l'accord électrique, si l'on veut concevoir une antenne fonctionnant sur l'ensemble des bandes décadiques, on devra accorder avec un LC série lorsque l'impédance au point d'alimentation est faible et avec un LC parallèle lorsque celle-ci est forte.

■ Effet de sol

De même l'effet du sol va entraîner une modification du diagramme de rayonnement et notamment de l'angle de rayonnement maximal par combinaison des rayonnements direct et réfléchi. Aussi, les radioamateurs font très attention à maîtriser la hauteur de l'antenne au-dessus du sol de façon atteindre l'angle optimal de rayonnement en fonction de la fréquence utilisée, comme cela est montré sur la figure 6.8. L'angle optimal étant celui qui évite les multibonds et donc les pertes par réflexion sur le sol qui n'est pas parfaitement conducteur.

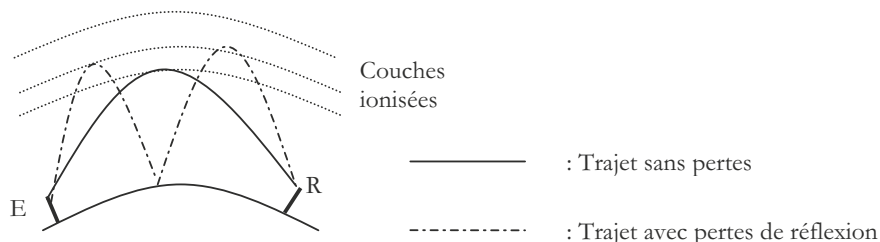


Figure 6.8 – Importance de l'angle de rayonnement.

■ Exemple d'antennes en bandes décadiques

- **L'antenne *ground plane*** : c'est une antenne quart d'onde omnidirectionnelle à polarisation verticale avec sol artificiel. Comme elle est asymétrique, elle est idéalement utilisée avec une ligne asymétrique (coaxiale). L'adaptation se fait par transformateur quart d'onde de coaxial $50\ \Omega$ entre l'antenne d'impédance $36\ \Omega$ et une ligne d'amenée de $75\ \Omega$, ou alors par inclinaison du sol artificiel (l'impédance variera entre les bornes 36 et $73\ \Omega$) et donc une adaptation possible à une ligne d'amenée $50\ \Omega$. Le rayonnement est maximum dans un plan quasi horizontal et donc se prête bien au trafic VHF en émission ou en transmission ionosphérique.
- **L'antenne *Zeppelin*** : elle doit comporter un nombre entier de demi-ondes le long de son brin rayonnant constitué par une ligne accordée alimentée en bout et est plus difficile à mettre au point que l'antenne Levy.

- **L'antenne Levy (double Zepp)** : elle est l'antenne de prédilection du radioamateur qui veut émettre ou recevoir sur l'ensemble des bandes décadiques de 3,5 à 30 MHz. Avec un circuit d'accord variable, on peut l'accorder sur l'ensemble de la plage de fréquences mentionnée.

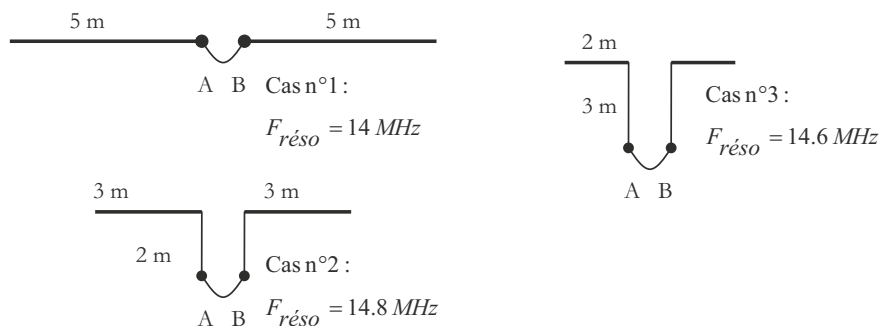


Figure 6.9 – Principe de fonctionnement de l'antenne Levy.

En connectant en entrée d'antenne une boucle, on autorise la mesure de la fréquence de résonance à l'aide d'un grid-dip (sorte de pince ampérométrique). On constate que la fréquence de résonance n'est pas modifiée si la somme des longueurs horizontales (rayonnantes) et de la ligne parallèle reste constante (demi-onde). Cela signifie que le mode de vibration est celui de la demi-onde, et cela malgré le fait que la partie repliée parallèle ne rayonne pas puisque les fils annulent mutuellement leur rayonnement. On considère qu'il y a réduction du rayonnement dû à ce raccourcissement à partir d'une longueur rayonnante totale inférieure au quart d'onde.

En combinant au pied de l'antenne une combinaison de circuits d'accord LC série ou parallèle, il est possible de travailler sur l'ensemble des bandes. En effet, le circuit d'accord permet de « retomber » sur un nombre impair de demi-ondes et donc de considérer le point de jonction avec la ligne (point AB) comme un point d'intensité maximale lorsque l'on est dans un fonctionnement multiple de demi-ondes.

Le dernier avantage de cette antenne est lié à sa symétrie qui permet d'établir un régime d'ondes stationnaires parfaitement équilibré sur les deux moitiés de la structure.

On pourrait encore citer comme autre type d'antennes décadiques, l'antenne « Trap » qui est constituée d'un dipôle dans lequel ont été insérés des circuits résonants parallèles à des positions précises (qui jouent le rôle de pièges) et qui permettent de reconstituer le fonctionnement sur des bandes particulières qui ne sont pas harmoniques (fréquences basses de la bande HF). On peut enfin citer l'antenne cadre qui s'inspire à la fois de l'antenne Quad dans sa forme et de l'antenne à trappes résonnantes.

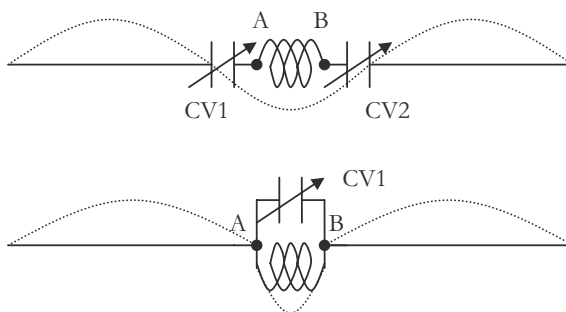


Figure 6.10 – Circuits d'accord pour fonctionnement multibandes.

Dans les bandes de fréquences supérieures, les techniques utilisées par les radioamateurs sont plus classiques. On peut citer les antennes Quad et Yagi, qui ont été utilisées et souvent mises au point de façon empirique pour les émissions de Télévision amateur (ATV). De nombreuses techniques d'adaptation d'antennes (coupleurs d'antennes) ont été étudiées, testées et éprouvées, qui permettent de prendre en compte les limites de réalisations pratiques.

6.2.3 Les réseaux de communications

On a pris l'habitude de distinguer les réseaux hertziens des réseaux mobiles. Les réseaux hertziens sont les réseaux qui utilisent le canal de propagation constitué par l'air entre deux points fixes. Les réseaux mobiles utilisent aussi le canal aérien. Ils sont destinés plus spécifiquement aux communications mobiles entre un point fixe (la station de base) et un point variable dans l'espace (le mobile). Certains systèmes sont développés de façon à communiquer entre des stations de bases variables dans le temps et l'espace. Ce sont les réseaux ad hoc.

■ Les réseaux hertziens

Ces réseaux sont utilisés pour les communications fixes et la diffusion de la radio ou de la télévision. Il y a plusieurs bandes de fréquences utilisées pour la télédiffusion terrestre qui se trouvent dans les bandes UHF et VHF. Pour les réseaux de télécommunications fixes, les fréquences sont comprises entre 1 et 40 GHz. Jusqu'à 20 GHz, les réseaux sont très encombrés.

Les faisceaux hertziens sont utilisés pour effectuer une transmission point à point grâce au transport de l'information sur une porteuse modulée. Actuellement la tendance s'oriente vers le choix des modulations numériques qui sont moins sensibles aux perturbations et permettent des traitements du signal plus performants.

Les liaisons par faisceaux hertziens servent à établir des communications dans des endroits non équipés en réseaux filaires. Ils permettent aussi des installations rapides et éventuellement provisoires.

Le réseau hertzien de diffusion est constitué d'un maillage régulier sur un territoire. Il permet de transporter la même information, au même instant sur ce territoire. Le réseau de télévision français, par exemple, est constitué d'environ 14 000 émetteurs répartis en 4 000 sites pour les six chaînes de base. Le réseau de la télévision numérique terrestre a démarré en 2005 avec 17 émetteurs couvrant 35 % du territoire français. La couverture, ensuite étendue en 2006 à 60 % du territoire, devrait couvrir l'ensemble de la France ultérieurement.

Les réseaux hertziens sont constitués d'antennes placées en vue directe l'une par rapport à l'autre. La portée des faisceaux est comprise entre 10 et 60 km, pouvant aller exceptionnellement jusqu'à 100 km. Des stations relais sont utilisées si la distance s'avère trop grande. Le relief, les bâtiments et la végétation constituent des obstacles potentiels à la propagation des faisceaux.

Les antennes sont placées en hauteur de façon à assurer une meilleure liaison. Nous avons vu précédemment que l'essentiel de l'énergie se trouve concentré dans le premier ellipsoïde de Fresnel. Tout obstacle se trouvant dans cet ellipsoïde abaisse la qualité de la liaison. Les hauteurs des antennes sont donc calculées en tenant compte des dimensions de celui-ci.

Les antennes de réseaux hertziens doivent être très directives, puisqu'elles assurent une transmission point à point. Les antennes à réflecteur parabolique (paragraphe 8.6) sont très utilisées, pour diverses raisons : leur grande directivité, leur robustesse, leur puissance d'émission. Elles sont souvent recouvertes d'un radôme de protection contre les intempéries et les différences de températures.

■ Les réseaux de télécommunications mobiles

Les réseaux de télécommunications permettent de transmettre des signaux entre une station de base et des terminaux mobiles. Les stations sont réparties sur le territoire et reliées entre elles. La

distance qui les sépare dépend de la densité des communications à transmettre et de l'environnement. Un environnement urbain nécessite des stations plus rapprochées en raison des nombreux obstacles que représentent les bâtiments, mais aussi en raison d'un trafic plus important. Chaque station de base dessert une zone l'entourant appelée cellule. Les cellules sont plus ou moins serrées selon le type de zone concernée. On parle de réseau microcellulaire dans une zone urbaine et de réseau picocellulaire plutôt à l'intérieur des bâtiments. La taille d'une microcellule est comprise entre 100 et 300 m. Lorsque la taille des cellules augmente, on passe au réseau cellulaire avec des tailles allant de 500 m à 2 km. Les macrocellules ont des dimensions comprises entre 2 km et 35 km.

Le réseau cellulaire s'appuie sur les deux éléments que sont les stations et les terminaux mobiles qui possèdent chacun la fonction d'émission et de réception. Les puissances mises en jeu à l'émission par les stations de base sont importantes. Elles dépendent du type de station. Les mobiles émettent une puissance plus faible, au maximum de 2 W pour les terminaux portables. Du fait des rôles différents des stations de base et des mobiles, leurs antennes sont différentes.

□ Antennes de terminaux mobiles

Les antennes de terminaux sont souvent très proches d'un plan de masse imparfait (le sol, le toit d'un véhicule ou le corps humain : la main ou la tête) qui modifie le diagramme de rayonnement. Les antennes pour les terminaux mobiles évoluent sans cesse. Ainsi les techniques utilisées pour les téléphones portables sont passées en quelques années d'antennes fouets, aux antennes hélicoïdales, puis aux antennes planaires. Certaines antennes sont repliées à des fins de miniaturisation (antennes PIFA, paragraphe 9.1.3). Les antennes utilisées sur les véhicules sont très souvent des antennes fouets qui sont constituées d'une antenne filaire de longueur égale à un quart d'onde au-dessus d'un plan de masse. Il est possible d'utiliser des antennes filaires de différentes tailles. Afin d'adapter les antennes, il est souvent nécessaire d'ajouter des selfs ou des capacités à l'entrée.

Nous allons montrer l'influence du plan de masse.

Prenons l'exemple d'une antenne filaire placée à une hauteur h au-dessus d'un plan conducteur, constitué soit par le sol, soit par le toit d'un véhicule. Nous allons donc étudier ce cas représenté par la figure 6.11, sur laquelle, le problème initial d'une antenne en présence d'un plan réflecteur, est représenté à gauche. Ce problème est remplacé par le problème équivalent de deux antennes, parcourues par un courant de même sens, distantes de $2h$ (voir paragraphe 2.4.4, Théorie des images).

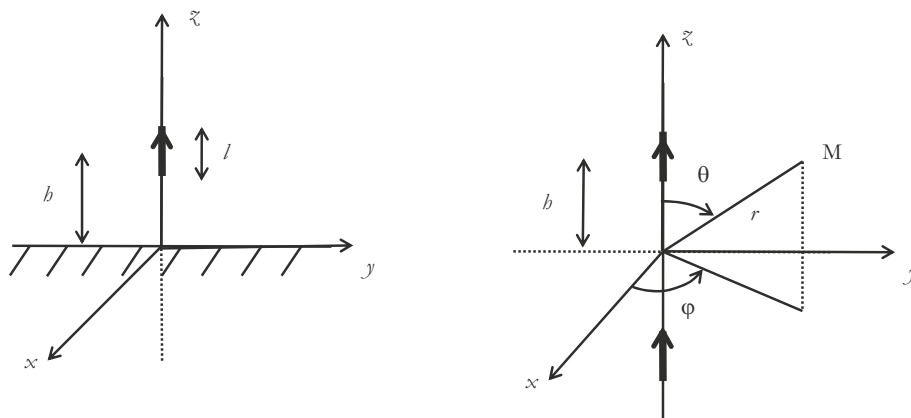


Figure 6.11 – Rayonnement d'une antenne filaire au-dessus d'un plan de masse et son problème équivalent.

On rappelle la fonction caractéristique d'une antenne filaire rectiligne :

$$F(\theta) = \left(\frac{\cos(n\pi \cos \theta) - \cos n\pi}{\sin \theta} \right)^2$$

Les deux antennes du problème équivalent sont décalées de h de chaque côté du plan. D'après la théorie des réseaux (paragraphe 5.1.2), le facteur de réseau introduit est égal à :

$$F_R = 4 (\cos(kh \cos \theta))^2$$

Le rayonnement global a pour fonction caractéristique de rayonnement le produit de ces deux fonctions. La symétrie axiale est maintenue. Le rayonnement du réseau constitué des deux antennes est maximum pour $\theta = \pi/2$. Le gain de cette antenne est double de ce qu'il serait pour une antenne seule, en absence de sol réflecteur, puisque la puissance n'est ici rayonnée que dans un demi-espace.

Pour des valeurs de θ différentes de $\pi/2$, le diagramme de rayonnement dépend de la hauteur du dipôle. En particulier, certaines directions peuvent présenter un rayonnement nul selon les valeurs de h .

□ Antennes de stations de base

Les stations de base doivent être puissantes (20 à 30 W). Les antennes sont donc associées sous forme de réseau. Le diagramme en résultant doit être omnidirectionnel autour de la station afin d'assurer la liaison avec les terminaux dont la position est quasiment aléatoire. Le diagramme dans le plan vertical ne doit pas présenter de zéros. L'ouverture verticale varie de 10° à 70° . La meilleure solution pour réaliser ce type de diagramme est l'association d'antennes verticalement, alimentées en amplitude et en phase afin de créer le diagramme de rayonnement recherché (voir le chapitre 5 sur les réseaux). On trouve ainsi des réseaux d'antennes filaires rectilignes (souvent des antennes demi-ondes), alignées et placées sur un mat.

À la réception, le réseau permet d'utiliser les techniques de diversités pour améliorer la qualité du signal (voir paragraphe 7.2).

La souplesse apportée par le contrôle de l'alimentation des réseaux permet aussi de rendre les antennes adaptatives : les faisceaux, à la réception, comme à l'émission, sont alors orientables dans des directions correspondant à celles d'un fort trafic. Il est aussi possible d'abaisser la valeur de la fonction caractéristique de l'antenne dans une direction où se trouve une source perturbatrice. L'ensemble de ces situations est géré grâce au traitement d'antennes (paragraphe 7.1)

Afin de renforcer le rayonnement, l'antenne peut être placée devant un plan de masse. Le principe des images permet de remplacer le problème initial par le problème équivalent de deux antennes antisymétriques (figure 6.12).

Calculons le champ rayonné par deux antennes filaires parallèles et parcourues par des courants d'intensité I , de sens opposés. Ces deux antennes constituent un réseau dont nous allons calculer le facteur de réseau F_R , avec les notations de la figure 6.12. La somme vectorielle de la contribution des deux antennes conduit au calcul de la fonction :

$$f_R(\theta, \phi) = e^{jkh \sin \theta \sin \phi} - e^{-jkh \sin \theta \sin \phi}$$

Soit :

$$f_R(\theta, \phi) = 2j \sin(kh \sin \theta \sin \phi)$$

La fonction réseau est donc donnée par :

$$F_R(\theta, \phi) = 4 \sin^2(kh \sin \theta \sin \phi)$$

Dans le cas $h = \lambda/4$:

$$F_R(\theta, \phi) = 4 \sin^2 \left(\frac{\pi}{2} \sin \theta \sin \phi \right)$$

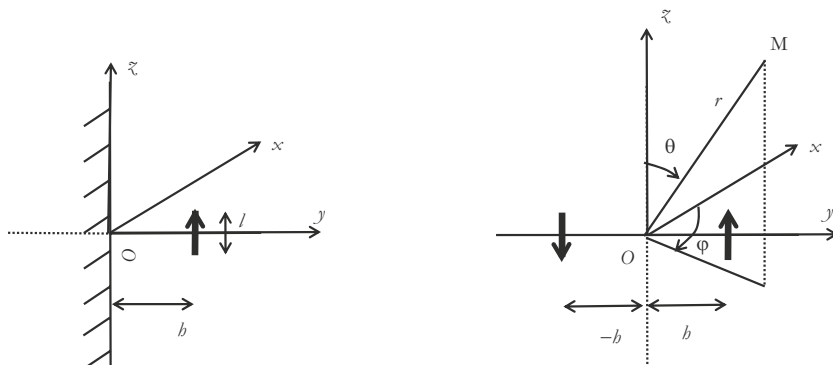


Figure 6.12 – Antenne filaire parallèle au plan de masse et son problème équivalent.

Cette fonction est maximum pour : $\theta = \pi/2$ et $\varphi = \pi/2$, c'est-à-dire dans la direction Oy , ce qui est normal, puisque les deux antennes présentent un déphasage nul dans la direction Oy . En effet les deux antennes sont séparées de $\lambda/2$. Donc la différence de marche introduit un déphasage de π , et du fait de l'opposition des deux courants elles présentent un déphasage supplémentaire de π . Le déphasage total est donc 2π . Il n'y a pas d'autre zéro de la fonction réseau.

Dans ce cas $h = \lambda/2$:

$$F_R(\theta, \phi) = 4 \sin^2(\pi \sin \theta \sin \phi)$$

Il existe, dans ce cas un zéro pour $\theta = \pi/2$ et $\varphi = \pi/2$, c'est-à-dire dans la direction Oy . Cela s'explique par la différence de marche dans la direction Oy qui correspond à un déphasage de 2π . Au final, en prenant en compte le déphasage dû au courant, le déphasage total est de π dans la direction Oy .

Donc le choix de la distance de l'antenne au plan de masse est fondamental. Afin d'imposer un diagramme de rayonnement visant dans la direction perpendiculaire au réflecteur, la distance h est choisie égale au quart de la longueur d'onde. Le diagramme de rayonnement est alors obtenu en multipliant la fonction caractéristique de l'antenne filaire par le facteur de réseau des deux antennes, représenté sur la figure 6.13. Bien entendu, le diagramme de rayonnement du problème initial se limite au demi-plan supérieur, le plan de masse jouant le rôle d'écran pour la partie inférieure.

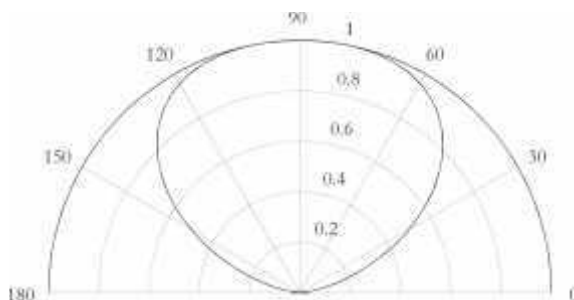


Figure 6.13 – Facteur de réseau de deux antennes parallèles, situées à $\lambda/4$ d'un plan de masse et alimentées par des courants de sens opposés.

L'ouverture de rayonnement est supérieure à 120° , d'où le nom d'**antenne sectorielle**, donné à ce type d'antenne.

Les antennes de stations de base les plus couramment utilisées pour le GSM sont des **antennes panneaux**. Elles sont constituées d'antennes filaires (quelquefois planaires) associées en un réseau linéaire vertical, placées devant un plan de masse. Chaque réseau constitue une antenne sectorielle. Les panneaux sont placés selon un triangle équilatéral, de façon à couvrir les 360° de l'espace autour d'un pylône (figure 6.14), en choisissant des antennes sectorielles d'ouverture 120°.

Le réseau est alimenté de façon à imposer un léger déphasage linéaire qui entraîne un angle de visée vers le bas, appelé *tilt* de l'antenne.

Un autre type d'antenne sectorielle utilise la réflexion sur un dièdre métallique. Le principe en est le même que celui de la réflexion sur un plan métallique. Il existe alors une image par rapport à chaque plan métallique. Le réseau équivalent comporte quatre éléments.

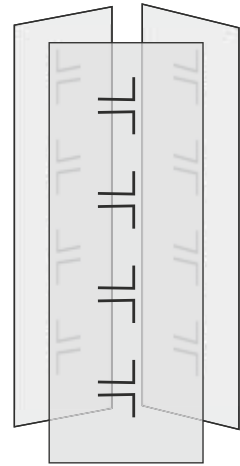


Figure 6.14 – Antenne panneaux.

6.2.4 Liaisons par satellites

Les techniques d'antennes ont évolué rapidement au moment de l'avènement des radars. Un nouvel essor a relancé leur développement avec le développement des antennes satellitaires.

Les caractéristiques recherchées pour les antennes satellites sont un grand gain (supérieur à 40 dB), un diagramme de rayonnement présentant si possible des lobes secondaires faibles. Nous verrons que les paraboles simples permettent d'obtenir des gains compris entre 40 et 50 dB, associés à des angles d'ouverture inférieurs au degré. Les antennes à réflecteurs secondaires permettent d'obtenir des gains de l'ordre d'une soixantaine de décibels, associés à des ouvertures de l'ordre de quelques minutes d'angle.

■ Différents types de satellites

Les satellites ont évolué au cours du temps depuis les premiers lancements dans les années soixante. Les premiers étaient passifs et utilisés comme réflecteurs du signal. Très rapidement les satellites ont eu des fonctions de réception et d'émission. Ils ont embarqué des charges utiles comportant des dispositifs électroniques de plus en plus complexes, rendant leur fonctionnement intelligent et autonome pour une durée de quelques années. Actuellement les satellites embarquent des récepteurs à très faible bruit, des émetteurs de puissance et un ensemble de fonctions électroniques qui dépendent du cahier des charges. Les satellites utilisent un dispositif de transposition de fréquence pour éviter les interférences entre les fréquences montantes et les fréquences descendantes, ces dernières étant en général plus basses.

Les satellites emportent de nombreuses antennes dont les fonctions sont différentes et qui utilisent différentes bandes de fréquences. Nous verrons leur utilisation au cours de ce chapitre. Les antennes satellitaires doivent présenter un grand gain et être qualifiées spatialement, c'est-à-dire qu'elles doivent répondre à certains critères essentiellement liés à la solidité, à la tenue en vibration, à la durée de vie et à la réponse en température. Ce dernier point est très sensible car les dilatations peuvent être très variables en raison des écarts de températures importantes dans l'espace. Ces écarts vont de -150 °C, côté espace, à plusieurs centaines de degrés, côté soleil.

□ Les satellites géostationnaires

Très rapidement l'utilité de satellites géostationnaires s'est imposée. Ces satellites, placés sur l'orbite circulaire équatoriale, à environ 36 000 km de la Terre, ont une période de rotation de 24 heures. Ils sont géosynchrones. L'orbite équatoriale abrite un peu moins de quatre cents satellites opérationnels. Elle est cependant encombrée d'anciens satellites et de débris de lancement. On y dénombre un peu plus d'un millier d'objets de plus d'un mètre.

La puissance est un problème crucial pour les antennes embarquées sur les satellites géostationnaires. En effet si l'on calcule l'atténuation en espace libre correspondant à la propagation de 36 000 km, on trouve :

$$\frac{1}{4\pi b^2} = \frac{1}{4\pi(3,6 \cdot 10^7)^2} \approx 6 \cdot 10^{-17} m^{-2}$$

C'est pourquoi toutes les antennes servant à la liaison satellite, aussi bien les antennes embarquées que les stations au sol, sont des antennes de grand gain. Le meilleur moyen d'obtenir un grand gain est, selon la formule [4.14], d'augmenter la surface de captation de l'antenne. Les antennes à réflecteur permettent, grâce à la surface importante du réflecteur, de capter le maximum de puissance. Certaines antennes utilisent deux réflecteurs. Nous verrons un peu plus loin le fonctionnement de ces antennes qui ont couramment des gains d'une cinquantaine de décibels. Dans la plupart des cas la source est un cornet. Selon l'état de polarisation attendue, le cornet est circulaire ou rectangulaire. Les cornets sont généralement corrugués pour limiter les lobes secondaires (voir paragraphe 9.1.1).

L'inconvénient majeur de ce type de satellite est de ne couvrir qu'une petite partie de la Terre. L'angle de vue de la Terre à partir d'un point de l'orbite géostationnaire étant de 17°, les latitudes élevées ne sont pas comprises dans l'ouverture de l'antenne.

□ Les satellites en orbite basse

Les satellites en orbite basse, aussi appelés LEO (*Low Earth Orbit*) ont une orbite inclinée par rapport au plan de l'équateur. Leur période est nettement plus faible que celle des satellites géostationnaires, puisqu'ils font le tour de la Terre plus d'une dizaine de fois par jour. Ces satellites défilent autour de la Terre et sont capables de couvrir l'ensemble de la planète grâce à un léger décalage de leur orbite à chaque tour. L'orbite des satellites LEO est de l'ordre de 800 à 1 000 km. Ils permettent d'atteindre des latitudes très élevées, contrairement aux satellites géostationnaires. Certaines orbites elliptiques permettent aux satellites de rester plus longtemps au-dessus d'une région. Leur vitesse est inversement proportionnelle à la distance au centre de la Terre. Par conséquent, ils vont plus vite au périégée qu'à l'apogée.

Comme leur distance à la Terre est beaucoup plus faible que pour les satellites géostationnaires, les contraintes de puissance peuvent être relâchées, ainsi les antennes ne sont pas nécessairement aussi puissantes. De plus, le coût du lancement est moins élevé. Cependant le suivi est plus délicat et du fait du défilement, les informations ne sont pas transmises en permanence. C'est pourquoi ils sont munis de systèmes de stockage de l'information.

Le suivi des satellites défilants est assuré par des antennes terrestres mobiles

□ Les satellites en orbite moyenne

Les satellites en orbite moyenne (MEO) ont des altitudes qui sont autour de 10 000 km. Ils ont globalement les mêmes caractéristiques que les satellites LEO. Cependant leur altitude plus élevée impose que la puissance d'émission soit plus importante. Le facteur d'atténuation en puissance est environ cent fois plus grand que pour les satellites LEO.

□ Les constellations de satellites

Afin d'améliorer la couverture globale de la surface terrestre, certains systèmes, appelés constellation de satellites, utilisent un ensemble de minisatellites. C'est le cas de la constellation « Iridium » qui comporte 66 satellites et de la constellation « GlobalStar » qui comporte 48 satellites sur des orbites basses (environ 1 400 km), inclinées à 50° par rapport à l'équateur. Les différents satellites prennent le relais en fonction de leur position. De ce fait, l'information peut être transmise en permanence et la zone de couverture est pratiquement mondiale. Le principe de fonctionnement des constellations est d'utiliser la diversité spatiale pour améliorer la qualité de service. En général un mobile est en vue de deux à quatre satellites

Certains satellites de la constellation peuvent communiquer entre eux. Cette fonction avait été conçue au départ pour assurer la confidentialité des communications.

Les constellations présentent de nombreux avantages :

- La masse de chaque satellite est faible (environ 500 kg alors que la masse d'un satellite géostationnaire est de plusieurs tonnes).
- La distance à la Terre est faible, donc il n'y a pas de décalage temporel dans la communication. Les puissances mises en jeu sont plus faibles.
- En cas de panne, un autre satellite peut prendre le relais.
- Les obstacles sont moins gênants puisqu'on peut utiliser le satellite qui est le mieux placé.

Le système GPS (*Global Positioning System*) est constitué de trente et un satellites sur des orbites situées à une altitude de 20 000 km. Les satellites sont répartis sur six plans orbitaux, inclinés de 55°. Les satellites reprennent la même position chaque jour car la période de rotation est de 11 h 58 min. Le système de localisation repose sur l'exploitation des signaux provenant de quatre satellites en vue.

Le futur système de positionnement européen GALILEO devrait être complètement opérationnel vers 2013. Il utilisera une trentaine de satellites.

□ Les stations de réception terrestres

Les stations de réception des signaux provenant des satellites peuvent être de plusieurs tailles.

Les stations les plus petites sont les stations GPS qui n'ont pas besoin de recevoir un débit important et doivent être mobiles et portables.

Les antennes de réception de télévision doivent recevoir un signal provenant d'une distance de 36 000 km avec un débit relativement élevé. Les antennes ont des diamètres de l'ordre du mètre. Il existe des stations à grand gain, conçues pour recevoir un débit élevé. Ces antennes sont utilisées à l'échelle nationale pour capter les informations en provenance de satellites scientifiques, par exemple. Dans ce cas il s'agit d'antennes paraboliques, comportant souvent un montage Cassegrain.

Les antennes de certaines stations sont mobiles afin de permettre la poursuite du satellite durant un certain intervalle de temps.

■ Les bandes de fréquences

Comme il a déjà été dit en introduction de cet ouvrage, le choix des bandes de fréquences s'effectue en fonction de l'atténuation due à la traversée de l'atmosphère. Les hyperfréquences sont bien adaptées à ces transmissions et ce d'autant plus qu'elles sont pratiquement insensibles à la pluie et qu'elles traversent les nuages. Les cristaux de glace ont quelquefois des effets dépolarisants, selon leur taille et selon les fréquences de transmission. On évite bien sûr pour les transmissions satellitaires les bandes d'absorption atmosphériques (chapitre 1).

Les systèmes utilisent une fréquence montante différente de la fréquence descendante pour des problèmes de compatibilité électromagnétique. Ces fréquences sont adoptées au niveau mondial lors des conférences de l'Union Internationale des Télécommunications. Voici quelques fréquences utilisées, respectivement la fréquence montante et la fréquence descendante :

- bande L : 1,6/1,4 GHz. Cette bande est très utilisée pour les communications avec les mobiles (marine, aviation, véhicules terrestres)
- bande C : 6/4 GHz
- bande X : 8/7 GHz, réservée aux applications militaires
- bande Ku : 14/12 GHz
- bande Ka : 30/20 GHz. Cette bande encore peu utilisée représente un potentiel important pour les liaisons à haut débit. Les contraintes techniques et les coûts associés sont relativement forts en raison de la fréquence élevée. Quelques satellites (en particulier américains, canadiens

et japonais) fonctionnent déjà dans cette bande de fréquences pour des transmissions de données et de la diffusion pour la télévision haute définition.

Les antennes sont donc conçues en fonction des fréquences d'utilisation qui peuvent varier d'une application à l'autre. Elles n'ont pas forcément besoin d'avoir une large bande passante car elles servent à une application donnée, généralement positionnée dans une bande de fréquences spécifique.

■ Les applications

Les satellites se sont peu à peu imposés comme moyen de transmettre de l'information car ils permettent de couvrir des zones larges qui peuvent être inaccessibles par voie filaire ou par voie hertzienne. Certaines zones montagneuses ne sont accessibles que par ce moyen.

□ Téléphonie fixe

Les satellites ont tout d'abord été utilisés pour la téléphonie fixe. La communication est envoyée au satellite qui la reçoit et la réémet. Lors de ce type de communication un léger retard est perçu qui est dû au temps de propagation de l'aller-retour du signal entre la Terre et le satellite. Maintenant, outre la voix, ces satellites transmettent aussi des données. La largeur des canaux est adaptée au débit à transmettre. La largeur de bande reste cependant raisonnable.

□ Télédiffusion

La télédiffusion fait partie des utilisations importantes des satellites. Tout le monde est maintenant familier avec les petites paraboles qui permettent de recevoir les émissions télévisées.

□ Localisation

Le GPS (*Global Positioning System*) est un système très répandu de localisation. Il fonctionne grâce à une petite station de réception qui détecte la position d'un mobile par rapport à la position de quatre satellites. Ce système utilisé par les militaires américains a été mis à la disposition des civils. Il donne une précision sur la localisation d'environ 20 m. Devant le succès du GPS, un autre système va être mis en place au niveau européen, le système GALILEO.

□ Balise

Les satellites émettent des signaux spécifiques sur une très faible bande qui servent de balises pour les mobiles.

□ Communications mobiles

La téléphonie mobile utilise aussi les services satellites. Les utilisateurs en sont la marine, l'aviation et les véhicules terrestres.

□ Observation de la Terre

Les satellites d'observations sont de plus en plus nombreux. Les observations peuvent porter sur des suivis de phénomènes naturels terrestres : météorologie, inondations, éruptions volcaniques ou bien des phénomènes liés aux activités humaines : pollutions, modifications des surfaces liées à l'agriculture... JASON, par exemple, est un satellite dédié plus particulièrement à l'observation océanographique. Les satellites d'observation peuvent être des satellites géostationnaires (METEOSAT) ou des satellites défilants (SPOT). Il existe aussi de nombreux satellites d'observation militaires. Quelques pays ont leurs propres satellites d'observation.

□ Expérimentations scientifiques

Certains satellites sont expérimentaux et permettent des essais sur de nouvelles technologies. D'autres, embarquant de nombreux appareils de mesures sont qualifiés de satellites scientifiques. Des dispositifs de télémétrie permettent de mesurer l'altitude du satellite. Certains, très précis, sont capables de mesurer la déformation de la surface des océans.

□ Télémétrie, télécommande

Enfin tous les satellites sont suivis très attentivement. Des mesures sont effectuées en permanence pour vérifier leur position. Dès le moindre écart à leur trajectoire ou à leur orientation, des dispositifs télécommandés depuis la Terre permettent le démarrage de moteurs afin de piloter à distance le satellite. Les dispositifs de télémétries et de télécommande passent par des antennes spécialisées.

□ Location de canaux

Les transmissions par satellites ont pris un tel essor et offrent une telle variété de services, de qualité, que des opérateurs ou des organisations, INMARSAT... EUTELSAT, ont créé un marché proposant de nombreux services.

L'organisation internationale INTELSAT, privatisée en 2001, gère actuellement une cinquantaine de satellites au niveau mondial.

Le consortium INMARSAT met à disposition quatre satellites géostationnaires, dont deux sont situés au-dessus de l'océan Atlantique et les deux autres au-dessus de l'océan Pacifique et au-dessus de l'océan Indien. Ces satellites servent à la transmission de la voix et à de données. Ces satellites sont relayés par des stations terrestres et la zone de couverture est ainsi très large. Des données peuvent mettre jusqu'à cinq minutes pour transiter.

Au niveau de l'Europe, il existe des sociétés telles que EUTELSAT qui proposent des services de transmission de la voix de données et de télédiffusion. Le programme ECS (*European Communication Satellite*) et TELECOM portent sur les mêmes services. Il faut signaler quelques programmes nationaux qui ont fonctionné avec succès tels que ITALSAT ou les programmes français TELECOM 1 et 2.

□ Évolutions

Les systèmes de satellites évoluent actuellement vers la transmission haut débit. Un certain nombre d'études sont en cours. Certains satellites supportent déjà des transmissions haut débit. On peut citer le satellite japonais KIZUNA, mis en orbite en 2008 qui fonctionne en bande Ka. Il propose un accès à ultra haut débit en tout point situé en Asie. Les débits visés sont de 1,2 Gbit/s avec une antenne de réception au sol de 5 m de diamètre. Une antenne plus petite permet de recevoir un débit moindre. C'est un système qui peut être mis en place en urgence pour assurer des moyens de communications lors d'un événement exceptionnel, par exemple.

Dans un avenir proche, il est envisagé de procéder à une identification par satellites. Cela pourrait permettre par exemple de suivre les navires de façon automatique.

Les utilisations des satellites s'orientent vers le haut débit. Une norme DVB-S (*Digital Video Broad Casting - Satellite*) a été décidée. Des systèmes sont mis en place pour transmettre du haut débit par satellites, utilisant la bande 27-31 GHz.

Les formes d'antennes utilisées à bord des satellites sont très variées. Elles dépendent des utilisations qui en sont faites et des bandes d'utilisation. Par exemple, une antenne altimétrique est très différente d'une antenne de diffusion.

Dans la suite nous présenterons les antennes à réflecteur(s) qui sont de loin les plus nombreuses. D'autres types d'antennes sont utilisés qui seront décrits plus loin dans l'ouvrage.

■ Les antennes à réflecteurs

Les antennes pour satellites doivent avoir une grande directivité. Comme cette propriété est obtenue pour une surface rayonnante grande, on utilise un ou deux réflecteurs qui agrandissent la surface effective de l'antenne.

□ Les antennes à un réflecteur

La forme des antennes à un réflecteur est généralement parabolique. On verra, au cours de ce chapitre, des variantes. Le réflecteur est métallique, donc pratiquement parfait. Dans son plan de coupe il est parabolique et a, la plupart du temps, une symétrie axiale (figure 6.15).

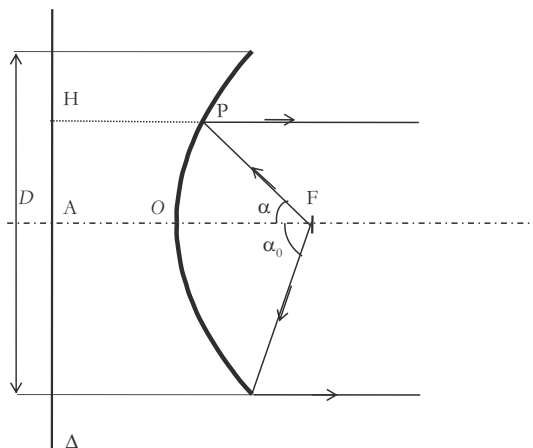


Figure 6.15 – Définition géométrique des paramètres de la parabole.

Principe de fonctionnement

Considérons un point source placé au foyer de la parabole. Par définition, tout rayon provenant du foyer (rayon primaire) se réfléchit et repart parallèlement à l'axe de la parabole (rayon secondaire).

Rappelons certaines propriétés d'un parabolôïde de révolution (3D) ou d'une parabole (2D). C'est par définition le lieu des points à égale distance d'un point, le foyer, et d'un plan en 3D, appelé le plan directeur ou d'une droite en 2D, la directrice (Δ). De ce fait, quel que soit le point P appartenant à la parabole, la relation suivante est vérifiée :

$$FP = PH$$

Les rayons secondaires issus de F sont parallèles entre eux et sont en phase. Donc tout plan perpendiculaire à l'axe du parabolôïde est un plan de phase.

L'équation de la parabole d'axe Oz est :

$$y^2 = 4fz$$

OF est la distance focale, notée f .

La surface géométrique de l'ouverture de l'antenne est ainsi égale à $\pi \frac{D^2}{4}$. On conçoit ainsi que cette antenne ait un grand gain selon la formule :

$$G = 4\pi \frac{A_{eff}}{\lambda^2}$$

Où A_{eff} représente l'aire effective de l'antenne qui dépend de la répartition du champ sur l'ouverture et qui est toujours inférieure ou égale à la surface géométrique.

Si le point source est légèrement décalé dans le plan focal, le rayon secondaire s'écarte légèrement de l'axe. En effet, le rayon réfléchi fait avec la normale un angle égal à l'angle d'incidence.

Or dans le plan focal se trouve placée une source ayant une certaine extension. Donc les rayons partant de l'ouverture ne sont pas tous rigoureusement parallèles entre eux et cela contribue à donner au rayonnement issu de la parabole une certaine ouverture angulaire. Cette explication est relativement sommaire car elle s'appuie sur la théorie des rayons qui est valable sous certaines approximations.

De nombreuses notions sur la conception d'antennes paraboliques s'appuient sur la théorie des rayons. On suppose, en particulier que l'onde provenant de la source est sphérique et que tous les rayons viennent se réfléchir sur la parabole. Ceci n'est vérifié que si le réflecteur est dans le champ lointain de l'antenne. Supposons que l'ouverture de l'antenne soit notée b . Nous verrons au paragraphe 9.1.1, portant sur les cornets que, si celui-ci a une longueur l , son ouverture optimale est donnée (cornet conique) par :

$$b = \sqrt{3l\lambda}$$

Rappelons que la condition de champ lointain est :

$$R > \frac{2b^2}{\lambda}$$

qui se traduit si l'on choisit une ouverture optimale, par :

$$R > 6l$$

Cette condition est à prendre en compte pour que les raisonnements qui suivent soient valables. Elle est généralement vérifiée sans difficulté. Rappelons que si le point d'observation est dans le champ proche, on ne peut plus considérer l'onde comme sphérique.

Pour déterminer le diagramme de rayonnement d'une parabole, il existe deux méthodes :

- L'une s'appuie sur le fait que le champ électromagnétique crée des courants induits sur la parabole. Ces courants sont la source du champ dans tout l'espace.
- L'autre utilise le rayonnement des ouvertures. La connaissance du champ sur le plan situé juste à la sortie du paraboloïde permet de connaître le champ dans tout l'espace.

Calcul du diagramme de rayonnement

Nous choisirons de donner les éléments de conception des antennes à partir de cette seconde méthode. D'après la théorie, le spectre d'onde plane sur la surface rayonnante donne la forme du diagramme de rayonnement [3.43] L'ouverture rayonnante est prise juste à la sortie du paraboloïde, de diamètre D .

Répartition d'amplitude sur l'ouverture

La répartition de puissance sur l'ouverture dépend à la fois de la répartition de puissance de la source primaire et de la distance à la source. Pour un cornet conique, cette répartition est souvent représentée par une puissance du cosinus de l'angle d'ouverture. La forme de la surface réflectrice modifie la répartition de puissance sur le plan de sortie du fait de son orientation locale en fonction du point de réflexion. En effet, considérons une parabole dans le plan (Oyz) (figure 6.16)

L'élément dS pour valeur :

$$dS = \sqrt{(dr)^2 + (rd\alpha)^2}$$

Par définition de la parabole, en notant f la distance focale :

$$r(\alpha) = \frac{2f}{1 + \cos \alpha} = \frac{f}{\cos^2(\alpha/2)}$$

La puissance élémentaire envoyée dans le cône élémentaire d'angle $d\alpha$ est égale à la puissance rayonnée dans dy :

$$P_\alpha(\alpha)d\alpha = P(y)dy$$

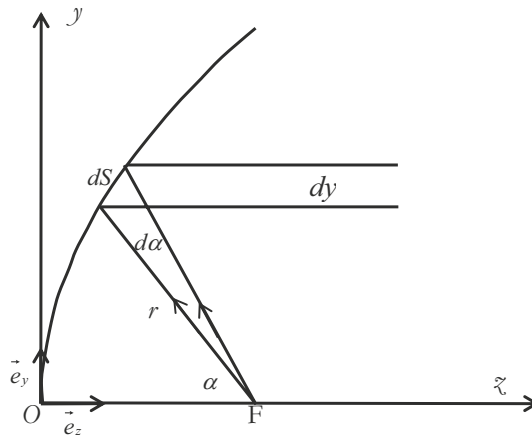


Figure 6.16 – Réflexion d'un rayon sur une parabole.

La tangente à la parabole s'exprime en fonction des vecteurs unitaires $\vec{e}_z$ et $\vec{e}_y$:

$$\vec{t} = \frac{\vec{e}_y dy + \vec{e}_z dz}{\sqrt{(dy)^2 + (dz)^2}}$$

Le calcul de l'élément différentiel selon y se calcule par :

$$dy = \vec{t} \cdot \vec{e}_y dS = \sqrt{(dr)^2 + (rd\alpha)^2} \frac{2f/y}{\sqrt{1 + (2f/y)^2}} d\alpha$$

En tenant compte de :

$$y = r \sin \alpha \quad \text{et de} \quad \frac{dr}{d\alpha} = r \tan \frac{\alpha}{2}$$

On obtient :

$$P(y) = P_\alpha(\alpha) \frac{\cos^2 \alpha/2}{f}$$

On constate que la répartition de puissance varie comme l'intensité de la source multipliée par le carré du cosinus $\alpha/2$. Donc l'amplitude imposée par la source est modulée par un terme en cosinus $\alpha/2$. Ceci induit une amplitude plus faible sur les bords du fait de la réflexion sur une parabole. Ce calcul est effectué en 2D. Il s'applique à une géométrie invariante dans la direction Ox . Il correspond à une ligne de sources placées selon Oz et rayonnant vers une surface de section parabolique, invariante par translation selon Oz .

Considérons maintenant un paraboloïde de révolution. L'intensité de la source dépend de l'angle azimutal φ (figure 6.17).

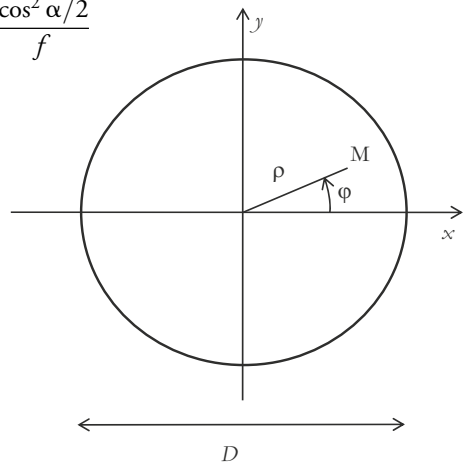


Figure 6.17 – Ouverture de la parabole.

La répartition de puissance sur l'ouverture dépend du rayon ρ sur la surface rayonnante et de ϕ . Un calcul en trois dimensions donne la répartition de puissance sur l'ouverture en fonction de l'intensité de la source $P_{\Omega}(\alpha, \phi)$:

$$P(\rho, \phi) = \frac{P_{\Omega}(\alpha, \phi)}{f^2} \cos^4 \frac{\alpha}{2}$$

L'amplitude, qui est donnée comme la racine carrée de la densité surfacique de puissance, varie comme le carré du cosinus $\alpha/2$. L'amplitude est donc affaiblie sur les bords de la surface rayonnante par rapport à ce qu'elle serait sans le paraboloïde. Ceci permet de réduire les lobes secondaires.

Une fois la répartition du champ connue pour chaque polarisation, il suffit d'effectuer la transformée du champ dans l'ouverture pour connaître le champ lointain. Ce calcul, difficile analytiquement, peut être effectué de façon numérique.

Blocage de l'ouverture

Le système utilisant une parabole présente un inconvénient majeur lié au fait que la source se trouve dans le champ de rayonnement. Elle fait donc écran, sur une certaine surface, au rayonnement incident ou au rayonnement émis (figure 6.18). C'est ce qu'on appelle le blocage de l'ouverture.

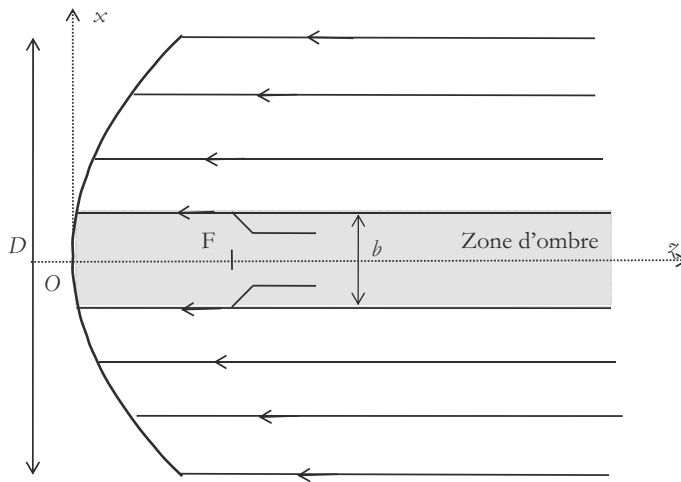


Figure 6.18 – Principe du blocage de l'ouverture.

Différentes considérations sur le rayonnement de la source vont permettre de donner certains éléments de conception du système rayonnant constitué de la source et de la parabole.

Tout d'abord, rappelons que si la source a des dimensions larges, son diagramme de rayonnement est fin, c'est-à-dire que l'angle d'ouverture est petit. Le lobe principal peut, dans ce cas, n'utiliser qu'une partie du réflecteur et le système n'est pas optimal.

Au contraire, si la source est de petite taille, le faisceau est large. Dans ce cas le lobe principal peut déborder de la parabole. C'est le phénomène de débordement (*spill-over*). Il existe alors un rayonnement arrière qui peut être gênant.

On admet que l'optimum de fonctionnement est obtenu lorsque la puissance au bord de la parabole est d'environ 10 dB en dessous de la puissance au maximum.

La distance focale de la parabole doit être choisie en fonction de l'ouverture de l'antenne. En effet selon la figure 6.15, la relation suivante est vérifiée :

$$\tan \alpha_0 = \frac{D/2}{f - D^2/16f} = \frac{0,5}{f/D - 0,0625D/f}$$

Une parabole plus aplatie a une plus grande distance focale, à ouverture constante et l'angle sous lequel est vue la parabole depuis la source est plus petit.

Prenons le cas d'un système pour lequel l'atténuation en puissance est de 10 dB sur les bords de la parabole. Si on augmente la distance focale, la parabole est plus aplatie, il faut alors augmenter la distance entre la source et la parabole qui sera vue sous un angle plus petit. On aura, dans ce cas, un phénomène de débordement. Pour éviter cela, il faut augmenter la taille de la source. Dans ce cas le faisceau est plus étroit, mais la source primaire écrante davantage le faisceau. La solution est donc de rapprocher la source. À l'inverse donc, lorsqu'on rapproche la source, la parabole est plus profonde et la source primaire doit être plus petite. Cependant ce raisonnement ne peut pas conduire à une distance focale trop petite car :

- La puissance rayonnée doit être suffisante.
- Il existe, de plus, des limitations de taille par rapport à la longueur d'onde émise.
- Si la source primaire est trop proche du réflecteur, des phénomènes de couplages existent

C'est pourquoi les distances choisies pour les systèmes à parabole sont en général tels que :

$$0,35 < \frac{f}{D} < 0,5$$

Le diagramme de rayonnement résulte de la combinaison des sources secondaires issues de l'ouverture du paraboloïde. Sans tenir compte de l'ombre de la source, le terme à calculer est le spectre d'onde plane qui est la transformée de Fourier de la répartition de l'amplitude dans l'ouverture sous la forme (paragraphe 3.7.1) :

$$\vec{f}_T(k_x, k_y) = \iint_S \vec{E}_o(x, y) e^{j(k_x x + k_y y)} dx dy$$

S représente la surface totale de l'ouverture. L'expression du champ contient l'atténuation selon le carré du cosinus $\alpha/2$ due à la parabole. En tenant compte du fait que l'ouverture n'est pas totalement illuminée et en appelant S_b , la surface correspondant au blocage des ondes par la source primaire, on constate que la transformée de Fourier doit prendre la forme :

$$\vec{f}_T'(k_x, k_y) = \iint_{S-S_b} \vec{E}_o(x, y) e^{j(k_x x + k_y y)} dx dy$$

L'efficacité de l'antenne s'en trouve donc nettement diminuée, puisque ce sont, en général, les ondes qui proviennent du centre de l'ouverture qui participent le plus au diagramme de rayonnement dans l'axe.

Un autre inconvénient de ce système est le niveau de puissance réfléchi vers la source primaire puisqu'elle se trouve directement en face du réflecteur et au point le plus proche. Cela induit sur la source un taux d'ondes stationnaires important.

Source décalée

Afin de s'affranchir du problème de blocage de l'ouverture, certains systèmes utilisent une source décalée selon le schéma de la figure 6.19.

Les propriétés de la parabole restant les mêmes quel que soit le point de réflexion, il est possible d'orienter la source de façon à utiliser un seul côté de la parabole. Le réflecteur métallique n'est alors plus symétrique et son contour n'est plus circulaire. L'ouverture peut être ainsi augmentée de façon significative.

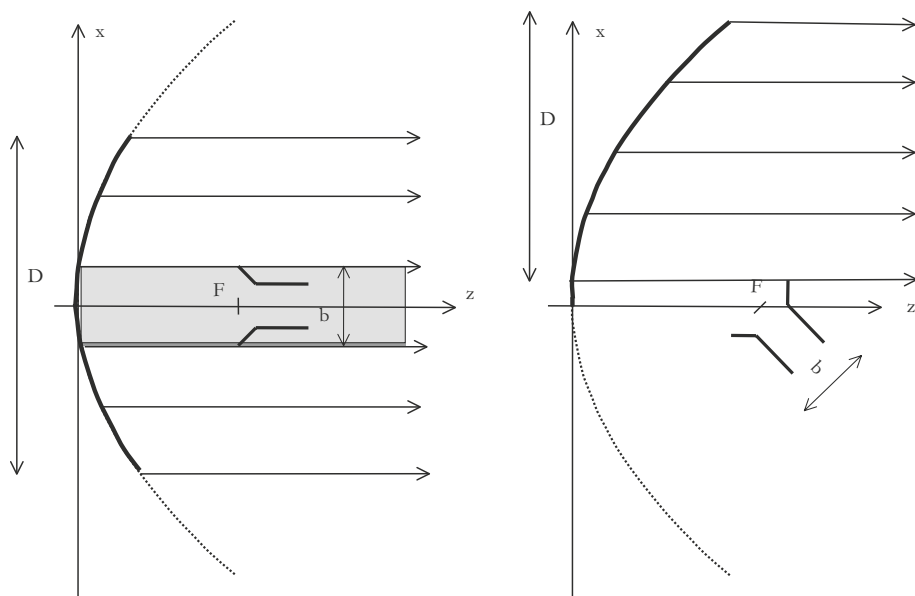


Figure 6.19 – Parabole à source décalée.

Sur la figure 6.19, les deux paraboles ont la même distance focale. La source est placée au même endroit, mais orientée différemment. L'encombrement est sensiblement le même si l'on considère que la partie métallique est représentée par les traits gras.

L'inconvénient de ce type d'antenne à source décalée est d'induire une polarisation croisée plus importante. L'avantage pour la source est évident car aucune puissance n'est réfléchie vers la source.

Diffraction

La diffraction par les bords est un phénomène à prendre en compte pour le calcul du rayonnement d'une antenne parabolique. En effet, les bords de la parabole se comportent comme une arête et diffractent le rayonnement. C'est un phénomène du second ordre qui perturbe légèrement le diagramme de rayonnement. Cependant pour s'en affranchir certains dispositifs présentent des paraboles à bords arrondis ou bien recouverts d'une matière absorbante.

La source primaire, par sa forme, peut également diffracter le rayonnement puisqu'elle est constituée de métal présentant aussi des arêtes.

Dans ce dispositif, les bras qui soutiennent la source provoquent aussi des phénomènes de diffraction.

□ Les antennes à deux réflecteurs

Afin d'améliorer les caractéristiques des antennes paraboliques à un réflecteur, on adjoint au système un deuxième réflecteur. C'est le cas pour le montage Cassegrain et pour le montage Grégorien. Nous verrons que la directivité du système antennaire est augmentée et que le rapport signal sur bruit est augmenté.

Antenne Cassegrain

Une antenne Cassegrain est représentée sur la figure 6.20

La source émet une onde qui se réfléchit une première fois sur un réflecteur hyperbolique convexe, puis une seconde fois sur le réflecteur parabolique. Les dimensions sont choisies de telle façon

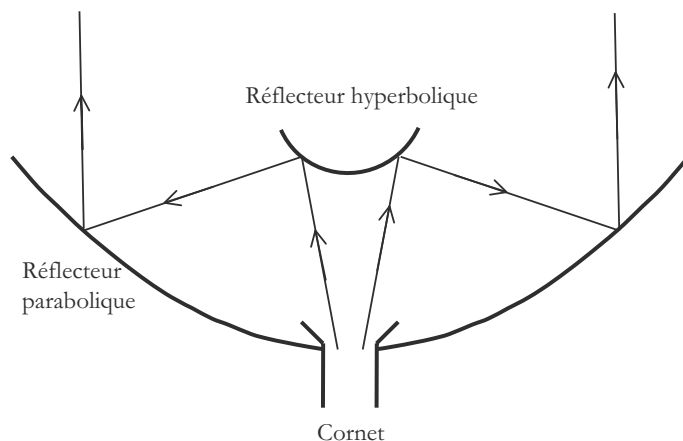


Figure 6.20 – Antenne Cassegrain.

que les rayons partent de l'ouverture parallèlement entre eux. Comme l'ouverture est plus grande grâce aux réflexions successives, le gain de cette antenne est plus grand que celui d'une antenne parabolique simple.

Ce montage présente l'avantage d'être moins sensible aux parasites. En effet, supposons que cette antenne soit utilisée en réception d'un signal venant du ciel et qu'une source parasite émette un signal provenant de la terre, la parabole étant dirigée vers le ciel, comme sur la figure. S'il n'y a qu'un seul réflecteur, ce signal parasite peut entrer directement, par débordement, sur le capteur (cornet) qui est alors dirigé vers le bas dans le montage conventionnel. Par contre s'il y a un second réflecteur, le cornet, orienté comme il est indiqué sur la figure, ne peut pas être atteint par le signal parasite. Le cornet orienté vers le ciel reçoit le rayonnement parasite provenant du ciel qui est plus faible que celui provenant de la terre dans les bandes de fréquences considérées. Ce système a une température de bruit faible.

La position du cornet, soit en émission, soit en réception, lui permet d'être très proche de la partie électronique, évitant ainsi des câbles trop longs. Ceci permet aussi de réduire le bruit du système.

Les deux remarques précédentes montrent pourquoi cette antenne a un meilleur rapport signal sur bruit.

Un autre avantage est de présenter un débordement (*spillover*) faible.

Dans la suite, nous allons montrer comment placer les réflecteurs. Considérons une source au point S qui émet une onde sphérique (figure 6.21)

Cette onde est interceptée par le réflecteur qui renvoie les rayons en respectant la loi de la réflexion : l'angle du rayon réfléchi par rapport à la normale est égal à l'angle d'incidence. Afin d'utiliser ensuite un paraboloïde pour la seconde réflexion, on désire se placer dans les conditions qui ont été décrites précédemment, à savoir que tous les rayons semblent venir d'une source émettant un front d'onde sphérique. Le point F joue le rôle de cette source. On doit donc vérifier :

$$FB = FB_0$$

$$\text{Et donc } \overline{SA} + \overline{AB} = \overline{SA_0} + \overline{A_0B_0}$$

Introduisons le point O à égale distance de la source et du point F :

$$\overline{SO} = \overline{OF}$$

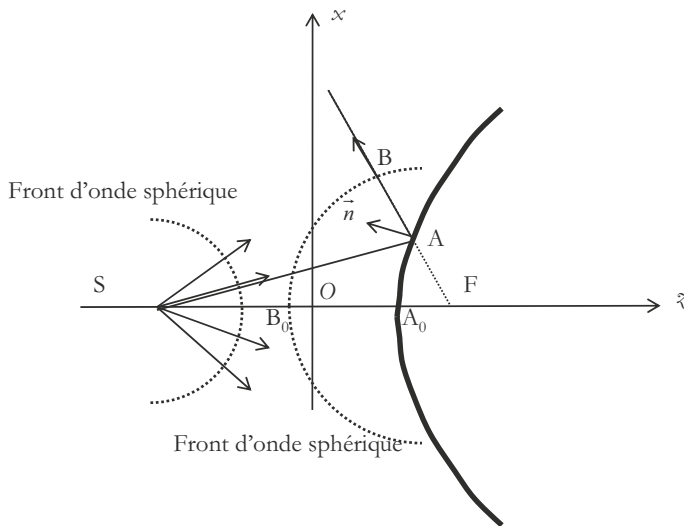


Figure 6.21 – Forme hyperbolique du réflecteur.

On obtient :

$$\overline{SA} + \overline{AB} = \overline{SA_0} + \overline{A_0B_0} = \overline{SO} + \overline{OA_0} + \overline{A_0B_0}$$

Et comme :

$$\overline{AB} = \overline{AF} + \overline{FB} = \overline{A_0F} + \overline{FB_0}$$

On déduit :

$$\overline{SA} - \overline{FA} = 2\overline{OA_0}$$

Le réflecteur doit donc être le lieu des points dont la différence de distance à deux points fixes est constante. C'est donc une hyperbole de foyers S et F.

La configuration idéale est de placer le sommet de la parabole en S et son foyer en F. Cependant l'ouverture de la source primaire représente une surface prise sur le paraboloïde et on retrouve le problème du blocage de l'ouverture. C'est pourquoi ce montage n'est utilisé que lorsque la parabole est grande :

$$D > 40\lambda$$

Ce type d'antenne n'est utilisé que lorsque le gain doit être très grand, comme dans les applications satellitaires ou radioastronomiques.

Un avantage de ce montage est d'augmenter la distance focale de la parabole équivalente, à ouverture égale. Posons $g = \frac{OS}{OF}$, appelé le grandissement du système. La parabole de focale f , associée au réflecteur hyperbolique est équivalente à une parabole de focale f' . La focale équivalente f' est donnée, en fonction de la focale f par :

$$f' = g.f$$

La courbure de la parabole équivalente à celle du montage Cassegrain peut donc être plus faible, puisqu'elle correspond à une plus grande focale. Cette qualité est utile lorsqu'on utilise plusieurs sources, car une parabole à forte courbure est plus sensible au décalage de position par rapport au foyer.

Réflecteur grégorien

Le principe du réflecteur grégorien repose aussi sur une double réflexion. Le système est constitué d'une parabole associée à un réflecteur elliptique concave.

De même que dans le cas d'un montage Cassegrain, le montage grégorien peut présenter une source décalée. La forme du réflecteur principal est alors dissymétrique.

□ Les réflecteurs formés

Pour de nombreuses applications la forme du faisceau doit être particulière. Une solution pour obtenir une répartition définie passe par la formation du ou des réflecteurs qui consiste à calculer la fonction que doit avoir sa surface.

Pour s'en convaincre, il suffit de se rapporter au calcul qui a été fait pour la répartition d'amplitude à la sortie de la parabole. Pour un paraboloïde, l'amplitude varie comme le carré de $\cos \alpha/2$. Cette variation vient pour une grande partie de l'atténuation en fonction de la distance séparant la source du réflecteur. Si le réflecteur s'écarte de cette position, il est possible d'obtenir une atténuation plus grande si la distance est plus grande ou inversement. Cette latitude au niveau de la conception du réflecteur permet d'obtenir un diagramme de rayonnement particulier.

Il est aussi possible de jouer sur la forme du premier réflecteur dans le cas d'une antenne Cassegrain ou grégorienne. Le diagramme de rayonnement est plus sensible à la forme de sa surface.

Le calcul analytique ne peut pas être conduit jusqu'au bout. Il faut avoir recours à des méthodes numériques basées dans un premier temps sur la théorie géométrique de la diffraction. La méthode intégrale donne des résultats plus précis, mais plus longs en temps de calcul. La méthode des éléments finis peut aussi être utilisée, à condition d'avoir de grosses ressources en termes de mémoire.

■ La zone de couverture

De nombreuses applications satellitaires reposent sur la définition de la zone de couverture des antennes. C'est la zone pour laquelle la puissance reçue au sol est suffisante pour faire fonctionner le système.

Cette zone est généralement divisée en plusieurs zones en fonction des valeurs de la puissance reçue. Il apparaît alors des cartes de niveau de puissance.

En principe la forme de cette zone doit résulter de l'intersection du faisceau rayonnant une puissance donnée avec la surface terrestre. La zone de couverture d'une antenne simple est donc approximativement constituée de la surface intérieure d'un cercle ou d'une ellipse, selon l'angle entre l'axe de l'antenne et la normale à la surface. Les différents niveaux qui apparaissent en partant du centre correspondent à la fois à une décroissance de la puissance due au facteur d'atténuation inversement proportionnel au carré de la distance Terre-satellite et à une décroissance dans le diagramme de rayonnement.

Les zones de couvertures sont en fait beaucoup plus complexes. Par exemple un des satellites du système INTELSAT couvre, grâce à 28 transpondeurs en bande C, toute la zone atlantique est et ouest, c'est-à-dire l'Amérique du Sud, une grande partie de l'Amérique du Nord, la majeure partie de l'Afrique (sauf l'est) et pratiquement toute l'Europe.

Ce type de couverture n'est possible que grâce à des combinaisons d'antennes (voir paragraphe 5.2).

Ainsi, on trouve couramment des cartes donnant la puissance isotrope rayonnée équivalente (PIRE) pour une région donnée. Rappelons que la PIRE est égale à la puissance d'alimentation de l'antenne multipliée par le gain de l'antenne d'émission. C'est la puissance qui serait émise par une antenne isotrope placée au même endroit que l'antenne réelle et qui apporterait la même densité de puissance à l'endroit de l'antenne de réception. Les niveaux de puissance sont indiqués en dBW. Ces niveaux diminuent de 1 dBW pour une augmentation du diamètre de la zone couverte d'environ de 12,25 %, pour des satellites géostationnaires.

Prenons l'exemple du satellite ASTRA qui diffuse des émissions de télévisions dans une bande de fréquence centrée sur 12,75 GHz. Les tubes à ondes progressives qui alimentent les antennes sont d'une centaine de Watt (120-130 W). Pour obtenir une PIRE de 50 dBW, il faut que le gain de l'antenne d'émission soit de l'ordre de 30 dB, ce qui est tout à fait réalisable avec des antennes à réflecteur du type de celles qui viennent d'être décrites.

La valeur classique de la PIRE nécessaire pour recevoir la télévision par satellite avec une petite parabole (diamètre 60 cm) est de l'ordre 50 dBW. Les cartes fournissent donc des niveaux de puissance autour de cette valeur. Pour des raisons pratiques, ces niveaux sont remplacés par des tailles de paraboles à utiliser. Plus on s'écarte de la zone centrale de couverture plus les diamètres d'antennes doivent être grands. Le tableau 6.5 donne une idée de cette variation.

Tableau 6.5 – Relation entre la PIRE et le diamètre de la parabole de réception

PIRE (dBW)	50	48	46	44	42	40
Diamètre (m)	0,6	0,75	0,95	1,2	1,5	1,90

Lorsque la PIRE est supérieure à 50 dBW, on serait tenté d'utiliser une antenne plus petite. Cependant, un rapide calcul nous donne l'ouverture de l'antenne approximativement circulaire, à une fréquence de 12,75 GHz :

$$1,22 \frac{\lambda}{D} = 1,22 \frac{3 \cdot 10^8}{12,75 \cdot 10^9 \cdot 0,6} \approx 0,05 \text{ rad}$$

Cette valeur qui est légèrement inférieure à 3° donne une idée du pouvoir de séparation de l'antenne entre deux satellites. Si on diminue la taille de l'antenne, son angle d'ouverture augmente. Il y a donc un risque de capter des signaux provenant de deux satellites différents, sur la même puissance et la même polarisation.

Un autre paramètre important est lié à la zone de couverture, c'est le facteur de mérite du récepteur, appelé aussi sensibilité. Ce facteur résulte de l'évaluation du rapport signal à bruit. Pour un récepteur porté à une température thermodynamique T , la puissance associée au bruit de l'électronique est donnée en fonction de la largeur de bande ΔB par :

$$N = kT\Delta B$$

La constante de Boltzmann k a pour valeur : $k = 1,38 \cdot 10^{-23} \text{ J} \cdot \text{K}^{-1}$

Le rapport signal sur bruit s'exprime pour l'antenne de réception par :

$$\frac{S}{N} = \frac{P_r}{N} = \frac{P_r}{kT\Delta B}$$

Or

$$P_r = P_t \frac{\lambda^2}{(4\pi r)^2} G_t G_r$$

Donc :

$$\frac{S}{N} = P_t G_t \frac{\lambda^2}{(4\pi r)^2} \frac{1}{k\Delta B} \frac{G_r}{T}$$

Le rapport G_r/T donne la sensibilité du récepteur. Pour une température élevée, ce rapport sera plus petit. La réception sera moins bonne. D'où l'idée d'utiliser des dispositifs refroidis. Une autre façon d'améliorer la réception est d'augmenter le gain de l'antenne de réception.

6.2.5 Communications courte distance

Sous cette dénomination traduite de l'anglais SRD (*Short Range Devices*), on inclut toutes les applications radio uni ou bidirectionnelles qui ont une faible capacité à interférer avec d'autres équipements. Ce sont en général des applications dont le débit de données reste faible (inférieur au Mbit/s) avec des schémas de codage et de modulation peu complexes en comparaison des techniques utilisées en communications de téléphonie étendue où les données (audio, vidéo et data) doivent être transmises à haut débit tout en économisant le spectre occupé. On les distingue donc en cela d'autres applications qui peuvent aussi être de nature courte distance comme les WBAN (*Wireless Body Area Network*), WPAN (*Wireless Personal Area Network*) ou WLAN (*Wireless Local Area Network*).

Nous pouvons lister quelques applications :

- Identification par radiofréquences (RFID)
- Télécontrôle pour systèmes domotiques
- Systèmes d'alarme et de prévention
- Télétransmission de capteurs distants
- Automobiles (immobilisation, démarrage de véhicules...)

Les bandes de fréquences allouées peuvent varier de région à région. Les fréquences d'utilisation sont listées dans le tableau 6.6.

Tableau 6.6 – Allocations des fréquences européennes pour les « srd »

Bande	Applications	Standards de régulation
13,553-13,567 MHz	RFID HF	ERC REC70-03 - EN300 220
40,66 – 40,70 MHz	Télémétrie, capteurs...	ERC REC70-03 - EN300 220
433,05 à 434,79 MHz	Domotique, capteurs...	ERC REC70-03 - EN300 220
863 à 870 MHz	Capteurs, contrôle...	ERC REC70-03 - EN300 220
2 400 à 2 485 MHz	Télésurveillance...	Non spécifique ERC REC70-03 - EN300 440

Dans les équipements SRD, on suit une tendance identique à celle des constructeurs d'antennes d'équipements de téléphonie mobile. C'est-à-dire que l'on cherche à remplacer l'antenne externe, en général des monopôles, par des antennes internes, en général à faible profil dont les formes sont dérivées de l'antenne imprimée ou patch. Aujourd'hui, une antenne interne bien conçue peut afficher un rendement presque équivalent à son homologue externe.

En plus des caractéristiques électriques et de rayonnement qui doivent être maintenues sur l'ensemble de la bande de fréquences (gain, TOS...) les antennes en communications courte distance doivent afficher des caractéristiques particulières :

- une grande intégrabilité
- un rendement élevé
- un rayonnement omnidirectionnel
- une insensibilité à l'environnement métallique proche.

Les puissances mises en jeu sont en général faibles (inférieures à la dizaine de mW), il est donc essentiel de préserver le rendement d'antenne, qui traduit la conversion de puissance électrique fournie en puissance rayonnée, afin de prolonger l'énergie de la batterie très fréquemment utilisée dans ces applications SRD.

Il existe quelques antennes très utilisées en communication SRD. Il est préférable d'utiliser une antenne de volume quand cela est possible pour des questions de rendement d'antenne. Toutefois si le faible profil devient une contrainte forte, alors il reste la possibilité d'utiliser une antenne de

type imprimée comme un monopole, un dipôle replié, une boucle ou une antenne de type CMS comme une hélice ou une antenne céramique. Nous allons traiter successivement les cas des antennes monopole, hélicoïdale et boucle à faible profil.

□ Monopole quart d'onde imprimé

Un dipôle replié peut être attractif car sa forte impédance d'entrée 292 Ohm (4 fois l'impédance d'un dipôle simple) permet de l'adapter facilement aux impédances des circuits « front-end » RF. Un autre avantage est qu'il s'agit d'une antenne symétrique et que la plupart des circuits intégrés fonctionnent en différentiel. Cela évite l'usage d'un symétriseur (balun) et donc des pertes supplémentaires. Ses caractéristiques de rayonnement sont similaires à celle du dipôle imprimé.

Toutefois, dans la plupart des cas, la structure dipôle se révèle trop grande et une structure monopôle ou dérivée est souvent retenue.

Comme l'antenne dipôle, l'antenne quart d'onde est une antenne électrique. À ce titre, son champ proche est influencé par les matériaux diélectriques se trouvant à proximité. Par exemple, elle va être désaccordée si on l'approche du corps d'une personne (le ϵ_r d'un corps humain est d'environ 80). Nous aurons donc à définir une constante diélectrique effective pour prendre en compte le fait que des lignes de champ se trouvent dans deux structures différentes :

$$\epsilon_{eff} = \frac{\epsilon_r + 1}{2} - \frac{\epsilon_r - 1}{2} \times \left[\frac{1}{\sqrt{1 + \frac{12h}{w}}} + 0.04 \left(1 - \frac{w}{h}\right)^2 \right]$$

où w est la largeur des traces du monopole et h la hauteur du substrat.

Cette constante permettra de calculer la longueur du monopole :

$$L_{mono_pcb} = \frac{\lambda}{4\sqrt{\epsilon_{eff}}}$$

Tous les éléments parasites, comme les capacités de couplage à la masse, les inductances supplémentaires à chaque coude du monopole ou l'influence de l'électronique de proximité et du boîtier, vont modifier l'impédance de l'antenne.

Lorsque l'antenne doit être raccourcie pour économiser l'espace occupé, il faut se rappeler que la bande passante va décroître comme le cube du rapport a/λ :

$$BW = \frac{1}{Q} = \frac{1}{Q_r} + \frac{1}{Q_p} = \frac{1}{Q_r} \left(\frac{Q_p + Q_r}{Q_p} \right) = \left(\frac{2\pi a}{\lambda} \right)^3 \times \frac{1}{\eta}$$

où a est le rayon de la sphère minimum entourant l'antenne

η est le rendement de l'antenne

Q_r est le coefficient de qualité lié au rayonnement

Q_p est le coefficient de qualité lié aux pertes Joule

La résistance de rayonnement décroît avec la longueur et le facteur de qualité augmente, entraînant une réduction de la bande :

$$R_r = 395 \times \left(\frac{L}{\lambda} \right)^2 \text{ pour } 0 < L < \frac{\lambda}{8}$$

$$R_r = 1\,220 \times \left(\frac{L}{\lambda} \right)^{2.5} \text{ pour } \frac{\lambda}{8} < L < \frac{\lambda}{4}$$

Pour une antenne monopole raccourcie, possédant un coude (figure 6.22), l'impédance d'entrée est capacitive et donc le circuit d'adaptation sera constitué d'une self série et d'une capacité

parallèle dont les pertes s'additionneront aux pertes du conducteur d'antenne et diminueront le facteur de qualité.



Figure 6.22 – Antenne monopole imprimée Texas Instruments.

□ Antenne hélicoïdale à mode transversal

En enroulant le monopole, on aboutit à la structure d'antenne hélicoïdale (figure 6.23).

Lorsque la circonférence et le pas de l'hélice sont faible devant λ , on obtient un mode de rayonnement dit transversal. Le rayonnement est similaire à celui d'un monopole, donc perpendiculaire à l'axe de l'antenne avec une polarisation elliptique. En général, pour une hélice verticale la composante de champ électrique axiale est plus forte que la composante radiale. Le calcul analytique d'une hélice est en général plus complexe que celui d'un monopole. Du point de vue des caractéristiques électriques, sa largeur de bande est faible et elle est très sensible aux tolérances des composants d'adaptation. Classiquement le gain vaut 0,5 à 1 dBi et le rendement de l'ordre de 60 à 70 %.

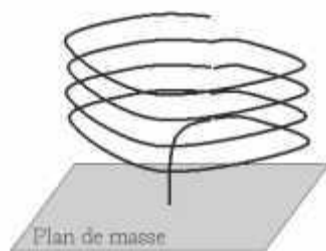


Figure 6.23 – Antenne hélicoïdale verticale.

□ Petites antennes boucle

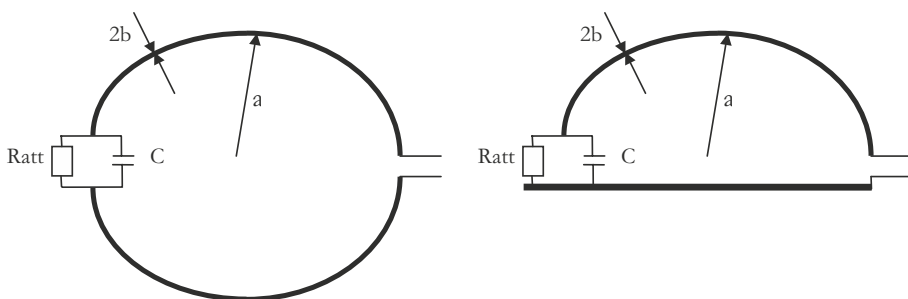


Figure 6.24 – Antenne boucle imprimée différentielle et « single-ended ».

Pour pouvoir faire l'hypothèse d'un courant quasiment constant, il est nécessaire que la circonférence soit inférieure à $\lambda/10$. Dans ce cas, l'antenne est assimilée à une self rayonnante dont la valeur est celle de la self du ruban sur le substrat. Électriquement, il faut donc disposer une capacité pour accorder l'antenne.

Pour le calcul de l'inductance L , la section du ruban est considérée comme étant circulaire de rayon $b = 0,35d + 0,24w$ où d est l'épaisseur de cuivre et w la largeur de la trace.

La résistance de rayonnement est faible, typiquement inférieure à 1Ω . Les pertes prennent en compte les pertes du conducteur mais aussi les pertes de la capacité d'accord C qui ne sont plus négligeables (ESR).

En raison du fort rapport L/C , le facteur de qualité est important, ce qui rend l'antenne sensible au désaccord, effet que l'on peut diminuer avec la résistance en parallèle R_{att} . Cette résistance transformée en résistance série vaut $R_{att_série} = \frac{2\pi f L}{Q} - R_r - R_p$ et est définie en fonction des autres pertes et du facteur de qualité utilisable.

Cela donne une efficacité d'antenne : $\eta = \frac{R_r}{R_r + R_p + R_{att_série}}$

Dans la plupart des cas, la résistance de rayonnement est plus faible que celle de pertes (surtout $R_{att_série}$), donnant une faible efficacité : $\eta = \frac{QR_r}{2\pi f L}$

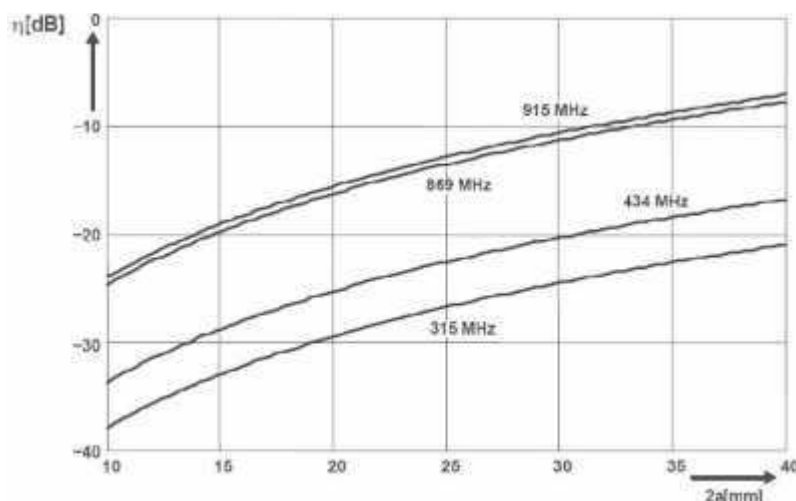


Figure 6.25

L'antenne boucle donne une polarisation linéaire dans le plan perpendiculaire à la boucle. Elle est du type magnétique et n'est donc pas désaccordée par les matériaux diélectriques environnants. C'est donc un choix judicieux pour une application où l'antenne doit être portée près du corps malgré des rendements inférieurs à 50 %.

6.3 Télévision et radiodiffusion FM

6.3.1 Performances requises

Les bandes de fréquences utilisées pour la télévision et la radiodiffusion FM occupent une grande partie du spectre VHF et UHF. Elles sont rappelées dans le tableau 6.7 pour la France.

La largeur d'un canal TV est de 6 MHz pour tous les canaux, ce qui implique une largeur de bande relative de 10,5 % pour le canal 2 et de 0,7 % pour le canal 69. En ce qui concerne la bande FM, un canal occupe 200 kHz de bande, représentant 0,2 % de largeur de bande.

Les performances requises dépendent de la classe de la station et du choix du canal. En France les stations émettent en polarisation horizontale la plupart du temps.

Tableau 6.7 – Bandes TV et FM et fréquences pour la France

Canaux	Fréquences (MHz)
TV E2 à E4	47-68 (VHF)
FM	88-108 (VHF)
Bande III : TV L5 à L10	174-223 (VHF)
Bande IV : TV C21 à C38	470-606 (UHF)
Bande V : TV C38 à C69	606-862 (UHF)

6.3.2 Caractéristiques générales des antennes

Il existe de nombreuses antennes utilisées aussi bien pour la transmission que la réception. Toutefois, pour la conception des antennes utilisées en transmission, il est impératif de respecter certaines caractéristiques et notamment l'adaptation de l'antenne à l'amplificateur de puissance :

- TOS maximum de 1.1 pour les canaux TV
- TOS maximum de 1.2 pour les canaux FM

La puissance ERP, exprimée comme le produit de la puissance électrique incidente par le gain d'antenne (référé au gain du dipôle demi-onde) a été établie pour permettre à toutes les stations, à une hauteur d'antenne au-dessus du sol moyen donnée, de couvrir des zones approximativement égales.

Dans les applications de radiodiffusion à modulation de fréquence (FM) et de télévision, la majorité des stations cherchent à afficher un diagramme de rayonnement omnidirectionnel dans le plan de l'azimut. On peut aussi trouver des diagrammes en forme de cardioïde, de cacahuète... ; cela afin de se protéger d'une source de rayonnement parasite ou de ne pas envoyer d'énergie dans une zone non habitée. De plus, il est intéressant de chercher à obtenir le plus grand gain, pour ne pas atteindre les limites de puissance électrique que permet la technologie.

Afin d'augmenter le gain dans le plan azimutal, il va être nécessaire d'utiliser une grande ouverture dans le plan vertical.

Il est donc impératif que le diagramme de rayonnement soit omnidirectionnel dans le plan horizontal et très directif dans le plan vertical. Cela nous mène à envisager l'utilisation de dipôles verticaux colinéaires. En effet, en raison de la symétrie de révolution d'un dipôle et de la théorie des groupements d'antennes, le diagramme va présenter un maximum dans le plan horizontal.

En général, les largeurs de faisceaux sont de l'ordre de 7, 4, 2 et 1°, respectivement aux bandes VHF basse, FM, VHF haute et UHF. On peut se rappeler que le gain est donné approximativement par $50\lambda/D$, référé au dipôle demi-onde. Un problème commun est d'incliner le diagramme en azimut et de remplir les trous de rayonnement (*null filling*), notamment lorsque l'antenne est située près d'une zone résidentielle.

La structure d'antennes consiste généralement en un réseau vertical constitué d'antennes ou de variantes d'antennes dipôles, boucles, hélices, fentes...

La plupart du temps, ces antennes sont montées le long d'un mât (*side-mount antennas*). L'inclinaison ou le remplissage des zones mortes sont obtenus en modifiant les phases et l'amplitude des courants d'alimentation des différentes baies d'antennes ou plus simplement par un tilt mécanique.

À l'exception de la largeur de bande relative, les caractéristiques d'antennes pour la radiodiffusion FM ou TV sont similaires.

6.3.3 Les antennes en transmission

Elles sont pratiquement tout le temps montées sur mât pour garantir les performances en s'éloignant du sol et utilisées en réseau.

6.3.4 Les antennes de types panneaux

On appelle antenne panneau une antenne qui comprend un écran réflecteur avec un élément simple rayonnant de type dipôle. Par rapport aux antennes Yagi, on peut citer comme avantages principaux :

- Gain, diagramme de rayonnement et TOS plus constant sur la bande de fréquences
- Construction physique plus compacte
- Couplage faible avec la structure porteuse.

On peut parfois utiliser des panneaux réflecteurs devant les antennes pour éviter les réflexions erratiques sur la tour ou utiliser le mât comme faisant parti du réflecteur.

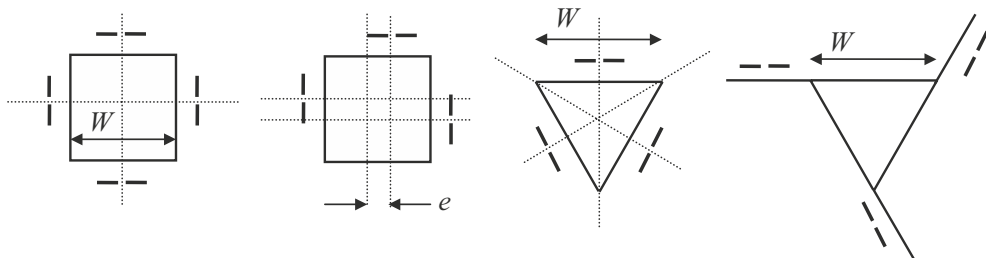


Figure 6.26 – Différents types d'antennes panneau.

En théorie, chaque couple antenne/panneau affiche un diagramme unidirectionnel dans une direction perpendiculaire au panneau. Toutefois, en pratique, il sera indispensable de concevoir l'antenne en ayant pris en compte le mât triangulaire ou carré de la tour support.

Pour obtenir un diagramme omnidirectionnel, chaque couple doit afficher une largeur de faisceau de 90° pour une tour carrée et 120° pour une tour triangulaire.

Pour un diagramme omnidirectionnel, la largeur W des tours (figure 6.26 gauche) ne doit pas excéder 1λ . On trouve aussi dans le cas des *skewed dipoles*, l'excentration des panneaux (figure 6.26 droite), cas utilisé lorsque la tour mesure plus de 4λ de côté (nécessaire pour ne pas dégrader le ratio $E_{\max}/E_{\min}$).

■ Les antennes en polarisation circulaire

Beaucoup d'antennes TV et FM polarisées circulairement sont réalisées d'après le concept d'un réseau circulaire de dipôles inclinés selon leur axe. Les dipôles pouvant être linéaires, en forme de V, courbés ou de forme plus complexe. Chaque dipôle rayonne un champ polarisé horizontalement mais en ajustant le diamètre du réseau circulaire et l'angle d'inclinaison ψ , il est possible d'obtenir un diagramme omnidirectionnel en polarisation circulaire. C'est une approche qui diffère donc de celle considérant des antennes affichant un diagramme unidirectionnel et déjà polarisées circulairement (antenne tourniquet, antennes dipôles croisés en cavité, hélice...).

□ L'antenne Tourniquet

Elle est constituée de deux dipôles demi-onde perpendiculaires et alimentés en quadrature de phase. Le champ créé par cet ensemble s'écrit :

$$E = \frac{\cos(90^\circ \cos \theta)}{\sin \theta} \cos \omega t + \frac{\cos(90^\circ \sin \theta)}{\cos \theta} \sin \omega t \quad [6.1]$$

Le champ est donc approximativement constant dans une direction normale aux dipôles et affiche une polarisation circulaire. La variation d'amplitude de diagramme étant d'environ $\pm 0,3$ dB. Elle est facilement montée sur un mât vertical et donc peu facilement être mise en réseau.

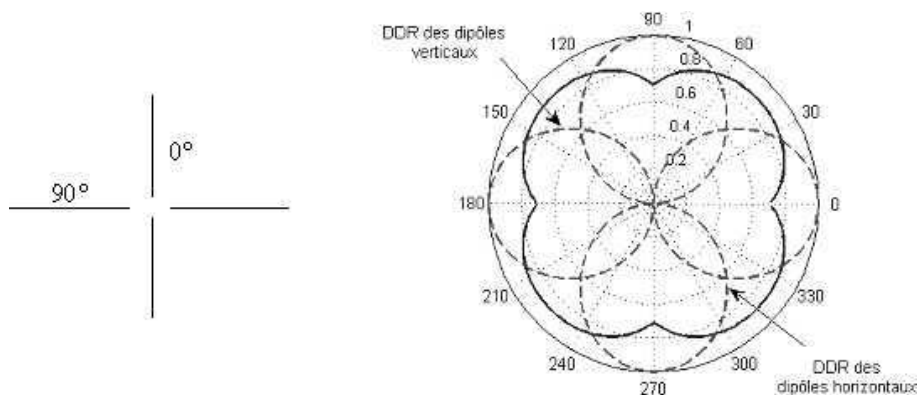


Figure 6.27 – Caractéristiques de l'antenne Tourniquet.

□ L'antenne « Skewed dipole »

Elle est constituée d'un réseau circulaire de dipôles tel que chaque dipôle rayonne en polarisation linéaire, mais placé en configuration telle que le réseau fonctionne en polarisation circulaire (figure 6.28). Pour des applications de radiodiffusion, on utilise un certain nombre de réseaux élémentaires placés autour d'un mât conducteur vertical :

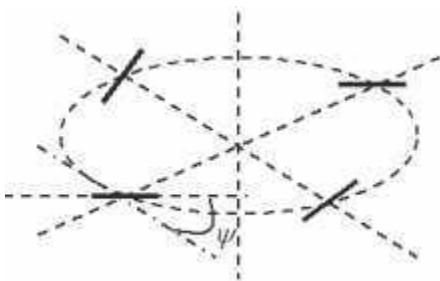


Figure 6.28 – Antenne « skewed dipole ».

Chaque réseau est constitué de trois antennes alimentées en phase et montées symétriquement autour du mât. L'angle de « skew » ψ est choisi pour produire des composantes horizontale et verticale égales. Quand on prend en compte l'ensemble des rayonnements élémentaires, on obtient un diagramme de rayonnement omnidirectionnel en azimut avec un faible rapport axial pour tous les angles d'azimut.

Dans la bande UHF, les antennes sont identiques à celles utilisées en VHF. On peut toutefois noter que, en raison d'une largeur de bande relative plus faible, il est permis d'utiliser plus facilement des antennes dites résonnantes comme les antennes coaxiales ou à guide d'ondes fendues.

En ce qui concerne les antennes utilisables aussi dans les bandes FM, on retrouve des antennes panneaux. On peut les utiliser dans une configuration où le mât supporte un certain nombre d'antennes, chacune dédiée à une station particulière. Dans le cas où l'on désire une antenne unique pour le multiplex des stations, on peut utiliser des antennes boucles en réseau ou non

dans le cas d'une polarisation horizontale. La largeur de bande reste faible mais suffisante pour la bande à couvrir en FM. En ce qui concerne la polarisation circulaire, on peut grouper les antennes en deux catégories : la combinaison d'antennes boucle horizontale/dipôle vertical et les « skewed dipôles ». La seconde catégorie permet d'obtenir de meilleures performances en TOS sur une plus grande largeur de bande (environ 5 MHz contre 1 MHz).

■ Les antennes en polarisation horizontale

□ L'antenne Super Tourniquet

Pour obtenir un très faible TOS sur une largeur de bande d'environ 50 %, on peut utiliser l'antenne Super Tourniquet où les dipôles sont remplacés par leurs structures duales c'est-à-dire des fentes rayonnantes. La polarisation est horizontale conformément à la configuration du champ électrique dans la fente. La longueur de la fente est légèrement supérieure à la demi-onde. Très souvent le plan métallique est ajouré pour diminuer la résistance au vent. En termes de performances, c'est une antenne de choix qui permet d'obtenir un TOS de 1,1 sur une largeur de bande de 50 % sur l'ensemble des canaux VHF. Cette antenne est aussi utilisée en réseau. En général, de 2 à 6 baies sont utilisées pour les canaux 2 à 6 et jusqu'à une quinzaine pour les canaux 7 à 13.

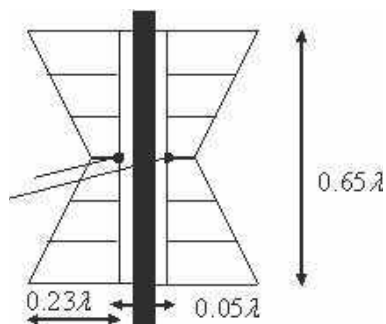


Figure 6.29 – Caractéristiques de l'antenne Super Tourniquet.

■ Les antennes en réception

Elles sont nombreuses. On pourrait citer notamment les antennes boucle, papillon et les antennes large bande comme l'antenne Yagi-Uda et log périodique. Il leur est demandé d'avoir comme caractéristiques :

- un gain suffisant pour ne pas dégrader le bilan de liaison
- une impédance d'entrée correcte
- un diagramme avec peu de lobes latéraux et arrière
- une très grande largeur de bande (de 50 à 1 000 MHz)

L'antenne extérieure la plus commune est celle qui combine une antenne VHF et une antenne UHF pour former une structure unique. Les antennes VHF les plus communes sont les antennes Yagi-Uda (voir le paragraphe 8.5) avec réflecteur, les antennes triangulaires avec réflecteur plat et les antennes LPDA (*Log periodic Dipole Array*).

La LPDA (figure 6.30) est une antenne large bande constituée d'une mise en parallèle de dipôles résonant demi-ondes. Chaque brin rayonne de façon optimale lorsque sa longueur correspond à la demi-onde. Nous avons donc un rayonnement des hautes fréquences sur les petits éléments et vice-versa. Elle est donc capable d'offrir un gain constant et une impédance d'entrée constante sur une largeur de bande de rapport 30 :1. Elle est conçue pour vérifier la loi géométrique :

$$\frac{x_{n+1}}{x_n} = \frac{l_{n+1}}{l_n} = \tau = cte \quad [6.2]$$

On appelle τ le facteur d'accroissement. Ses performances principales sont :

- Gain de 6 à 11 dB (en référence au dipôle demi-onde).
- Une déviation standard en gain de l'ordre de 1,5 dB sur la bande.
- Une largeur de bande donnée par le ratio longueur max sur longueur min de l'ordre de 3.

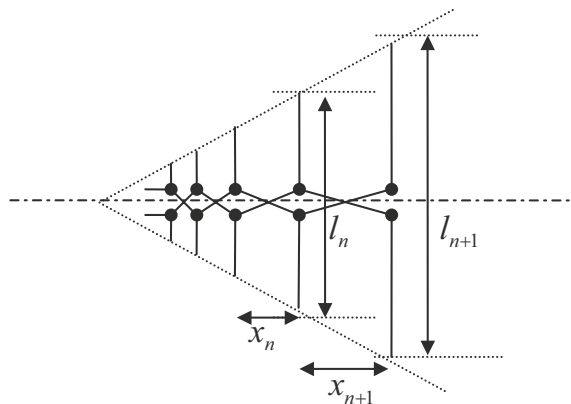


Figure 6.30 – Géométrie de l'antenne LPDA.

6.4 Radar

6.4.1 Généralités

Le RADAR est une technique de détection et de mesure de position d'objets réfléchissant les ondes électromagnétiques. Les caractéristiques d'antennes radar sont étroitement liées aux capacités de couverture, c'est-à-dire la capacité à voir « loin », et de résolution du radar, c'est-à-dire la capacité à discerner deux cibles proches.

En vertu de la loi liant l'ouverture rayonnante et la largeur du faisceau, il est assez facile d'imaginer que l'antenne doit afficher une grande surface de rayonnement ou de captation si l'on désire un grand pouvoir de résolution.

Pour une application donnée, le choix de la fréquence radar se fera donc en fonction de la taille et de la largeur de faisceau d'antenne possible mais aussi en fonction du bruit environnant. De façon générale, les caractéristiques d'antennes importantes sont données dans le tableau 6.8.

Tableau 6.8 – Caractéristiques importantes des antennes radar

1. Puissances max et moyenne	7. Surface de balayage
2. Gain	8. Précision de pointage
3. Largeur de faisceau	9. Volume, poids
4. Niveau de lobes secondaires	10. Capacité à être démontée
5. Largeur de bande en impédance	11. Sensibilité à l'environnement
6. Polarisation et cross polarisation	12. Coût

Les caractéristiques de l'antenne choisie vont être fortement dépendantes de l'environnement dans lequel le radar va opérer. Par exemple, s'il s'agit d'un radar aéroporté, le volume occupé par l'antenne est primordial alors qu'une antenne faible profil aura la préférence pour être plaquée sur un avion ou un missile.

Aussi, si le radar est fixe et opère dans des contrées à forts taux de précipitation comme dans les régions arctiques, l'antenne devra être conçue avec son radôme.

De la même façon, dans un environnement de type civil, le radar subit des interférences qui sont de nature involontaires. Au contraire, un radar militaire va subir des agressions intentionnelles

pour le mettre hors d'opération, il devra donc être capable de subir ces signaux de contre-mesures et souvent l'antenne devra afficher une grande largeur de bande (contre les interférentes bandes étroites) et un faible taux de lobes secondaires pour éviter l'aveuglement.

Contrairement aux antennes de communication, il est nécessaire d'utiliser de fortes puissances, en général pulsées avec un rapport cyclique de 0,1 à 10 %, pour localiser des cibles lointaines. Cela place une contrainte particulière sur la capacité de l'antenne à supporter des puissances allant de quelques dizaines de kW à quelques MW. Il peut être intéressant, et c'est la tendance actuelle, d'utiliser des antennes réseau qui distribuent la puissance sur les différents éléments rayonnants. Il est important aussi de noter que tous les radars doivent avoir une antenne capable d'effectuer un balayage de faisceau, de façon mécanique ou électronique. Actuellement, un grand nombre de radars sont à balayage électronique, ce qui permet d'améliorer les vitesses de balayage en évitant l'inertie mécanique.

Il est aussi important de réaliser que la précision de pointage d'antenne est une donnée fondamentale puisqu'elle a un impact direct sur la capacité de résolution. Dans le cas d'une antenne mécanique, cette caractéristique est liée à la mécanique elle-même tandis que pour une antenne électronique, le paramètre clé est la précision obtenue sur la valeur des déphasages de chaque antenne. Il est habituel de catégoriser les radars en radar de recherche ou de poursuite. La fonction de recherche implique que le radar doit balayer périodiquement le même volume et calculer les positions des objets rencontrés.

Dans un radar de poursuite, une ou plusieurs cibles sont sous surveillance, ce qui oblige le radar à avoir des caractéristiques dynamiques (changement brusque de la position du faisceau par exemple ou traitement des données plus rapide) meilleures que celles d'un radar de recherche. Parfois, certains radars à balayage électronique possèdent les deux fonctions de façon à ce qu'une cible particulière soit détectée parmi un grand nombre de cibles puis suivie.

Les fréquences utilisées en techniques radar peuvent aller de quelques MHz pour les radars OTH (*Over The Horizon*) à quelques dizaines de GHz pour les radars anticollisions.

Comme principales applications, nous pouvons citer :

- contrôle de trafic aérien (civil et militaire)
- aide à la navigation aérienne (aéroport...)
- aide à la navigation maritime
- aide au contrôle des satellites
- télédétection
- médicales (traitement des tumeurs...)
- de contrôle routier (vitesse).

6.4.2 Radar de recherche ou veille ou surveillance

Un radar de recherche ou veille ou surveillance scanne une zone à intervalle régulier. Il a été montré que pour un temps de balayage fixé et une taille de cible donnée, la capacité de couverture du radar est liée au produit de la puissance moyenne transmise par la surface d'ouverture de l'antenne et n'est donc pas seulement liée au choix de la fréquence. Bien sûr, celle-ci reste importante puisque l'ouverture est exprimée en λ^2 .

Si l'on s'en tient à la forme simple de l'équation du radar monostatique (même antenne d'émission et de réception) :

$$R_{\max} = \left[\frac{P_t G A_e \sigma}{(4\pi)^2 S_{\min}} \right]$$

Avec P_t : puissance de transmission

G : gain de l'antenne de transmission/réception

A_e : ouverture effective de l'antenne de transmission/réception

σ : surface équivalente radar

$S_{\min}$: signal détectable minimum

On voit que les paramètres importants de l'antenne sont le gain de transmission et son ouverture effective.

Mais cette équation ne décrit pas de façon adéquate la performance d'un radar pratique. En effet, beaucoup d'autres paramètres ont une influence notable qui aboutit à une portée moins grande que la portée théorique.

Il est important de noter que la probabilité de détection d'une cible sera différente si le radar utilise (cas pratique) l'intégration des pulses reçus lors de la période de rotation de l'antenne. Le nombre de pulses reçus N est :

$$N = \frac{\theta_B f_p}{\dot{\theta}_S} = \frac{\theta_B f_p}{6 f_S}$$

Avec θ_B : ouverture du faisceau (en °)

$\dot{\theta}_S$: vitesse angulaire de rotation (en °/s)

f_p : fréquence de répétition des pulses (en Hz)

f_S : fréquence de rotation de l'antenne (en tours par minute)

Ici aussi, nous notons que la largeur de faisceau doit être faible si l'on veut maximiser le nombre de pulses à traiter.

Intuitivement, il semble préférable que l'antenne rayonne dans un pinceau étroit pour avoir une bonne résolution de détection. En général, ce n'est pas la solution retenue, en raison des contraintes opérationnelles du système sur le temps de scan maximum, ce qui signifie que le radar ne peut pas rester longtemps sur un pixel de la zone. Cela est particulièrement vrai si la résolution recherchée est grande. Le nombre de cellules peut être réduit si le pinceau est large dans une direction et étroit dans une autre (figure 6.31), ce que l'on obtient avec une antenne à grande dimension horizontale et étroite en dimension verticale. On obtient cette forme en translatant sur un axe horizontal un ensemble d'arcs de paraboloïde verticaux.

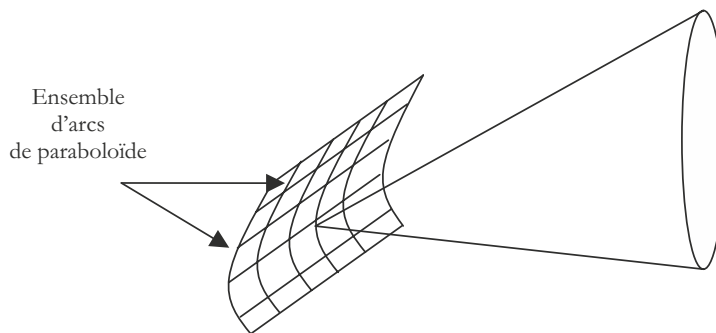


Figure 6.31 – Principe de la génération d'un faisceau vertical étroit.

La couverture possible avec ce type d'antenne est en général inadéquate pour la détection des cibles sur des directions de site élevées.

Pour le concepteur du système, le problème se résume donc à la synthèse d'un certain diagramme de rayonnement correspondant à la zone à couvrir. Cette zone est toujours tracée dans le plan vertical, sachant que dans le plan horizontal, on cherche à obtenir un diagramme qui est le plus étroit possible. La couverture dans ce plan est obtenue par rotation de l'antenne.

Le choix de la forme du diagramme va dépendre de l'application et notamment de la largeur de la zone à couvrir exprimée en angle.

■ Diagramme en cosécante carré

Une technique possible est de modifier le diagramme précédent pour obtenir un pinceau dont la forme est proportionnelle au carré de la cosécante de l'angle d'élévation.

La figure 6.32 illustre le problème de la détection de cibles basses ou hautes en fonction de la forme du diagramme de rayonnement.

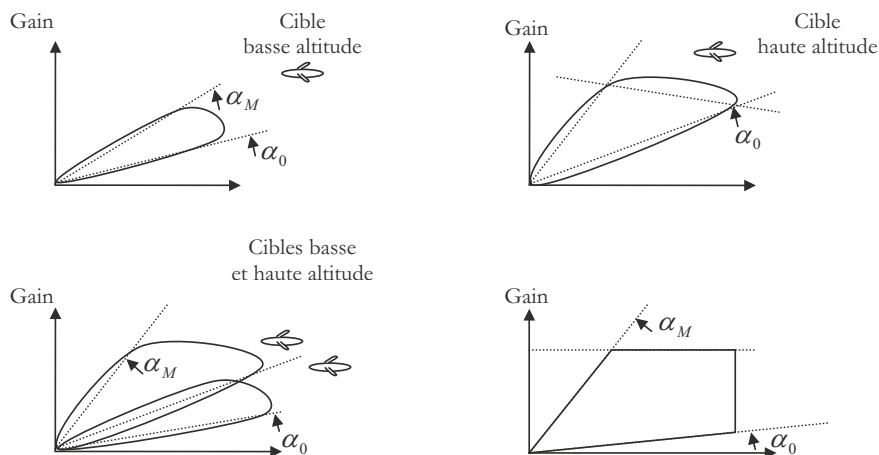


Figure 6.32 – Problème de la détection de cibles basses et hautes.

Le gain exprimé en fonction de l'angle d'élévation α s'écrit :

$$G(\alpha) = G(\alpha_0) \frac{\csc^2 \alpha}{\csc^2 \alpha_0} \quad \text{pour } \alpha_0 < \alpha < \alpha_M$$

avec $G(\alpha_0)$: gain dans la direction α_0

α_0 : angle d'élévation minimum

α_M : angle d'élévation maximum

En pratique, pour des angles inférieurs à α_0 , le gain est quelconque et suit une loi en cosécante carré seulement sur la partie $\alpha_0 < \alpha < \alpha_M$. Idéalement, il serait souhaitable d'avoir α_M aussi proche que possible de 90° mais en pratique, cela est difficile à obtenir.

Ce diagramme possède l'importante propriété suivante : la puissance renvoyée par des cibles situées à une hauteur constante H ne dépend pas de la direction de visée du radar.

Si l'on remplace le gain d'une antenne en cosécante carré dans l'équation du radar simplifiée, nous obtenons :

$$P_R = \frac{P_T G^2(\alpha_0) \csc^4(\alpha) \lambda^2 \sigma}{(4\pi)^3 \csc^4(\alpha_0) R^4}$$

Cette puissance est de la forme :

$$P_R = K \frac{\csc^4(\alpha)}{R^4}$$

D'un point de vue géométrique, on a : $\csc(\alpha) = \frac{R}{H}$

Nous pouvons donc constater que la puissance reçue est indépendante de la direction de visée si les cibles sont à des altitudes identiques.

En pratique, cela n'est pas exactement vrai car il n'est pas possible de synthétiser exactement un diagramme en cosécante carré.

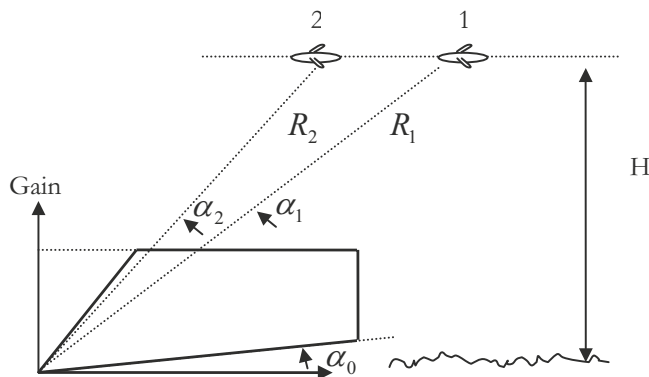


Figure 6.33 – Détection de cibles d'altitude identique et de direction différente.

D'un point de vue général, il est possible d'écrire que la portée R d'un radar dans une direction α pour une certaine probabilité de détection est donnée par :

$$R^4(\alpha) = a G^2(\alpha)$$

avec a : facteur qui dépend notamment de la puissance du transmetteur, surface d'ouverture d'antenne, sensibilité du récepteur...

$G^2(\alpha)$: gain de l'antenne dans cette direction α

Comme le gain est proportionnel à la puissance rayonnée et que celle-ci est proportionnelle au carré du champ, l'équation de la portée devient :

$$R(\alpha) = a_1 E(\alpha)$$

Nous voyons que portée et champ électrique rayonné sont liés. Il est habituel de tracer cette valeur de portée en coordonnées polaires ; la courbe obtenue n'est donc que le diagramme de rayonnement de l'antenne. Cette variation de portée s'appelle la couverture radar.

■ Conception d'antennes de diagramme en cosécante carré

D'un point de vue construction physique, ce diagramme en cosécante carré peut être obtenu :

- Par distorsion des arcs de parabole et une source primaire unique (figure 6.34) : la partie haute de l'antenne approxime une parabole et renvoie les rayons issus de la source primaire

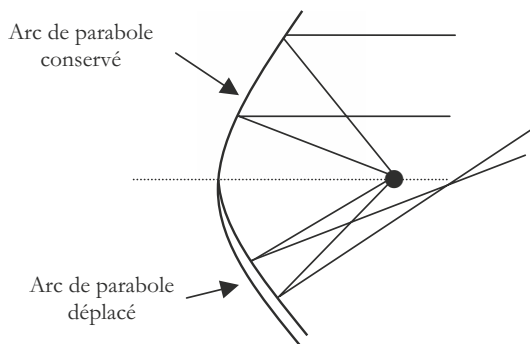


Figure 6.34 – Antenne cosécante carré produite par déformation du bas de la parabole.

parallèlement à l'axe de l'antenne tandis que la partie basse s'éloigne de la forme parabolique pour pouvoir diriger le rayonnement vers le haut.

- En conservant ces arcs mais en illuminant ce réflecteur par une source primaire constituée de plusieurs antennes souvent disposées en ligne (figure 6.35) : lorsque les cornets sont correctement espacés et alimentés, la somme des différents faisceaux secondaires permet de diriger correctement le rayonnement dans une certaine gamme d'angles.
- Avec une antenne réseau (par synthèse de diagramme).

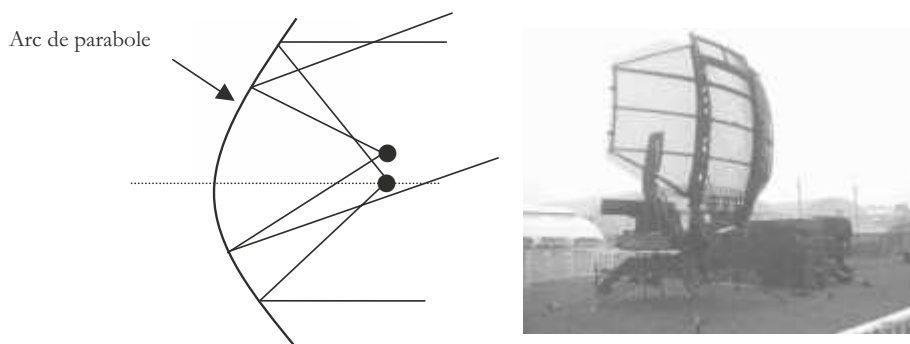


Figure 6.35 – Antenne cosécante carré produite par deux sources primaires en décalage – radar 3D TPS-43 de Westinghouse (source : <http://en.wikipedia.org/wiki/AN/TPS-43>).

Il est à noter, comme nous pouvons le voir sur la photo de la figure 6.35, que le réflecteur doit posséder une double courbure pour avoir la formation correcte du faisceau dans les directions définies par l'angle de site et des propriétés de focalisation en azimut.

Le radar 3D TPS-43 de Westinghouse affiche les caractéristiques suivantes :

- Bande de fréquences : 2,9 à 3,1 GHz
- Fréquence de répétition : 250 Hz
- Largeur de pulse : 6,5 microsecondes
- Puissance crête : 4,0 MW
- Puissance moyenne : 6,7 KW
- Largeur de faisceau (horizontal) : 1,1°
- Largeur de faisceau (vertical) : 1,5 à 8,1°, au total 20° de couverture par 6 faisceaux
- Vitesse de rotation : 6 tours par minute
- Portée maximale : 450 km
- Caractéristiques d'antenne : ouverture du réflecteur : 4,27 m de haut par 6,20 m de large

6.4.3 Radar de poursuite

Un radar de poursuite permet de donner à tout instant la distance et la direction d'une cible tout en se maintenant pointé sur elle de façon automatique. Cela se fait grâce à l'élaboration de signaux d'erreur en élévation et en azimut si la poursuite doit se faire en trois dimensions. Ces signaux permettent l'asservissement de l'antenne dans la direction de la cible.

Une technique de base utilisée est la technique d'écartométrie qui permet de délivrer deux signaux :

- Un signal appelé signal somme qui constitue l'écho radar et dont l'amplitude permettra d'obtenir la distance de la cible

- Un signal appelé signal différence qui sera l'image de l'écart d'angle entre la direction de pointage de l'antenne et la direction de la cible

Du point de vue de l'implantation de cette technique au système d'antenne, trois méthodes ont été étudiées mais seulement les deux dernières ont été réellement utilisées :

- Génération séquentielle de lobe
- Balayage conique du faisceau
- Comparaison des amplitudes et des phases des signaux monopulse

Nous reviendrons par la suite sur les principes de ces deux techniques et notamment le terme de monopulse. Sur la figure 6.36 sont représentés les signaux d'écartométrie en fonction de l'angle formé par la direction de visée de l'antenne et la direction de la cible.

Les deux signaux sont bien entendus liés et le signal d'écho doit être maximisé lorsque le signal d'erreur est nul.

La fonction d'erreur doit être impaire afin de déterminer le signe de l'écart angulaire et donc corriger le dépointage. Si la fonction est linéaire, cela permet d'avoir une estimation de l'angle du dépointage. Cette plage linéaire s'appelle la plage d'écartométrie. Afin d'obtenir une précision de pointage importante, il est nécessaire de maximiser la pente de ce signal d'erreur.

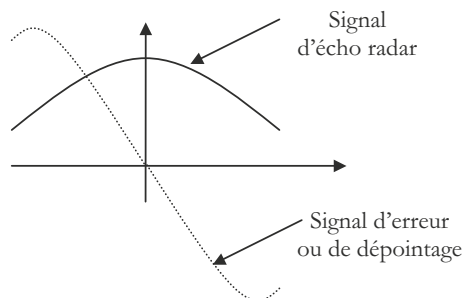


Figure 6.36 – Signaux générés en écartométrie.

■ Commutation séquentielle de lobe

Une méthode pour obtenir l'amplitude et la direction de l'erreur angulaire est de commuter le faisceau d'antenne entre deux positions angulaires (figure 6.37).

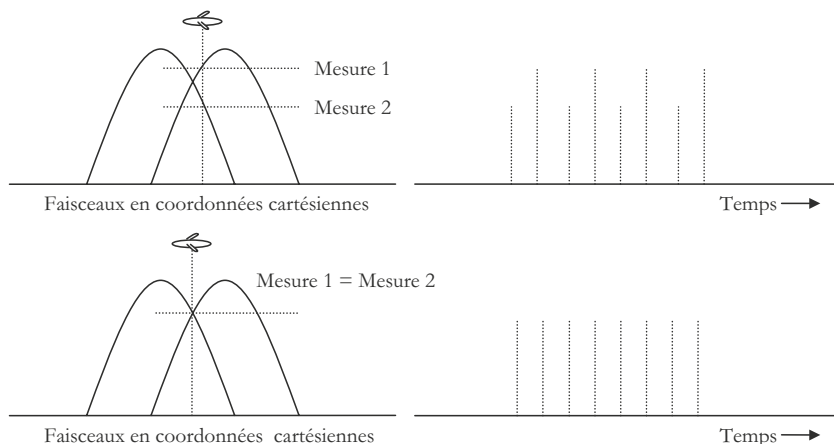


Figure 6.37 – Génération de signal d'erreur angulaire par commutation de direction de lobe.

Le signal d'erreur est fabriqué à partir de la différence des deux valeurs mesurées correspondant au gain du diagramme dans la direction de visée. Ce signal d'erreur dépend de l'écart angulaire de la cible par rapport à l'axe de l'antenne et sera nul lorsque la cible se confond avec l'axe de l'antenne. La différence entre les deux mesures sera d'autant plus forte que les lobes seront étroits et donc que le gain d'antenne sera grand.

D'un point de vue pratique, on procède par commutation de la source primaire du réflecteur. Bien sûr, la technique décrite ci-dessus permet la poursuite de la cible sur un seul axe. Par extension, il faudra ajouter deux autres sources primaires pour obtenir la poursuite dans la direction orthogonale à la première.

■ Système à balayage conique

C'est simplement une extension de la technique précédente qui consiste à faire varier continûment la position du faisceau plutôt que de façon discrète (quatre positions pour une détection en deux dimensions). En général, l'axe de rotation est confondu avec l'axe de l'antenne.

Un écho va être retourné dont la fréquence de modulation correspondra à la fréquence de rotation du faisceau. Cette modulation sera nulle lorsque l'axe de rotation et l'axe de la cible seront confondus.

L'amplitude de cette modulation va dépendre de l'angle de décalage (*squint-angle*), angle formé entre l'axe du faisceau et l'axe de rotation.

Pratiquement, il s'agit de faire varier la position du centre de phase de l'antenne primaire, qui va décrire un cercle de rayon r situé dans le plan focal. Le diagramme de révolution va donc tourner dans l'espace et la direction de rayonnement maximal va décrire un cône de demi-angle au sommet θ . La figure 6.38 représente deux positions extrêmes du faisceau. Quand une cible se présente dans une direction γ par rapport à l'axe de l'antenne, le système délivre un signal d'amplitude $M1$ en position 1 et $M2$ en position 2 et un signal compris entre $M1$ et $M2$ pour toutes les autres positions (en dehors du plan de la feuille). Le signal se retrouve modulé à la fréquence de rotation de la source primaire avec un taux :

$$m = \frac{M2 - M1}{M2 + M1}$$

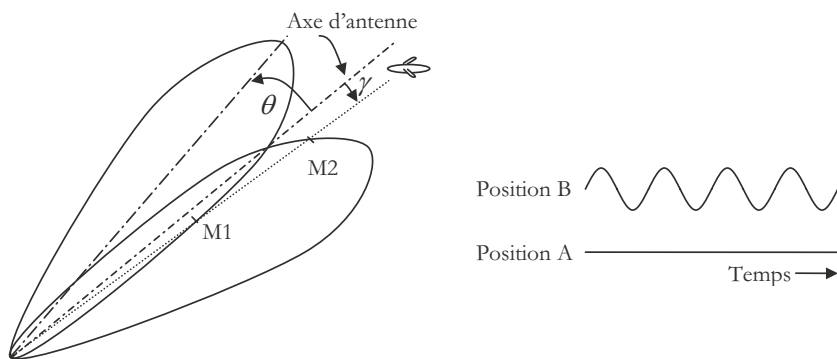


Figure 6.38 – Poursuite par balayage conique du faisceau.

La pente du signal d'écartométrie doit être maximisée ; cela a été démontré par L.Thourel pour un paraboloïde de diamètre D :

$$\left(\frac{dm}{d\varepsilon} \right)_{\varepsilon \rightarrow 0} = 1.87 \frac{D^2}{\lambda^2} \theta$$

avec $\left(\frac{dm}{d\varepsilon} \right)_{\varepsilon \rightarrow 0}$: pente du signal d'erreur

ϕ : excentration de la direction du faisceau par rapport à l'axe de l'antenne

Cette pente affiche les caractéristiques suivantes :

- Elle est maximisée lorsque l'excentration est forte
- Elle varie en fonction du carré du diamètre

Bien sûr, cela est en conflit avec la maximisation de portée si l'excentration devient trop forte. Dans ce cas, une maximisation du couple portée-pente donne pour la pente :

$$\left(\frac{dm}{d\varepsilon} \right)_{\varepsilon \rightarrow 0} = 0.945 \frac{D}{\lambda}$$

Nous pouvons donc conclure qu'il y a intérêt à choisir des antennes de grand diamètre.

■ Système à radar monopulse

Il s'agit toujours de mesures angulaires mais cette fois en traitant séparément chaque impulsion de retour, d'où le nom de monopulse.

On peut procéder par monopulse d'amplitude ou de phase.

Dans le cas du monopulse d'amplitude, l'antenne primaire est généralement constituée de deux cornets placés symétriquement de part et d'autre du foyer de la parabole (séparés d'une distance d) et dont l'angle d'ouverture à 3 dB vaut l'écart angulaire entre les faisceaux. La technique de traitement est similaire à celle par balayage conique.

Dans le cas du monopulse de phase, les deux antennes fournissent des signaux identiques en amplitude mais différents en phase :

$$\varphi = \frac{2\pi d}{\lambda} \sin \theta \approx \frac{2\pi d}{\lambda} \theta$$

La distance entre les centres de phase des antennes est égale au diamètre de chaque cornet (cornets joints). À l'émission, les deux antennes sont alimentées en phase et le faisceau de rayonnement est donc unique. En réception, les deux signaux somme et différence sont formés et utilisés comme précédemment :

$$\left(\frac{\vec{\Delta}}{\vec{\Sigma}} \right) = \frac{\vec{s}_1 - \vec{s}_2}{\vec{s}_1 + \vec{s}_2} = \tan \frac{\varphi}{2} \cong \frac{2\pi}{\lambda} \frac{\theta}{\theta_0}$$

avec θ_0 : angle d'ouverture à 3 dB

θ : angle entre la direction de l'antenne et la direction de la cible

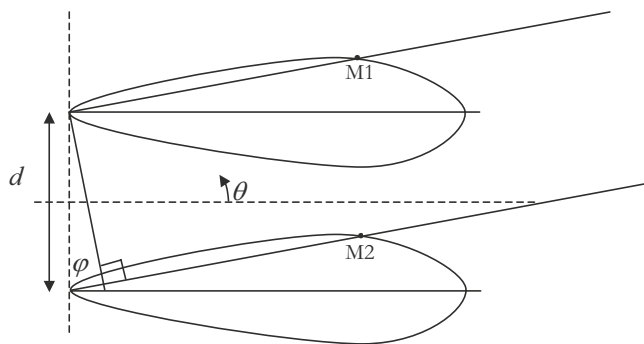


Figure 6.39 – Poursuite par monopulse de phase.

En conclusion, il faut aussi noter l'arrivée massive, depuis une dizaine d'années, des antennes réseau à commande de phase. Elles sont de deux types : actives (antennes indépendantes les unes des autres) ou passives (une seule source divisée puis déphasée à chaque antenne). Leur versatilité

en termes d'application, leur faible profil et leur rapidité de configuration pour suivre des cibles rapides sont des qualités recherchées.

6.5 Télédétection

La télédétection est par définition le domaine qui consiste à détecter à distance. Les applications dans ce domaine sont actuellement très nombreuses. Plusieurs techniques sont utilisées pour cela. Notons très globalement les techniques reposant sur des méthodes acoustiques (dont nous ne parlerons pas ici) ou électromagnétiques. Le large spectre offert par ces dernières est utilisé dans le domaine visible, infrarouge, des rayons X et des hyperfréquences. Nous nous intéresserons, dans ce chapitre à ce dernier domaine pour lequel l'observation de la terre constitue une des applications privilégiées. Les différents types d'antennes dont nous avons parlé sont choisis en fonction de leurs caractéristiques selon les applications visées que nous allons décrire.

Il n'est pas question ici d'approfondir toutes les méthodes de la télédétection mais de donner certains principes qui aident à comprendre les objectifs à atteindre et le choix du segment correspondant aux antennes.

Les systèmes de télédétection travaillent à distance. Les distances concernées sont très variables, allant de valeurs inférieures au mètre jusqu'à plusieurs milliers de kilomètres et même plus pour ce qui concerne la radioastronomie qui sera abordée dans un autre paragraphe. Les méthodes utilisées sont donc bien différentes en fonction des distances.

Le principe des méthodes en télédétection est de capter l'onde électromagnétique associée à la présence d'un objet. On distingue alors deux cas :

- le système envoie une onde électromagnétique qui interagit avec la matière constituant cet objet. Ce système capte l'onde émise en retour et l'analyse pour donner les caractéristiques de l'objet qu'on appelle souvent sa signature. Il s'agit de systèmes de type **radar**.
- l'objet lui-même émet naturellement une onde électromagnétique liée à sa température thermodynamique, qui est captée par le système. On parle alors de **radiométrie**.

Dans le premier cas, la télédétection est dite active et dans le second cas elle est dite passive.

La télédétection recouvre un ensemble de techniques d'observation de la Terre dans les domaines de la géologie, l'océanographie, la glaciologie, l'environnement, la climatologie, etc.

Les plates-formes utilisées en télédétection sont, outre les stations au sol, les ballons atmosphériques, les avions et les satellites. Ces derniers sont très utilisés et font l'objet de missions regroupant souvent plusieurs nations. À l'échelle européenne, de nombreux programmes de télédétection sont en cours. Les satellites embarquent de nombreux appareils de mesures. Les uns fonctionnent dans le domaine optique ou proche infrarouge, les autres dans le domaine des hyperfréquences. Par exemple, le satellite ENVISAT, de l'Agence spatiale européenne, embarque à son bord plusieurs radiomètres et radars, un radar à synthèse d'ouverture et un radar altimétrique.

6.5.1 Historique

Historiquement, il n'est pas très facile de dire à quand remontent les premières expériences en télédétection. On peut faire remarquer que Heinrich Hertz a démontré très vite l'intérêt des ondes électromagnétiques pour ce type d'application à la fin du XIX^e siècle en prouvant que différents objets avaient la propriété de réfléchir les ondes. Ses expériences ont été conduites sur un domaine de fréquences relativement basses (quelques centaines de MHz), mais cela n'enlève rien à la généralité du principe. Le premier brevet permettant de détecter des navires a été déposé par Hülsmeier en 1904. À partir de ce moment les recherches se sont accélérées surtout dans le domaine militaire, mais aussi dans le domaine des transports. Les radars ont utilisé des fréquences de plus en plus élevées jusqu'à atteindre les bandes hyperfréquences. Ceci a été rendu possible

grâce au développement des sources hyperfréquences. Avec la montée en fréquence la résolution devenait de plus en plus petite et l'imagerie micro-onde a pu se développer. Les premières images issues d'un radar aéroporté ont été produites durant la période de la Seconde Guerre mondiale. Ce n'est vraiment que vers les années 1960, avec la déclassification des images obtenues à l'aide de radars, que la télédétection s'est imposée pour des applications scientifiques d'observation de la Terre. Jusqu'alors les systèmes étaient aéroportés. Les performances des systèmes augmentant et l'utilisation de satellites se développant, il a été possible d'embarquer des systèmes complets dans l'espace.

Les techniques de radiométrie, quant à elles, sont issues des travaux menés en radioastronomie dans lesquels les antennes étaient pointées vers le ciel. Ce n'est que vers la fin des années 1950 que la radiométrie a commencé à se développer pour l'observation terrestre.

Actuellement, les expériences de télédétection sont aussi bien aéroportées qu'embarquées à bord de satellite. La différence essentielle tient à la distance d'observation. Pour des distances de l'ordre du kilomètre, la zone couverte est bien moins grande que pour des observations à plusieurs centaines, voire plusieurs milliers de kilomètres. Par contre dans ce dernier cas la résolution est plus faible et les appareils à embarquer doivent être plus puissants.

6.5.2 Pourquoi les micro-ondes en télédétection

Les micro-ondes se sont imposées comme support pour la télédétection pour différentes raisons :

- Les ondes correspondant à ce spectre se propagent bien à travers l'atmosphère, avec peu d'atténuation sauf autour de certaines fréquences (20, 60, 120, 180 GHz). Ces plages d'atténuation sont dues aux résonances des ondes avec les molécules de l'atmosphère. En dehors de ces zones, l'atténuation due à la traversée verticale de l'atmosphère est inférieure à 10 % jusqu'à environ 40 GHz.
- Les nuages sont traversés par les ondes sans grande perturbation. On distingue les nuages de glace à travers lesquels la perturbation des ondes est inférieure à quelques pourcents jusqu'à des fréquences de l'ordre de 30 GHz. Seul un effet dépolarisant existe. Les nuages de vapeur d'eau ont un peu plus d'effet. Jusqu'à 10 GHz l'atténuation y est inférieure à 15 %. Pour des fréquences plus élevées, l'atténuation augmente significativement pour atteindre entre 30 % et 60 % selon le type de nuages.
- La pluie a aussi un effet sur la propagation des ondes électromagnétiques. Il est difficile de donner des chiffres précis car il faut prendre en compte non seulement l'épaisseur de la zone de pluie traversée, mais aussi la taille des gouttes de pluie. Ainsi une pluie fine a moins d'effet aux fréquences usuelles. Par contre les pluies constituées de grosses gouttes affectent sensiblement les ondes au-delà de 20 GHz. On peut cependant conclure que pour des fréquences inférieures à 10 GHz, la pluie donne une atténuation faible.
- Les hyperfréquences peuvent pénétrer le couvert végétal car les longueurs d'ondes concernées sont comprises entre 1 m et 1 mm. Ceci permet, selon les fréquences utilisées, soit d'étudier les propriétés de ce couvert, soit de le traverser pour atteindre les propriétés du sol. Cependant il faut savoir que les ondes électromagnétiques sont très sensibles à l'humidité des milieux traversés. Cela peut être un avantage pour étudier, par exemple, l'état de stress des végétaux. Par contre cela peut être un inconvénient lorsqu'on cherche à atteindre les propriétés d'un sol à travers une végétation humide.
- Les fréquences concernées étant éloignées du spectre solaire ne sont pas perturbées par celui-ci. Elles permettent d'obtenir des informations complémentaires sur les propriétés des milieux observés. En effet, on attribue les interactions onde/surface à :
 - des résonances moléculaires pour les fréquences optiques et dans le proche infrarouge

- la forme de la surface ainsi qu'aux propriétés diélectriques en volume pour les hyperfréquences.

On comprend ainsi le grand intérêt des micro-ondes en télédétection qui en fait un outil indispensable, utilisable de nuit comme de jour et ce malgré la couverture nuageuse. Les techniques développées sont celles du sondage et de l'imagerie.

6.5.3 Instruments utilisant des antennes en télédétection

Nous présentons les principaux instruments utilisés en télédétection et surtout nous donnons quelques éléments en vue du choix des antennes. Certaines applications sont basées, soit sur des méthodes actives, soit sur des méthodes passives. On donnera quelques exemples de réalisations d'expériences.

Parmi les applications de la télédétection utilisant les radars, il faut distinguer plusieurs modes de fonctionnement :

- le mode simple d'émission d'onde en continu
- le mode impulsionnel selon lequel on mesure le temps d'aller-retour du signal
- le mode FMCW (*Frequency Modulated Continuous Wave*) dans lequel l'onde envoyée est modulée linéairement en fréquence.

■ Altimètres

L'application qui s'appuie sur le fonctionnement le plus simple du radar est liée à la mesure de distance. Le temps d'aller-retour d'une impulsion est lié à la distance entre le radar et la surface réfléchissante. Les altimètres utilisent différents modes de fonctionnement du radar. Ce principe est utilisé pour connaître la forme d'un terrain. Il est utilisé en océanographie pour connaître la hauteur de la mer et déterminer la forme du géoïde. Les courants marins créent des différences de hauteur de la surface qui peuvent être détectées par ce moyen. Un matériel de ce type a été embarqué à bord du satellite SEASAT avec une résolution de 10 cm. Il est ainsi possible de réaliser une cartographie de la hauteur de la mer. On peut citer les missions des satellites de l'ESA de la série ERS, situé sur des orbites à environ 800 km d'altitude. Le satellite TOPEX-POSEIDON a permis de mettre en évidence le phénomène climatique El Nino. JASON-1 (2001) et JASON-2 (2008) ont pris la relève de celui-ci. Ces deux satellites sont placés sur une orbite inclinée à 66° et à 1 336 km d'altitude. Cela leur permet une couverture presque complète de la surface marine, en repassant tous les dix jours au-dessus d'un même point. Les missions suivantes, vers 2015, réaliseront des mesures de toutes les hauteurs des surfaces aquatiques du globe, y compris des rivières, grâce à des radars fournissant une précision de quelques centimètres avec une résolution de 10 à 30 mètres.

■ Radars météorologiques

Le principe du fonctionnement du radar météorologique repose sur la mesure du signal rétrodiffusé par une zone contenant des gouttes de pluie. L'onde électromagnétique émise par l'antenne du radar interagit avec les gouttes de pluie par un phénomène de diffusion, appelée la diffusion de Mie. Ce phénomène est caractérisé par le fait que la longueur d'onde est du même ordre de grandeur que la taille des gouttes de pluie. Une partie du signal rétrodiffusé revient à l'antenne du radar et donne une information sur la taille des gouttes et leur densité.

La distance de la zone pluvieuse peut être déterminée en utilisant plusieurs radars. Grâce au déphasage entre des signaux successifs, on peut calculer sa vitesse de déplacement. Le signal rétrodiffusé vient d'une zone étendue en volume. On parle de sondeur en volume. La portée de tels systèmes est limitée par le fait que la largeur de la région sondée devient de plus en plus grande au fur et à mesure qu'elle s'éloigne du radar.

Les radars météorologiques sont utilisés à terre, pour déterminer les zones de pluies. Utilisant ce principe, on trouve la catégorie des sondeurs à bord de satellite, comme celui embarqué à bord de CloudSat (CPR, *Cloud Profiling Radar*) qui fonctionne à 94,05 GHz. Il a pour mission de déterminer le profil en eau liquide ou solide des nuages. Sa résolution est de 1,4 km transversalement, de 3,5 km le long de la trajectoire et 500 m verticalement.

Le sondeur embarqué à bord du satellite NOAA-18 mesure le contenu en vapeur d'eau en utilisant cinq fréquences dont trois sont dans la bande d'absorption de la vapeur d'eau.

■ Sondeurs ionosphériques

Le principe des sondeurs ionosphériques a été découvert très tôt après que les premières expériences sur les ondes électromagnétiques aient été faites. L'ionosphère est une région entourant la Terre à une altitude située entre 60 et 800 km d'altitude. Elle est composée de particules chargées qui interagissent très fortement avec les ondes électromagnétiques dans certaines bandes de fréquences. Ainsi les ondes décimétriques sont réfléchies complètement par l'ionosphère. Les ondes hectométriques (ondes moyennes) sont absorbées par l'ionosphère. La constitution de l'ionosphère varie considérablement selon le moment de la journée. Les ondes qui ont interagi avec l'ionosphère sont analysées pour en recueillir des informations sur sa constitution.

■ Diffusomètres

Le terme de diffusomètre (*scatterometer*) recouvre les systèmes radars qui, de façon très générale, donnent une information d'amplitude et de phase sur le signal rétrodiffusé. Il existe plusieurs types de diffusomètres. Certains sont mobiles autour d'un axe, permettant de faire varier les angles d'observation.

Le diffusomètre mesure, en fait, la surface équivalente radar $\sigma(\theta, \phi)$

Certains diffusomètres sont utilisés pour la mesure de la vitesse du vent. Par exemple celui qui est utilisé sur le satellite MetOp de l'agence spatiale européenne mesure la vitesse du vent sur les surfaces marines, avec une résolution standard de $25 \text{ km} \times 25 \text{ km}$, en utilisant une fréquence de 5,255 GHz. Le principe de la mesure de vitesse du vent repose sur l'utilisation, dans le diffusomètre Doppler, de quatre faisceaux orthogonaux pour lesquels on mesure le coefficient de rétrodiffusion.

■ Radars imageurs

Les radars imageurs permettent d'obtenir des images de la surface terrestre. Ils sont embarqués à bord d'avions ou de satellites. Leur résolution est meilleure s'ils sont situés à une altitude basse, par contre le champ observé est plus petit.

Selon la longueur d'onde, les images contiennent des informations de nature différente. Les longueurs d'onde plus élevées pénètrent bien la végétation et les sols, alors que les longueurs d'onde plus courtes sont sensibles à la forme de la surface. Les sols humides ont tendance à absorber l'onde électromagnétique. Il en résulte une réflexion faible du signal. D'une façon générale, l'onde rétrodiffusée est sensible à la valeur de la permittivité des sols. L'interprétation des images permet donc d'obtenir des informations sur les propriétés de la surface. Elles sont aussi utilisées pour les études de glaciologie. Certains radars dont les ondes pénètrent dans le sol permettent de détecter des inhomogénéités dans le sol, avec des applications à l'hydrologie, à la pétrologie et même à l'archéologie.

■ Radar à ouverture de synthèse (SAR)

La synthèse d'ouverture est de plus en plus utilisée pour les observations satellites, car elle permet une grande résolution, pour tout type d'observation : les océans, la terre et les glaces. Les satellites ERS1 et ERS-2 embarquent de tels dispositifs, fonctionnant à 5,3 GHz. À titre d'exemple, citons la résolution du radar à ouverture synthétique du satellite RadarSat2 qui peut aller jusqu'à 3 m.

■ Interférométrie radar

L'interférométrie radar est une méthode récente qui repose sur le fait que les radars enregistrent la phase des signaux. Il est ainsi envisageable de comparer les phases de signaux obtenus à des instants différents, au-dessus d'une même zone. Si les surfaces observées n'ont pas changé, les signaux ne feront apparaître aucune différence de phase. Dès qu'un détail de taille supérieure à la résolution de l'appareil est modifié, les signaux font apparaître une différence de phase qui doit être interprétée. Cette méthode demande une précision très grande sur chacun des paramètres de la mesure. On a pu montrer des mouvements de terrains, par ce procédé. L'illustration la plus célèbre est l'observation du mouvement de la faille de San Andrea. Les applications sont nombreuses dans le domaine de la géologie et de la volcanologie.

■ La radiométrie hyperfréquence

Les radiomètres observent le signal incohérent émis par les surfaces et permettent aussi de réaliser des images radiométriques. Les radiomètres hyperfréquences sont utilisés dans différents buts. Certains, comme celui qui est embarqué sur le satellite JASON-2 est destiné à mesurer le contenu en vapeur d'eau de l'atmosphère afin de corriger les mesures issues de l'altimètre radar associé. Sa résolution est de 25 km. D'autres sont utilisés pour la mesure de température, comme le MWTS (*Microwave Temperature Sounder*), à bord des satellites polaires chinois de la série Feng Yun à vocation météorologique, lancés depuis 2004 jusqu'en 2008. Le radiomètre fonctionne autour de 50 GHz, avec une précision de 0,5 K.

Certains radiomètres permettent de réaliser des images, sur le même principe que les radars imageurs. Le radiomètre AMSR-E embarqué à bord du satellite EOS-Aqua, dans le cadre de la mission d'observation globale de la Terre, TERRA, fonctionne sur six fréquences, de 6,9 à 89 GHz. L'antenne est d'un diamètre de 1,6 m et supporte les polarisations H et V. La précision sur la température observée va de 0,3 K (pour la plus basse fréquence) à 1,1 K (pour la plus élevée). Les pixels sont de l'ordre de $10 \text{ km} \times 10 \text{ km}$.

6.5.4 Télédétection active

Dans ce paragraphe, nous présenterons les principes de fonctionnement des systèmes radars utilisés uniquement pour la télédétection. Ils reposent sur la mesure du signal renvoyé par la cible observée. Dans un système radar un dispositif électronique, appelé émetteur est couplé à une antenne d'émission qui rayonne l'onde vers une cible. L'onde réfléchie est captée par une antenne de réception couplée à la partie électronique, appelée récepteur. Un dispositif permet de comparer le signal émis et le signal reçu. Les propriétés recherchées sont extraites grâce à un traitement de signal effectué dans un dispositif électronique basse fréquence aboutissant soit à un enregistrement, soit à un affichage.

On distingue les radars bi-statiques qui ont deux antennes différentes pour l'émission et la réception et les radars mono-statiques qui ne fonctionnent qu'avec une seule antenne. Dans ce dernier cas, le mode d'émission et le mode de réception fonctionnent alternativement dans le temps. Le radar doit commuter rapidement entre ces deux modes. Il est aussi muni de dispositifs d'isolation pour éviter que le signal d'émission à forte puissance n'affecte le signal reçu dont la puissance est généralement très faible.

La polarisation des ondes reçues par le radar constitue une source d'information supplémentaire. On distingue en général les deux polarisations : horizontale (H), parallèle au sol et verticale (V) dans le plan vertical et perpendiculaire à la polarisation H.

La forme du signal de l'onde émise permet de distinguer plusieurs types de radars :

- le radar pulsé
- le radar Doppler
- le radar modulé continûment en fréquence

Nous allons exposer rapidement leur principe de fonctionnement.

Le radar apparaît comme une source d'onde électromagnétique. L'onde renvoyée par la cible a une phase bien déterminée et mesurable. Le principe des mesures avec un radar repose sur la cohérence de cette onde réfléchie qui fait de la phase une grandeur portant une information.

■ Le radar pulsé

Le radar pulsé est muni d'un dispositif de synchronisation qui permet de déclencher une impulsion. Celle-ci est envoyée vers un modulateur qui génère des impulsions courtes sur une onde porteuse haute fréquence, elle-même créée par un oscillateur. L'onde est ensuite rayonnée par l'antenne. Ce dispositif, souvent monostatique, possède un dispositif de commutation et un isolateur.

L'onde reçue, après avoir été captée par l'antenne est dirigée vers le récepteur, composé d'un préamplificateur, d'un mélangeur avec un oscillateur local à la fréquence de l'onde émise. Il résulte du mélange entre l'onde reçue (f_r) et l'onde à fréquence de la porteuse (f_p), deux ondes aux fréquences somme et différence : $f_r + f_p$ et $f_r - f_p$. Cette dernière est la fréquence intermédiaire. Grâce à un dispositif de filtrage, on ne conserve que la fréquence intermédiaire qui est amplifiée. Le dispositif de synchronisation permet de commander les commutateurs entre l'état de réception et d'émission. Il permet aussi de déclencher l'impulsion une fois le signal complètement reçu. Si la largeur de l'impulsion est bien maîtrisée à l'émission, elle ne l'est pas à la réception, car elle dépend de la forme de la cible. La synchronisation permet aussi de déclencher le dispositif d'acquisition.

Le temps entre l'émission et la réception est accessible à la mesure et permet de remonter à la distance parcourue. En effet, considérant que l'onde se propage à la vitesse de la lumière c , la distance R entre le radar et la cible est donnée en fonction du temps mesuré T par :

$$R = 2cT$$

Le radar donne accès à la distance à laquelle se trouve la cible.

Notons $P_e(t)$, la puissance à l'émission. Supposons la cible très petite, l'onde réfléchie aura une forme proche de celle de la puissance émise. En introduisant un facteur α dû à l'atténuation et à la réflexion de l'onde, on obtient la forme de l'onde reçue $P_r(t)$:

$$P_r(t) \approx \alpha P_e(t - T)$$

Si la cible est étendue, le signal reçu est plus complexe. Le coefficient de réflexion α dépend de la distance R définie sur la figure 6.40.

Le coefficient $\alpha(R)$ inclut les variations de l'atténuation due à la distance et au coefficient de réflexion qui peut être différent selon le point de la surface et du diagramme de rayonnement de l'antenne qui dépend de l'angle de vue de la surface élémentaire.

La puissance reçue apparaît alors comme :

$$P_r(t) = \int P_e\left(t - \frac{2R}{c}\right) \alpha(R) dR$$

Le signal reçu résulte d'une convolution et sa forme est complexe. Tout le travail de l'analyse de ce signal repose sur des techniques de déconvolution. En particulier la largeur temporelle du signal est toujours plus grande que la largeur du signal émis.

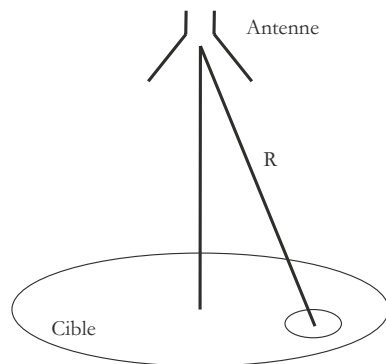


Figure 6.40 – Position de la cible par rapport au radar.

■ Le radar Doppler

L'effet Doppler repose sur la différence entre la fréquence de l'onde reçue et la fréquence de l'onde émise lorsque le signal se réfléchit sur une cible en mouvement.

Considérons une onde se propageant dans le sens des x positifs. Son équation dans le référentiel (Oxy) , lié au radar est donnée par :

$$y(x, t) = a \cos 2\pi f_r \left(t - \frac{x}{c} \right)$$

Les coordonnées (X, Y) , dans le référentiel de la cible (figure 6.41) se déplaçant à la vitesse $\vec{u}$ se déduisent par une translation des coordonnées liées au radar :

$$X = x - ut$$

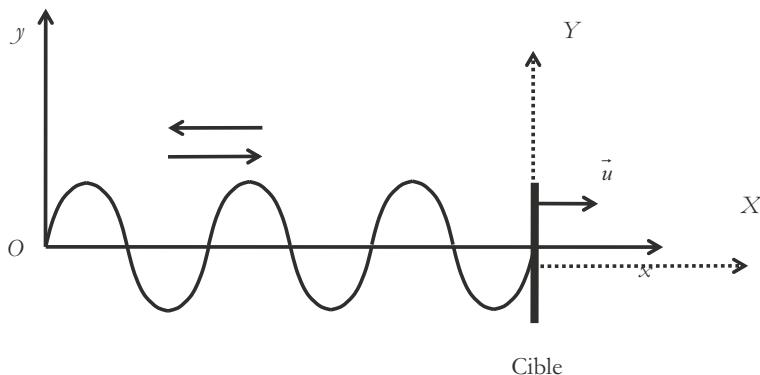


Figure 6.41 – Référentiel de la cible en mouvement.

L'équation de l'onde dans le référentiel de la cible est donc :

$$Y(X, t) = a \cos 2\pi f_r \left(t - \frac{X + ut}{c} \right)$$

Soit encore :

$$Y(X, t) = a \cos 2\pi f_r \left(\left(1 - \frac{u}{c} \right) t - \frac{X}{c} \right)$$

L'onde est réfléchi dans ce référentiel avec un coefficient de réflexion r . L'équation de l'onde réfléchi est donc :

$$Y_r(X, t) \approx ar \cos 2\pi f_r \left(\left(1 - \frac{u}{c} \right) t + \frac{X}{c} \right)$$

Cette onde a pour équation dans le référentiel du radar :

$$Y_r(X, t) \approx ar \cos 2\pi f_r \left(\left(1 - \frac{u}{c} \right) t + \frac{x - ut}{c} \right)$$

Soit encore, dans le référentiel initial :

$$y_r(x, t) \approx ar \cos 2\pi f_r \left(\left(1 - 2\frac{u}{c} \right) t + \frac{x}{c} \right)$$

La fréquence qui revient est donc :

$$f = f_r + f_D = f_r - \frac{2u}{\lambda} \quad [6.3]$$

u représente la composante de la vitesse selon la direction allant du radar à la cible.

Le signal est créé par un oscillateur à une fréquence f_r . Il passe ensuite dans un circulateur avant d'être rayonné par l'antenne. L'onde revient décalée en fréquence. Elle est captée par l'antenne, passe dans le circulateur puis dans un mélangeur qui mélange les signaux émis et réfléchis. Le résultat du mélange donne deux signaux ayant des fréquences correspondant à la somme et à la différence des fréquences mélangées. La fréquence la plus élevée est éliminée par filtrage. L'autre fréquence est la fréquence Doppler f_D . Il est donc possible d'avoir accès à la vitesse d'une cible, selon l'expression [6.3]. Dans la plupart des cas la cible est complexe et les vitesses ne sont pas constantes. Le traitement de signal permet d'extraire les informations sur les vitesses.

■ Le radar modulé continûment en fréquence

Le système radar le plus commun est de ce type. Le signal est modulé en fréquence selon une rampe linéaire (figure 6.42). La période de la modulation est T .

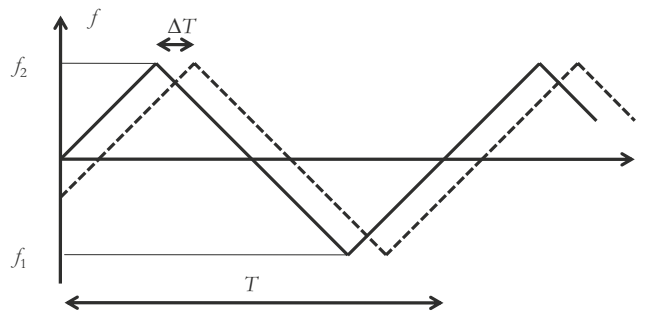


Figure 6.42 – Modulation linéaire en fréquence.

Le trait plein représente la fréquence du signal émis. Le signal reçu est décalé du temps d'aller-retour ΔT de l'onde entre le radar et la cible. Le trait pointillé représente la fréquence du signal de retour. Lorsque la pente de la rampe est positive, la fréquence reçue f_R est plus petite que la fréquence émise f_E . C'est l'inverse lorsque la pente de la rampe est négative.

Un générateur de signaux triangulaires module une porteuse. Le signal est rayonné par une antenne d'émission. Au retour le signal est reçu par une antenne, souvent différente de l'antenne d'émission. Après une pré-amplification, le signal, de fréquence f_R , est envoyé sur l'entrée d'un mélangeur. L'autre entrée du mélangeur reçoit, au même instant, un signal proportionnel au signal qui est en train d'être émis, c'est-à-dire à la fréquence f_E . À l'issue du mélange on ne garde que la fréquence différence $\Delta f = f_E - f_R$, par filtrage.

La bande de fréquence de fonctionnement est par définition :

$$B = f_2 - f_1$$

La pente p de la rampe est donnée par :

$$p = \frac{B}{T/2} \quad \text{d'où :} \quad \Delta f = \frac{2B}{T} \Delta T$$

On obtient la relation entre la distance R de la cible au radar et la différence de fréquence mesurée :

$$\Delta f = \frac{4BR}{cT}$$

■ Le radar à ouverture latérale

Les radars à ouverture latérale sont très utilisés car ils permettent un large champ d'observation. Ce sont des radars pulsés. L'antenne est dépointée sur le côté. Elle est conçue de façon à avoir une ouverture plus grande dans le plan perpendiculaire à l'axe de propagation (rayonnement vertical) que dans le plan parallèle au déplacement (rayonnement horizontal).

Le rayonnement dans le plan transverse a une ouverture β (figure 6.43). L'impulsion envoyée par le radar, de largeur temporelle ΔT , rencontre le sol, est réfléchie, puis revient au radar. L'onde touche tout d'abord le sol à l'endroit le plus proche (A), puis un peu plus loin sur la surface vers B. Le point B est touché en dernier. Donc en enregistrant le signal en fonction du temps, on a des informations sur les points situés entre A et B.

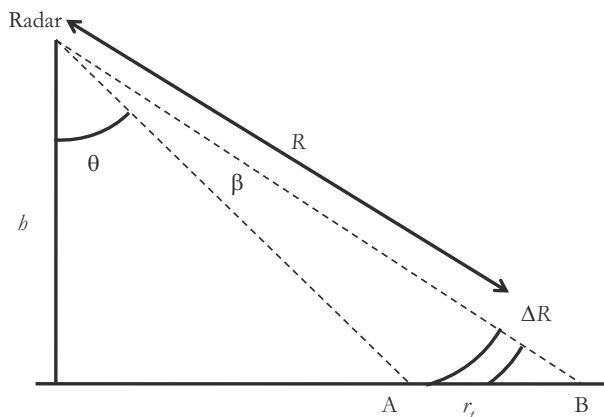


Figure 6.43 – Rayonnement du radar dans le plan transverse.

Si T représente le temps parcouru entre l'émission et la réception, la distance R est donnée par :

$$R = \frac{cT}{2}$$

La largeur de l'impulsion ΔT correspond donc à une distance parcourue ΔR :

$$\Delta R = \frac{c\Delta T}{2}$$

On ne peut résoudre des objets au sol que si leur distance est supérieure à r_t , projection au sol de ΔR :

$$r_t = \frac{c\Delta t}{2 \sin \theta}$$

Le rayonnement dans le plan parallèle au déplacement est beaucoup plus fin angulairement que dans le plan transverse. Si la dimension longitudinale de l'antenne est L , on sait que son ouverture angulaire γ est fonction de la dimension de l'antenne et de la longueur d'onde λ :

$$\gamma \approx \frac{\lambda}{L}$$

La largeur r_l balayée sera :

$$r_l = \frac{\lambda}{L} R$$

Prenons le cas d'un avion volant à 6 km d'altitude, embarquant un radar muni d'une antenne inclinée de 45° , de dimension longitudinale égale à 2 m, envoyant des impulsions de $0,1 \mu s$,

avec une porteuse de longueur d'onde 2 cm. Les résolutions le long de l'axe de propagation et transversalement prennent les valeurs :

résolution longitudinale : $r_l = 84$ m, résolution transversale : $r_t = 21$ m.

On constate d'après ce qui a été dit que la résolution longitudinale ne dépend que de la taille de l'antenne dans la direction longitudinale, alors que la résolution transverse dépend de la largeur temporelle de l'impulsion.

La dimension transversale l_t de l'antenne donne l'ouverture β de l'antenne qui détermine la trace du radar. L'ouverture est donnée par :

$$\beta = \frac{\lambda}{l_t}$$

La trace a pour dimension :

$$AB \approx \frac{\beta R}{\cos \theta}$$

Avec les valeurs numériques déjà utilisées, pour une antenne de dimension transverse $l_t = 0,2$ m, on trouve : $AB = 1\,200$ m. Dans ce cas la fauchée a une taille de $84 \text{ m} \times 1\,200 \text{ m}$.

Le signal reçu par le radar est enregistré temporellement. Il faut que l'envoi des impulsions soit synchronisé avec la vitesse du radar afin de ne pas perdre d'informations au passage au-dessus du sol. Le signal rétrodiffusé est ensuite analysé afin d'en extraire les propriétés physiques du sol : humidité, nature du sol, végétation...

■ Le radar imageur

Le radar imageur est conçu sur le principe qui vient d'être exposé. Le signal est enregistré ligne par ligne. Chaque ligne correspond à l'enregistrement des informations contenues durant la largeur de l'impulsion, c'est-à-dire celles qui concernent l'ouverture transversale. La ligne reflète donc les propriétés d'une région perpendiculaire à la vitesse du radar, de largeur r_l et dont la longueur est la trace transversale de l'ouverture de l'antenne.

Dans ce système, le signal haute fréquence est démodulé. La puissance de modulation, basse fréquence, extraite est envoyée sur un dispositif qui permet un affichage en niveaux de gris proportionnellement à la valeur de la puissance. Une ligne d'images est donc affichée durant la durée de l'impulsion. Chaque pixel a une dimension $r_l \times r_t$. La ligne a une largeur r_l . L'impulsion suivante permet d'afficher une autre ligne, etc. On obtient donc une image en deux dimensions qui constitue une image radar.

Le dispositif de synchronisation est un organe fondamental de ce type de radar, puisque la vitesse de l'avion doit imposer l'envoi des impulsions au bon moment pour couvrir tout le champ et que le dispositif d'affichage doit passer d'une ligne à l'autre lorsque les informations d'une ligne sont complètes.

On conçoit que si, l'antenne est fortement inclinée ($> 10^\circ$), les images seront déformées en raison de l'effet du cosinus dans la projection. Certains dispositifs d'affichage tiennent compte de cet effet pour redresser l'image. Lors des traitements ultérieurs, cet effet est toujours corrigé. Selon l'inclinaison, le relief peut créer un effet d'ombre. Certains pixels ne recevant alors aucune information.

Les antennes des radars imageurs émettent et reçoivent des signaux dans les deux polarisations H et V. La possibilité existe donc de réaliser des images de type HH, HV, VH, VV, correspondant à la forme de la polarisation de réception en fonction de la polarisation d'émission.

Les images radar sont ensuite traitées afin d'en extraire des informations, soit de relief, soit physiques, après correction à l'aide d'un modèle numérique de terrain. Les images obtenues ont une grande dynamique. Elles présentent un chatoiement (speckle) constituant un bruit qui doit être éliminé par traitement. Ce phénomène vient de la taille des objets du même ordre de grandeur que la longueur d'onde.

Le principe du radar imageur vient d'être exposé en prenant comme exemple le radar pulsé. D'autres modes de fonctionnement peuvent être utilisés pour réaliser des images, comme le radar Doppler.

■ Le radar à synthèse d'ouverture

Le radar à synthèse d'ouverture appelé aussi SAR (*Synthetic Aperture Radar*) repose sur le principe de la réalisation d'une antenne artificielle de dimension plus grande que l'antenne réelle. On obtient alors une résolution plus grande dans la direction de propagation du radar. Pour cela, les différentes positions que prend l'antenne réelle au cours du temps sont utilisées. On enregistre au cours du temps les signaux reçus avec leur phase, puis ils sont analysés par un traitement de signal approprié, complexe et nécessitant une connaissance exacte de la position du radar.

Ce principe a bien sûr des limites. L'antenne virtuelle reconstituée ne peut pas prendre une taille très grande. Elle est limitée par l'angle de vue de l'antenne et par le bruit du traitement lors de l'addition de signaux.

6.5.5 Radiométrie hyperfréquence

Dans ce paragraphe, nous donnerons les principes de la radiométrie et les caractéristiques des antennes utilisées.

■ Quelques grandeurs radiométriques

La radiométrie est le domaine d'étude du rayonnement électromagnétique dans la partie du spectre utilisant des antennes, c'est-à-dire les radiofréquences et les hyperfréquences. Les grandeurs radiométriques sont analogues à celles de la photométrie dans le domaine optique. L'usage leur a donné des noms différents.

Par opposition à la télédétection active, la télédétection passive utilise un rayonnement incohérent.

□ La brillance

Considérons une antenne d'émission dont l'aire effective est A_e et une antenne de réception d'aire effective A_r . Elles sont placées face à face, de façon à ce que leurs axes soient communs, à une distance R suffisamment grande pour qu'elles soient en champ lointain l'une de l'autre. Elles sont correctement orientées de façon à être sensibles à la même polarisation (figure 6.44).

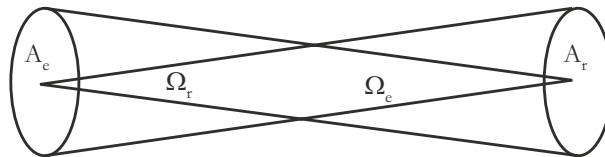


Figure 6.44 – Positions des antennes d'émission et de réception.

Soit Ω_e l'angle solide sous lequel on voit l'antenne d'émission depuis l'antenne de réception et Ω_r l'angle solide sous lequel on voit l'antenne de réception depuis l'antenne d'émission. Pour le raisonnement, on suppose les antennes telles qu'on puisse considérer la puissance comme constante sur leur surface. Alors la puissance reçue P_r sur l'antenne de réception est fonction de son aire et de la densité surfacique de puissance S_e émise :

$$P_r = S_e A_r$$

Exprimons la densité surfacique de rayonnement en fonction de la distance :

$$P_r = \frac{F_e}{r^2} A_r$$

Supposons maintenant que l'antenne d'émission soit une surface émettrice de rayonnement incohérent comme le ciel ou le sol. La puissance reçue par l'antenne de réception se caractérise de la même façon. Le terme F_e qui est une densité de puissance par unité d'angle solide est par définition l'intensité de la source. Cette intensité rapportée à la surface de la source est par définition la brillance B qui s'exprime en $\text{W.sr}^{-1}.\text{m}^{-2}$:

$$B = \frac{F_e}{A_e}$$

On déduit :

$$P_r = B \frac{A_r A_e}{r^2}$$

Introduisons l'angle solide sous lequel on voit la surface émettrice :

$$P_r = BA_r \Omega_e$$

Cette expression est établie en supposant que la surface émettrice est équivalente à une antenne, dont on a moyenné les propriétés de rayonnement sur sa surface. Considérant maintenant le même raisonnement avec des surfaces élémentaires d'émission, la valeur de la puissance élémentaire reçue se généralise. On tient alors compte du fait que la surface émettrice est grande et que la puissance émise présente des variations angulaires. La puissance élémentaire reçue par l'antenne de réception, dont on considère l'aire de réception suffisamment petite, se met alors sous la forme :

$$dP_r = A_r B(\theta, \phi) F_n(\theta, \phi) d\Omega \quad [6.4]$$

$B(\theta, \phi)$ est la brillance de la surface. Elle dépend de l'angle d'observation. Dans cette dernière expression, la fonction de répartition $F_n(\theta, \phi)$ prend en compte la réponse angulaire de l'antenne de réception.

Généralement le rayonnement émis par une surface dépend de la fréquence. La brillance spectrale $B_f(\theta, \phi)$ représente la brillance par unité de fréquence.

La puissance reçue par l'antenne est finalement :

$$P_r = \frac{1}{2} A_r \int_f^{f+\Delta f} \iint_{4\pi} B_f(\theta, \phi) F_n(\theta, \phi) d\Omega df \quad [6.5]$$

Le facteur $\frac{1}{2}$ placé devant les intégrales vient du fait que le rayonnement reçu étant incohérent, il présente toutes les polarisations. L'antenne n'est sensible qu'à une polarisation et donc ne reçoit en moyenne que la moitié de la puissance.

□ Rayonnement du corps noir

Le principe fondamental de l'étude qui suit, repose sur le fait qu'un objet porté à une température thermodynamique non nulle émet un rayonnement électromagnétique incohérent fonction de sa température.

Nous allons étudier le rayonnement des différentes surfaces émettrices rencontrées dans la réalité par rapport à la surface de référence d'un corps noir.

Le corps noir est défini comme un objet en équilibre thermodynamique à une température T , qui absorbe toute la puissance qu'il reçoit. Une bonne représentation du corps noir est celle d'une enceinte fermée en équilibre thermodynamique dans la paroi de laquelle est pratiquée une ouverture très petite qui permet au rayonnement d'entrer, mais aussi de s'échapper. Toute puissance entrant par cette ouverture est piégée dans l'enceinte. En retour cette enceinte émet de la puissance. À l'équilibre thermodynamique, toute la puissance absorbée est aussi émise.

La brillance spectrale du corps noir est calculable en introduisant à la fois les principes de physique statistique et ceux de physique quantique. La démonstration aboutit à la loi de Planck dans

laquelle la brillance spectrale s'exprime en fonction de la fréquence f (Hz) et de la température thermodynamique T (K) :

$$B_f = \frac{2hf^3}{c^2} \frac{1}{e^{hf/kT} - 1} \quad [6.6]$$

La vitesse de la lumière est : $c = 3.10^8 \text{ m.s}^{-1}$

La constante de Planck est : $h = 6,62.10^{-34} \text{ J.s}$

La constante de Boltzmann est : $k = 1,38.10^{-23} \text{ J.K}^{-1}$

On constate que la brillance spectrale du corps noir ne présente pas de dépendance angulaire. En effet le rayonnement du corps noir est isotrope.

Les propriétés du corps noir sont aussi bien connues des opticiens. L'équivalent de la brillance en optique est la luminance.

La brillance spectrale est représentée sur la figure 6.45 en prenant comme paramètre la température. Elle passe par un maximum par rapport à la fréquence qui est fonction de la température selon une loi de température à la puissance trois. On connaît mieux cette loi en optique sous le nom de loi de Wien qui fait intervenir un terme équivalent à la brillance par longueur d'onde, appelée émittance spectrale dont le maximum, en fonction de la longueur d'onde, varie selon une loi de température à la puissance cinq.

Les fréquences jusqu'à 10^{10} Hz sont des fréquences radio ou hyperfréquences qui concernent les antennes. Au-delà, on trouve les infrarouges puis les fréquences optiques et plus loin, après les rayons ultraviolets, les rayons X. La courbe correspondant à une température de 6 000 K représente approximativement le rayonnement du soleil. Le maximum de brillance est obtenu pour les longueurs d'onde correspondant au jaune. La courbe à 300 K correspond à la brillance de surfaces portées à des températures usuelles. Le rayonnement émis est, dans ce cas, relativement faible. Le principe de la radiométrie est de capter ce rayonnement.

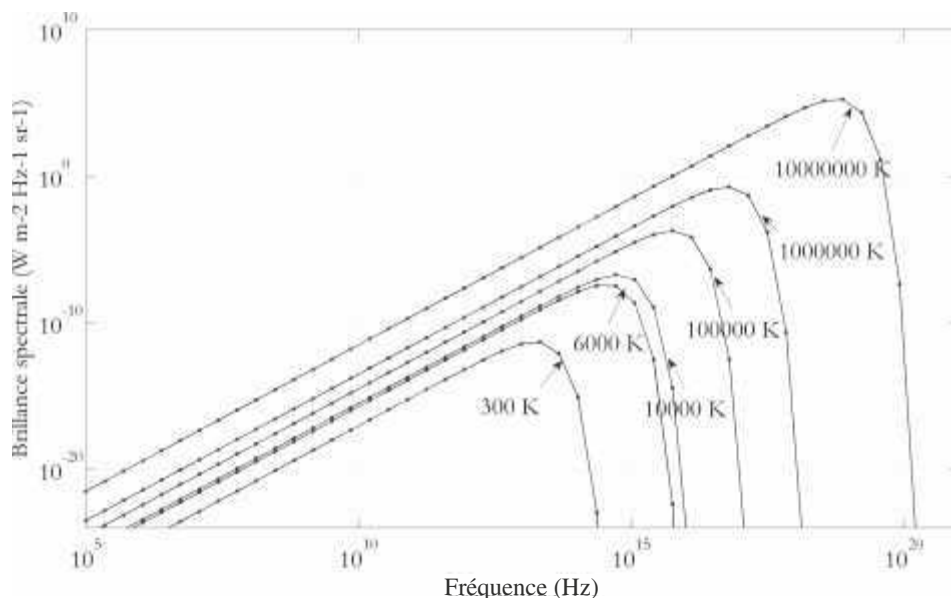


Figure 6.45 – Brillance spectrale du corps noir.

Pour les fréquences les plus basses, la variation de la brillance spectrale est linéaire. En effet pour :

$$\frac{hf}{kT} \ll 1, \text{ l'approximation suivante } e^{\frac{hf}{kT}} - 1 \approx \frac{hf}{kT} \text{ est valable.}$$

On en déduit :

$$B_f = 2f^2 \frac{kT}{c^2}$$

soit encore :

$$B_f = \frac{2kT}{\lambda^2} \quad [6.7]$$

Cette loi est appelée la loi de Rayleigh-Jeans. À la température de 300 K, l'erreur entre la loi de Planck et son approximation est inférieure à 1 % jusqu'à 120 GHz. La linéarité entre la brillance et la température thermodynamique va permettre de définir la température de brillance.

□ Relation entre la brillance et la température

Comparons deux expériences reposant sur les propriétés du corps noir. Dans un cas le corps noir, porté à la température T , contient une antenne qui recueille le rayonnement existant à l'intérieur et transmet la puissance à une ligne (figure 6.46). Dans l'autre cas l'antenne est remplacée par une résistance aux bornes de laquelle la tension est mesurée.

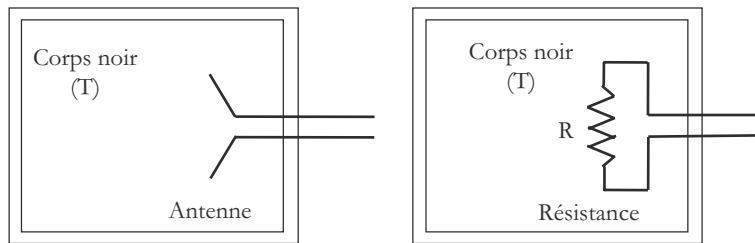


Figure 6.46 – Comparaison du comportement d'une antenne et d'une résistance placées dans une enceinte assimilable à un corps noir.

L'antenne reçoit le rayonnement existant dans l'enceinte. Supposons que l'on reste dans le domaine de température et de fréquence pour lesquelles la loi de Rayleigh-Jeans reste valable. La puissance reçue s'exprime par :

$$P_r = \frac{1}{2} A_r \int_f^{f+\Delta f} \iint_{4\pi} \frac{2kT}{\lambda^2} F_n(\theta, \phi) d\Omega df$$

Après simplification :

$$P_r = A_r \frac{kT}{\lambda^2} \Delta f \Omega_p$$

Rappelons que l'angle solide Ω_p d'une antenne (défini au chapitre 4) est lié à l'aire effective de l'antenne de réception par la relation :

$$\Omega_p = \frac{\lambda^2}{A_r}$$

La puissance reçue est donc :

$$P_r = kT \Delta f \quad [6.8]$$

La puissance reçue par l'antenne est proportionnelle à la température thermodynamique.

On retrouve le même comportement pour une résistance. Le théorème de Nyquist montre que la puissance du bruit reçue aux bornes de la résistance dépend uniquement de la température thermodynamique de celle-ci et est donnée par la même formule. Cette remarque conduit à justifier la définition de la résistance de rayonnement d'une antenne.

□ Rayonnement d'un corps non noir

Le corps noir est un corps idéal. Certains objets peuvent être assimilés à des corps noirs, la Terre globalement, par exemple. Plus les objets sont réfléchissants, moins ils absorbent de puissance et plus ils s'écartent du comportement d'un corps noir. On les dit non noirs ou bien gris. Pour caractériser leur émission, on conserve la notion de brillance.

La brillance d'une surface grise est alors liée à la température de brillance de la surface $T_B(\theta, \phi)$ par analogie à celle d'un corps noir par :

$$B(\theta, \phi) = \frac{2kT_B(\theta, \phi)}{\lambda^2} \Delta f$$

Remarquons que la brillance, de même que la température de brillance, dépendent de l'angle sous lequel on voit la surface, contrairement au corps noir.

L'émissivité d'une surface est définie par le rapport entre la brillance de la surface et la brillance du corps noir, portés à la même température thermodynamique :

$$e(\theta, \phi) = \frac{B(\theta, \phi)}{B_{\text{corps noir}}} \quad [6.9]$$

C'est une mesure du rayonnement propre émis par la surface, à une température thermodynamique donnée.

L'émissivité s'écrit donc comme le rapport entre la température de brillance et la température thermodynamique. Elle est inférieure ou égale à l'unité.

$$e(\theta, \phi) = \frac{T_B(\theta, \phi)}{T} \leq 1 \quad [6.10]$$

Pour des fréquences d'observation de l'ordre du GHz, la température de brillance d'un sol sec est en général inférieure d'une dizaine de pourcents à sa température thermodynamique. Lorsque le sol possède une humidité d'environ 20 %, sa température de brillance représente environ 70 % de sa température thermodynamique. Elle n'est plus que de 50 % de sa température thermodynamique pour un sol gorgé d'eau à 50 % de son volume.

La température de brillance des surfaces océaniques est de l'ordre d'une centaine de Kelvin, pour des fréquences de l'ordre du giga-hertz. Elle augmente pour des fréquences plus élevées : entre 160 K et 180 K pour 34 GHz. L'étude de la température de brillance des océans permet d'accéder à des mesures de salinité, puisque celle-ci modifie la conductivité de l'eau de mer et donc sa réflectivité.

□ Température apparente

D'après ce qui vient d'être expliqué, on conçoit qu'une puissance reçue soit exprimée comme une température. Cependant cette notion est complexe car on a vu que la température de brillance a une variation angulaire. De ce fait le rayonnement que reçoit l'antenne peut avoir plusieurs sources qui s'ajouteront (figure 6.47). Pour une antenne satellite d'observation de la Terre, par exemple, le rayonnement reçu est celui émis par le sol. Celui-ci n'étant pas homogène, il faut tenir compte des émissivités ou des températures de brillance des différentes surfaces le constituant, qui sont vues par l'antenne. D'autres sources existent comme la diffusion par l'atmosphère correspondant au rayonnement de celle-ci vers le haut (T_{atm}), la diffusion par le sol des divers rayonnements qu'il reçoit en provenance de l'atmosphère, du soleil, etc. T_{AS} est la température de brillance du rayonnement provenant des sources éclairant le sol et T_{dif} correspond à la puissance diffusée par

le sol. Tous ces rayonnements perturbent la mesure de la température de brillance du sol T_S . Ils doivent être évalués afin de faire une mesure précise.

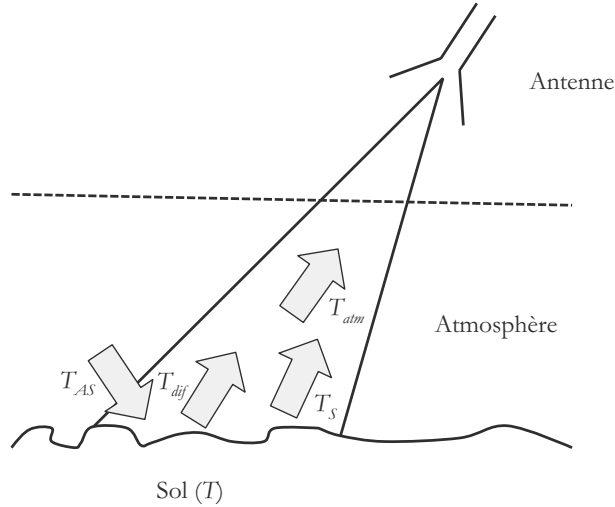


Figure 6.47 – Différentes sources vues par l'antenne d'un satellite.

De nombreux exemples existent en télédétection et les radiomètres mesurent les températures des surfaces aquatiques aussi bien que terrestres. Ces températures permettent d'accéder à la mesure de paramètres physiques.

Soit $B_{Total}(\theta, \phi)$, la puissance reçue au niveau de l'antenne, obtenue en sommant tous les rayonnements. La température apparente $T_{AP}(\theta, \phi)$ est définie par :

$$B_{Total}(\theta, \phi) = \frac{2kT_{AP}(\theta, \phi)}{\lambda^2} \Delta f$$

En général la température de brillance de la surface observée est différente de la température apparente, à moins qu'il n'y ait aucune atténuation, ni aucune source parasite de rayonnement.

La température apparente ne dépend pas de l'antenne utilisée. Elle dépend cependant de sa position.

□ Température d'antenne

La puissance reçue par l'antenne est constituée des puissances élémentaires correspondant à chaque direction partant de la surface observée vers l'antenne. Cependant l'antenne agit sélectivement selon chaque direction, selon sa fonction caractéristique de rayonnement. L'expression de la puissance reçue est donc :

$$P_r = \frac{1}{2} A_r \int_f^{f+\Delta f} \iint_{4\pi} \frac{2k}{\lambda^2} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega df$$

Définissons la température d'antenne par :

$$T_A = \frac{\iint_{4\pi} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega}{\iint_{4\pi} F_n(\theta, \phi) d\Omega} \quad [6.11]$$

La température d'antenne apparaît comme la valeur moyennée de la température apparente, pondérée par la fonction caractéristique de rayonnement.

La puissance reçue par l'antenne s'exprime simplement par :

$$P_r = kT_A \Delta f$$

Prenons l'exemple d'une antenne d'angle solide Ω_p visant une surface de température apparente T_{app} . Supposons la surface telle que l'angle solide sous laquelle on la voit Ω_s soit plus petit que Ω_p . Supposons que la température de cette surface soit uniforme, le calcul conduit à la température d'antenne :

$$T_A = \frac{\Omega_s}{\Omega_p} T_{app}$$

Ce calcul est vrai à condition qu'il n'y ait aucune autre émission de puissance en dehors de la surface observée.

□ Qualité d'une antenne radiométrique

L'antenne a pour but de capter la puissance de façon à obtenir une mesure physique la plus précise possible. Sur l'exemple qui vient d'être évoqué, la température d'antenne est différente de la température apparente parce que le rapport des angles solides est différent de 1.

Si l'angle solide de l'antenne est plus petit que l'angle solide sous lequel on voit la surface, la température d'antenne est égale à la température apparente.

Finalement, pour un objet étendu de température apparente variable, l'antenne qui donne le plus de précision est celle qui a l'ouverture la plus fine. L'ouverture de l'antenne est le facteur limitant à la précision si la température apparente varie.

Cette première approche intuitive nous permet de saisir les qualités de l'antenne d'un radiomètre. Avant de quantifier cette étude il faut remarquer que le rayonnement capté par une antenne correspond à de la puissance entrant selon le lobe principal, mais aussi par les lobes secondaires.

Ainsi la puissance provenant d'une source assez puissante dans la direction des lobes secondaires peut modifier considérablement la mesure en s'ajoutant à la puissance provenant de la surface visée dans la direction du lobe principal. Une autre qualité d'une antenne radiométrique est donc de posséder des lobes secondaires très bas.

Revenons à la température d'antenne. Décomposons son expression [6.11] sous forme d'une intégrale sur le lobe principal ($l.p.$) plus une intégrale sur le reste de l'espace, incluant les lobes secondaires ($4\pi - l.p.$).

$$T_A = \frac{\iint_{l.p.} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega}{\iint_{4\pi} F_n(\theta, \phi) d\Omega} + \frac{\iint_{4\pi - l.p.} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega}{\iint_{4\pi} F_n(\theta, \phi) d\Omega} \quad [6.12]$$

Définissons la température apparente effective ($T_{l.p.}$) dans le lobe principal par :

$$T_{l.p.} = \frac{\iint_{l.p.} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega}{\iint_{l.p.} F_n(\theta, \phi) d\Omega}$$

Rappelons que l'efficacité dans le lobe principal a été définie dans le chapitre 4 par :

$$\eta_M = \frac{\iint_{l.p.} F_n(\theta, \phi) d\Omega}{\iint_{4\pi} F_n(\theta, \phi) d\Omega}$$

Définissons de la même façon la température apparente effective ($T_{l.s.}$) dans les lobes secondaires par :

$$T_{l.s.} = \frac{\iint_{4\pi-l.p.} T_{app}(\theta, \phi) F_n(\theta, \phi) d\Omega}{\iint_{4\pi-l.p.} F_n(\theta, \phi) d\Omega}$$

La température d'antenne apparaît reliée à ces différentes grandeurs par :

$$T_A = \eta_M T_{l.p.} + (1 - \eta_M) T_{s.l.} \quad [6.13]$$

Soit encore :

$$T_{l.p.} = \frac{T_A - (1 - \eta_M) T_{s.l.}}{\eta_M}$$

Cette expression montre l'erreur théorique que l'on fait lorsqu'on assimile la température effective dans le lobe principal à la température d'antenne. L'erreur est d'autant plus faible que l'efficacité dans le lobe principal est proche de l'unité. Cela revient à dire que les lobes secondaires doivent être très faibles.

Afin d'affiner cette étude, nous devons considérer les pertes de l'antenne, liées à son efficacité η_l qui correspond au rendement de l'antenne. Si T_A est la température à l'entrée de l'antenne, la température vue du récepteur T'_A sera plus faible :

$$T'_A = \eta_l T_A$$

Cette expression de la température vue du récepteur est incomplète, car l'antenne possède une température thermodynamique T_0 . Elle va donc émettre aussi son propre rayonnement qui sera reçu par le récepteur. Cela constitue la puissance de bruit de l'instrument qui arrive au récepteur et s'exprime sous la forme de la température de bruit T_B :

$$T_B = (1 - \eta_l) T_0$$

L'expression correcte de la température vue du récepteur est donc :

$$T'_A = \eta_l T_A + (1 - \eta_l) T_0$$

En utilisant [6.13], cette température s'exprime, par rapport aux températures à mesurer par :

$$T'_A = \eta_l \eta_M T_{l.p.} + \eta_l (1 - \eta_M) T_{s.l.} + (1 - \eta_l) T_0$$

On constate qu'en plus d'avoir une bonne efficacité dans le lobe principal, une antenne pour radiomètre doit avoir une efficacité d'antenne très grande. La puissance du bruit dépendant de la température physique de l'antenne, certaines antennes sont refroidies pour améliorer la précision de mesure.

□ Imagerie radiométrique

De la même façon que les radars, les radiomètres peuvent être utilisés pour réaliser des images. Il suffit alors d'orienter leur antenne. Deux cas sont possibles :

- la variation de l'orientation du faisceau transversalement au mouvement du radiomètre s'effectue de façon mécanique, soit par un mouvement de l'antenne, soit par un mouvement d'un réflecteur intermédiaire. Le balayage dans la direction perpendiculaire est obtenu grâce au mouvement de la plate-forme.
- l'autre solution est obtenue par balayage électronique.

Le système d'imagerie consiste ensuite à capter les informations et à les restituer sous forme d'images.

6.6 Radioastronomie

La radioastronomie est le domaine d'observation des corps célestes dans les bandes de fréquences micro-ondes et millimétriques. Différents objets émettent à ces fréquences : quasar, supernovae, fond diffus cosmologique

L'atmosphère terrestre est transparente pour des fréquences assez basses, mais la vapeur d'eau absorbe les ondes pour des fréquences élevées. En particulier aux fréquences millimétriques, cet effet est extrêmement gênant, d'autant plus que la puissance du rayonnement est très faible, les objets observés étant très loin. C'est pourquoi, les radiotélescopes sont construits en altitude, dans des régions sèches. Une autre possibilité consiste à envoyer des radiotélescopes dans l'espace.

Une autre caractéristique des radiotélescopes est d'avoir une extension très grande. Pour situer précisément des objets lointains, une grande surface de captation est nécessaire, impliquant une ouverture angulaire très faible.

Les antennes sont souvent composées de deux réflecteurs. Le récepteur se trouve juste derrière le réflecteur primaire qui est percé de façon à laisser passer le rayonnement qui vient se focaliser (montage Cassegrain ou grégorien). Les antennes peuvent être associées en réseau.

Les qualités d'une antenne pour la radioastronomie sont :

- sa sensibilité
- son pouvoir de résolution

La sensibilité est obtenue en augmentant la taille de l'antenne et la qualité du récepteur. L'électronique est refroidie pour diminuer la sensibilité au bruit.

Le pouvoir de résolution est la capacité d'un télescope à distinguer deux objets. Il est proportionnel à la longueur d'onde et inversement proportionnel au diamètre de l'antenne. Donc une surface maximale est aussi recherchée pour ce point.

Certains systèmes augmentent artificiellement la surface du radiotélescope en recevant les signaux en phase sur plusieurs antennes identiques. Les principes d'interférométrie sont alors utilisés pour augmenter le pouvoir séparateur et la sensibilité du système. C'est le cas du projet e-VLBI (*Very Long Baseline Interferometer*) qui recueillera, à partir de 2009, les données de seize radiotélescopes dans le monde. Le diamètre du télescope équivalent sera de 11 000 km. Tous les observatoires seront calés sur la même longueur d'onde et synchronisés sur une horloge atomique.

La construction des radiotélescopes est souvent le résultat de coopérations internationales.

Nous allons préciser les caractéristiques recherchées des antennes des radiotélescopes en donnant quelques exemples des plus célèbres.

6.6.1 Le radiotélescope de Nançay

Les ondes venant de l'espace se réfléchissent sur une surface métallique grillagée, dont la maille est de 1,25 cm. Ce réflecteur, de taille totale égale à 200 m × 40 m est inclinable verticalement. Le rayonnement est ensuite envoyé sur un miroir, grillagé, de même constitution, fixe de 300 m × 35 m, situé à 460 m du premier et de rayon 560 m. Il concentre les rayons au foyer situé entre les deux miroirs, où se trouvent les antennes. Celles ci sont mobiles selon deux directions, horizontale et verticale, de façon à suivre la source. Les antennes sont des cornets qui reçoivent des fréquences de 1,4 GHz, 1,6 GHz, 3,3 GHz. C'est le troisième plus grand radiotélescope du monde. Il a été achevé en 1965.

Il a permis d'étudier certaines raies spectrales de radicaux fondamentaux pour la matière inter-stellaire, comme H, OH, CH. De nombreuses observations de corps célestes ont été obtenues. La première observation d'une comète dans le domaine des ondes radio a eu lieu en 1973.

6.6.2 Le radiotélescope d'Arecibo

Il est constitué d'un réflecteur métallique formé d'une portion de sphère horizontale, de 305 m de diamètre, dont la concavité est dirigée vers le ciel. Sa distance focale est de 132 m. Il est construit sur le principe du montage grégorien comprenant le miroir primaire sphérique, renvoyant les ondes vers un miroir secondaire concave qui focalise le rayonnement vers un miroir tertiaire. Ces deux derniers miroirs sont maintenus par un dispositif de câbles permettant de positionner l'ensemble très précisément au-dessus du miroir primaire.

Les antennes sont reliées à un récepteur refroidi à l'hélium liquide. Elles fonctionnent de 10 MHz à 10 GHz.

Il a été construit en 1963. Depuis sa construction, de nombreuses découvertes ont été faites, dont celles de certains pulsars. Une carte détaillée de la distribution des galaxies dans l'Univers a pu être établie. Des données sur l'atmosphère et l'ionosphère terrestres ont été mesurées. Des astéroïdes ont été observés.

6.6.3 Le radiotélescope de Green Bank

Ce radio télescope a la particularité de présenter une surface active pour son réflecteur principal. Celui-ci est composé de plus de 2 000 panneaux ajustables. Des actuateurs agissent sur leurs positions afin de maintenir la forme de la surface réfléchissante. La surface totale est de 7 854 m². Ses antennes cornets circulaires permettent d'analyser les signaux de 300 MHz à 50 GHz. Une extension est prévue à 100 GHz. C'est le plus grand radio télescope orientable du monde. Il est situé en Virginie, dans une zone isolée de toute pollution électromagnétique.

6.6.4 Le radiotélescope ALMA

(Atacama Large Millimeter-submillimeter Array)

Prévu pour fonctionner à partir de 2012, c'est un observatoire pour les longueurs d'ondes allant de 0,3 à 9,6 mm. Il permettra d'étudier les rayonnements émis par les objets froids, entre 3 K et 100 K, comme des nuages froids, des étoiles en formation ou les premières galaxies, éloignées d'une dizaine de milliards d'années. Le rayonnement micro-ondes n'étant pas absorbé par ces corps, il peut nous transmettre des informations sur leur constitution.

Situé à 5 000 m d'altitude, c'est un réseau géant de radiotélescopes qui comprendra au plus 80 télescopes constitués d'antennes de 12 m de diamètre. Les positions de ces antennes seront variables permettant des espacements allant de 150 m à 18 km. On attend une résolution très fine, de 0,005" pour les plus faibles longueurs d'onde.

6.6.5 Le satellite Planck

Le satellite Planck est conçu pour embarquer de nombreux appareils de mesures qui permettront d'apporter des éléments aux théories sur la naissance de l'Univers. Les bandes d'observation sont comprises entre 30 GHz et 850 GHz

Son lancement est prévu en 2009. Il observera le rayonnement fossile (2,7 K)

Refroidi à 0,1 K, le récepteur est relié à des antennes multifaisceaux, multifréquences, munies d'un système de deux réflecteurs, avec une source décentrée. Ce système permettra de mesurer la polarisation pour certaines fréquences.

La mission de Planck fait suite à celles des deux satellites Cobe (*Cosmic Background Explorer*), WMAP (*Wilkinson Microwave Anisotropy Probe*) qui ont permis d'analyser le rayonnement micro-onde existant il y a dix milliards d'années et de réaliser une carte de l'univers pour ce rayonnement.

Cobe a été placé sur une orbite de 900 km autour de la Terre. Parmi les appareils embarqués se trouvaient trois radiomètres (31 GHz, 52 GHz, 81 GHz) ayant pour but de mesurer les intensités du rayonnement.

WMAP, plus ancien, a pu confirmer l'existence de matière noire. Il a apporté des précisions sur l'âge de l'Univers et sa géométrie. Le système contient des réflecteurs de 1,4 m × 1,6 m qui renvoient les ondes vers des antennes dans cinq bandes de fréquences, de 22 à 90 GHz.

Bibliographie

http://www.anfr.fr/pages/presentation/gestion_spectre.pdf

ARRL – *The ARRL Antenna book*, 2002.

BALANIS C.A. – *Antenna theory, Analysis and Design*, Wiley, 1996, ISBN 0471 592684.

BONN F., ROCHON G. – *Précis de télédétection, vol. 1.*, Paris, Presses de l'Université du Québec, 1996.

BRAULT R. – *Les antennes*, Paris, Éditions ETSE, 1987.

CHU L.J. – *Physical limitations of Omni-Directional Antennas*, J. Appl. Phys., vol. 19, décembre 1948, pp. 1163-1175.

COMBES P.F. – *Micro-ondes, vol.2 Circuits passifs, propagation, antennes*, Paris, Dunod, 1997.

DARRICAUT J. – *Radars - Concepts et fonctionnalités*, E6 660, Les Techniques de l'Ingénieur.

ELLIOTT R.S.. – *Antenna Theory and Design*, New Jersey, IEEE Press John Wiley & Sons , 2003.

GUILBERT C. – *Pratique des antennes - TV, FM, Réception, Emission*, Paris, Radio, 1989.

JOHNSON R.C., JASIK H. – *Antenna Engineering Handbook*, McGraw Hill, 2^e édition, 1984.

LAVERGNAT J., SYLVAIN M. – *Propagation des ondes radioélectriques*, Paris, Masson, 1997.

Le VINE D. M., LAGERLOEF G.S., COLOMB F.R., YUEH S.H., PELLERANO F.A., *Aquarius : An Instrument to Monitor Sea Surface Salinity From Space*, IEEE Transactions On Geoscience And Remote Sensing, vol. 45, n° 7, juillet 2007.

LO Y.T., LEE S.W. – *Antenna Handbook*, Van Nostrand, 1988.

LOY M., IBOUN S. – *ISM bands and Short Range Devices Antennas*, Texas Instruments, Application Report SWRA046A.

NELSON R.A. – *ATI courses*. ATIcourses@aol.com.

REMY J. G., CUEUGNIET J., SIBEN C. – *Systèmes de radiocommunications avec les mobiles*, Paris, Eyrolles, CENT-ENST, 1992.

SKOLNIC M.I. – *Introduction to radar systems*, MacGraw-Hill, 2^e édition, 1980.

SKOLNIC M.I. – *Radar Handbook*, MacGraw-Hill, 2^e édition, 1990.

THOUREL L. – *Les antennes*, Dunod, 1^{re} édition, 1971.

ULABY F., MOORE R., FUNG A. – *Microwave remote sensing Active and passive, vol.1.*, Norwood, MA, Artech House, 198197.

WHEELER H.A. – *Fundamental Limitations of Small Antennas*, Proc. IRE, vol. 35, décembre 1947, pp. 1479 – 1484..

7 • SYSTÈMES MULTI-ANTENNES

7.1 Éléments de traitement d'antenne

Pour accroître la sensibilité d'une antenne, on est conduit à augmenter sa taille. Ceci induit éventuellement des difficultés de réalisation ou de manipulation (l'orientation mécanique devient difficile). Il est naturel, lorsque l'on dispose de plusieurs antennes, de chercher à exploiter conjointement les mesures. À partir d'un réseau de capteurs, on peut former une antenne à plusieurs éléments, où les signaux délivrés par chacun des capteurs sont combinés entre eux de manière à orienter « électroniquement » l'antenne ou à réaliser d'autres fonctions. Cette orientation électronique est réalisée par des circuits spécialisés ou par des algorithmes. On parle alors de traitement. Ces quelques pages forment une introduction très simplifiée aux méthodes de traitement d'antenne, l'objectif essentiel étant de sensibiliser aux principes de base de quelques méthodes (on ne discutera notamment pas des bornes de performances, de la robustesse au bruit, etc.). Dans un premier temps, on présentera un modèle simple et les paramétrisations possibles du champ d'onde reçu. Dans un second temps, on donnera l'analogie entre l'imagerie d'un champ de source et l'estimation spectrale, dans le cas d'une propagation par ondes planes. On présentera ensuite les méthodes classiques pour un champ continu : formation de voies, méthode de Capon, modélisation AR. Enfin, sous l'hypothèse de sources ponctuelles, on présentera quelques méthodes à « haute résolution ». La présentation sera appuyée par des résultats de simulation.

7.1.1 Modèle et paramétrisation

Lorsque le champ est stationnaire, le milieu homogène et isotrope, alors le milieu se comporte comme un filtre, et le champ au point $\mathbf{r}$ est :

$$e(\mathbf{r}, t) = \int_{\sigma} \int_{t'} G(\sigma - \mathbf{r}, t - t') s(\sigma, t') d\sigma dt'$$

où G désigne la fonction de Green du milieu, et s l'amplitude du champ émis. Cette relation, qui est une convolution spatiale et temporelle, exprime que le champ $e(\mathbf{r}, t)$, au point $\mathbf{r}$ et au temps t , résulte de la superposition de toutes les émissions aux points σ , et au cours du temps.

Pour un capteur placé en $\mathbf{r}_i$, en tenant compte de ses propres caractéristiques (fonction de transfert en fréquence, directivité), on a :

$$x_i(t) = x(\mathbf{r}_i, t) = \int_{\sigma} \int_{t'} h_i(\sigma - \mathbf{r}_i, t - t') s(\sigma, t') d\sigma dt' + b_i(t)$$

On a supposé ici que les sources étaient immobiles (l'effet Doppler est un effet non linéaire).

L'hypothèse « bande étroite » permet de simplifier l'écriture précédente : en effet, la transformée de Fourier (TF), à la fréquence f , de la relation précédente est :

$$x_i(f) = \int_{\sigma} h_i(\sigma - \mathbf{r}_i, f) s(\sigma, f) d\sigma$$

Si le champ est à bande étroite autour d'une fréquence f_0 , on a :

$$s(\sigma, t') \simeq s(\sigma, f_0) e^{-j2\pi f_0 t} = s_0(\sigma) e^{-j2\pi f_0 t}$$

En collectant les mesures sur M capteurs placés en $\mathbf{r}_1, \dots, \mathbf{r}_M$, on a :

$$\begin{bmatrix} x_1(t, f_0) \\ \vdots \\ x_M(t, f_0) \end{bmatrix} = e^{-j2\pi f_0 t} \int s_0(\sigma) \begin{bmatrix} h_1(\sigma - \mathbf{r}_1, f_0) \\ \vdots \\ h_M(\sigma - \mathbf{r}_M, f_0) \end{bmatrix} d\sigma + \begin{bmatrix} b_1(t, f_0) \\ \vdots \\ b_M(t, f_0) \end{bmatrix}$$

soit

$$\mathbf{x}(t, f_0) = e^{-j2\pi f_0 t} \int s_0(\sigma) \mathbf{a}(\sigma) d\sigma + \mathbf{b}(t, f_0)$$

où $s_0(\sigma)$ est l'amplitude émise en σ , à la fréquence f_0 , et $\mathbf{a}(\sigma)$ est le vecteur directionnel, qui intègre les informations sur le modèle et le milieu de propagation ainsi que les caractéristiques des capteurs. L'hypothèse bande étroite est très importante, car elle permet de découpler les contributions spatiales et temporelles. Lorsque les signaux ne sont pas bande étroite, on se ramène à ce cas de figure, soit par un filtrage bande étroite en réception, soit en effectuant une transformée de Fourier des données en début de traitement.

On notera τ_{ik} le retard de propagation entre deux capteurs :

$$\tau_{ik} = T(\sigma - \mathbf{r}_i, f_0) - T(\sigma - \mathbf{r}_k, f_0)$$

où $T(\sigma - \mathbf{r}_i, f_0)$ désigne le temps de propagation de σ à $\mathbf{r}_i$, à la fréquence f_0 .

On peut omettre les dépendances en temps et fréquence en ne conservant que les amplitudes complexes :

$$\mathbf{x} = \int s_0(\sigma) \mathbf{a}(\sigma) d\sigma + \mathbf{b}$$

La connaissance du modèle de propagation, des caractéristiques du réseau (gain, directivité des capteurs, géométrie...) permet de relier $\mathbf{a}(\sigma)$ à un ensemble de paramètres physiques, θ .

Le vecteur paramétré $\mathbf{a}(\theta)$ est le *vecteur directionnel* ou *steering vector*. En enregistrant les données recueillies sur les capteurs à différents instants $t_1, \dots, t_K$ (ou à partir du calcul sur différents blocs de données), on obtient une série de réalisations ou *snapshots* $\mathbf{x}_1, \dots, \mathbf{x}_K$, qui sont utiles pour l'estimation des quantités statistiques.

■ Exemples

Dans ce chapitre, les différentes méthodes seront illustrées pour un modèle de propagation sphérique, et un modèle de propagation par ondes planes, ces ondes illuminant un réseau de M capteurs répartis sur une droite, les capteurs étant uniformément répartis, avec une distance intercapteurs notée d .

□ Modèle sphérique sur une antenne linéaire uniforme (ALU)

On utilisera la modélisation suivante : la source S est située à une distance H de l'axe de l'antenne et on prend un capteur de référence à une distance D de la perpendiculaire vers la source. Dans ces conditions, le vecteur directionnel prend la forme :

$$\mathbf{a}(\theta)^T = \begin{bmatrix} \frac{1}{R_1} e^{-j2\pi f_0 R_1 / c} & \dots & \frac{1}{R_M} e^{-j2\pi f_0 R_M / c} \end{bmatrix}$$

avec $R_m = (H^2 + (D + md)^2)^{\frac{1}{2}}$. Le paramètre θ regroupe ici la hauteur H de la source à l'axe de l'antenne et la distance D à un capteur de référence : $\theta^T = \begin{bmatrix} D & H \end{bmatrix}$.

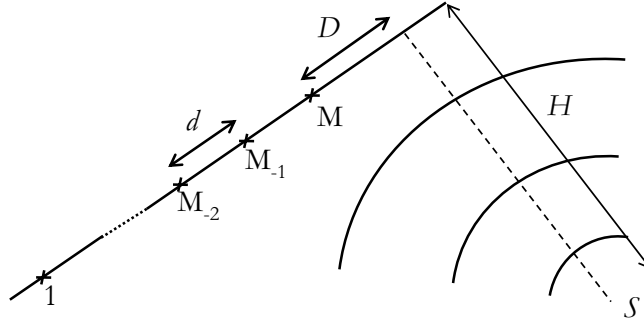


Figure 7.1 – Ondes à propagation sphérique illuminant une ALU.

Remarque : le vecteur directionnel $a(\theta)$ est invariant par rotation autour de l'axe de l'antenne : ceci correspond à une ambiguïté sur la localisation des sources situées sur un cercle de rayon H , pour une antenne linéaire.

□ Modèle plan sur une antenne linéaire uniforme (ALU)

Lorsque R devient très grand devant les dimensions de l'antenne, le front d'onde devient plan, et le retard de propagation entre deux capteurs s'exprime comme

$$2\pi f_0 \tau_{ik} = \mathbf{r}_{ik}^+ \mathbf{k} = (\mathbf{r}_i - \mathbf{r}_k)^+ \mathbf{k}$$

avec $\mathbf{k}$ le vecteur d'onde.

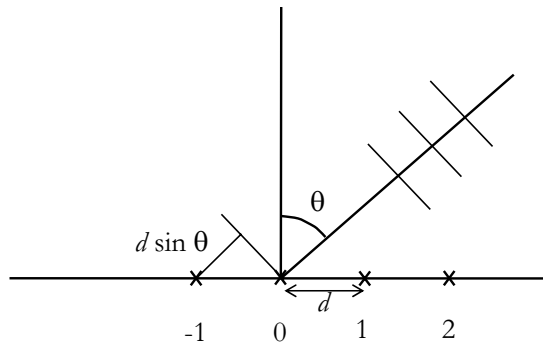


Figure 7.2 – Onde plane illuminant une ALU.

Le paramètre θ se réduit à l'azimut θ , et en notant τ le retard de propagation entre deux capteurs consécutifs, on a :

$$\tau = \frac{\mathbf{r}^+ \mathbf{k}}{2\pi f_0} = -\frac{d \sin \theta}{c}$$

ce qui correspond au temps mis par l'onde plane pour parcourir la distance $d \sin \theta$ à la vitesse c (le signe moins provient de la direction de $\mathbf{k}$).

Le vecteur directionnel peut alors s'exprimer comme :

$$\begin{aligned} a(\theta)^T &= \left[1 \quad e^{j2\pi f_0 \frac{d \sin \theta}{c}} \quad \dots \quad e^{j2\pi(M-1)f_0 \frac{d \sin \theta}{c}} \right] \\ &= \left[1 \quad e^{j2\pi u(\theta)} \quad \dots \quad e^{j2\pi(M-1)u(\theta)} \right] \end{aligned}$$

avec $u(\theta) = d \sin \theta / \lambda$.

Le vecteur $\mathbf{a}(\theta)$ est indépendant des effets de la propagation jusqu'au premier capteur. On ne pourra alors pas remonter à l'amplitude ou aux caractéristiques énergétiques des sources situées en θ , mais simplement remonter aux caractéristiques mesurées au niveau de l'antenne en provenance de la direction θ .

■ Illustrations

Nous illustrerons les différentes techniques présentées à l'aide des deux exemples suivants

- (A) Un signal résultant de la superposition de trois ondes planes, d'amplitudes respectives 2, 2 et 1, et d'angles d'incidence respectifs $\theta_1 = -20^\circ$, $\theta_2 = 38^\circ$, et $\theta_3 = 40^\circ$.
- (B) Un signal résultant de la superposition de trois ondes à propagation circulaire, les trois sources, également d'amplitudes 2, 2 et 1, étant caractérisées par les paramètres $[D, H]$ suivants : $[D_1, H_1] = [8, 35]d$, $[D_2, H_2] = [3, 40]d$ et $[D_3, H_3] = [-4, 14]d$, où d est la distance intercapteurs.

Dans les deux cas, on prend $d/\lambda = 1/4$.

7.1.2 L'imagerie d'un champ continu

L'imagerie consiste à étudier la répartition spatiale de telle ou telle propriété, en particulier la distribution spatiale d'intensité $I_s(\mathbf{n})$, dans la direction $\mathbf{n}$. Cela permet d'établir une cartographie du milieu. On trouve des applications en radar, sonar (cartographie sous-marine), astronomie.

■ Généralités

On débute avec la fonction d'intercorrélation :

$$\gamma(\mathbf{r}_{ik}) = E [x_i(t)x_k(t)^*]$$

entre les mesures $x_i(t)$ et $x_k(t)$ prises sur les capteurs positionnés en $\mathbf{r}_i$ et $\mathbf{r}_k$, et où $E[\bullet]$ désigne l'espérance mathématique.

Supposons que les capteurs soient identiques et de même gain sur la région D étudiée. De plus, on suppose que le signal est bande étroite et le modèle de propagation plan. Alors :

$$\begin{aligned} \gamma(\mathbf{r}_{ik}) &= E \left[\int_D f(v_0)s(\mathbf{n})e^{-j2\pi r_i^+ \mathbf{n}/\lambda} d\mathbf{n} \int_D f(v_0)^*s(\mathbf{n}')^* e^{j2\pi r_k^+ \mathbf{n}'/\lambda} d\mathbf{n}' \right] \\ &= \iint_D |f(v_0)|^2 E [s(\mathbf{n})s(\mathbf{n}')^*] e^{-j2\pi(r_i^+ \mathbf{n} - r_k^+ \mathbf{n}')/\lambda} d\mathbf{n}d\mathbf{n}' \end{aligned}$$

Si le champ est incohérent, c'est-à-dire si $E [s(\mathbf{n})s(\mathbf{n}')^*] = I_s(\mathbf{n})\delta(\mathbf{n} - \mathbf{n}')$ alors :

$$\gamma(\mathbf{r}_{ik}) = \int_D |f(v_0)|^2 I_s(\mathbf{n})e^{-j2\pi r_{ik}^+ \mathbf{n}/\lambda} d\mathbf{n}$$

Ce lien entre la distribution d'intensité $I_s(\mathbf{n})$ et les corrélations $\gamma(\mathbf{r}_{ik})$ constitue un résultat connu comme le théorème de Van Cittert-Zernicke. Il indique que le lien entre $\gamma(\mathbf{r}_{ik})$ et $I_s(\mathbf{n})$ est une simple **transformée de Fourier (inverse) multidimensionnelle**. Ce résultat est bien sûr analogue au théorème de Wiener Kintchine en analyse spectrale qui indique que la densité spectrale de puissance et la fonction de corrélation forment une paire de transformées de Fourier.

Dès lors, la recherche de la distribution d'intensité $I_s(\mathbf{n})$ correspond à un problème d'estimation spectrale : à partir des échantillons de la fonction de corrélation spatiale, il s'agit de reconstruire la distribution d'intensité qui vérifie

$$I_s(\mathbf{n}) \propto \sum_r \gamma(\mathbf{r})e^{j2\pi \mathbf{r}^+ \mathbf{n}/\lambda}$$

L'inversion directe, telle que fournie par la formule précédente, souffre du fait qu'il faut estimer la fonction d'autocorrélation spatiale $\gamma(r)$ à partir des données (erreurs d'estimation) ainsi que du fait que les dimensions limitées de l'antenne entraînent nécessairement un support fini de la fonction de corrélation calculée, ce qui limite alors la résolution.

Une autre solution pour estimer la puissance consiste à utiliser un **filtre spatial** $\mathbf{w}(\theta_0)$ afin d'isoler la portion de $\mathbf{x}$, $\mathbf{x}(\theta_0)$, caractérisée par le paramètre θ_0 du modèle :

$$x(\theta_0) = \mathbf{w}(\theta_0)^+ \mathbf{x} = \sum_{k=1}^M w_k(\theta_0) x_k(t, f_0)$$

Ceci consiste ainsi à effectuer une combinaison linéaire des signaux recueillis sur chacun des capteurs, en choisissant les coefficients de pondération de manière à favoriser l'entrée de paramètre θ_0 . Ce principe est illustré sur la figure 7.3

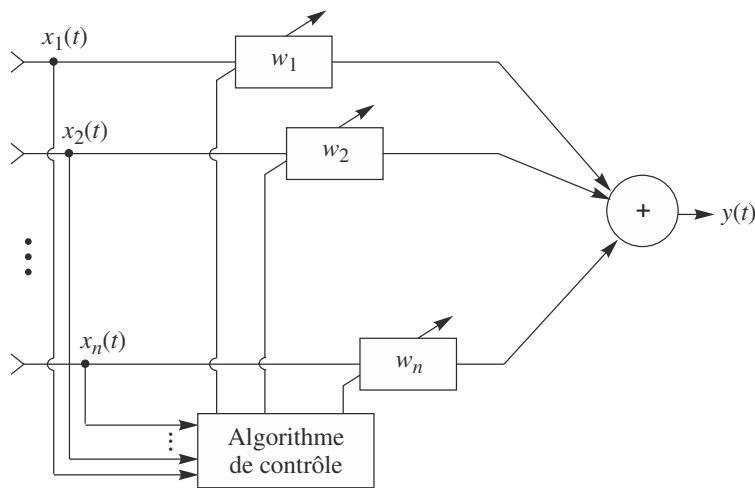


Figure 7.3 – Traitement d'antenne par filtrage spatial

La puissance en sortie du filtre spatial s'écrit :

$$E \left[|x(\theta_0)|^2 \right] = E \left[|\mathbf{w}(\theta_0)^+ \mathbf{x}|^2 \right] = \mathbf{w}(\theta_0)^+ E \left[\mathbf{x} \mathbf{x}^+ \right] \mathbf{w}(\theta_0) = \mathbf{w}(\theta_0)^+ \mathbf{R}_x \mathbf{w}(\theta_0)$$

avec $\mathbf{R}_x$ la matrice de corrélation des observations.

Dans le cas d'un champ continu, on a vu que

$$\mathbf{x} = \int s(\theta) \mathbf{a}(\theta) d\theta + \mathbf{b}$$

on obtient

$$E \left[|x(\theta_0)|^2 \right] = \int I_s(\theta) |\mathbf{w}(\theta_0)^+ \mathbf{a}(\theta)|^2 d\theta$$

en supposant à nouveau le champ incohérent, et le bruit centré.

On constate que cette puissance sera d'autant plus proche de $I_s(\theta_0)$ que $\mathbf{w}(\theta_0)^+ \mathbf{a}(\theta)$ sera sélectif si $|\mathbf{w}(\theta_0)^+ \mathbf{a}(\theta)|^2 \sim \delta(\theta_0 - \theta)$ alors $E \left[|x(\theta_0)|^2 \right] \sim I_s(\theta_0)$. Il convient donc de bien choisir le filtre spatial, en relation avec le modèle de réception $\mathbf{a}(\theta)$. La section suivante présente un premier exemple de filtrage spatial, avec la formation de voies, et nous examinerons ensuite le filtre spatial associé à la méthode de Capon.

■ La formation de voies

La technique de formation de voies vise à estimer la puissance émise pour une valeur particulière du paramètre θ .

L'idée de la formation de voies est de sommer les contributions recueillies sur les différents capteurs de manière constructive, et d'augmenter ainsi le rapport signal sur bruit. On introduit ainsi une série de gains ajustables, qui servent à privilégier une composante du signal (caractérisée par un paramètre θ) : on parle alors d'orientation électronique. On balaie ensuite les différentes valeurs possibles du paramètre pour reconstituer une carte de la distribution de puissance.

Considérons le cas d'une seule source de paramètre θ_0 . Les observations correspondantes sont $\mathbf{x} = s\mathbf{a}(\theta_0) + \mathbf{b}$, où le vecteur $\mathbf{b}$ collecte les contributions du bruit sur les différents capteurs ; on suppose que le bruit est blanc spatialement, c'est-à-dire décorrélé de capteur à capteur et de puissance σ_b^2 . On combine alors les observations recueillies sur les différents capteurs par $\mathbf{w}^+ \mathbf{x} = s \mathbf{w}^+ \mathbf{a}(\theta_0) + \mathbf{w}^+ \mathbf{b}$. La puissance en sortie du filtre spatial w s'écrit alors

$$\sigma_{FV} = E[|s|^2] |\mathbf{w}^+ \mathbf{a}(\theta_0)|^2 + |\mathbf{w}|^2 \sigma_b^2$$

Dans ces conditions, il est facile de constater que le rapport signal sur bruit :

$$r = \frac{E[|s|^2] |\mathbf{w}^+ \mathbf{a}(\theta_0)|^2}{|\mathbf{w}|^2 \sigma_b^2}$$

est maximum lorsque $\mathbf{w}$ est colinéaire au vecteur directionnel $\mathbf{a}(\theta_0)$: $\mathbf{a}(\theta_0) = \lambda \mathbf{w}$. Comme ceci doit être prolongé pour chaque paramètre θ possible, on prendra donc $\mathbf{w}$ comme une fonction de θ ; $\lambda \mathbf{w}(\theta) = \mathbf{a}(\theta)$. On choisit généralement la normalisation pour que la puissance d'un bruit blanc en entrée ne soit pas modifiée : si on prend un bruit blanc en entrée de matrice de corrélation $\mathbf{R}_x = \sigma_b^2 \mathbf{I}$, avec $\mathbf{I}$ la matrice identité, on a

$$\sigma_{FV}(\theta) = E[|x(\theta)|^2] = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta) = \sigma_b^2 \lambda^2 \mathbf{a}(\theta)^+ \mathbf{a}(\theta) = \sigma_b^2$$

ce qui fournit λ^2 , et :

$$\mathbf{w}(\theta) = \frac{\mathbf{a}(\theta)}{\sqrt{\mathbf{a}(\theta)^+ \mathbf{a}(\theta)}}$$

La puissance en sortie de la formation de voies devient :

$$E[|x(\theta)|^2] = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta) = \frac{\mathbf{a}(\theta)^+ \mathbf{R}_x \mathbf{a}(\theta)}{\mathbf{a}(\theta)^+ \mathbf{a}(\theta)}$$

On notera ainsi que la technique de formation de voies peut s'appliquer quel que soit le modèle de propagation : il suffit de connaître l'expression du vecteur directionnel.

Dans le cas d'un modèle plan, la formation de voie consiste simplement à compenser les retards de propagation entre les différents capteurs, en ajustant des déphasages par l'intermédiaire de gains sur chaque voie de manière à effectuer une sommation constructive du signal d'intérêt.

En pratique, la puissance sera estimée de la manière suivante : filtrage, quadrature, intégration. Ceci est illustré sur la figure 7.4.

Les signaux recueillis à la sortie de chacun des capteurs sont affectés d'un gain, puis sommés. La puissance est ensuite estimée par quadrature puis intégration temporelle.

Considérons maintenant le cas particulier d'un modèle de propagation par ondes planes et d'une antenne linéaire uniforme.

Dans ce cas, on a :

$$\mathbf{w}^+ = \frac{1}{\sqrt{M}} \begin{bmatrix} 1 & e^{-j2\pi u(\theta_0)} & \dots & e^{-j2\pi(M-1)u(\theta_0)} \end{bmatrix}$$

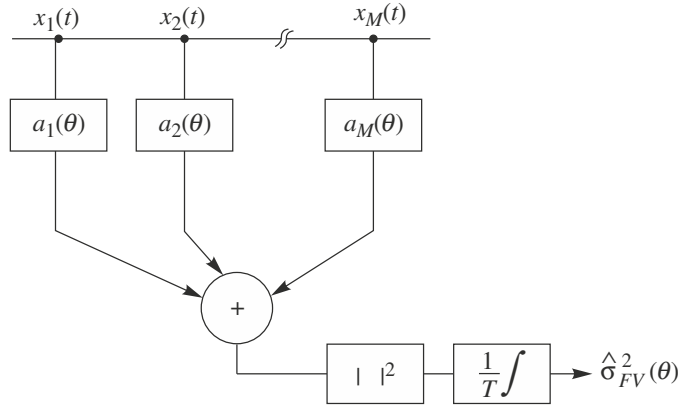


Figure 7.4 – Principe de la formation de voies.

et

$$\begin{aligned}
 \sigma_{FV}^2(\theta) &= \frac{1}{M} \mathbb{E} \left[|a(\theta)^+ x|^2 \right] = \frac{1}{M} \mathbb{E} \left[\left| \sum_{m=1}^M x(m) e^{-j2\pi m u(\theta)} \right|^2 \right] \\
 &= \frac{1}{M} \mathbb{E} \left[\sum_{m_1=1}^M \sum_{m_2=1}^M x(m_1) x(m_2)^* e^{-j2\pi(m_1 - m_2)u(\theta)} \right] \\
 &= \frac{1}{M} \sum_{m_1=1}^M \sum_{m_2=1}^M \gamma_x(m_1 - m_2) e^{-j2\pi(m_1 - m_2)u(\theta)} \\
 &= \frac{1}{M} \sum_{m=-M+1}^{M-1} (M - |m|) \gamma_x(m) e^{-j2\pi m u(\theta)}
 \end{aligned}$$

soit finalement :

$$\sigma_{FV}^2(\theta) = \sum_{m=-M+1}^{M-1} g(m) \gamma_x(m) e^{-j2\pi m u(\theta)}$$

avec $g(m)$ la fonction d'autocorrélation du réseau et $\gamma_x(m)$ la séquence de corrélation des observations. On reconnaît dans cette dernière relation l'expression d'une transformée de Fourier, où $u(\theta)$ joue le rôle d'une fréquence spatiale normalisée. Ainsi, la formation de voies correspond à la transformée de Fourier de la fonction de corrélation, pondérée par une fenêtre de Bartlett (fenêtre triangulaire). Lorsque la fonction d'autocorrélation spatiale $\gamma_x(m)$ est estimée à partir des données, on obtient alors un simple périodogramme ; si la matrice de corrélation $\mathbf{R}_x$ est estimée par un moyennage sur K réalisations $\mathbf{x}_k$:

$$\hat{\mathbf{R}}_x = \frac{1}{K} \sum_{k=1}^K \mathbf{x}_k \mathbf{x}_k^+$$

alors la puissance correspondante vaut

$$\begin{aligned}
 \hat{\sigma}_{FV}^2 &= \mathbf{w}(\theta)^+ \hat{\mathbf{R}}_x \mathbf{w}(\theta) = \frac{1}{MK} \sum_{k=1}^K |\mathbf{a}(\theta_0)^+ \mathbf{x}_k|^2 \\
 &= \frac{1}{K} \sum_{k=1}^K \frac{1}{M} \left| \sum_{m=1}^M x_k(m) e^{-j2\pi m u(\theta)} \right|^2
 \end{aligned}$$

et prend exactement la forme d'un périodogramme moyenné.

Le calcul de la puissance théorique σ_{FV}^2 en sortie de la formation de voies nous montre qu'elle s'exprime comme la transformée de Fourier du produit $g(m)\gamma_x(m)$. Ainsi, la puissance s'exprimera aussi comme le produit de convolution de la transformée de la fenêtre et de la distribution d'intensité I_s (à un terme de bruit près) :

$$\sigma_{FV}^2(\theta) = G(u(\theta)) * I_s(u(\theta)) + \sigma_b^2(u(\theta))$$

avec
$$G(u(\theta)) = \frac{1}{M} \left(\frac{\sin(\pi M u(\theta))}{\sin(\pi u(\theta))} \right)^2$$

Pour une source ponctuelle, équivalente à une « impulsion » dans le domaine spatial, par exemple dans la direction θ_0 , on a $I_s(u(\theta)) = \sigma_s^2 \delta(u(\theta) - u(\theta_0))$, et, en l'absence de bruit,

$$\sigma_{FV}^2(\theta) = G(u(\theta)) * I_s(u(\theta)) = \sigma_s^2 G(u(\theta) - u(\theta_0))$$

Ainsi, la fonction $G(u(\theta))$ correspond à la réponse du réseau à une impulsion, et exhibe une limitation en résolution, liée à la largeur du lobe principal de $G(u(\theta))$. D'autre part, les lobes secondaires, importants, collectent de la puissance hors de la direction visée. La fonction $G(u(\theta))$ définit alors le diagramme de directivité de l'antenne.

La figure 7.5 représente la fonction $G(u(\theta))$ en fonction de $u(\theta)$: le lobe principal est d'autant plus large que l'antenne est courte et la résolution en $u(\theta)$ varie en $1/M$; la résolution en fréquences spatiales varie pour sa part en $1/Md$, c'est-à-dire de façon inversement proportionnelle à la dimension de l'antenne formée par le réseau de capteurs. Cette résolution est la résolution de Fourier, et nous avons vu que la formation de voie s'interprète comme une méthode de Fourier. Pour accroître cette capacité de résolution, il n'y a guère d'autre moyen que d'augmenter le nombre de capteurs ou la distance intercapteurs d . Il sera aussi possible, comme on le verra plus loin, d'envisager d'autres méthodes, à super ou à haute résolution.

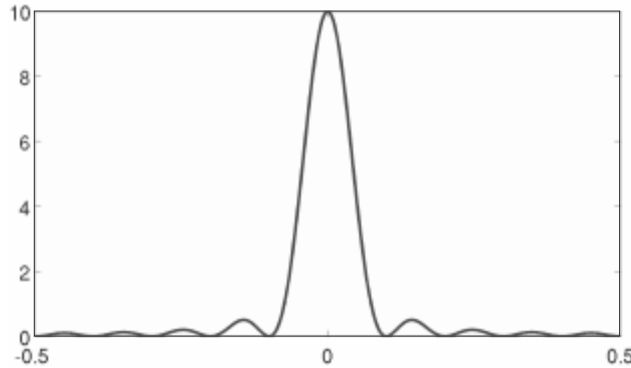


Figure 7.5 – Diagramme directivité en fréquences spatiales normalisées $u(\theta)$, avec $M = 10$.

Il est possible de pondérer cette formation de voie de manière à obtenir un diagramme de directivité avec des lobes secondaires plus bas, mais au détriment d'un élargissement du lobe principal.

Le traitement par formation de voies est ainsi caractérisé par son diagramme de directivité, fixé pour chaque direction d'arrivée.

■ Illustrations

Considérons pour commencer le cas (A) de la somme de trois contributions de sources situées en -20° , 38° et 40° . La figure 7.6 donne le résultat obtenu avec $M = 10$. On observe que les sources ponctuelles sont très étalées, et qu'il est absolument impossible de séparer les sources en 38° et 40° .

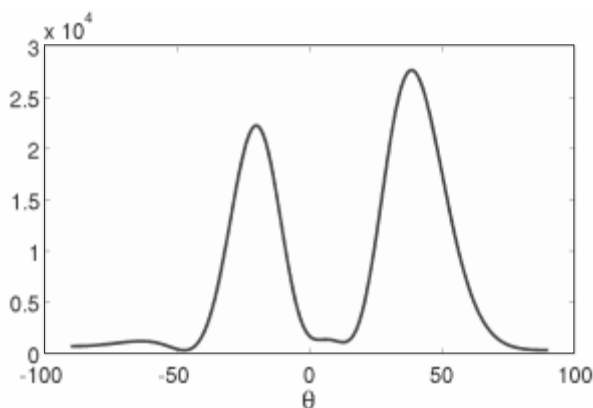


Figure 7.6 – Sortie de la formation de voies dans l'exemple (A), avec $M = 10$.

La largeur, en $u(\theta)$, de la fonction $G(u(\theta))$ est en $1/M$. Aussi, on pourra résoudre les deux sources si $u(\theta_3) - u(\theta_2) \geq 1/M$, ce qui fournit $M \geq \lambda / (d (\sin(\theta_3) - \sin(\theta_2)))$, soit $M \geq 148$. La figure 7.7 fournit le résultat obtenu pour $M = 200$: on observe que cette fois-ci les sources sont bien séparées, mais bien sûr au prix d'un nombre important de capteurs.

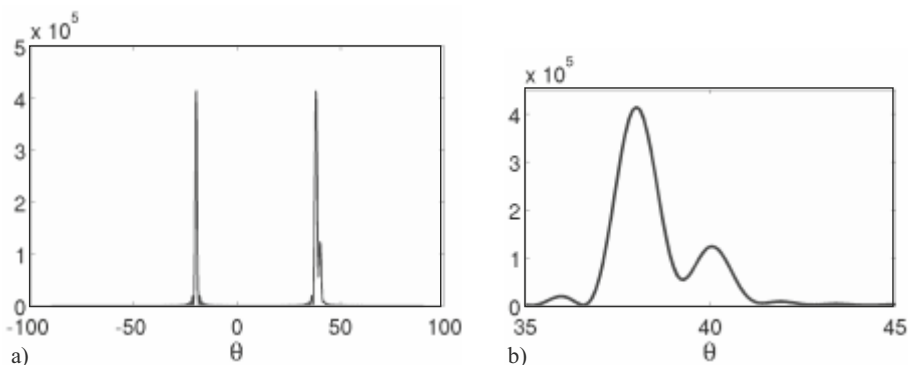


Figure 7.7 – Sortie de la formation de voies dans l'exemple (A), avec $M = 200$.
Sortie complète (a) et détail (b).

Considérons maintenant le cas du modèle de propagation circulaire, exemple (B). La technique de formation de voies est bien sûr applicable, avec l'expression du vecteur directionnel correspondant à la paramétrisation que nous avons donnée dans ce cas. Il s'agit alors de calculer la puissance en sortie du filtre spatial, en parcourant une grille de valeurs $[D, H]$ possibles. Les maxima du résultat donnent alors une indication sur les positions des sources. Le résultat obtenu

ainsi avec $M = 100$ est donné figure 7.8. Comme précédemment, la localisation des sources s'affine et la résolution s'améliore lorsqu'on augmente le nombre de capteurs.

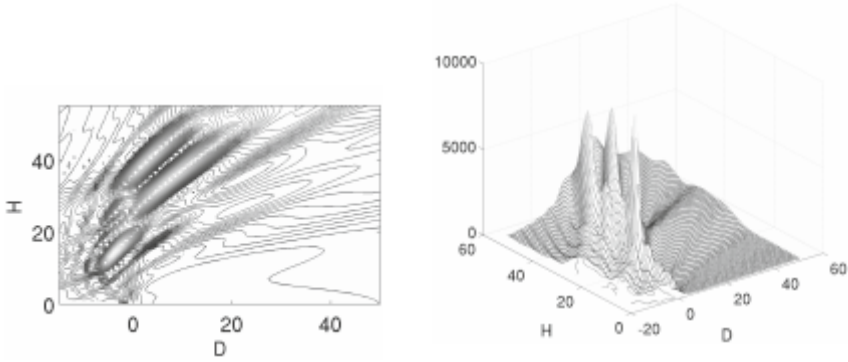


Figure 7.8 – Sortie de la formation de voies pour un modèle circulaire, avec $M = 100$.

À l'issue de cette section, nous pouvons retenir qu'il est possible de construire une antenne synthétique à partir d'un réseau de capteurs, chaque capteur du réseau n'ayant pas nécessairement de propriétés de directivité. C'est le traitement des signaux recueillis sur les capteurs qui permet de transformer un réseau inerte en une antenne. Plus encore, la modification des coefficients du filtre spatial $\mathbf{w}(\theta)$ permet d'orienter électroniquement l'antenne, de manière similaire à l'orientation mécanique d'une antenne standard. Le diagramme de directivité de l'antenne obtenue est lié à la forme du réseau (nous n'avons ici considéré qu'une antenne linéaire) et au nombre de capteurs le composant. D'autres techniques de traitement du signal permettent d'améliorer les performances en résolution.

■ La méthode de Capon

Du fait de la limitation en résolution et de la présence des lobes secondaires importants, des sources situées en dehors de la direction d'intérêt, ou plus généralement pour un paramètre θ différent, peuvent « renvoyer » de la puissance et ainsi corrompre les mesures, voire, même masquer des sources de faible puissance. La méthode introduite par Capon permet de tenir compte de l'ensemble des sources présentes pour pallier la limitation en résolution de la formation de voies et réduire l'amplitude des lobes secondaires : c'est une méthode adaptative. Le filtre spatial est ajusté afin d'orienter l'antenne tout en minimisant la contribution des sources qui ne sont pas situées dans la direction scrutée.

On cherche donc à *minimiser la puissance globale* recueillie pour le paramètre θ :

$$\sigma^2(\theta) = \mathbf{E} [\mathbf{w}(\theta)^+ \mathbf{x}] = \mathbf{w}(\theta)^+ \mathbf{E} [\mathbf{x}^+ \mathbf{x}] \mathbf{w}(\theta) = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta)$$

avec $\mathbf{R}_x$ la matrice d'autocorrélation, sans modifier la puissance dans la direction visée. Ce *desideratum* s'écrit sous la forme de la contrainte :

$$\mathbf{w}(\theta)^+ \mathbf{a}(\theta) = 1$$

ce qui signifie simplement que l'on impose un gain unitaire dans la direction θ . Par suite, la méthode sélectionnera un filtre spatial $\mathbf{w}(\theta)$ qui minimisera nécessairement la puissance renvoyée en θ par les « brouilleurs ».

On écrit alors le Lagrangien correspondant à ce problème de minimisation sous contrainte :

$$L(\mathbf{w}(\theta); \lambda) = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta) - \lambda (\mathbf{w}(\theta)^+ \mathbf{a}(\theta) - 1)$$

Le minimum de $L(\mathbf{w}(\theta); \lambda)$ suivant le premier argument est atteint pour $\mathbf{w}(\theta)$ tel que :

$$\frac{\partial L(\mathbf{w}(\theta); \lambda)}{\partial \mathbf{w}(\theta)} = 2\mathbf{R}_x \mathbf{w}(\theta) - \lambda \mathbf{a}(\theta) = \mathbf{0}$$

soit

$$\mathbf{w}(\theta) = \frac{\lambda}{2} \mathbf{R}_x^{-1} \mathbf{a}(\theta)$$

La contrainte $\mathbf{w}(\theta)^+ \mathbf{a}(\theta) = 1$ fournit alors :

$$\frac{\lambda}{2} = \frac{1}{\mathbf{a}(\theta)^+ \mathbf{R}_x^{-1} \mathbf{a}(\theta)}$$

soit finalement :

$$\mathbf{w}(\theta) = \frac{\mathbf{R}_x^{-1} \mathbf{a}(\theta)}{\mathbf{a}(\theta)^+ \mathbf{R}_x^{-1} \mathbf{a}(\theta)}$$

Le spectre de Capon correspondant est quant à lui :

$$\sigma_{CAP}^2 = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta) = \frac{1}{\mathbf{a}(\theta)^+ \mathbf{R}_x^{-1} \mathbf{a}(\theta)}$$

Le filtre est adaptif à la direction visée et prend en compte l'ensemble de l'environnement, par l'intermédiaire de la matrice de corrélation, afin de minimiser les contributions des brouilleurs. Ce traitement est aussi caractérisé par un pouvoir de résolution dépendant du rapport signal sur bruit (RSB) ; plus le RSB est élevé, meilleure est la résolution.

La technique mise en œuvre dans la méthode de Capon peut-être étendue au cas où l'on souhaite prendre en compte plus de contraintes. Typiquement, on pourra rechercher non seulement à assurer un gain unité dans la direction d'intérêt, mais aussi à imposer un gain nul dans les directions de brouilleurs connus. On introduit ainsi une matrice de contrainte $\mathbf{C}$ telle que $\mathbf{C}^+ \mathbf{w}(\theta) = \mathbf{c}$. Un exemple simple est le cas où l'on veut annuler un brouilleur dans une direction prédéterminée θ_b , soit $\mathbf{a}(\theta_b)^+ \mathbf{w}(\theta) = 0$, en plus de la condition $\mathbf{a}(\theta)^+ \mathbf{w}(\theta) = 1$. Dans ce cas on a $\mathbf{C} = \begin{bmatrix} \mathbf{a}(\theta) & \mathbf{a}(\theta_b) \end{bmatrix}$ et $\mathbf{c} = \begin{bmatrix} 1 & 0 \end{bmatrix}^T$. On cherche comme précédemment à minimiser la puissance globale $\sigma^2(\theta) = \mathbf{w}(\theta)^+ \mathbf{R}_x \mathbf{w}(\theta)$ sous la contrainte $\mathbf{C}^+ \mathbf{w}(\theta) = \mathbf{c}$. La résolution conduit alors à la solution

$$\mathbf{w}(\theta) = \mathbf{R}_x^{-1} \mathbf{C} (\mathbf{C}^+ \mathbf{R}_x \mathbf{C})^{-1} \mathbf{C}$$

qui se réduit à la formule de Capon standard lorsque $\mathbf{C} = \mathbf{a}(\theta)$ et $\mathbf{c} = 1$. Notons qu'il est possible de calculer le filtre de Capon $\mathbf{w}(\theta)$ de manière itérative (par un algorithme de minimisation de type gradient), afin de réduire la charge en calcul.

■ Illustrations

Pour l'exemple (A) du modèle plan, avec trois sources situées en -20° , 38° et 40° , on obtient les résultats présentés sur la figure 7.9, avec $M = 10$. Ces résultats sont équivalents à ceux obtenus pour la formation de voies, mais avec 200 capteurs !

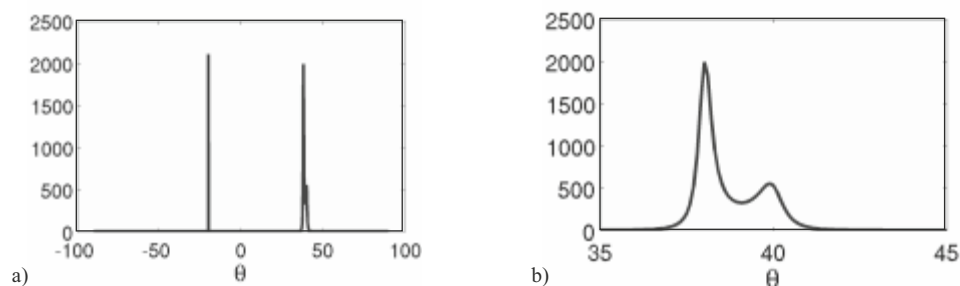


Figure 7.9 – Sortie de la méthode de Capon dans l'exemple (A), avec $M = 10$. Sortie complète a) et détail b).

Dans le cas du modèle circulaire, exemple (B), les performances sont également meilleures qu'avec la simple formation de voies. La figure 7.10 fournit les résultats avec 15 capteurs.

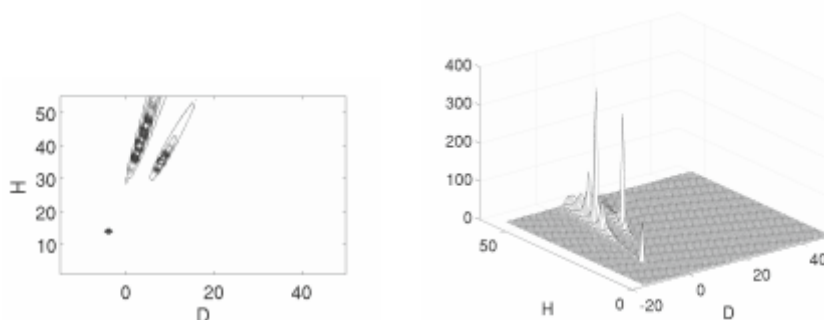


Figure 7.10 – Sortie de la méthode de Capon dans l'exemple (B), avec $M = 15$.

■ Filtrage

Pour terminer cette section, nous présenterons une application de filtrage spatial. Dans l'exemple (A), on suppose que les amplitudes sont maintenant variables : on prend un bruit en -20° et un autre en 40° et un signal binaire est émis par la source située en 38° . On suppose que la fréquence porteuse est suffisamment haute pour que l'hypothèse bande étroite soit vérifiée. Les trois sources illuminent l'antenne. On mesure donc sur chacun des capteurs un mélange, avec différents déphasages, des signaux émis par les trois sources. La figure 7.11 donne ainsi le signal recueilli sur le premier capteur.

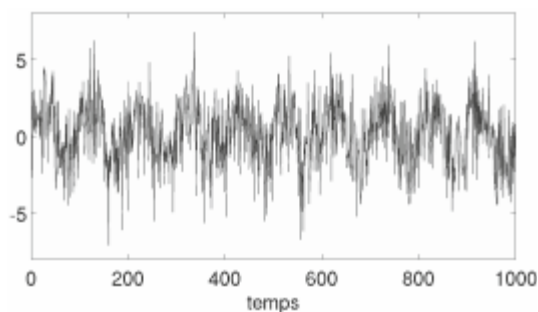


Figure 7.11 – Signal temporel recueilli sur le premier capteur (mélange de deux bruits et d'un train d'impulsions).

On utilise alors le filtre spatial de la méthode de Capon, orienté dans la direction 38° et on observe la sortie de la somme pondérée des signaux recueillis sur les capteurs. Le résultat est rapporté sur la figure 7.12 : on est capable d'extraire du mélange le signal binaire, avec une faible interférence résiduelle, essentiellement liée à l'émission proche en 40° . De telles applications du filtrage spatial sont maintenant mises en œuvre dans le cadre des communications cellulaires, en équipant les stations de plusieurs antennes de réception, ce qui permet de focaliser sur certains utilisateurs, et ainsi d'augmenter la capacité du système.

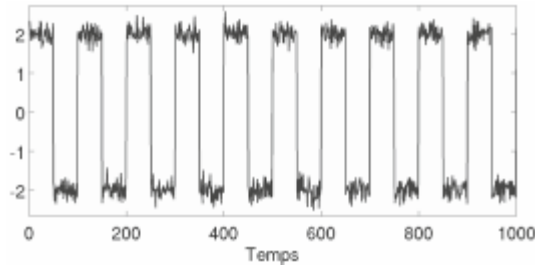


Figure 7.12 – Signal temporel obtenu après filtrage spatial de Capon dans la direction 38° .

■ Prédiction linéaire

On a noté que, sous l'hypothèse d'ondes planes et d'antenne linéaire uniforme, le problème de reconstruction de la distribution d'intensité est équivalent à un problème d'estimation spectrale. Il est donc possible d'utiliser les méthodes paramétriques standard appliquées en analyse spectrale, et notamment les méthodes autorégressives (AR). On sait que ces méthodes qui reposent sur un modèle paramétrique de la densité spectrale de puissance permettent d'aller au-delà de la limite de résolution de Fourier.

Rappelons ici simplement que l'idée est d'utiliser un modèle paramétrique de la densité spectrale et d'estimer les paramètres de ce modèle. Parmi les différentes familles possibles, le modèle AR est particulièrement utilisé, car les paramètres AR $\alpha^T = [\alpha_1 \ \alpha_2 \ \dots \ \alpha_q]$ sont obtenus en résolvant système linéaire :

$$\mathbf{R}_x \alpha = \sigma^2 \mathbf{e}$$

avec $\mathbf{R}_x$ la matrice de corrélation, $\mathbf{e}^+ = [1 \ 0 \ \dots \ 0]$ et $\mathbf{a}^+ \mathbf{e} = 1$.

Enfin, le spectre AR s'exprime alors comme :

$$\sigma_{AR}^2(\theta) = \frac{\sigma^2}{\left| \sum_{k=1}^q \alpha_k e^{-j2\pi k u(\theta)} \right|^2}$$

7.1.3 Méthodes sous l'hypothèse de sources ponctuelles

Bien que les méthodes présentées précédemment soient toujours applicables, et disposent d'une certaine robustesse, lorsque le champ est composé d'un ensemble de sources ponctuelles, il est possible d'utiliser des méthodes spécifiques très performantes. Ces méthodes sont dites à haute résolution.

Pour P sources ponctuelles, on a :

$$s(\sigma) = \sum_{p=1}^P s_p \delta(\sigma - \sigma_p)$$

Dans ce cas, le vecteur d'observation recueilli sur le réseau de capteurs s'écrit :

$$\mathbf{x} = \sum_{p=1}^P s_p \mathbf{a}(\theta_p) + \mathbf{b}$$

où θ_p est le paramètre du modèle correspondant à la position σ_p .

Cette dernière relation indique que l'observation non bruitée appartient à un espace vectoriel de dimension P (si les vecteurs $\mathbf{a}(\theta_p)$ sont linéairement indépendants, ce qui est raisonnable). Cette remarque fondamentale est à l'origine des méthodes à haute résolution.

Le vecteur des observations peut également s'exprimer comme :

$$\mathbf{x} = \mathbf{A}\mathbf{s} + \mathbf{b}$$

avec $\mathbf{A} = [\mathbf{a}(\theta_1) \ \dots \ \mathbf{a}(\theta_M)]$ et $\mathbf{s} = [s_1 \ \dots \ s_P]$.

La matrice de corrélation Γ_x des observations s'écrit

$$\begin{aligned} \mathbb{E} [\mathbf{x}\mathbf{x}^+] &= \mathbf{A}\mathbb{E} [\mathbf{s}\mathbf{s}^+] \mathbf{A}^+ + \mathbb{E} [\mathbf{b}\mathbf{b}^+] \\ \Gamma_x &= \mathbf{A}\Gamma_s \mathbf{A}^+ + \Gamma_b = \Gamma_y + \Gamma_b \end{aligned}$$

avec des notations évidentes.

La matrice Γ_s , de dimension $P \times P$ est de rang P , si le nombre de capteurs est supérieur à P .

La matrice Γ_y est également de rang P et son noyau est de dimension $(M - P)$. Cette matrice, la matrice de corrélation de l'observation non bruitée, peut être décomposée en éléments propres selon :

$$\Gamma_y = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^+$$

où $\mathbf{U}$ est la matrice des vecteurs propres et $\mathbf{\Lambda}$ la matrice diagonale des valeurs propres. Les P premiers vecteurs propres définissent l'espace signal et les $M - P$ suivants, l'espace bruit. Les matrices de corrélation étant définies non négatives, toutes les valeurs propres sont positives ou nulles. La matrice des vecteurs propres, $\mathbf{U}$ est unitaire, et vérifie ainsi $\mathbf{U}\mathbf{U}^+ = \mathbf{I}_M$.

Lorsque le bruit est spatialement blanc, Γ_b est de la forme $\Gamma_b = \sigma^2 \mathbf{I}_M$, avec $\mathbf{I}_M$ la matrice identité de dimension $M \times M$. On peut alors écrire Γ_x sous la forme :

$$\Gamma_x = \Gamma_y + \Gamma_b = \Gamma_y = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^+ + \sigma^2 \mathbf{U}\mathbf{U}^+$$

ou encore

$$\Gamma_x = \mathbf{U} (\mathbf{\Lambda} + \sigma^2 \mathbf{I}_M) \mathbf{U}^+$$

Ceci fournit la décomposition en éléments propres de Γ_x , qui montre que les vecteurs propres sont inchangés, et que les valeurs propres sont augmentées de σ^2 . Les valeurs propres correspondant à un « espace bruit », qui étaient nulles, valent maintenant σ^2 .

Par contre, si le bruit est corrélé, il est nécessaire de soustraire à Γ_x une estimée Γ_b de la matrice de corrélation du bruit.

La remarque fondamentale est alors que les sous-espaces « signal » et « bruit » sont orthogonaux. Tout vecteur $\mathbf{v}_b$ de l'espace bruit est ainsi orthogonal à un vecteur de l'espace signal. Or l'espace signal est engendré par les P vecteurs directionnels $\mathbf{a}(\theta_i)$. On a donc :

$$\mathbf{a}(\theta_i)^+ \mathbf{v}_b = 0$$

Dès lors, il suffit de sélectionner un ou plusieurs vecteurs du sous-espace bruit pour tester cette orthogonalité pour tous les paramètres θ . Pour cela, plutôt que de rechercher les passages par zéro du produit scalaire $\mathbf{a}(\theta_i)^+ \mathbf{v}_b$ on peut utiliser la fonction de détection, ou de discrimination :

$$d(\theta) = \frac{1}{|\mathbf{a}(\theta_i)^+ \mathbf{v}_b|^2}$$

qui est maximale (infinie) lorsque θ prend pour valeur l'un des vecteurs paramètres des sources. En pratique, la fonction de détection est seulement maximale (et non infinie) au voisinage des solutions.

Bien sûr, il est nécessaire de connaître le nombre de sources, qui définit la dimension de l'espace signal, et partant, de l'espace bruit. Ce nombre n'est le plus souvent pas connu, mais doit être déterminé à partir des données. Une méthode simple est la suivante : on a vu qu'en présence de bruit blanc, les valeurs propres sont augmentées de σ^2 . Comme les valeurs propres initiales sont positives (la matrice de corrélation est définie non négative), les valeurs propres de l'espace bruit sont les plus faibles et égales à σ^2 . Il suffit donc de rechercher les $M - P$ valeurs propres les plus faibles qui soient toutes égales à une même valeur. En pratique, on n'obtient pas un plancher de valeurs toutes égales entre elles, mais les valeurs propres sont toujours légèrement décroissantes ; on recherche donc plutôt une rupture de pente des valeurs propres classées par ordre décroissant. D'autres techniques fondées sur les critères AIC ou MDL peuvent également être utilisées.

Lorsque les deux sous-espaces sont définis, il reste à choisir un vecteur $\mathbf{v}_b$ du sous-espace bruit. Ce choix, arbitraire, donne lieu à plusieurs variantes.

■ Le vecteur orthogonal

On peut choisir comme représentant une combinaison linéaire particulière de tous les vecteurs du sous-espace bruit.

Tufts et Kumaresan ont préconisé de choisir un vecteur de norme minimale, car celle-ci conduit à une estimée de variance minimale. Cette méthode est aussi appelée MIN-NORM.

On sélectionne ainsi la solution du problème de minimisation suivant :

$$\begin{cases} \min \mathbf{v}^+ \mathbf{v} & \text{norme minimale} \\ \text{avec } \Pi \mathbf{v} = \mathbf{v} & \mathbf{v} \text{ appartient au sous espace bruit} \\ \text{et } \mathbf{v}^+ \mathbf{e} = 1 & 1^{\text{re}} \text{ composante à 1} \end{cases}$$

où l'on a utilisé $\Pi = \sum_{i=P+1}^M \mathbf{v}_i \mathbf{v}_i^+$ avec $\mathbf{v}_i$ les vecteurs propres du sous-espace bruit, et $\mathbf{e}^+ = \begin{bmatrix} 1 & 0 & \cdots & 0 \end{bmatrix}$. Le Lagrangien du problème de minimisation s'écrit :

$$L(\mathbf{v}, \lambda) = \mathbf{v}^+ \mathbf{v} + \lambda (\mathbf{v}^+ \mathbf{e} - 1)$$

En tenant compte de $\mathbf{v}^+ \mathbf{e} = \mathbf{v}^+ \Pi \mathbf{e}$, on obtient $\mathbf{v} = \lambda \Pi \mathbf{e} / 2$, et la vérification de la contrainte $\mathbf{v}^+ \mathbf{e} = 1$ conduit à

$$\mathbf{v} = \frac{\Pi \mathbf{e}}{\mathbf{e}^+ \Pi \mathbf{e}}$$

■ Le projecteur orthogonal, ou goniomètre, ou MUSIC

Dans cette célèbre méthode introduite indépendamment par Schmidt (1979) sous l'acronyme de MUSIC (*MUltiple SIgnal Classification*) et Bienvenu et Kopp (1980) sous le nom de goniomètre, on utilise l'ensemble des vecteurs du sous-espace bruit pour tester l'orthogonalité :

$$\sum_{i=P+1}^M |\mathbf{v}_i^+ \mathbf{a}(\theta)|^2 = 0$$

En introduisant à nouveau une fonction de détection, on a :

$$d(\theta) = \frac{1}{\sum_{i=P+1}^M |\mathbf{v}_i^+ \mathbf{a}(\theta)|^2} = \frac{1}{\mathbf{a}(\theta)^+ \Pi \mathbf{a}(\theta)}$$

Dans le cas d'une antenne rectiligne uniforme et d'un modèle de propagation par ondes planes, le vecteur directionnel est de la forme :

$$\mathbf{a}(\theta)^T = \begin{bmatrix} 1 & e^{j2\pi u(\theta)} & \dots & e^{j2\pi(M-1)u(\theta)} \end{bmatrix}$$

On peut alors donner une solution explicite à la relation d'orthogonalité, puisque :

$$\mathbf{a}(\theta)^+ \mathbf{v} = \sum_{i=0}^{M-1} v_i z^{-i} = 0$$

pour les $z_k = e^{j2\pi u(\theta_k)}$ correspondant à des sources, et avec $\mathbf{v}$ un vecteur du sous-espace bruit. Ainsi il suffit de rechercher les P racines d'un polynôme de degré M (on retient les racines de module 1), pour en déduire les valeurs des paramètres θ des P sources. Ceci s'étend bien sûr au cas d'une combinaison des vecteurs du sous-espace bruit, et si on recherche par exemple les racines du polynôme $\begin{bmatrix} 1 & z & \dots & z^{M-1} \end{bmatrix} \Pi \begin{bmatrix} 1 & z & \dots & z^{M-1} \end{bmatrix}^+$ associé à MUSIC, on obtient la méthode Root-MUSIC.

Dans le cas où le sous-espace bruit est de dimension 1, soit $M = P + 1$, la méthode du vecteur orthogonal est la même que celle du projecteur, et est appelée Méthode de Pisarenko (1973).

■ Illustrations

Pour le modèle plan, exemple (A), on obtient les résultats présentés sur la figure 7.13, avec $M = 10$.

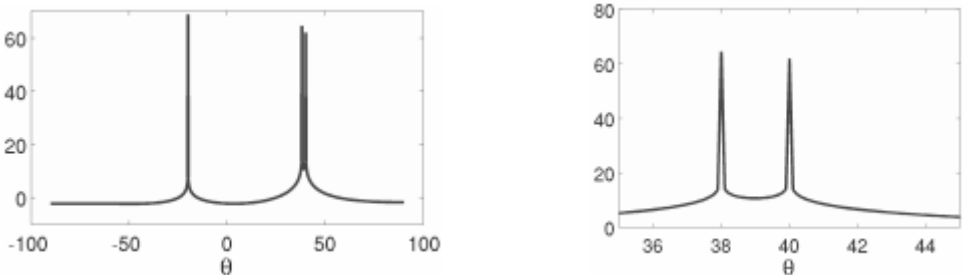


Figure 7.13 – Sortie de la méthode MUSIC pour l'exemple (A), avec $M = 10$. Sortie complète et détail.

Pour l'exemple (B), la méthode MUSIC est également applicable, en prenant pour vecteur directionnel $\mathbf{a}(\theta)$ l'expression obtenue pour le modèle de propagation par ondes circulaires. Les

résultats correspondants sont donnés figure 7.14. On obtient une très bonne localisation, mais au prix, il est vrai, d'une complexité accrue ainsi que d'une sensibilité aux hypothèses.

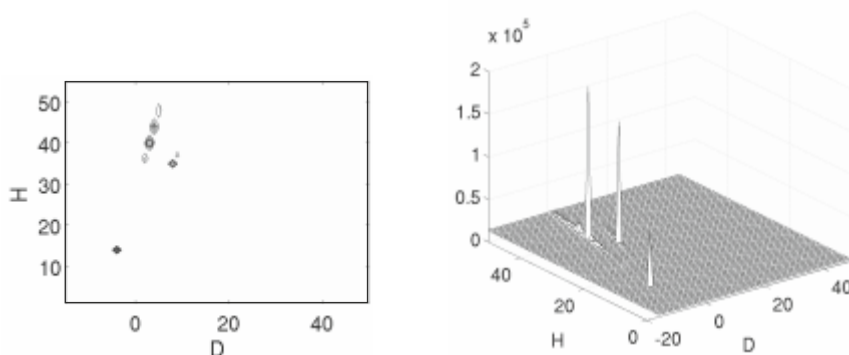


Figure 7.14 – Sortie de la méthode MUSIC pour l'exemple (B), avec $M = 15$.

7.2 Diversité

La propagation d'une onde électromagnétique peut être soumise à de multiples contraintes. L'absence d'obstacle entre l'émetteur et le récepteur constitue le cas le plus favorable et la propagation s'effectue alors en ligne droite. On ne retrouve pas cette situation à l'extérieur ou à l'intérieur des bâtiments où la présence d'obstacles crée des phénomènes de réflexion et de diffraction. Ces conditions conduisent ainsi à la propagation par trajets multiples. Ceci se traduit par un évanouissement rapide de la puissance reçue et par la dispersion temporelle du signal au niveau de l'antenne réceptrice. La conséquence directe en est la limitation de la portée et du débit de la transmission, alors qu'aujourd'hui, la demande de hauts débits exigés dans les communications sans fils de dernières générations se fait de plus en plus pressante. Pour contrer les effets des trajets multiples, des techniques de diversité ont été mises au point.

La figure 7.15 synthétise l'ensemble des contraintes que subit le signal lors de la transmission.

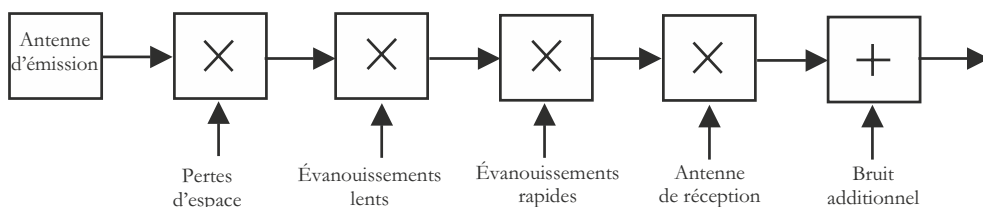


Figure 7.15 – Modélisation du canal de propagation.

On constate sur cette figure que trois modèles complémentaires existent et permettent de prédire la puissance captée par le récepteur. Il s'agit :

- de l'atténuation de l'onde due à la distance (pertes d'espace ou affaiblissement de parcours), qui traduit le fait que la puissance captée par le récepteur dépend de la distance qui le sépare de l'émetteur,

- des variations d'amplitudes dues aux obstacles sur le trajet (effet de masquage ou évanouissements lents)
- des variations d'amplitude et de phase dues aux trajets multiples (évanouissements rapides) qui sont liées au fait que l'onde émise peut suivre plusieurs chemins jusqu'au récepteur.

Pour limiter le brouillage qui résulte des effets ci-dessus cités, on utilise des antennes dites à diversité. Il s'agit d'éléments rayonnants reconfigurables dont le principe consiste à multiplier le nombre de canaux indépendants disponibles en réception. Ainsi, à partir de ces signaux et à l'aide d'une méthode de recombinaison basée sur un critère de puissance, il est possible de reconstituer un signal à plus fort rapport signal sur bruit ce qui permet d'améliorer la fiabilité et la qualité de la transmission. Nous verrons que la diversité dans une antenne peut prendre de multiples formes que nous allons détailler un peu plus loin.

On peut noter sur la figure 7.15 que l'association des termes due aux pertes d'espace et aux effets de masquage constitue le *Long Term Fading* compte tenu des fluctuations lentes du signal résultant. On l'appelle aussi *Local Mean* ou Moyenne Locale. Ces effets sont mis en évidence sur la figure 7.16 à l'aide d'un signal prélevé dans un environnement typiquement indoor. Sur cette figure sont superposés à la fois le signal original, d'où l'on peut constater la présence des évanouissements rapides et profonds, et la moyenne locale à variation plus lente déduite de ce signal par une opération de filtrage.

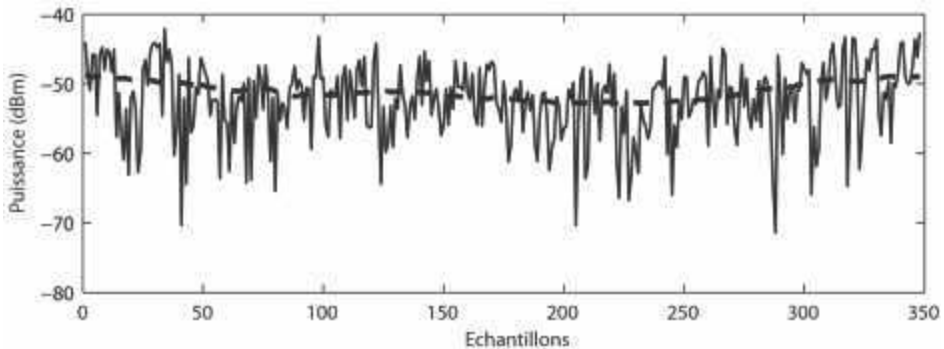


Figure 7.16 – Mise en évidence des phénomènes physiques sur un signal prélevé en environnement indoor.

Si les propriétés du terme dû aux pertes d'espace sont déterministes et peuvent être approchées par de nombreux modèles, en revanche, les évanouissements lents et rapides sont de nature aléatoires. Ils peuvent être décrits d'un point de vue statistique par des lois de type Rayleigh, Rice ou Nakagami selon la position relative de l'émetteur et du récepteur et les caractéristiques de l'environnement.

Enfin, l'évaluation des performances des antennes à diversité est effectuée par la détermination du coefficient de corrélation entre les différents signaux reçus sur les canaux et du gain en diversité du système qui, en première approximation, donne une indication pertinente sur l'amélioration du rapport signal sur bruit.

Pour un nombre de canaux donnés et une méthode de recombinaison choisie, les performances du système seront d'autant meilleures que la corrélation entre les voies (ou branches) sera faible. Il faudra également s'assurer que les puissances moyennes reçues sur chacune des branches sont équivalentes. Un troisième paramètre permet de quantifier cette différence : le *Power Imbalance* ou déséquilibre de puissance entre les branches.

Nous proposons dans ce qui suit une description des différents types de diversité que l'on peut rencontrer dans les dispositifs de transmission.

7.2.1 Méthode de diversité

■ Diversité d'espace

En diversité d'espace, le message est transmis par un émetteur unique vers plusieurs récepteurs distincts dont les aériens sont suffisamment espacés pour que la transmission respecte dans chaque cas une statistique d'évanouissement indépendante. Il est cependant nécessaire de choisir judicieusement la distance de séparation entre les deux antennes du récepteur de façon à obtenir deux signaux non corrélés. Cet espacement optimal est obtenu par l'évaluation du coefficient de corrélation en fonction de la séparation entre les antennes. À ce titre, R.H. Clarke a établi la relation simplifiée entre l'enveloppe du coefficient de corrélation ρ_e et la distance entre deux antennes dipôles demi-onde. Elle s'exprime de la façon suivante :

$$\rho_e \cong J_0^2 \left(\frac{2 \pi d}{\lambda} \right)$$

où J_0 représente la fonction de Bessel de première espèce et d'ordre 0, d la distance entre antennes et λ la longueur d'onde.

La relation ci-dessus ne donne qu'une approximation dans la mesure où l'on suppose que la distribution des angles d'arrivées des ondes uniformément répartis, d'égales polarisations et toutes contenues dans l'azimut.

On remarque sur la figure 7.17 que pour $d > \lambda/4$, le coefficient de corrélation théorique ne dépasse pas 0,2.

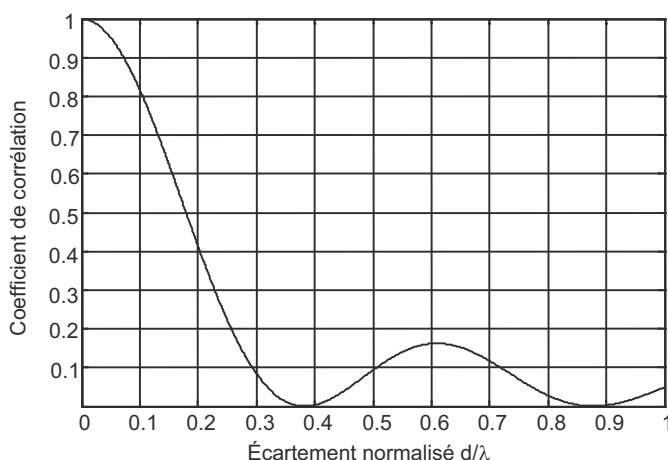


Figure 7.17 – Coefficients de corrélation théorique et mesurés en fonction de l'écartement entre antennes.

Notons enfin que cette distance peut subir des fluctuations sensibles car elle est fortement dépendante des caractéristiques sur la distribution statistique des signaux dans l'environnement de mesures et du niveau moyen de puissance reçue sur chacune des branches.

■ Diversité de polarisation

Nous venons de voir que la diversité d'espace nécessite d'observer une distance d'au moins $\lambda/4$ entre antennes. D'un point de vue pratique, pour assurer l'indépendance entre branches, cet espacement doit être d'autant plus grand que la dispersion angulaire des signaux venant de l'émetteur sera faible. On retrouve cette situation dans les antennes pour station de base où

la distance entre éléments peut atteindre la dizaine de longueurs d'onde ce qui rend les systèmes peu discrets et difficilement intégrables.

La diversité de polarisation constitue une alternative intéressante à la diversité d'espace. On remarque en effet qu'au cours de la propagation entre le terminal et la station de base, la polarisation de l'onde subit des modifications plus ou moins importantes selon la nature de l'environnement. En particulier, une composante orthogonale à la polarisation principale apparaît et elle possède une statistique d'évanouissement indépendante de la polarisation principale, ce qui représente déjà un avantage en termes de gain en diversité. Enfin, la diversité de polarisation évite l'espacement physique entre antennes. En effet, l'idée consiste à concevoir deux antennes co-situées, sensibles chacune aux deux polarisations.

La figure 7.18 illustre un tel système ici obtenu à partir d'une antenne patch carrée alimentée par deux fentes orthogonales en croix. Chacune des fentes est couplée à une ligne d'alimentation microruban. Une excitation sur l'une des deux lignes imprimées favorisera le rayonnement d'un champ électrique orienté parallèlement à la ligne excitée. Il est ainsi possible, au moyen d'un dispositif de contrôle, de commuter électroniquement la polarisation de verticale à horizontale. D'autres antennes permettent des commutations entre les deux sens de polarisation circulaire.

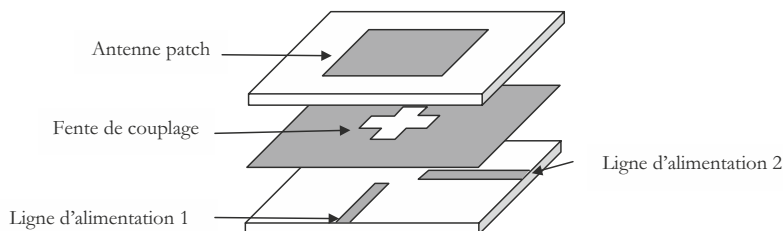


Figure 7.18 – Exemple d'antennes co-situées à polarisations orthogonales

Au même titre que la diversité d'espace, les performances en termes de gain en diversité seront d'autant meilleures que les puissances moyennes reçues sur chacune des branches seront équilibrées. La caractérisation préalable de l'environnement de mesure demeure donc incontournable. En effet, l'expérience montre qu'un environnement riche en réflexion permet, par trajets multiples, une dépolarisation suffisante ce qui améliore sensiblement le gain en diversité.

■ Diversité de diagramme ou diversité angulaire

Nous avons vu que les contraintes d'espacement entre antennes dans les systèmes de communications pouvaient constituer un handicap dans le cas de la diversité d'espace. La diversité de polarisation constitue une solution au même titre que la diversité angulaire.

Dans ce dernier cas, on déterminera un certain nombre de directions d'arrivée que l'on pourra isoler à la réception en utilisant des antennes directives. Pour couvrir l'ensemble des trajets, une solution consiste à multiplier le nombre de récepteurs, chacun étant associé à un trajet. Cela se traduit par une augmentation de l'encombrement. On préfère utiliser des antennes intelligentes capables de commuter de façon électronique les diagrammes de rayonnement émis/reçus pour les orienter sur la direction qui présente le plus fort rapport signal sur bruit. La diversité angulaire permet donc d'améliorer la qualité de la transmission et la sécurité en réduisant l'interférence des sources. Enfin, cette technique peut être appliquée aussi bien au niveau de l'émetteur (station de base) que du terminal mobile.

D'un point de vue pratique, la diversité angulaire peut être réalisée en utilisant le concept des antennes à éléments parasites. La figure 7.19 illustre cette méthode. L'idée consiste à utiliser un réseau d'antennes (deux antennes dans cet exemple) dont l'espacement d sera choisi de façon à

obtenir un couplage important. Une seule des antennes est alimentée tandis qu'une charge réactive (capacitive ou inductive) (Z_1 , Z_2) commutable électroniquement sera placée sur l'élément couplé de façon à rendre l'élément rayonnant principal réflecteur ou directeur.

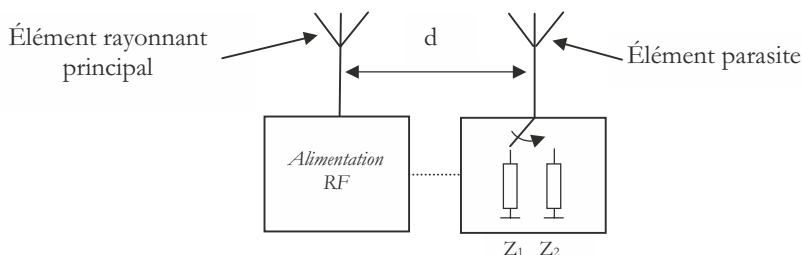


Figure 7.19 – Principe de l'antenne à élément parasite couplé

La commutation de charge permet ainsi de réaliser la commutation sectorielle de diagramme comme le montre la figure 7.20.

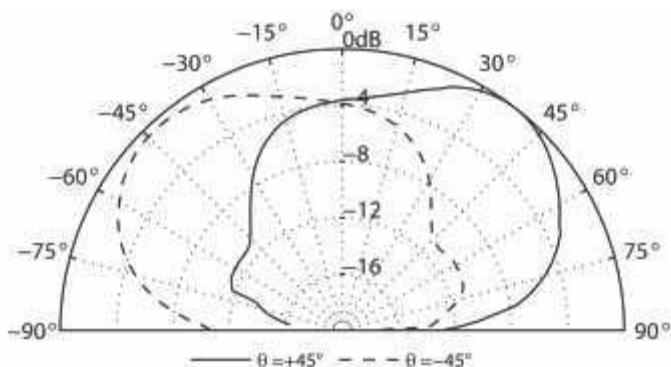


Figure 7.20 – Champ électrique rayonné pour deux valeurs de charge (Z_1 , Z_2)

Un des intérêts de ce concept vient du fait que la recherche d'un couplage élevé impose un écartement réduit entre antennes ($< 0,2\lambda_0$) ce qui rend la structure compacte et plus facilement intégrable.

■ Diversité en fréquence

En diversité fréquentielle, le signal d'émission est envoyé à l'aide de deux fréquences porteuses distinctes de sorte que les statistiques d'évanouissement des signaux reçus soient indépendantes. Appliquée tout d'abord aux radars pour résoudre le problème de la fluctuation des cibles illuminées, la diversité de fréquence a ensuite connu un fort développement dans les radiocommunications cellulaires avec des stations mobiles. Le standard GSM 900 et la norme DCS 1800 utilisent la diversité de fréquence pour lutter contre les effets d'évanouissements.

Plus précisément, les communications dans la norme GSM actuelle s'effectuent à partir d'une liste de fréquences autour de la fréquence porteuse. Lors de la communication, la fréquence de la porteuse varie à un rythme rapide « sautant » ainsi d'une fréquence à l'autre à l'intérieur de la liste. On réalise ainsi une transmission par saut de fréquence (*frequency hopping*) ce qui conduit de fait à un procédé de transmission par paquets d'information de longueur limitée.

■ Diversité en temps

Dans la diversité en temps, les mêmes signaux sont envoyés sur le canal à des intervalles de temps différents. Cet intervalle étant fonction du taux d'affaiblissement et de la vitesse de déplacement du mobile. Cependant, l'efficacité de la diversité en temps est limitée lorsque le mobile n'est pas en mouvement ce qui n'est pas le cas des autres techniques de diversité.

7.2.2 Techniques de recombinaison des signaux en réception

Les méthodes de diversité précédemment citées permettent l'obtention de deux ou plusieurs signaux. À partir de ces signaux, pour reconstituer un signal dont les propriétés en termes de rapport signal sur bruit sont plus élevées, il est nécessaire d'utiliser des techniques de recombinaison adéquates. Cette recombinaison s'effectue selon deux principes :

- la combinaison basée sur la commutation entre les signaux reçus par le récepteur. Cette commutation s'effectuant selon un algorithme basé sur le niveau du signal reçu sur chacune des branches,
- la combinaison basée sur la sommation des différents signaux, système plus complexe mais cependant plus performant en termes de qualité du signal obtenu.

Il existe actuellement trois techniques de recombinaisons qui peuvent être classées de la façon suivante.

□ Selection Combining (SC)

C'est la plus simple des techniques de diversité. Elle consiste, à partir de M antennes en réception, à sélectionner instantanément la branche qui présente le plus fort niveau du signal sur bruit instantané et à la connecter au récepteur. Ainsi, lorsqu'un signal s'évanouit, un autre est sélectionné avec un niveau supérieur. La qualité du signal résultant sera d'autant meilleure que le nombre d'antennes M est élevé, multipliant ainsi le nombre de branches à la réception. Une variante de cette méthode appelée *Switched Combining* consiste à ne commuter d'une branche à l'autre que lorsque le rapport signal sur bruit instantané de la branche sélectionnée est inférieur à un seuil prédéterminé.

□ Maximal Ratio Combining (MRC)

Il s'agit ici de la technique la plus performante. Un traitement adaptatif est ici appliqué à l'ensemble des signaux issus des M antennes. Chacun des signaux est pondéré à l'aide d'un gain complexe, lui-même calculé à partir du rapport signal sur bruit instantané relevé sur chacune des branches en réception. Les signaux obtenus sont ensuite mis en phase avant sommation afin de maximiser le niveau du signal recombinaison. Ce type de recombinaison nécessite cependant une architecture matérielle complexe. En effet, une mise à jour du gain complexe est nécessaire et il faut s'assurer que les signaux sont correctement traités en phase.

□ Equal Gain Combining (EGC)

Cette technique se situe entre les deux méthodes précédemment décrites en termes de performances et de complexité. C'est une alternative au Maximal Ratio Combining sans traitement adaptatif. Les gains sur chacune des branches sont ici constants et ne dépendent pas du niveau du signal reçu. La procédure consiste en une remise en phase et une sommation simple des signaux à la réception. Les performances sont cependant proches de celles obtenues avec le Maximal Ratio Combining avec une complexité architecturale moindre.

7.2.3 Évaluation des performances des systèmes à diversité : gain en diversité

Le gain en diversité est le paramètre qui permet de quantifier l'amélioration, en termes de rapport signal sur bruit, obtenu avec l'utilisation d'un système à diversité. Nous avons précisé plus haut

que, pour une technique de recombinaison choisie (SC, MRC, EGC), cette amélioration était d'autant meilleure que le coefficient de corrélation entre les branches était faible et que le niveau de puissance moyen reçu sur chacune des branches était équivalent.

Le gain en diversité est défini à partir des fonctions de répartition des i branches et de la branche issue de la recombinaison des signaux.

Pour déterminer le gain en diversité, on introduit pour les i branches, le rapport signal sur bruit instantané (noté γ_i) ainsi que la valeur moyenne de ce rapport (noté Γ_i). On fait de même pour la branche issue de la recombinaison (respectivement γ_c et Γ_c). Le gain est alors déterminé par l'écart entre la branche qui présente le rapport signal sur bruit normalisé, γ_i/Γ_i , le plus élevé et le rapport γ_c/Γ_c , ceci pour un niveau de probabilité donné γ_s/Γ (généralement 1 ou 10 %).

$$DG \text{ (dB)} = \left[\frac{\gamma_c}{\Gamma_c} \text{ (dB)} - \frac{\gamma_i}{\Gamma_i} \text{ (dB)} \right]_{P(\gamma_c < \frac{\gamma_i}{\Gamma_i})}$$

À titre d'exemple, la figure 7.21 représente l'évolution des fonctions de répartition pour une branche et celles obtenues pour deux, trois et quatre branches recombinaison par Selection Combining. En supposant une distribution des enveloppes de type Rayleigh et non corrélées entre elles, on remarque que pour une branche, 1 % des échantillons ont des valeurs du rapport signal sur bruit normalisé inférieures à -20 dB. Si l'on considère maintenant deux branches pour le même niveau de probabilité, cette valeur passe à -10 dB ce qui représente un gain en diversité de 10 dB.

Les performances des systèmes à diversité dépendent, entre autres, du type de recombinaison et du nombre de branches. On remarque sur la figure 7.21 que le gain atteint respectivement 14 et 16 dB pour trois et quatre branches et marquera une stabilisation au-delà de six branches. Enfin, le coefficient de corrélation entre branches, idéalement nul, n'affectera les performances du système que si sa valeur devient supérieure à 0,5.

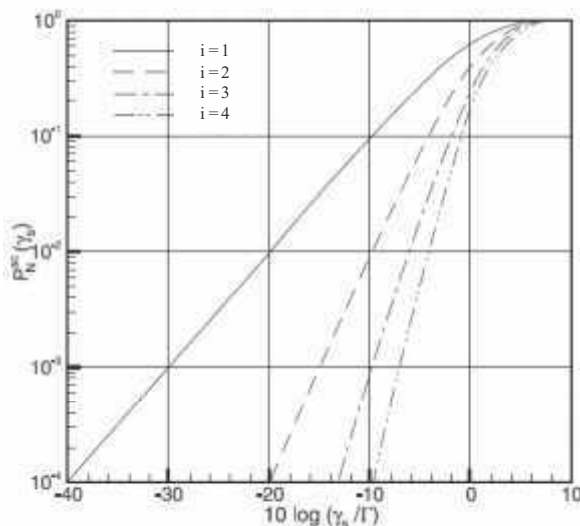


Figure 7.21 – Fonctions de répartition obtenues pour 1, 2, 3 et 4 branches recombinaison par Selection Combining.

7.3 Systèmes MIMO

7.3.1 Introduction

Cette section est principalement consacrée aux systèmes MIMO, le sigle MIMO (*Multiple Input Multiple Output*) représentant les systèmes dans lesquels on utilise plusieurs antennes en émission et plusieurs antennes en réception.

On commence par exposer les principales caractéristiques des canaux radio-mobiles ainsi que les conséquences des trajets multiples et de la mobilité en termes d'évanouissements (on dira par la suite fading) et d'effet Doppler. On présente les performances d'une liaison numérique en présence de fading et on montre l'intérêt des techniques de diversité. On décrit ensuite le multiplexage spatial et les principes des systèmes MIMO, puis on donne quelques résultats sur la capacité et la probabilité de coupure de ces systèmes. On termine par les méthodes permettant de se rapprocher des performances théoriques en abordant succinctement le codage spatio-temporel et les récepteurs optimaux ou sous-optimaux. On conclut en expliquant le compromis entre le gain de diversité et le gain de multiplexage (ou de débit).

S'agissant d'une introduction aux systèmes MIMO, certains aspects comme les systèmes multi-utilisateurs ne sont pas abordés.

L'utilisation de plusieurs antennes en réception est mise en œuvre depuis longtemps afin de tirer partie de la diversité ainsi obtenue (on peut parler dans ce cas de systèmes SIMO : *Single Input Multiple Output*). L'utilisation de plusieurs antennes en émission (systèmes MISO : *Multiple Input Single Output*) est appliquée de manière classique pour la formation de faisceau qui associe plusieurs éléments rayonnants dans une antenne réseau de façon à augmenter le gain de rayonnement dans une direction fixe ou variable. Elle peut aussi apporter une diversité en émission par l'utilisation de codes spatio-temporels par exemple.

Par la suite on note n_t le nombre d'antennes en émission et n_r le nombre d'antennes en réception. L'utilisation de plusieurs antennes en émission et en réception (MIMO) peut apporter deux types de gains pour la transmission dans un canal présentant du fading, un gain en débit et un gain en diversité.

- **Gain en débit et multiplexage spatial** : en émission, on peut créer un ensemble de canaux, chacun transportant un flux élémentaire de données. On peut ainsi augmenter le débit. On parle de multiplexage spatial. Les flux élémentaires peuvent provenir d'un flux global original qui est décomposé à l'aide d'un convertisseur série-parallèle à l'émission puis qui est recomposé par un convertisseur parallèle-série en réception. Pour séparer les flux en réception, il faut au moins autant d'antennes qu'en émission. Le gain sur le débit est limité par le minimum du nombre d'antennes en émission et en réception $\min(n_t, n_r)$.
- **Gain en diversité** : dans un canal avec fading, les signaux se propagent sur différents trajets et le signal reçu est une combinaison de différentes versions, éventuellement distordues et bruitées, d'un même signal original. Un même signal peut au maximum être transmis par $n_t n_r$ chemins différents et le gain de diversité est au plus égal à $n_t n_r$.

Il existe un compromis entre le gain en débit et le gain en diversité. Les systèmes MIMO tirent profit de ces deux types de gains. Leurs performances dépendent des caractéristiques du canal et du nombre d'antennes utilisées en émission et en réception.

Le développement des techniques MIMO a constitué une innovation majeure pour les systèmes de communications mobiles. La technologie MIMO s'est imposée dans de nombreuses applications comme les réseaux locaux sans fil ou WLAN (standard IEEE 801.11n pour WIFI), les réseaux métropolitains ou WMAN (standards IEEE 802.16-2004 et IEEE 802.16e pour WiMAX ou IEEE 802.20), les réseaux sans fil régionaux ou WRAN (standard IEEE 802.22), les systèmes cellulaires de troisième génération et le standard 3GPP LTE (*Long Term Evolution*) par exemple.

Le développement des travaux sur les systèmes MIMO remonte au milieu des années quatre-vingt-dix, avec en particulier les travaux de Telatar (*Bell Laboratories*, 1995) sur le calcul de la capacité d'un canal MIMO, l'introduction de la notion de multiplexage spatial par A. Paulraj (université de Stanford, 1994), les travaux de l'équipe de Foschini (*Bell Laboratories*, 1996) sur les récepteurs BLAST (*Bell Laboratories Layered space time architecture*) et les travaux de l'équipe de Tarokh (*Bell Laboratories*, 1998) sur le codage spatio-temporel qui ont généralisé les concepts introduits précédemment par Alamouti (AT&T Wireless services 1998).

7.3.2 Canaux radio-mobiles : évanouissements et effet Doppler

Dans le cas d'une transmission sur un canal radio-mobile, le signal est dégradé par différents phénomènes : l'ajout de bruit mais aussi les phénomènes de fading dus aux trajets multiples et au déplacement relatif de l'émetteur et du récepteur (effet Doppler).

Dans la section sur la diversité, on a présenté les notions de pertes d'espace dues à la distance, d'évanouissements lents (ou à grande échelle) dus aux obstacles sur le trajet et d'évanouissements rapides (ou à petite échelle) dus aux trajets multiples. On parle de fading à grande échelle (*large scale fading*) pour désigner les pertes d'espace et les évanouissements dus aux obstacles et de fading à petite échelle (*small scale multipath fading*) pour désigner le fading rapide dû aux trajets multiples. Pour un mobile en déplacement (piéton, voiture...), les constantes de temps associées aux évanouissements lents (*large scale fading*) sont très importantes (de l'ordre de plusieurs minutes) et ce phénomène joue peu de rôle sur la conception de la couche physique de communication. Aussi nous centrons-nous ici sur les évanouissements rapides. Par la suite, on utilisera souvent le terme canal pour indiquer le canal radio-mobile.

■ Paramètres du canal

Pour caractériser le comportement temporel et fréquentiel du canal en lien avec les évanouissements rapides, on utilise des paramètres tels que :

- Bande de cohérence : B_c ,
- Temps de cohérence : T_c
- Étalement temporel (*delay spread*) : T_m ,
- Étalement Doppler (*Doppler spread*) : B_d .

On va voir que ces paramètres ne sont pas indépendants.

La bande de cohérence est liée à l'étalement temporel $B_c \propto 1/T_m$ (où le symbole $\propto$ signifie proportionnel à).

De même l'étalement Doppler est liée au temps de cohérence $B_d \propto 1/T_c$.

Les caractéristiques du canal sont données de façon relative aux caractéristiques du signal à transmettre, en particulier sa largeur de bande B et la durée des symboles T_S (avec généralement $B \approx 1/T_S$).

On verra par la suite comment on définit un canal sélectif en fréquence ou en temps selon les valeurs de B , B_c , B_d , T_S , T_c , T_m .

L'existence de trajets multiples entre un émetteur et un récepteur, typiques des communications radios à l'intérieur des bâtiments et des communications cellulaires, conduit à des interférences constructives ou destructives entre les signaux reçus.

Lors de son déplacement, le mobile rencontre successivement des nœuds et des ventres d'interférence qui sont séparés de distances de l'ordre de la demi-longueur d'onde de la porteuse utilisée dans le système de communication. Le canal est non stationnaire parce que le mobile se déplace et que son environnement change. En première approximation, on peut considérer que le canal est linéaire.

Dans un premier temps, on s'intéresse au canal SISO puis on étend les résultats au canal MIMO. On commence par présenter les paramètres du canal à partir d'exemples élémentaires.

■ Étalement temporel et bande de cohérence, sélectivité en fréquence

L'étalement temporel mesure l'écart maximal entre les retards $\tau_i(t)$.

Prenons l'exemple d'un canal linéaire stationnaire à deux trajets caractérisés par une atténuation réelle positive a_i , un retard τ_i et un déphasage ϕ_i (avec $i = 0$ ou 1). Pour une entrée $x(t)$ cosinusoidale à la fréquence f_c , d'amplitude A , la sortie $y(t)$ vaut :

$$y(t) = a_0 A \cos(2\pi f_c(t - \tau_0) + \phi_0) + a_1 A \cos(2\pi f_c(t - \tau_1) + \phi_1).$$

L'étalement temporel vaut $T_m = |\tau_1 - \tau_0|$.

La fonction de transfert $H(f)$ du canal s'écrit :

$$H(f) = a_0 e^{-j2\pi f \tau_0 + j\phi_0} + a_1 e^{-j2\pi f \tau_1 + j\phi_1}.$$

Le module carré de $H(f)$ est donné par :

$$|H(f)|^2 = a_0^2 + a_1^2 + 2a_0 a_1 \cos(2\pi f T_m + \phi_0 - \phi_1).$$

Cette fonction est de période $1/T_m$. La fonction de transfert fréquentielle présente donc périodiquement des maxima d'amplitude $|a_0 + a_1|$ et des minima d'amplitude $|a_0 - a_1|$ (voir figure 7.22).

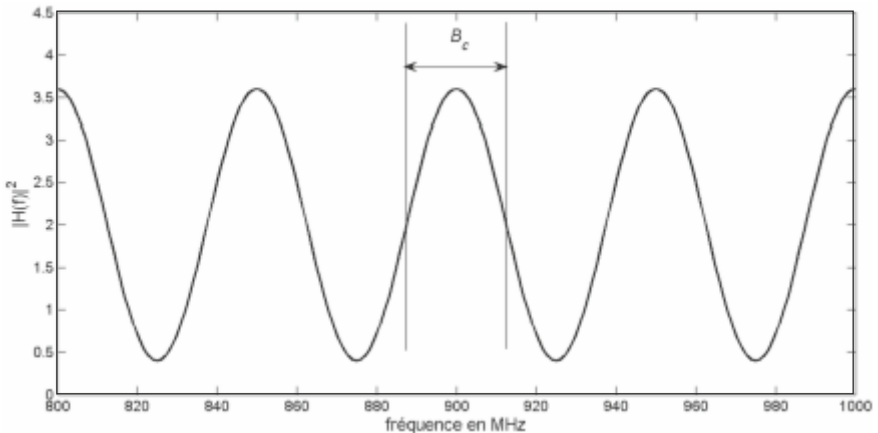


Figure 7.22 – Fonction de transfert fréquentielle d'un canal à deux trajets. Bande de cohérence.

Cet exemple simple permet de comprendre le lien entre la bande cohérence B_c et l'étalement temporel T_m , et de vérifier que $B_c \approx 1/(2T_m)$. En effet, on peut considérer que la fonction de transfert du canal est à peu près constante sur une largeur de bande (bande de cohérence) inférieure à $1/(2T_m)$ (qui est dans l'exemple l'écart entre un minimum et un maximum successifs de la fonction de transfert).

On dit que le canal est sélectif en fréquence si la bande de cohérence B_c est inférieure à la largeur de bande du signal B ($B_c \ll B$) ou de façon équivalente si la durée des symboles est plus courte que l'étalement temporel du canal ($T_s \ll T_m$). On parle aussi de canal dispersif en temps. Dans ce cas, la fonction de transfert fréquentielle du canal fluctue sur la plage de fréquence du signal et les différents symboles reçus interfèrent entre eux. Pour les transmissions à haut débit, on doit généralement considérer que le canal est sélectif en fréquence.

■ Étalement Doppler et temps de cohérence

On considère maintenant un exemple simple de canal variant en temps pour introduire les notions de temps de cohérence T_c et d'étalement Doppler B_d .

Soit un mobile portant un récepteur se déplaçant à la vitesse v entre un émetteur fixe et un mur sur lequel se réfléchit l'onde radio. Le signal reçu par le mobile est la somme du signal de l'émetteur et du signal réfléchi sur le mur (figure 7.23). La réflexion sur le mur déphase le signal de 180° .

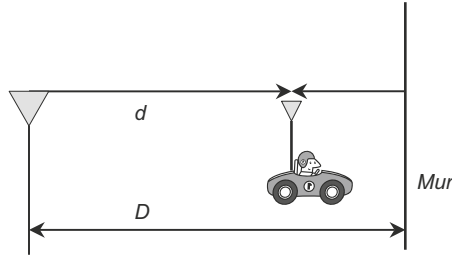


Figure 7.23 – Cas d'une propagation avec une onde directe et une onde réfléchie entre une antenne fixe et une antenne mobile.

On suppose que le signal $x(t)$ émis par l'antenne est cosinusoidal :

$$x(t) = A \cos(2\pi f t)$$

Pour une propagation en espace libre, l'amplitude du champ reçu décroît de façon inversement proportionnelle à la distance, et le signal $y(t)$ reçu par le mobile s'écrit :

$$y(t) = \frac{A}{d(t)} \cos\left(2\pi f \left(t - \frac{d(t)}{C}\right)\right) - \frac{A}{2D - d(t)} \cos\left(2\pi f \left(t - \frac{2D - d(t)}{C}\right)\right)$$

Dans cette équation, C représente la vitesse de la lumière.

En remplaçant $d(t)$ par sa valeur $d(t) = d_0 + vt$, on obtient :

$$\begin{aligned} y(t) &= \frac{A}{d_0 + vt} \cos\left(2\pi f \left(1 - \frac{v}{C}\right) t - 2\pi f \frac{d_0}{C}\right) \\ &\quad - \frac{A}{2D - d_0 - vt} \cos\left(2\pi f \left(1 + \frac{v}{C}\right) t - 2\pi f \left(\frac{2D - d_0}{C}\right)\right) \\ y(t) &= \frac{A}{d_0 + vt} \cos\left(2\pi (f - \Delta f) t - 2\pi f \frac{d_0}{C}\right) \\ &\quad - \frac{A}{2D - d_0 - vt} \cos\left(2\pi (f + \Delta f) t - 2\pi f \left(\frac{2D - d_0}{C}\right)\right) \end{aligned}$$

On observe le changement de la fréquence caractérisant l'effet Doppler et dont le signe est négatif ou positif selon que le mobile s'éloigne ou se rapproche de l'émetteur :

$$\Delta f = \pm \frac{v}{C} f$$

On appelle étalement Doppler B_d l'écart maximal entre les deux fréquences générées par l'effet Doppler. Ici il vaut :

$$B_d = 2\Delta f = 2 \frac{v}{C} f$$

La différence de phase entre les deux termes de $y(t)$ vaut :

$$2\pi f \left(\frac{2D - d(t)}{C} \right) + \pi - 2\pi f \left(\frac{d(t)}{C} \right) = 4\pi f \left(\frac{D - d(t)}{C} \right) + \pi$$

Lorsque les deux ondes sont en phase, l'amplitude de $y(t)$ est maximale et quand les deux ondes sont en opposition de phase l'amplitude de $y(t)$ est minimale. La distance entre un maximum et un minimum successifs est égale à :

$$\Delta d = \frac{C}{4f} = \frac{\lambda}{4}$$

Cette distance est appelée la distance de cohérence. Le temps de cohérence est lié à cette distance divisée par la vitesse du mobile. Il est donc inversement proportionnel à l'étalement Doppler. Si on définit le temps de cohérence comme l'écart entre un maximum et un minimum consécutifs, on a :

$$T_c \approx \frac{\lambda}{4v} = \frac{C}{4f} = \frac{1}{2B_d}$$

Par exemple, pour une porteuse à 900 MHz, la distance entre deux maxima (ou minima) vaut 16,6 cm. Et pour un mobile se déplaçant à 60 km/h, le mobile passe par un minimum (ou un maximum) toutes les 10 ms. Le rythme d'occurrence des minima est de 100 Hz. Par ailleurs, l'étalement Doppler est égal à 100 Hz.

Pour un signal ayant une largeur de bande $B = 1/T_S$, où T_S est la durée des symboles, on dit que le canal est sélectif en temps (ou que le fading est rapide) si la durée des symboles est supérieure au temps de cohérence du canal $T_S > T_c$. Pendant la durée d'un symbole, l'amplitude du signal reçu fluctue de façon significative. Dans le cas contraire, on parle de fading lent.

Le canal peut être sélectif en temps et en fréquence si $T_S > T_c$ et $B > B_c$. Cette situation peut se produire dans le cas d'une transmission de données à bas débit avec étalement de spectre.

■ Modélisation du canal SISO

L'étude précédente de certaines situations élémentaires a permis d'introduire quelques termes de vocabulaire décrivant un canal radio. On peut préciser la définition de ces termes en généralisant l'analyse à l'aide d'une modélisation des effets des trajets multiples par un modèle linéaire variant avec le temps. Comme le canal est linéaire, on peut le caractériser par sa réponse impulsionnelle à l'instant t notée $h(\tau, t)$. Cette fonction est la réponse du canal à l'instant t à une impulsion de Dirac à l'instant $t - \tau$. Cette réponse impulsionnelle varie avec le temps. La relation entrée-sortie s'écrit :

$$y(t) = \int_{-\infty}^{+\infty} h(\tau, t) x(t - \tau) d\tau$$

Pour un canal multitrajets, si on suppose que les atténuations et les retards ne dépendent pas de la fréquence, la relation entrée-sortie, ramenée en bande de base, peut s'exprimer par :

$$y_e(t) = \sum_i a_i(t) x_e(t - \tau_i(t)) = \int_{-\infty}^{+\infty} h_e(\tau, t) x_e(t - \tau) d\tau$$

où l'indice e indique les signaux ramenés en bande de base. Le signal $y_e(t)$ reçu à l'instant t est la somme des différentes répliques du signal x_e provenant des différents trajets de retard $\tau_i(t)$ et de gain $a_i(t)$. On en déduit la réponse impulsionnelle à l'instant t , ramenée en bande de base :

$$h_e(\tau, t) = \sum_i a_i(t) \delta(\tau - \tau_i(t))$$

Généralement, cette réponse impulsionnelle varie lentement avec t et on peut la considérer stable sur une durée supérieure à l'étalement temporel du canal.

Dans la suite de ce paragraphe, on néglige la notation avec l'indice e .

Par transformée de Fourier, on déduit une fonction de transfert à l'instant t :

$$H(f, t) = \int_{-\infty}^{+\infty} h(\tau, t) e^{-2j\pi f\tau} d\tau = \sum_i a_i(t) e^{-2j\pi f\tau_i(t)}$$

La fonction $h(\tau, t)$ est une fonction aléatoire. Par définition, on dit que la réponse impulsionnelle du canal est stationnaire au sens large (WSS *Wide Sense Stationary*) par rapport au temps, si et seulement si :

$$E(h(\tau, t_1) h^*(\tau, t_2)) = R_h(\tau, t_1 - t_2)$$

Par ailleurs, on dit que la diffusion est non corrélée (*uncorrelated scattering*) si et seulement si :

$$\forall \tau_1 \neq \tau_2 \quad E(h(\tau_1, t) h^*(\tau_2, t)) = 0$$

L'hypothèse d'un canal WWS et non corrélé est souvent vérifiée. Elle peut être étendue en supposant que le canal est stationnaire dans le domaine fréquentiel :

$$\begin{aligned} E(h(\tau_1, t_1) h^*(\tau_2, t_2)) &= R_h(\tau_1, t_1 - t_2) \delta(\tau_1 - \tau_2) \\ E(H(f_1, t_1) H^*(f_2, t_2)) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} E(h(\tau_1, t_1) h^*(\tau_2, t_2)) e^{-2j\pi f_1\tau_1 + 2j\pi f_2\tau_2} d\tau_1 d\tau_2 \\ &= \int_{-\infty}^{+\infty} R_h(\tau_1, t_1 - t_2) e^{-2j\pi(f_1 - f_2)\tau_1} d\tau_1 \\ &= R_H(f_1 - f_2, t_1 - t_2) \end{aligned}$$

On définit la fonction de diffusion (*scattering function*) du canal, notée S , par une double transformée de Fourier de la fonction d'autocorrélation précédente :

$$S_H(\alpha, \beta) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} R_H(\Delta f, \Delta t) e^{-2j\pi(\alpha\Delta f + \beta\Delta t)} d\Delta f d\Delta t$$

En intégrant cette fonction de diffusion par rapport au retard α ou par rapport à l'écart en fréquence β , on obtient :

$$\begin{aligned} S_H(\beta) &= \int_{-\infty}^{+\infty} S_H(\alpha, \beta) d\alpha \\ S_H(\alpha) &= \int_{-\infty}^{+\infty} S_H(\alpha, \beta) d\beta \end{aligned}$$

Et on définit les paramètres caractérisant le canal de la façon suivante :

- L'étalement temporel T_m est la longueur de l'intervalle temporel où $S_H(\alpha)$ est « significativement non-nul ».
- L'étalement Doppler B_d est la longueur de l'intervalle fréquentiel où $S_H(\beta)$ est « significativement non-nul ».

L'expression « significativement non-nul » peut être définie de différentes façons. On peut considérer par exemple les écarts-types des fonctions.

Par ailleurs, ce modèle est complété par l'ajout d'un bruit blanc gaussien $n(t)$ indépendant du signal utile et de densité spectrale de puissance $N_0/2$. La relation entrée-sortie devient :

$$y(t) = \sum_i a_i(t) x(t - \tau_i(t)) + n(t)$$

■ Modèle discret équivalent en bande de base

La plupart du temps, on s'intéressera au modèle de canal discret équivalent en bande de base au canal original.

Le modèle équivalent en bande de base est obtenu en prenant l'enveloppe complexe¹ de la réponse impulsionnelle du canal multipliée par 0,5.

On obtient le modèle discret à partir du modèle équivalent en bande de base en le faisant suivre d'un filtre passe-bas de largeur de bande $W/2$ puis d'un échantillonneur à la fréquence W .

Les enveloppes complexes des signaux émis et reçu : $x_e(t)$ et $y_e(t)$ étant de largeur de bande $W/2$, elles peuvent être échantillonnées à la fréquence W et s'écrire (par exemple pour $y_e(t)$) :

$$y_e(t) = \sum_n y_e(n) \operatorname{sinc}(Wt - n)$$

où :

$$\operatorname{sinc}(x) = \frac{\sin(\pi x)}{\pi x}$$

Soit un canal multitrajets défini par la réponse impulsionnelle :

$$h(\tau, t) = \sum_i a_i(t) \delta(\tau - \tau_i(t))$$

Le modèle équivalent en bande de base du canal, pour une fréquence porteuse f_c , a pour réponse impulsionnelle :

$$h_e(\tau, t) = \sum_i a_i(t) e^{-i2\pi f_c \tau_i(t)} \delta(\tau - \tau_i(t)) = \sum_i b_i(t) \delta(\tau - \tau_i(t))$$

$$b_i(t) = a_i(t) e^{-i2\pi f_c \tau_i(t)}$$

La phase des coefficients $b_i(t)$ change rapidement.

La relation entrée-sortie exprimée en bande de base est la suivante :

$$\begin{aligned} y_e(t) &= \sum_i b_i(t) x_e(\tau - \tau_i(t)) = \sum_i b_i(t) \sum_n x_e(n) \operatorname{sinc}(W(t - \tau_i(t)) - n) \\ &= \sum_n x_e(n) \sum_i b_i(t) \operatorname{sinc}(W(t - \tau_i(t)) - n) \end{aligned}$$

Les échantillons de $y_e(t)$ obtenus à la fréquence d'échantillonnage W sont donnés par :

$$\begin{aligned} y_e(n) &= \sum_j x_e(j) \sum_i b_i(n/W) \operatorname{sinc}(n - j - W\tau_i(n/W)) \\ &= \sum_{k=n-j} x_e(n-k) \sum_i b_i(n/W) \operatorname{sinc}(k - W\tau_i(n/W)) \end{aligned}$$

Or :

$$y_e(n) = \sum_k h(k, n) x_e(n - k)$$

D'où l'on déduit que :

$$h(k, n) = \sum_i b_i(n/W) \operatorname{sinc}(k - W\tau_i(n/W))$$

¹ Pour mémoire, l'enveloppe complexe $x_e(t)$ d'un signal $x(t)$ est définie par : $x_e(t) = x(t) + jTH(x(t))$, où $TH(x(t))$ est la transformée de Hilbert de $x(t)$.

Pour des canaux à fading rapide non sélectifs en fréquence, cette relation se simplifie en :

$$y_e(n) = h(n)x_e(n)$$

Enfin, il faut ajouter à cette relation entrée-sortie l'enveloppe complexe du bruit blanc gaussien $n_e(t)$ filtré par le filtre passe-bas de largeur de bande $W/2$ et échantillonné à la fréquence W . En notant $b(n)$ les échantillons de bruit on obtient :

$$y_e(n) = \sum_k h(k, n)x_e(n - k) + b(n)$$

Le signal complexe $b(n)$ est gaussien. Ses parties réelles et imaginaires sont indépendantes et gaussiennes de même moyenne et variance $N_0/2$. Le bruit gaussien complexe $b(n)$ présentant une symétrie circulaire, on note sa loi $CN(0, \sigma^2)$ et on utilisera l'expression « loi gaussienne complexe circulaire ».

■ Modèles statistiques pour les coefficients du canal

Les coefficients de la réponse impulsionnelle du canal discret $h(k, n)$ correspondent à la somme de plusieurs trajets dont les retards sont proches de k/W . Différents modèles statistiques peuvent être utilisés pour ces coefficients selon l'environnement considéré.

Dans un environnement très riche en trajets et pour une communication entre un émetteur et un récepteur sans ligne de propagation directe (NLOS, *Non Line Of Sight*), on peut considérer que chaque coefficient est la somme d'un grand nombre de variables aléatoires complexes indépendantes à symétrie circulaire. Les coefficients résultant suivent donc une loi gaussienne complexe circulaire et de moyenne nulle. L'amplitude des coefficients de la réponse impulsionnelle suit alors une loi de Rayleigh, et la phase une loi uniforme entre 0 et 2π . On parle de fading de Rayleigh pour désigner cette situation.

Lorsqu'il existe des trajets multiples et une ligne de propagation directe entre l'émetteur et le récepteur (LOS, *Line Of Sight*) d'amplitude plus ou moins importante, le modèle de Rayleigh est remplacé par le modèle de Rice. En notant K le rapport entre la puissance sur le trajet direct et la puissance sur les autres trajets, la loi de probabilité pour le module r d'un coefficient $h(k, n)$ s'écrit (en supposant une variance égale à 1) :

$$p(r) = 2r(1 + K)e^{-(1+K)r^2 - K} I_0 \left(2r\sqrt{K(1 + K)} \right) \quad \forall r \geq 0$$

$$p(r) = 0 \quad \forall r < 0$$

Dans cette relation I_0 représente la fonction de Bessel modifiée d'ordre 0.

Quand K tend vers 0, on passe de la situation LOS à la situation NLOS et on retrouve la loi de Rayleigh :

$$p(r) = 2re^{-r^2} \quad \forall r \geq 0$$

$$p(r) = 0 \quad \forall r < 0$$

■ Modélisation des canaux MIMO

Dans un canal MIMO, on doit considérer n_t antennes en émission et n_r antennes en réception, ce qui conduit à $n_t n_r$ canaux élémentaires entre les antennes d'émission et les antennes de réception. Si on suppose que les canaux sont non sélectifs en fréquence, on a vu précédemment que la relation entrée-sortie pour le modèle du canal SISO discret équivalent en bande de base peut s'écrire (en notant x l'entrée et y la sortie) :

$$y(n) = h(n)x(n)$$

Si de plus le canal est non sélectif en temps, la relation devient :

$$y(n) = h x(n)$$

En notant $\mathbf{x}$ le vecteur de longueur n_t formé des signaux émis par les n_t antennes d'émission et $\mathbf{y}$ le vecteur de longueur n_r formé des signaux reçus par les n_r antennes de réception, le modèle discret équivalent en bande de base, pour un canal MIMO non sélectif en fréquence et à fading rapide, est représenté par la relation matricielle :

$$\mathbf{y}(n) = \mathbf{H}(n)\mathbf{x}(n) + \mathbf{b}(n)$$

Dans cette expression, $\mathbf{b}(n)$ représente le vecteur de bruit gaussien complexe circulaire de longueur n_r . La matrice du canal $\mathbf{H}(n)$ est une matrice rectangulaire de dimension $n_t \times n_r$. On parle de fading ergodique pour représenter cette situation.

Dans le cas d'un canal MIMO non sélectif en fréquence et à fading lent cette relation devient :

$$\mathbf{y}(n) = \mathbf{H}\mathbf{x}(n) + \mathbf{b}(n)$$

On parle alors de fading non ergodique. Cette situation correspond par exemple à un déplacement à la vitesse d'un piéton.

Quand le canal peut être considéré comme constant sur la durée d'un bloc de données, on parle de fading par bloc (*block-fading channel model*). Chaque bloc de données est affecté d'un gain aléatoire constant sur la durée du bloc (les blocs sont de durée inférieure au temps de cohérence du canal).

On peut rendre ergodique un canal non ergodique en introduisant un entrelacement en émission. Cet entrelacement est compensé par un désentrelacement en réception. Pour les applications dans lesquelles le retard doit rester limité, l'entrelacement peut devenir inefficace.

On peut détailler ces relations matricielles en faisant apparaître les gains sur chaque canal élémentaire reliant deux antennes :

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_{n_r} \end{bmatrix} = \begin{bmatrix} h_{1,1} & h_{1,2} & \cdots & h_{1,n_t} \\ h_{2,1} & h_{2,2} & \cdots & h_{2,n_t} \\ \vdots & \vdots & \ddots & \vdots \\ h_{n_r,1} & h_{n_r,2} & \cdots & h_{n_r,n_t} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_{n_t} \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_{n_r} \end{bmatrix}$$

Dans cette relation $h_{i,j}$ représente le gain complexe sur le canal reliant l'antenne d'émission j avec l'antenne de réception i . Par souci de simplicité, on n'a pas indiqué la variation avec le temps dans cette relation.

Sous l'hypothèse d'une matrice de canal constante pendant la durée d'un bloc de données (ou d'un mot de code), on peut définir la loi de probabilité conjointe de l'ensemble des éléments $h_{i,j}$ de la matrice du canal $\mathbf{H}$. Selon les caractéristiques du canal (en particulier la disposition des antennes et des diffuseurs de l'environnement), on pourra considérer différents cas pour la matrice $\mathbf{H}$ qui demanderont un nombre plus ou moins important de paramètres pour caractériser la loi de probabilité. Dans le cas d'un fading de Rayleigh, la valeur moyenne des éléments de la matrice est nulle. Dans le cas d'un fading de Rice, la moyenne n'est pas nulle. Le cas le plus simple est celui d'une matrice $\mathbf{H}$ formée d'éléments indépendants gaussiens complexes circulaires de moyenne nulle (*rich scattering*). Le cas d'une matrice $\mathbf{H}$ formée d'éléments gaussiens complexes circulaires, de moyenne nulle et corrélés, nécessite la connaissance des coefficients de corrélation de toutes les paires d'éléments $h_{i,j}$. Un cas moins complexe est celui d'une matrice $\mathbf{H}$ formée d'éléments à corrélation séparable, c'est-à-dire telle que :

$$E \left(h_{i,j} h_{i',j'}^* \right) = R_{i,i'} T_{j,j'}$$

La matrice de canal dans ce dernier cas se factorise en :

$$\mathbf{H} = \mathbf{R}^{1/2} \mathbf{H}_u \mathbf{T}^{1/2}$$

Dans cette expression, $\mathbf{H}_u$ est formée d'éléments non corrélés gaussiens complexes circulaires de moyenne nulle et de variance unité. Les matrices $\mathbf{R}$ et $\mathbf{T}$ sont hermitiennes définies non-négatives. Par la suite on considérera que la puissance reçue est constante et que $\text{trace}(\mathbf{T}) = n_t$ et $\text{trace}(\mathbf{R}) = n_r$.

De manière générale, le rang de la matrice $\mathbf{H}$ va conditionner les performances du système. Le rang de $\mathbf{H}$ est toujours inférieur au minimum de n_t et de n_r . Un cas très défavorable à l'approche MIMO est appelé « trou de serrure » (*uncorrelated keyhole*). Dans ce cas, la propagation est « canalisée » par exemple par un passage entre des murs ou par un tunnel. La matrice $\mathbf{H}$ s'écrit comme le produit d'un vecteur colonne et d'un vecteur ligne :

$$\mathbf{H} = \mathbf{h}_r \mathbf{h}_t^\dagger$$

La notation $\dagger$ représente une transposition plus une conjugaison de matrice.

Le rang de $\mathbf{H}$ vaut 1.

Un cas limite est celui où la distance séparant les antennes d'émission et de réception est grande par rapport à l'espacement entre les antennes d'émission et à l'espacement entre les antennes de réception. Tous les « rayons » suivent à peu près le même trajet et les différents gains constituant la matrice $\mathbf{H}$ sont tous quasiment égaux.

Dans le cas d'un canal sélectif en fréquence, la relation matricielle entrée-sortie s'écrit :

$$\mathbf{y}(n) = \sum_k \mathbf{H}(k, n) \mathbf{x}(n - k) + \mathbf{b}(n)$$

Par la suite, on se limite aux cas des canaux MIMO non sélectifs en fréquence (*flat fading*), c'est-à-dire des signaux à bande étroite par rapport à la bande de cohérence du canal. Cette hypothèse est vérifiée même pour des débits symboles élevés dans le cas des transmissions OFDM pour lesquelles chaque porteuse est modulée par un signal à bande étroite. Cette remarque explique pourquoi les systèmes MIMO sont souvent utilisés avec des modulations OFDM.

7.3.3 Canal SISO non sélectif et fading de Rayleigh

Les phénomènes de fading dans le cas d'une transmission sur un canal radio-mobile dégradent fortement les performances par rapport au cas d'un canal additif blanc gaussien (CABG).

On illustre cette affirmation avec un canal non sélectif en fréquence dont le modèle discret équivalent en bande de base est décrit par la relation entrée-sortie :

$$y(n) = h(n) x(n) + b(n)$$

Dans le cas d'un canal additif blanc gaussien CABG, h est constant et $b(n)$ est un bruit gaussien complexe circulaire de variance N_0 .

$$y(n) = x(n) + b(n)$$

Pour une modulation à deux états de phase (BPSK) $x = \pm c$, la probabilité d'erreur bit (BER) pour un récepteur cohérent est liée au rapport signal à bruit en réception (par convention on note $RSB = E_b/N_0$) par :

$$BER_{CABG} = Q\left(\sqrt{\frac{2E_b}{N_0}}\right) = Q\left(\sqrt{2RSB}\right)$$

Dans cette expression, la fonction $Q(x)$ est définie par :

$$Q(x) = \frac{1}{\sqrt{2\pi}} \int_x^{+\infty} e^{-\frac{u^2}{2}} du$$

On appelle RSB le rapport entre l'énergie moyenne par bit E_b et la densité spectrale monolatérale de bruit N_0 . Dans cet exemple, $E_b = c^2$. Quand le RSB est grand, la probabilité d'erreur décroît exponentiellement avec le RSB .

Dans le cas d'un canal de Rayleigh non sélectif en temps, h est une variable aléatoire gaussienne complexe circulaire de moyenne nulle.

$$\begin{aligned} y(n) &= h x(n) + b(n) \\ h &\sim CN(0,1) \end{aligned}$$

La figure 7.24 illustre les variations du module de l'amplitude du signal reçu au cours du temps.

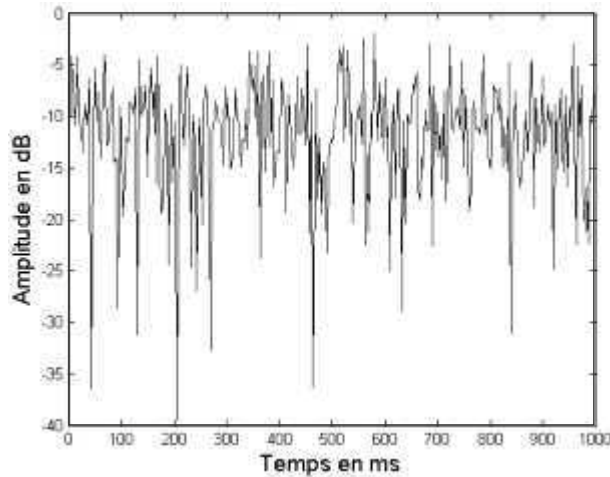


Figure 7.24 – Variation du module de l'amplitude du gain du canal dans un canal de Rayleigh.

Pour une modulation BPSK ($x = \pm c$) dans un canal de Rayleigh avec un récepteur cohérent, pour une valeur de h donnée, la probabilité d'erreur vaut :

$$BER_{\text{Rayleigh}}(h) = Q\left(\sqrt{2|h|^2 RSB}\right)$$

En moyennant cette expression sur les valeurs de h , on obtient la probabilité d'erreur pour le canal de Rayleigh :

$$E(BER_{\text{Rayleigh}}(h)) = BER_{\text{Rayleigh}} = \frac{1}{2} \left(1 - \sqrt{\frac{RSB}{1 + RSB}}\right)$$

Lorsque le RSB est grand, on peut considérer que la probabilité d'erreur bit décroît de façon inversement proportionnelle au RSB en réception :

$$BER_{\text{Rayleigh}} \approx \frac{1}{4RSB} \propto \frac{1}{RSB}$$

La figure 7.25 illustre les relations entre les probabilités d'erreur et le rapport E_b/N_0 dans les cas du canal CABG et du canal de Rayleigh avec une modulation BPSK.

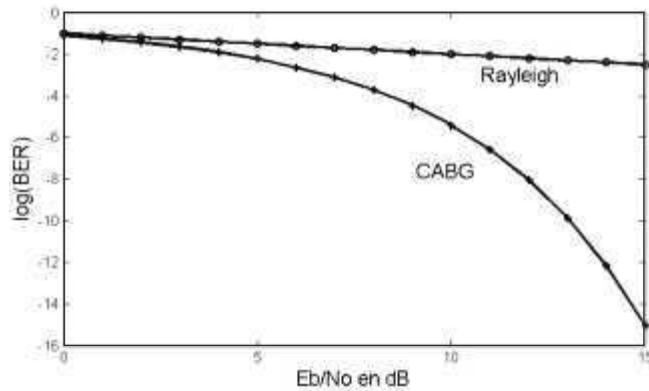


Figure 7.25 – Probabilités d'erreur en BPSK pour un canal CABG et un canal de Rayleigh

Cette mauvaise performance pour les canaux de Rayleigh s'explique par la probabilité importante d'évanouissement de l'amplitude. La probabilité d'erreur est liée à la probabilité que $|h|^2$ soit inférieur à l'inverse du RSB. En effet quand $|h|^2 \gg (1/RSB)$, la probabilité d'erreur est faible. Mais quand $|h|^2 < (1/RSB)$, le canal subit un évanouissement profond et la probabilité d'erreur est grande. La loi de probabilité de $|h|^2$ est une exponentielle décroissante et :

$$p\left(|h|^2 < \frac{1}{RSB}\right) \approx \frac{1}{RSB}$$

La probabilité d'erreur dans un canal de Rayleigh est donc directement liée à la probabilité d'évanouissement profond.

7.3.4 Utilisation de la diversité dans les canaux radio-mobiles

Pour améliorer les performances dans le cas d'un canal de Rayleigh ou plus généralement dans le cas d'un canal radio, on peut utiliser la diversité.

On parle d'une diversité d'ordre n pour indiquer que l'information est reçue par n canaux indépendants.

Dans cette section, on va montrer que pour une diversité d'ordre n , et de grands RSB, la probabilité d'erreur dans un canal de Rayleigh devient inversement proportionnelle au RSB à la puissance n .

$$BER_{\text{Rayleigh}} \propto \frac{1}{RSB^n} \quad [7.1]$$

Les techniques de diversité permettent donc de diminuer la probabilité d'erreur.

Dans la section sur la diversité, différentes techniques de diversité ont été présentées, avec pour mémoire, la diversité en temps, en espace, en fréquence.

L'ordre de diversité dépend de la technique de diversité employée. Par exemple :

- La diversité d'espace en réception utilise plusieurs antennes pour recevoir plusieurs versions d'un même signal. Les antennes doivent être suffisamment espacées (quelques longueurs d'onde) pour que ces versions soient peu corrélées. L'ordre de diversité est au maximum égal au nombre d'antennes.

- La diversité de trajets est utilisée dans les récepteurs Rake (utilisés dans les systèmes DS-CDMA). L'ordre de diversité dépend alors du nombre de trajets qui peuvent être séparés.
- La diversité temporelle consiste à envoyer r fois le même signal à des instants séparés d'une durée supérieure au temps de cohérence du canal pour que les signaux soient peu corrélés. L'ordre de diversité est alors inférieur ou égal à r .
- La diversité par code utilise un code de rendement R pour transmettre plusieurs versions non corrélées de l'information. Elle se rapproche de la diversité temporelle, cette dernière pouvant être considérée comme un code à répétition. Mais la diversité par code cherche à optimiser l'usage des ressources et l'efficacité spectrale. L'ordre de diversité n dépend alors de la distance minimale du code $d_{\min}$ et de l'ordre de diversité physique n_p : $n = \min(d_{\min}, n_p)$. On combine généralement un code et un entrelaceur.
- Quand plusieurs techniques de diversité sont combinées, l'ordre de diversité résultant est égal au produit des différents ordres. On peut par exemple combiner, dans les systèmes MIMO, des diversités d'espace et de temps ou d'espace et de trajet.

Dans le cas d'un système MIMO avec n_t antennes en émission et n_r antennes en réception, l'ordre maximal de diversité est égal à $n_t n_r$.

Pour illustrer le résultat donné dans l'équation [7.1] on étudie d'abord la diversité temporelle.

■ Diversité temporelle

On commence par analyser la diversité temporelle obtenue en utilisant un code à répétition et un entrelacement.

On considère une donnée BPSK x ($x = \pm c$) qui est codée par un code à répétition de longueur L . Grâce à un entrelaceur, les L répétitions sont dispersées dans L intervalles séparés d'une durée supérieure au temps de cohérence et donc soumis à des évanouissements indépendants entre eux caractérisés par les gains de canal h_l avec $l = 1, \dots, L$. Les L échantillons reçus y_l s'écrivent donc :

$$\begin{aligned} y_l &= h_l x + b_l, \quad l = 1, \dots, L \\ \mathbf{y}^T &= (y_1, y_2, \dots, y_L) \\ \mathbf{h}^T &= (h_1, h_2, \dots, h_L) \\ \mathbf{b}^T &= (b_1, b_2, \dots, b_L) \\ \mathbf{y} &= x\mathbf{h} + \mathbf{b} \end{aligned}$$

À partir du vecteur d'observations $\mathbf{y}$, le récepteur doit décider quelle est la donnée x reçue. La solution optimale au sens de la minimisation de la probabilité d'erreur consiste à projeter le vecteur d'observation $\mathbf{y}$ sur le vecteur $\mathbf{h}$ des gains du canal et à décider la valeur de x selon le signe de cette projection. En notant Y cette projection, on a :

$$Y = \frac{\mathbf{h}^\dagger}{\|\mathbf{h}\|} \mathbf{y} = x \frac{\mathbf{h}^\dagger \mathbf{h}}{\|\mathbf{h}\|} + \frac{\mathbf{h}^\dagger \mathbf{b}}{\|\mathbf{h}\|} = x \|\mathbf{h}\| + \frac{\mathbf{h}^\dagger \mathbf{b}}{\|\mathbf{h}\|} = x \|\mathbf{h}\| + B_p$$

La probabilité d'erreur dépend de la projection du bruit B_p et de la norme du vecteur $\mathbf{h}$.

Par rapport au cas scalaire (sans diversité), le terme xh est remplacé par $x \|\mathbf{h}\| = x \sqrt{\sum_{l=1}^L h_l^2}$. Pour une valeur de $\mathbf{h}$ donnée la probabilité d'erreur vaut :

$$BER_{\text{Rayleigh}} = Q\left(\sqrt{2 \|\mathbf{h}\|^2 RSB}\right)$$

Pour un fading de Rayleigh, la loi de probabilité de la norme au carré de $\mathbf{h}$ est une loi du χ^2 à L degrés de liberté. La probabilité d'erreur s'obtient en moyennant l'expression précédente sur les valeurs possibles de la norme carrée de $\mathbf{h}$. Pour les grandes valeurs du RSB , on peut faire la même hypothèse que précédemment et supposer que la probabilité d'erreur dépend surtout de la probabilité d'évanouissement profond. La probabilité qu'il y ait un évanouissement profond, c'est-à-dire que la norme carrée de $\mathbf{h}$ soit inférieure à $1/RSB$ est donnée par :

$$p\left(\|\mathbf{h}\|^2 < \frac{1}{RSB}\right) = p\left(\sum_{l=1}^L |h_l|^2 < \frac{1}{RSB}\right) \approx \frac{1}{L! (RSB)^L}$$

La probabilité d'erreurs pour un canal de Rayleigh et une diversité temporelle d'ordre L et pour les grandes valeurs du RSB varie donc en $1/RSB^L$. Pour les petites valeurs de RSB l'approximation n'est plus correcte car la probabilité d'erreur dépend alors autant du bruit additif que de la probabilité d'occurrence des évanouissements profonds.

On notera que la mise en œuvre de ce récepteur suppose connu le vecteur de gain du canal $\mathbf{h}$. On parle de CSI (*Channel State Information*) pour désigner la connaissance du canal disponible au niveau de l'émetteur ou du récepteur.

L'utilisation de la diversité temporelle avec un code à répétition et un entrelaceur améliore la probabilité d'erreur au détriment de l'efficacité spectrale puisque chaque symbole est transmis L fois. On peut améliorer l'efficacité spectrale en utilisant des codes plus performants. Une autre approche possible est d'utiliser la diversité d'espace obtenue en utilisant plusieurs antennes.

■ Diversité d'espace ou d'antennes

La diversité d'espace est obtenue en utilisant plusieurs antennes soit en émission, soit en réception, soit en émission et en réception. L'ordre de diversité dépend du nombre de chemins différents utilisables pour transmettre le signal. Pour n_t antennes en émission et n_r antennes en réception, il vaut au maximum $n_t n_r$.

□ Diversité d'antennes en réception

L'approche la plus classique utilise plusieurs antennes de réception (voir figure 7.26).

Pour n_r antennes de réception et 1 antenne d'émission, le vecteur de signal reçu s'écrit :

$$\begin{aligned} y_i &= h_i x + b_i, \quad i = 1, \dots, n_r \\ \mathbf{y}^T &= (y_1, y_2, \dots, y_{n_r}) \\ \mathbf{h}^T &= (h_1, h_2, \dots, h_{n_r}) \\ \mathbf{b}^T &= (b_1, b_2, \dots, b_{n_r}) \\ \mathbf{y} &= x\mathbf{h} + \mathbf{b} \end{aligned}$$

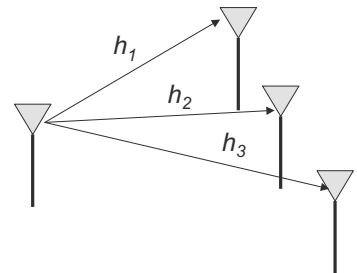


Figure 7.26 – Diversité d'antennes en réception.

On retrouve la même expression que pour la diversité temporelle avec l'avantage d'un gain de puissance. En effet pour une même puissance moyenne P , dans le cas du code à répétition chaque symbole x ne porte qu'une puissance P/L . Si le récepteur connaît le vecteur de gain du canal $\mathbf{h}$, il projette le signal reçu $\mathbf{y}$ sur $\mathbf{h}$. Il calcule donc :

$$Y = \frac{\mathbf{h}^\dagger}{\|\mathbf{h}\|} \mathbf{y} = x \frac{\mathbf{h}^\dagger \mathbf{h}}{\|\mathbf{h}\|} + \frac{\mathbf{h}^\dagger \mathbf{b}}{\|\mathbf{h}\|} = x \|\mathbf{h}\| + \frac{\mathbf{h}^\dagger \mathbf{b}}{\|\mathbf{h}\|} = x \|\mathbf{h}\| + B$$

Ce traitement porte le nom de *Maximum Ratio Combining* (MRC).

□ Diversité d'antennes en émission

Il est aussi possible d'utiliser plusieurs antennes en émission et d'effectuer une diversité en émission (figure 7.27).

Le signal reçu y est alors le produit scalaire du vecteur de gains du canal $\mathbf{h}$ et du vecteur émis $\mathbf{x}$.

$$\mathbf{x}^T = (x_1, x_2, \dots, x_{n_t})$$

$$\mathbf{h}^T = (h_1, h_2, \dots, h_{n_t})$$

$$y = \mathbf{h}^T \mathbf{x} + b = \sum_{i=1}^{n_t} h_i x_i + b$$

Si les n_t antennes transmettent le même signal x on obtient :

$$y = \mathbf{h}^T \mathbf{x} + b = \sum_{i=1}^{n_t} h_i x_i + b = x \sum_{i=1}^{n_t} h_i + b$$

$$y = xh + b, \quad \text{où } h = \sum_{i=1}^{n_t} h_i$$

On notera que, pour une comparaison équitable, il faut considérer une puissance de x de valeur n_t fois inférieure à celle utilisée pour la diversité de réception.

Pour le récepteur optimal et une valeur de h donnée, la probabilité d'erreur vaut :

$$BER_{\text{Rayleigh}} = Q\left(\sqrt{2|h|^2 RSB}\right) = Q\left(\sqrt{2 \left|\sum_{i=1}^{n_t} h_i\right|^2 RSB}\right)$$

Pour de grands RSB, la probabilité d'erreur est à peu près égale à la probabilité d'évanouissement profond, c'est-à-dire à :

$$BER \approx p\left(|h|^2 < \frac{1}{RSB}\right) = p\left(\left|\sum_{i=1}^{n_t} h_i\right|^2 < \frac{1}{RSB}\right)$$

Les éléments h_i étant gaussiens complexes circulaires, de moyenne nulle et de variance σ^2 , le module carré de h_i suit une loi exponentielle décroissante de moyenne σ^2 . De même, le module carré de h suit une loi exponentielle décroissante de moyenne $n_t \sigma^2$ (on suppose les h_i indépendants). Finalement, les performances avec ou sans diversité sont les mêmes avec cette approche. Autrement dit, émettre directement le même signal sur toutes les antennes n'apporte rien.

Une solution possible pour obtenir une diversité d'ordre n_t , est d'émettre le signal x alternativement sur chaque antenne. Le récepteur dispose alors des valeurs successives $h_i x$ qu'il peut recombinaison de manière optimale comme dans le cas de la diversité de réception. L'interférence éventuelle entre les symboles successifs est traitée par des techniques d'égalisation ou par un détecteur de séquences à maximum de vraisemblance MLSE (*Maximum Likelihood Sequence Estimator*). Cette approche est équivalente à la diversité temporelle avec un code à répétition. Elle n'utilise pas le temps de manière optimale.

Une solution plus efficace est possible si l'émetteur connaît le vecteur de gains du canal $\mathbf{h}$. Il peut alors pondérer le signal x émis par chaque antenne pour tenir compte du gain du canal sur le trajet entre cette antenne d'émission et le récepteur. L'émetteur effectue alors un traitement équivalent à ce que fait le récepteur dans le cas d'une diversité en réception (MRC).

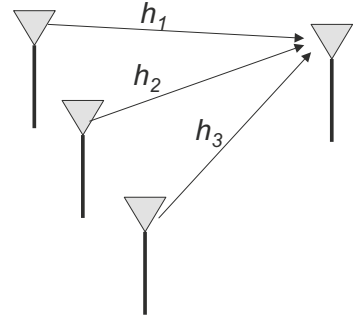


Figure 7.27 – Diversité d'antennes en émission.

Le vecteur émis s'écrit :

$$\mathbf{x} = x \frac{\mathbf{h}^*}{\|\mathbf{h}\|}$$

$$y = \mathbf{h}^T \mathbf{x} + b = x \|\mathbf{h}\| + b$$

On retrouve alors le même résultat qu'avec la diversité en réception. Cette solution nécessite que le système dispose d'une voie de retour sur laquelle le récepteur transmet à l'émetteur les caractéristiques du canal.

Enfin, même si l'émetteur ne dispose pas de la connaissance du canal, on peut obtenir un gain de diversité en émission par l'utilisation d'un codage spatio-temporel. Un exemple de code spatio-temporel est le code d'Alamouti.

□ Code d'Alamouti pour la diversité en émission

On considère le cas de deux antennes en émission ($n_t = 2$) :

Soit deux valeurs successives x_1 et x_2 à transmettre.

À l'instant n , on émet le vecteur $[x_1, x_2]$, c'est-à-dire qu'on envoie x_1 sur une antenne et x_2 sur l'autre antenne.

À l'instant suivant $n + 1$ on émet $[-x_2^*, x_1^*]$.

L'antenne de réception reçoit les signaux émis par chaque antenne, multipliés par des gains complexes h_1 et h_2 . Le vecteur $\mathbf{y}$ formé des échantillons reçus correspondant aux deux instants successifs s'écrit :

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \end{bmatrix}$$

$$y_1 = h_1 x_1 + h_2 x_2 + b_1$$

$$y_2 = -h_1 x_2^* + h_2 x_1^* + b_2 \text{ soit } y_2^* = -h_1^* x_2 + h_2^* x_1 + b_2^*$$

D'où on déduit :

$$\begin{bmatrix} y_1 \\ y_2^* \end{bmatrix} = \begin{bmatrix} h_1 & h_2 \\ h_2^* & -h_1^* \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2^* \end{bmatrix} = \mathbf{H}\mathbf{x} + \mathbf{b}$$

En supposant que le récepteur connaît le vecteur $\mathbf{h}$ de gains du canal, le traitement optimal consiste à projeter sur les deux vecteurs colonnes de la matrice $\mathbf{H}$, c'est-à-dire à multiplier le vecteur par la transposée hermitienne $\mathbf{H}^\dagger$ de la matrice de canal :

$$\begin{aligned} \tilde{\mathbf{y}} &= \mathbf{H}^\dagger \begin{bmatrix} y_1 \\ y_2^* \end{bmatrix} = \mathbf{H}^\dagger \mathbf{H}\mathbf{x} + \mathbf{H}^\dagger \mathbf{b} \\ &= \begin{bmatrix} h_1^* & h_2 \\ h_2^* & -h_1 \end{bmatrix} \begin{bmatrix} h_1 & h_2 \\ h_2^* & -h_1^* \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} h_1^* & h_2 \\ h_2^* & -h_1 \end{bmatrix} \begin{bmatrix} b_1 \\ b_2^* \end{bmatrix} \\ &= \begin{bmatrix} |h_1|^2 + |h_2|^2 & 0 \\ 0 & |h_1|^2 + |h_2|^2 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} h_1^* b_1 + h_2 b_2^* \\ h_2^* b_1 - h_1 b_2^* \end{bmatrix} \\ \text{Soit} \quad &= \left(|h_1|^2 + |h_2|^2 \right) \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} \tilde{b}_1 \\ \tilde{b}_2 \end{bmatrix} \end{aligned}$$

Les bruits $\tilde{b}_1$ et $\tilde{b}_2$ sont indépendants et le nouveau RSB vaut :

$$RSB_{\text{Alamouti}} = \left(|h_1|^2 + |h_2|^2 \right) \frac{RSB}{2}$$

Le code d'Alamouti permet donc de doubler le débit par rapport à la méthode consistant à émettre alternativement sur chaque antenne. Le *RSB* obtenu est deux fois plus petit que celui obtenu en diversité de réception.

7.3.5 Multiplexage spatial et principe des systèmes MIMO

Les techniques MIMO font appel à deux techniques pour optimiser la transmission sur les canaux radio-mobiles : la diversité et le multiplexage spatial.

On a vu que la diversité en réception (systèmes SIMO) ou en émission (systèmes MISO) permet d'améliorer la probabilité d'erreur.

Le multiplexage spatial a pour objectif d'augmenter le débit en profitant de la présence de plusieurs antennes d'émission pour transmettre des signaux différents sur ces antennes et ainsi augmenter potentiellement le débit sans augmenter la largeur de bande.

Pour expliquer le principe du multiplexage spatial, on considère un canal déterministe non sélectif en fréquence (voir figure 7.28).

La relation entrée-sortie pour le système MIMO à canal non sélectif déterministe s'écrit :

$$\mathbf{y} = \mathbf{H}\mathbf{x} + \mathbf{b}$$

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_{n_r} \end{bmatrix} = \begin{bmatrix} h_{1,1} & h_{1,2} & \cdots & h_{1,n_t} \\ h_{2,1} & h_{2,2} & \cdots & h_{2,n_t} \\ \vdots & \vdots & \ddots & \vdots \\ h_{n_r,1} & h_{n_r,2} & \cdots & h_{n_r,n_t} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_{n_t} \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_{n_r} \end{bmatrix}$$

Dans cette équation, $\mathbf{b}$ est un bruit blanc gaussien.

La matrice du canal $\mathbf{H}$ est une matrice rectangulaire de dimension $n_r \times n_t$. On peut factoriser $\mathbf{H}$ par une décomposition en valeurs singulières :

$$\mathbf{H} = \mathbf{U}\mathbf{D}\mathbf{V}^\dagger$$

Les matrices $\mathbf{U}$ et $\mathbf{V}$ sont des matrices unitaires de dimensions respectives $n_r \times n_r$ et $n_t \times n_t$, où n est le minimum de n_r et n_t . La matrice $\mathbf{D}$ est carrée de dimension $n \times n$ et diagonale. Elle est formée des valeurs singulières de $\mathbf{H}$.

On peut multiplier les deux membres de l'équation entrée-sortie par $\mathbf{U}^\dagger$:

$$\mathbf{U}^\dagger \mathbf{y} = \mathbf{U}^\dagger \mathbf{U} \mathbf{D} \mathbf{V}^\dagger \mathbf{x} + \mathbf{U}^\dagger \mathbf{b} = \mathbf{D} \mathbf{V}^\dagger \mathbf{x} + \mathbf{U}^\dagger \mathbf{b}$$

Cette nouvelle écriture indique que l'on peut récupérer les données $\mathbf{x}$ à l'aide des deux opérations linéaires suivantes, une en émission et l'autre en réception :

En émission, on effectue un précodage du signal en multipliant $\mathbf{x}$ par $\mathbf{V}$

$$\tilde{\mathbf{x}} = \mathbf{V}\mathbf{x}$$

En réception, on reçoit alors :

$$\mathbf{y} = \mathbf{U}\mathbf{D}\mathbf{V}^\dagger \tilde{\mathbf{x}} + \mathbf{b} = \mathbf{U}\mathbf{D}\mathbf{V}^\dagger \mathbf{V}\mathbf{x} + \mathbf{b} = \mathbf{U}\mathbf{D}\mathbf{x} + \mathbf{b}$$

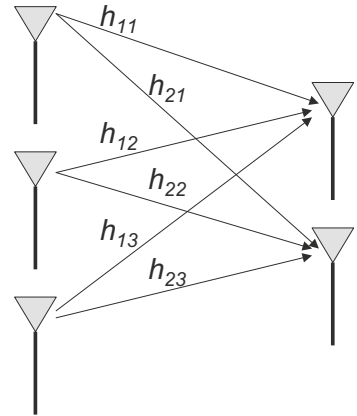


Figure 7.28 – Système MIMO avec $n_t = 3$ antennes d'émission et $n_r = 2$ antennes de réception.

En multipliant le signal reçu par $\mathbf{U}^\dagger$, on récupère le signal $\mathbf{x}$ bruité :

$$\tilde{\mathbf{y}} = \mathbf{U}^\dagger \mathbf{y} = \mathbf{U}^\dagger \mathbf{U} \mathbf{D} \mathbf{x} + \mathbf{U}^\dagger \mathbf{b} = \mathbf{D} \mathbf{x} + \mathbf{U}^\dagger \mathbf{b} = \mathbf{D} \mathbf{x} + \tilde{\mathbf{b}} \quad [7.2]$$

En notant λ_i les éléments diagonaux de $\mathbf{D}$, les coordonnées de $\tilde{\mathbf{y}}$ s'écrivent :

$$\begin{aligned} \tilde{y}_i &= \lambda_i x_i + \tilde{b}_i, \quad i = 1, \dots, r \\ \tilde{b}_i &\sim CN(0, N_0) \end{aligned}$$

Dans cette expression, r représente le nombre de valeurs singulières non nulles, c'est-à-dire le rang de la matrice $\mathbf{H}$. On peut noter que cette approche suppose le canal connu de l'émetteur et du récepteur.

La figure 7.29 résume ces opérations.

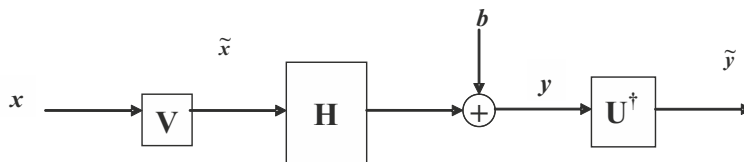


Figure 7.29 – Précodage et réception MIMO.

La décomposition en valeurs singulières met en évidence la décomposition du canal en canaux orthogonaux indépendants correspondant aux vecteurs singuliers. Le nombre de vecteurs singuliers orthogonaux et de valeurs singulières non nulles dépend des caractéristiques du canal qui déterminent les propriétés de la matrice $\mathbf{H}$ (en particulier son rang). Il est au maximum égal à n le nombre minimal d'antennes en émission ou en réception.

On peut donc au plus transmettre sur n canaux virtuels indépendants et augmenter ainsi le débit par un facteur n par rapport au cas SISO.

La question est de savoir quelle est la capacité de ce système et quelles techniques mettre en œuvre pour s'en rapprocher.

7.3.6 Capacité d'un canal MIMO

Discuter précisément des calculs de capacités des canaux sort du champ de cet ouvrage. On ne donne ici que quelques notions d'intérêt général.

On rappelle que la capacité d'un canal représente le débit maximal d'information qu'il est possible de transmettre sur le canal avec une probabilité d'erreur arbitrairement faible.

Pour mémoire, la capacité du canal additif blanc gaussien à temps continu, pour une largeur de bande W et un RSB donnés, exprimée en bits par seconde par Hz (bps/Hz), vaut :

$$C_{CABG} = \log_2(1 + RSB) \text{ bps/Hz}$$

Lorsque le canal a un gain déterministe, la capacité a la valeur donnée par l'équation précédente. Lorsque le gain du canal est une variable aléatoire, la capacité est elle-même aléatoire et peut dans certains cas devenir très faible. On introduit alors la notion de coupure du canal et on définit une capacité de coupure à p % qui est la valeur de capacité réalisée dans au moins $(100 - p)$ % des cas.

On présente dans cette section quelques résultats sur la capacité des canaux MIMO, en se limitant aux canaux non sélectifs en fréquence et en distinguant les cas des canaux déterministes, des canaux à fading rapide (canaux ergodiques) et des canaux à fading lent (canaux quasi-statiques ou à fading par bloc).

■ Canal déterministe

On commence par la présentation de la capacité du canal MIMO non sélectif en fréquence et de matrice $\mathbf{H}$ déterministe (vu précédemment, éq. 7.2) connue de l'émetteur et du récepteur. Ce canal peut être vu comme r canaux parallèles, où r est le rang de la matrice $\mathbf{H}$ (nombre de valeurs singulières non nulles), chacun de ces canaux étant associé à une valeur de gain λ_i et à un bruit additif gaussien complexe (voir figure 7.30). La capacité C de ce canal est obtenue pour des signaux x_i gaussiens circulaires, indépendants, de moyenne nulle et dont la variance est obtenue par la méthode appelée « *water filling* ». Cette méthode se caractérise par le fait de transmettre plus de puissance dans les canaux les meilleurs.

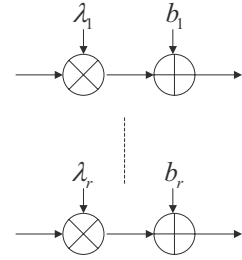


Figure 7.30 – Canal MIMO vu comme r canaux parallèles

$$E(|x_i|^2) = \begin{cases} (\mu - N_0 \lambda_i^{-2}) & \text{si } (\mu - N_0 \lambda_i^{-2}) > 0 \\ 0 & \text{si } (\mu - N_0 \lambda_i^{-2}) \leq 0 \end{cases}$$

La constante μ est définie pour tenir compte de la contrainte de puissance moyenne :

$$P(\mu) = \sum_{i=1}^r \max((\mu - N_0 \lambda_i^{-2}), 0)$$

La capacité s'écrit finalement :

$$C(\mu) = \sum_{i=1}^r \max\left(\log_2\left(\frac{\mu \lambda_i^2}{N_0}\right), 0\right)$$

On montre que, pour de grands RSB , la capacité est obtenue pour une quasi-équirépartition des puissances et elle vaut approximativement :

$$\begin{aligned} C &= \sum_{i=1}^r \log_2\left(1 + \frac{P \lambda_i^2}{r N_0}\right) = \sum_{i=1}^r \log_2\left(1 + \frac{RSB \lambda_i^2}{r}\right) \\ C &\approx r \log_2(RSB) + \sum_{i=1}^r \log_2\left(\frac{\lambda_i^2}{r}\right) \end{aligned}$$

La capacité dépend du rang de la matrice $\mathbf{H}$ et de son conditionnement, c'est-à-dire du rapport entre les deux valeurs singulières extrêmes. La matrice $\mathbf{H}$ est d'autant mieux conditionnée et la capacité est d'autant plus grande que ce rapport est proche de 1. Des trajets correspondant à des directions d'arrivée proches seront d'autant mieux séparables que la longueur du réseau d'antennes de réception est importante.

Dans le cas particulier d'un canal de matrice $\mathbf{H}$ égale à la matrice identité avec $n_t = n_r = n$, le canal est formé de n canaux parallèles identiques indépendants et la capacité globale est la somme de n capacités élémentaires correspondant à des canaux gaussiens avec une puissance égale à $1/n^{\text{ème}}$ de la puissance totale. La capacité totale vaut donc :

$$C = \sum_{i=1}^n \log_2\left(1 + \frac{P}{n N_0}\right) = n \log_2\left(1 + \frac{RSB}{n}\right)$$

Par rapport à un système SISO, on observe un gain en débit possible lié au terme n qui multiplie le logarithme.

■ Canaux aléatoires, capacité ergodique, capacité de coupure

Pour les canaux radio aléatoires non stationnaires, on définit plusieurs notions de capacité du canal correspondant à des hypothèses différentes sur la vitesse de variation du canal (et ceci de façon relative au rythme des symboles transmis).

□ Canaux ergodiques, capacité ergodique pour un canal de Rayleigh

Si le canal varie assez vite, la moyenne temporelle permet d'approcher la moyenne stochastique. On peut considérer une capacité « au sens de Shannon » en effectuant des moyennes statistiques. On parle alors de *capacité ergodique*. On illustre ici cette approche avec le cas d'un canal de Rayleigh ergodique.

On suppose que chaque utilisation du canal correspond à une nouvelle réalisation de la matrice aléatoire $\mathbf{H}$. On fait l'hypothèse d'un canal de Rayleigh, c'est-à-dire que les éléments de $\mathbf{H}$ suivent une loi gaussienne circulaire, de moyenne nulle et de variance unité. On suppose aussi que le récepteur connaît $\mathbf{H}$. L'information mutuelle moyenne est maximale entre l'émetteur et le récepteur (ce qui correspond à la capacité du canal) pour un vecteur signal $\mathbf{x}$ de loi gaussienne circulaire de moyenne nulle et dont la matrice de covariance est la matrice identité I_{n_t} multipliée par RSB/n_t . On montre que la capacité ainsi obtenue vaut :

$$C = E \left[\log_2 \left(\det \left(I_{n_r} + \frac{RSB}{n_t} \mathbf{H} \mathbf{H}^\dagger \right) \right) \right] \text{ bps/Hz}$$

Quand le nombre d'antennes d'émission tend vers l'infini, la capacité tend vers

$$C = \log_2 \left(\det \left((1 + RSB) I_{n_r} \right) \right) = n_r \log_2 (1 + RSB)$$

Une borne supérieure de C est donnée par :

$$C \leq \min \left[n_r \log_2 (1 + RSB), n_t \log_2 \left(1 + \frac{n_r}{n_t} RSB \right) \right]$$

La figure 7.31 représente la capacité d'un canal SISO additif gaussien et la borne supérieure de la capacité d'un canal de Rayleigh ergodique MIMO, pour différentes valeurs du couple n_t, n_r .

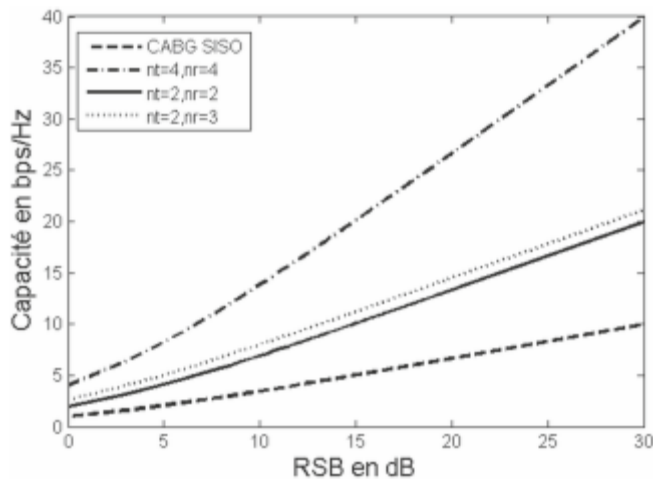


Figure 7.31 – Capacité (borne supérieure) d'un canal de Rayleigh ergodique pour différents couples n_t, n_r .

On constate que, pour les RSB élevés, la pente de la capacité est déterminée par $\min(n_t, n_r)$, la technique MIMO permet d'augmenter le nombre de degrés de liberté. Pour les petits RSB , le rapport entre la capacité MIMO et la capacité SISO est de l'ordre de n_r .

□ Canaux à fading lent, capacité de coupure

Si le canal varie lentement (canal quasi-statique ou à évanouissement par bloc), il reste constant pendant une longue durée (par exemple toute la durée de la transmission). On ne peut plus calculer la capacité comme une moyenne. On considère la capacité $C(\mathbf{H})$ comme une variable aléatoire dépendant de la réalisation du canal aléatoire $\mathbf{H}$. Et pour un débit de transmission R (efficacité spectrale) donné, on définit la probabilité de coupure (*outage probability*) qui est la probabilité que la capacité $C(\mathbf{H})$ soit inférieure à R :

$$p_{\text{outage}}(R) = p(C(\mathbf{H}) < R)$$

Pour un canal de Rayleigh, on a :

$$p_{\text{outage}}(R) = p \left[\log_2 \left(\det \left(I_{n_r} + \frac{RSB}{n_t} \mathbf{H} \mathbf{H}^\dagger \right) \right) < R \right]$$

On définit aussi une capacité de coupure C_ϵ (ϵ -*outage capacity*) qui est le débit maximal permettant une probabilité de coupure inférieure à ϵ :

$$C_\epsilon = p_{\text{outage}}^{-1}(\epsilon)$$

7.3.7 Codage spatio-temporel

Le codage spatio-temporel (STC, *Space Time Code*) effectue un codage à la fois en temps et en espace. Il existe différents types de STC, par exemple des codes en bloc ou des codes treillis. Pour un traitement approfondi du sujet, le lecteur intéressé pourra se reporter à Tarokh, Jafarkhani, Gesbert et A. Paulraj.

On se limite ici aux codes en bloc.

On considère un code en bloc de longueur N et n_t antennes d'émission. Le code est utilisé pour transmettre, pendant chaque période de N intervalles de temps, un ensemble de K symboles utiles x_i . Le code est formé d'un ensemble de matrices $\mathbf{X}$ de n_t lignes et N colonnes.

Le signal reçu sur les n_r antennes de réception pendant N intervalles de temps peut s'exprimer sous la forme d'une matrice $\mathbf{Y}$ formée de n_r lignes et N colonnes, la ligne i contenant les échantillons $y_{i,j}$ reçus par l'antenne i aux N instants successifs j du bloc. La matrice $\mathbf{Y}$ est reliée à la matrice $\mathbf{X}$ par la relation :

$$\mathbf{Y} = \mathbf{H}\mathbf{X} + \mathbf{B}$$

Dans cette équation, $\mathbf{B}$ est une matrice de bruit dont les éléments sont supposés gaussiens, complexes circulaires, de moyenne nulle.

Les performances (en termes de gain de diversité et de multiplexage) obtenues avec un code dépendent de la matrice stochastique $\mathbf{D}$ formée des différences entre deux matrices de code $\mathbf{X}$, en particulier de la trace et du déterminant de la matrice $\mathbf{D}\mathbf{D}^\dagger$.

On présente maintenant quelques codes simples.

■ Code d'Alamouti avec deux antennes en émission et en réception

On a vu qu'il était possible d'obtenir un gain de diversité en émission en utilisant un code d'Alamouti. Ce code a été présenté dans le cas $n_t = 2, n_r = 1$. Il s'agit d'un code spatio-temporel

(codage en espace et en temps) défini par la matrice :

$$\begin{array}{c} \text{Temps} \rightarrow \\ \begin{array}{cc} t_n & t_{n+1} \\ \text{Espace} \downarrow \begin{array}{l} \text{Antenne 1} \\ \text{Antenne 2} \end{array} & \begin{pmatrix} x_1 & -x_2^* \\ x_2 & x_1^* \end{pmatrix} \end{array} \end{array}$$

Dans le cas où on dispose de deux antennes d'émission et de deux antennes de réception, on peut reprendre les équations vues précédemment mais en considérant les deux antennes de réception. On note $y_{i,j}$ le signal reçu par l'antenne i au temps j avec $i = 1$ ou 2 et $j = 1$ ou 2 (représentant deux instants successifs). On peut écrire :

$$\begin{bmatrix} y_{11} & y_{12} \\ y_{21} & y_{22} \end{bmatrix} = \begin{bmatrix} h_{11} & h_{12} \\ h_{21} & h_{22} \end{bmatrix} \begin{bmatrix} x_1 & -x_2^* \\ x_2 & x_1^* \end{bmatrix} + \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix}$$

Cette expression peut-être réécrite sous la forme :

$$\tilde{\mathbf{y}} = \begin{bmatrix} y_{11} \\ y_{21} \\ y_{12}^* \\ y_{22}^* \end{bmatrix} = \begin{bmatrix} h_{11} & h_{12} \\ h_{21} & h_{22} \\ h_{12}^* & -h_{11}^* \\ h_{22}^* & -h_{21}^* \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} b_{11} \\ b_{21} \\ b_{12}^* \\ b_{22}^* \end{bmatrix} = \tilde{\mathbf{H}}\mathbf{x} + \tilde{\mathbf{b}}$$

La matrice $\tilde{\mathbf{H}}$ a la propriété suivante :

$$\tilde{\mathbf{H}}^\dagger \tilde{\mathbf{H}} = \left(|h_{11}|^2 + |h_{12}|^2 + |h_{21}|^2 + |h_{22}|^2 \right) \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \left(|h_{11}|^2 + |h_{12}|^2 + |h_{21}|^2 + |h_{22}|^2 \right) \mathbf{I}_2$$

En supposant que le récepteur connaît le vecteur $\mathbf{h}$ de gains du canal, le traitement optimal consiste à multiplier le vecteur $\tilde{\mathbf{y}}$ par la transposée hermitienne $\tilde{\mathbf{H}}^\dagger$ de la matrice $\tilde{\mathbf{H}}$:

$$\begin{aligned} \tilde{\mathbf{H}}^\dagger \tilde{\mathbf{y}} &= \tilde{\mathbf{H}}^\dagger \tilde{\mathbf{H}}\mathbf{x} + \tilde{\mathbf{H}}^\dagger \tilde{\mathbf{b}} \\ \tilde{\mathbf{H}}^\dagger \tilde{\mathbf{y}} &= \begin{bmatrix} h_{11}^* y_{11} + h_{21}^* y_{21} + h_{12} y_{12}^* + h_{22} y_{22}^* \\ h_{12}^* y_{11} + h_{22}^* y_{21} - h_{11} y_{12}^* - h_{21} y_{22}^* \end{bmatrix} \\ &= \left(|h_{11}|^2 + |h_{12}|^2 + |h_{21}|^2 + |h_{22}|^2 \right) \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} h_{11}^* b_{11} + h_{21}^* b_{21} + h_{12} b_{12}^* + h_{22} b_{22}^* \\ h_{12}^* b_{11} + h_{22}^* b_{21} - h_{11} b_{12}^* - h_{21} b_{22}^* \end{bmatrix} \\ &= \left(|h_{11}|^2 + |h_{12}|^2 + |h_{21}|^2 + |h_{22}|^2 \right) \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} n_1 \\ n_2 \end{bmatrix} \end{aligned}$$

On peut donc supprimer en réception les interférences spatiales entre x_1 et x_2 et les performances obtenues correspondent à celles d'un système de diversité en réception avec une seule antenne d'émission et quatre antennes de réception (diversité d'ordre 4) utilisant une méthode MRC (*Maximum Ratio Combining*), à une perte de 3 dB près, si on utilise la même puissance totale. Dans la dernière équation, n_1 et n_2 sont des bruits gaussiens complexes de moyenne nulle. Ce code permet le maximum de diversité, mais le nombre de symboles émis à chaque utilisation du canal est égal à 1 seulement.

■ Codes spatio-temporels linéaires

Un intérêt du code d'Alamouti est de conduire avec un décodage simple à un système sans interférences spatiales en réception tout en multipliant le RSB par un facteur $\sum_{i,j} |h_{ij}|^2$. On parle

de schéma orthogonal.

On peut écrire en réception un système :

$$\tilde{\mathbf{y}} = \tilde{\mathbf{H}}\mathbf{x} + \tilde{\mathbf{b}}$$

avec :

$$\tilde{\mathbf{H}}^\dagger \tilde{\mathbf{H}} = \left(\sum_{i,j} |h_{i,j}|^2 \right) \mathbf{I}_{n_t}$$

On peut définir d'autres codes spatio-temporels possédant le même type de simplicité de décodage. En particulier les codes blocs spatio-temporels linéaires sont définis de la manière suivante. On considère n_t antennes d'émission et une durée de N intervalles de temps, pour transmettre K symboles x_i . La matrice code $\mathbf{X}$ prend la forme :

$$\mathbf{X} = \sum_{k=1}^K (\Re(x_k) \mathbf{A}_k + j \Im(x_k) \mathbf{B}_k)$$

La matrice de réception $\mathbf{Y}$ vaut :

$$\mathbf{Y} = \mathbf{H}\mathbf{X} + \mathbf{B}$$

Les matrices $\mathbf{A}_k$ et $\mathbf{B}_k$ sont des matrices complexes de dimension $n_t \times N$.

Par exemple pour le code d'Alamouti :

$$\mathbf{X} = \Re(x_1) \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} + j \Im(x_1) \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} + \Re(x_2) \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix} + j \Im(x_2) \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} = \begin{bmatrix} x_1 & -x_2^* \\ x_2 & x_1^* \end{bmatrix}$$

À partir des K valeurs complexes x_i , on définit un vecteur colonne $\widehat{\mathbf{x}}$ regroupant les parties réelles et imaginaires des composantes :

$$\widehat{\mathbf{x}} = \begin{bmatrix} \Re(x_1) & \Im(x_1) & \cdots & \Re(x_K) & \Im(x_K) \end{bmatrix}^T$$

On définit une matrice de canal $\widehat{\mathbf{H}}$ formée de $n_t N$ lignes et de $2K$ colonnes et obtenue en regroupant les matrices vect ($\mathbf{H}\mathbf{A}_k$) et j vect ($\mathbf{H}\mathbf{B}_k$) pour les K valeurs de k :

$$\widehat{\mathbf{H}} = \begin{bmatrix} \text{vect}(\mathbf{H}\mathbf{A}_1) & j \text{vect}(\mathbf{H}\mathbf{B}_1) & \cdots & \text{vect}(\mathbf{H}\mathbf{A}_K) & j \text{vect}(\mathbf{H}\mathbf{B}_K) \end{bmatrix}$$

À partir des matrices $\mathbf{B}$ et $\mathbf{Y}$, on définit des vecteurs colonnes :

$$\widehat{\mathbf{b}} = \text{vect}(\mathbf{B}), \quad \widehat{\mathbf{y}} = \text{vect}(\mathbf{Y})$$

et on peut écrire l'équation entrée-sortie sous la forme d'un système d'équations linéaires dont les inconnues sont les K symboles x_i :

$$\begin{aligned} \widehat{\mathbf{y}} &= \text{vect}(\mathbf{Y}) = \text{vect}(\mathbf{H}\mathbf{X} + \mathbf{B}) \\ \widehat{\mathbf{y}} &= \sum_{k=1}^K \Re(x_k) \text{vect}(\mathbf{H}\mathbf{A}_k) + \Im(x_k) \text{vect}(j \mathbf{H}\mathbf{B}_k) + \widehat{\mathbf{b}} \\ \widehat{\mathbf{y}} &= \widehat{\mathbf{H}}\widehat{\mathbf{x}} + \widehat{\mathbf{b}} \end{aligned}$$

Ce système comprend $n_r N$ équations linéaires et K inconnues complexes. Les K inconnues x_i ne pourront donc être récupérées que si $K \leq n_r N$.

7.3.8 Récepteurs optimaux et sous-optimaux

Cette section décrit très succinctement quelques récepteurs utilisés pour les systèmes MIMO. Le récepteur optimal au sens du maximum de vraisemblance (récepteur ML) choisit la matrice de code $\mathbf{X}$ qui minimise la distance avec le signal reçu $\mathbf{Y}$:

$$\min_{\mathbf{X}} \|\mathbf{Y} - \mathbf{H}\mathbf{X}\|^2$$

Puis à partir de $\mathbf{X}$, il récupère les symboles x_i .

Le récepteur optimal au sens du maximum de vraisemblance peut être trop complexe quand le nombre d'antennes est grand. Plusieurs solutions sous-optimales moins complexes ont été proposées dont certaines obtiennent des performances très proches de la solution optimale.

On peut distinguer les récepteurs linéaires et les non-linéaires. Les premiers effectuent une transformation linéaire (caractérisée par une matrice $\mathbf{A}$) sur le signal reçu de façon à réduire les interférences spatiales et à simplifier les calculs du récepteur ML, le critère devenant :

$$\min_{\mathbf{X}} \|\mathbf{A}\mathbf{Y} - \mathbf{X}\|^2$$

Dans les récepteurs non-linéaires, le récepteur effectue un prétraitement linéaire puis soustrait du résultat une estimation des interférences spatiales obtenue par un premier décodage (SIC : *Successive Interference Cancellation*).

La question de la réception dans un système MIMO présente une grande similarité avec celle de l'égalisation dans les canaux SISO dispersifs et avec les techniques de détection multi-utilisateurs utilisées en CDMA. Les techniques de réception utilisées sont très proches. Parmi les récepteurs on peut citer les récepteurs « *zero-forcing* », MMSE (*Minimum Mean Square Error*), et les architectures d'émission et de réception VBLAST, DBLAST.

■ Algorithme linéaire « zero-forcing » ou ZF

Un exemple simple de récepteur linéaire utilise comme matrice $\mathbf{A}$ la matrice pseudo-inverse de la matrice du canal $\mathbf{H}$. Dans ce cas, l'opération linéaire de multiplication par $\mathbf{A}$ a pour effet de supprimer les interférences spatiales. On parle alors d'algorithme « *zero-forcing* ». Comme en égalisation, cet algorithme cherche à annuler les interférences des autres émetteurs sans se soucier des conséquences sur le bruit, c'est-à-dire au prix d'une dégradation du *RSB*.

$$\mathbf{A} = (\mathbf{H}^\dagger \mathbf{H})^{-1} \mathbf{H}^\dagger$$

$$\mathbf{A}\mathbf{Y} = \mathbf{X} + \mathbf{A}\mathbf{B}$$

De meilleurs récepteurs linéaires sont possibles, tels que le récepteur MMSE qui effectue un filtrage linéaire optimal.

■ Algorithmes MMSE, MMSE-DFE, MMSE-SIC

L'algorithme MMSE cherche à minimiser les interférences et le bruit, c'est-à-dire le rapport signal sur bruit plus interférences *RSBI*. C'est un algorithme linéaire plus efficace que l'algorithme ZF en présence de faibles *RSB*. La matrice $\mathbf{A}$ qui minimise le critère MMSE est :

$$\mathbf{A} = \arg \min E (|\mathbf{A}\mathbf{y} - \mathbf{x}|^2)$$

$$\mathbf{A} = \mathbf{H}^\dagger (\mathbf{H}\mathbf{H}^\dagger + \mathbf{R}_b)^{-1}$$

Dans cette expression $\mathbf{R}_b$ représente la matrice de covariance du bruit.

Mais pour de grandes valeurs du RSB , ce critère laisse subsister de l'interférence entre les différentes voies. Pour améliorer les performances et se rapprocher de la capacité théorique, on peut utiliser des récepteurs non-linéaires cherchant à supprimer les interférences résiduelles tels que les récepteurs MMSE-DFE (*Minimum Mean Square Error Decision Feedback Equalizer* ou en français récursif à retour de décision) et MMSE-SIC.

L'algorithme MMSE-SIC (*MSE Successive Interference Cancellation*) est un algorithme non-linéaire qui supprime les interférences de manière itérative. Il décode l'un des trains de données (disons le train n° 1) par un récepteur MMSE puis il soustrait du signal reçu les interférences générées par le train n° 1 et décode le train de données n° 2 par un récepteur MMSE. Il itère ensuite ce processus de décodage d'un des trains de données en soustrayant les interférences générées par les trains déjà décodés.

L'algorithme MMSE-SIC permet d'atteindre la capacité théorique du canal MIMO ergodique.

■ Architecture VBLAST et fading rapide

L'approche MMSE-SIC a été proposée par les *Bell Laboratories* pour les systèmes MIMO dans l'architecture VBLAST (*Vertical Bell Laboratories Layered Space Time Architecture*).

Dans cette architecture, le train de symboles x_i est décomposé en n_t trains de symboles chacun d'eux aiguillé vers l'une des n_t antennes d'émission. Ces trains sont codés indépendamment les uns des autres sans codage spatio-temporel.

Le récepteur effectue un prétraitement linéaire du signal reçu. Puis il fonctionne de manière itérative avec une méthode de type MMSE-SIC. Le décodage commence par le train de données présentant le meilleur RSB .

La lettre V de VBLAST indique que l'encodeur organise les données successives en couches verticales (figure 7.32) par opposition à des couches diagonales utilisées dans l'architecture DBLAST qui est présentée plus loin.

La solution VBLAST est intéressante pour les canaux à fading rapide.

Dans le cas d'un fading rapide et d'un grand RSB , le gain de performance est égal au minimum des nombres d'antennes d'émission et de réception (gain dû au multiplexage spatial). Dans le cas d'un fading rapide et d'un RSB faible, le gain de performance est égal au nombre d'antennes de réception (gain de diversité).

■ Architecture DBLAST (Diagonal BLAST) et fading lent

L'architecture VBLAST ne permet pas d'atteindre la capacité de coupure d'un canal à fading lent car elle n'effectue pas de diversité en émission, les différents trains de données étant transmis indépendamment sur chaque antenne. Et à la différence du cas d'un fading rapide, il n'y a pas non plus de diversité temporelle. Aussi est-il nécessaire de mettre en œuvre un codage spatio-temporel ou de diversité temporelle pour réaliser une diversité en émission dans le cas du fading lent et pouvoir atteindre la capacité de coupure du canal.

L'architecture DBLAST (*Diagonal Bell Laboratories Layered Space Time Architecture*) est intéressante pour les canaux à fading lent. L'encodeur DBLAST utilise une disposition spatio-temporelle en couches diagonales.

Le train de symboles x_i est décomposé en n_t trains de symboles chacun d'eux aiguillés vers l'une des n_t antennes d'émission. Ces trains sont « dispersés » de façon diagonale sur les différentes antennes d'émission (figure 7.32) et sur le temps.

En réception, il faut combiner cette structure diagonale avec l'approche d'annulation successive des interférences.

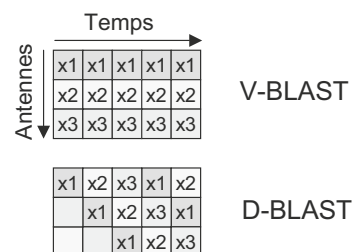


Figure 7.32 – Architectures VBLAST et DBLAST.

7.3.9 Compromis diversité-multiplexage spatial (gain en débit)

Dans un système MIMO, on cherche à optimiser le compromis entre le gain de diversité et le gain de multiplexage [A. Paulraj].

Le gain de multiplexage maximal $r_{\max}$ vaut :

$$r_{\max} = \min(n_t, n_r).$$

La diversité maximale $d_{\max}$ vaut :

$$d_{\max} = n_t n_r.$$

La diversité permet d'améliorer la probabilité d'erreur dans les canaux avec fading. Le multiplexage spatial permet d'augmenter le débit (l'efficacité spectrale) en augmentant le nombre de degrés de liberté de la communication.

Il existe un compromis entre le gain de diversité et le gain de multiplexage. L'objectif du codage spatio-temporel est d'optimiser ce compromis.

On définit le gain de diversité d comme la valeur de l'exposant du RSB dans la relation donnant la probabilité d'erreur p en fonction du RSB (pour de grandes valeurs de RSB) :

$$p(RSB) \approx RSB^{-d},$$

$$d = - \lim_{RSB \rightarrow \infty} \frac{\log(p(RSB))}{\log(RSB)}.$$

De même, on définit le gain en multiplexage r comme :

$$R(RSB) \approx r \log_2(RSB) \text{ bits/s/Hz}.$$

$$r = \lim_{RSB \rightarrow \infty} \frac{R(RSB)}{\log_2(RSB)}.$$

Dans cette expression $R(RSB)$ représente le débit possible pour un RSB donné avec une probabilité d'erreur donnée. Si on augmente de 3 dB le RSB , dans le cas SISO, pour une probabilité d'erreur donnée et une modulation QAM, on peut augmenter d'un bit la taille de la constellation QAM. Dans le cas MIMO avec multiplexage spatial, on peut augmenter de r bits.

Bibliographie

AÏSSAT H., CIRIO L., GRZESKOWIAK M., LAHEURTE J.M., PICON O. — *Reconfigurable circularly polarized antenna for short-range communication systems*, IEEE Transactions on Microwave Theory and Techniques, vol. 54, n°6, pp. 2856-2863, June 2006

ALAMOUTI S. — *A simple transmit diversity technique for wireless communications*, IEEE Journal on Selected Areas in Communications, Special Issue on Signal Processing for Wireless Communications, vol. 16, Issue 8, pp. 1451-1458, 1998.

BIGLIERI E., TARICCO G. — *Transmission and Reception with Multiple Antennas : Theoretical Foundations*, Now Publishers Inc. the essence of knowledge, USA, 2004.

COMPTON Jr., R.T. — *Adaptive Antennas : Concepts and Performances*, New York, Prentice Hall, 1988.

DIETRICH C. B., DIETZE K., NEALY J. R., STUTZMAN W. L. — *Spatial, polarisation, and pattern diversity for wireless handheld terminals*, IEEE Transactions on Antennas and Propagation, vol. 49, n° 9, sept. 2001.

- FOSCHINI G.J., GANS M.J. — *On Limits of Wireless Communications in a Fading Environment when Using Multiple Antennas*, Wireless Personal Communications, vol. 6, pp. 311-335, 1998.
- GESBERT D., SHAFI M., SHIU D. S., SMITH P., AND NAGUIB A. — *From theory to practice : An overview of mimo space-time coded wireless systems*, IEEE Journal on Selected Areas in Communications, 21(3) :281– 302, avril 2003.
- GODARA, L. C. — *Smart antennas*, Boca Raton, CRC Press, 2004.
- HAYKIN, S. — *Array Signal Processing*, New York, Prentice Hall, 1985.
- JAFARKHANI H. — *Space-Time Coding Theory and Practice*, Cambridge, 2000.
- JAKES W. C., — *Microwave mobile communications*, IEEE Press, 1993.
- LARSSON E. G., STOICA P. — *Space-Time Block Coding for Wireless Communications*, UK, Cambridge University Press, 2003.
- LINDMARK B., NILSSON M. — *On the available diversity gain from different dual-polarized antennas*, IEEE Journal on Selected Areas in Communications, vol. 19, n°2, février 2001.
- MATTHEIJSEN P., HERBEN M.H.A.J., DOLMANS G., LEYTEN L. — *Antenna-pattern diversity versus space diversity for use at handhelds*, IEEE Transactions on Vehicular Technology, vol. 53, n° 4, juillet 2004.
- MONZINGO, R. A., MILLER T. W. — *Introduction to Adaptive Arrays*, Raleigh, SciTech Publishing, Inc., 2004.
- NARAYANAN R. M., ATANASSOV K., STOILJKOVIC V., KADAMBI G. R. — *Polarization diversity measurements and analysis for antenna configurations at 1 800 MHz*, IEEE Journal on Transactions on Antennas and Propagation, vol. 52, n° 7, juillet 2004.
- PAULRAJ A., NABAR R., AND GORE D., — *Introduction to Space-Time Wireless Communications*, Cambridge University Press, mai 2003.
- PAULRAJ A.J., KAILATH T. — *Increasing Capacity in Wireless Broadcast Systems Using Distributed Transmission/Directional Reception*, U.S. Patent, n°5345500, 1994.
- POUSSOT B., LAHEURTE J-M., CIRIO L., PICON O., DELCROIX D., DUSSOPT L. — *Diversity measurements of a reconfigurable antenna with switched polarization and patterns*, IEEE Transactions on Antennas and Propagation, vol. 56, n°1, janvier 2008.
- STEELE R., — *Mobile radio communications : second and third generation cellular and WATM systems*, John Wiley, 1999.
- TAROKH V., SESHADRI N., CALDERBANK A.R. — *Space-Time codes for high data rate wireless communication : performance criterion and code construction*, IEEE Transactions on Information Theory, vol. 44, n°2, pp. 744-765, March 1998.
- TAROKH V., SESHADRI N., CALDERBANK A. R. — *Space-Time codes for high data rates wireless communications*, IEEE Trans. Information Theory, vol. 44, issue 2, pp. 744-765, mars 1998.
- TELATAR E. — *Capacity of Multi-Antenna Gaussian Channels*, Eur. Trans. Telecomm. , vol. 10, n° 6, pp. 585-596, novembre 1999.
- TELATAR E. — *Capacity of Multi-Antenna Gaussian Channels*, Technical report, AT & T Bell Labs, 1995.

TSE D., VISWANATH P. — *Fundamentals of Wireless Communication*, USA, Cambridge University Press, 2005.

Van TREES H. L. — *Optimum Array Processing*, New York, John Wiley & Sons, 2002.

WEINER M. M. — *Adaptive Antennas and Receivers*, Boca Raton, CRC Press, 2005.

ZHENG L., TSE D. — *Diversity and Multiplexing : a fundamental trade-off in multiple antennas channels*, IEEE Trans. on Information Theory, vol. 49, n°5, pp. 1073-1096, mai 2003.

8 • ANTENNES DÉFINIES SELON LE PHÉNOMÈNE PHYSIQUE CRÉANT LE RAYONNEMENT

La forme des antennes peut être très variée. Il en découle un fonctionnement différent pour chacune d'elles. Nous présentons dans ce chapitre une forme de classement qui n'est pas le seul possible. Il consiste à mettre en avant la caractéristique fondamentale que l'on recherche dans un type d'antenne donné. Celui-ci est choisi en fonction de contraintes particulières, soit liées à la physique, soit imposées par le système.

Selon cette classification, nous verrons, dans ce chapitre, quelles sont les antennes dont le fonctionnement repose sur des principes physiques. Dans le chapitre suivant, seront abordées les antennes dont la forme particulière est imposée par le système. : taille, forme, fréquence, etc.

Les antennes étudiées précédemment ne sont pas reprises dans ce chapitre, même si elles appartiennent à l'une des catégories présentées. Nous ne prétendons pas ici faire un catalogue exhaustif de toutes les antennes. Certaines ont été laissées de côté car elles sont relativement peu utilisées. Certaines reposent sur des phénomènes physiques très particuliers qui nous entraîneraient trop loin dans les formulations théoriques. C'est le cas des antennes fractales ou des antennes à ferrites. Pour terminer cette introduction, précisons qu'une antenne peut appartenir à plusieurs types. Par exemple, une antenne planaire peut être large bande. Cette remarque permet de comprendre la difficulté à proposer un classement unique.

8.1 Antennes boucles

Ces antennes se présentent sous la forme d'une boucle unique ou de boucles multitours. Elles ont été très utilisées dès les premiers âges de la radio, notamment par H. Hertz. Il est classique de distinguer l'étude de la boucle en fonction de son diamètre ou de sa circonférence exprimés en longueur d'onde.

En réception, lorsqu'elle est de petite dimension, on la retrouve dans les applications de radio-communications (réception en grandes ondes), de localisation en navigation radio ou de sonde en mesures de champ magnétique. Lorsqu'elle est de dimension comparable à la longueur d'onde, elle est souvent utilisée comme élément d'un groupement de source ou par les radioamateurs (*cubical-quad*).

8.1.1 Diagramme de rayonnement de l'antenne boucle à courant uniforme

L'étude théorique de la boucle de petite dimension, appelée aussi antenne magnétique, a été faite au paragraphe 3.1.5. Le calcul du champ émis, et donc du diagramme de rayonnement, peut être fait en assimilant l'antenne à un doublet magnétique (dual du doublet électrique). Il est montré que les champs lointains des boucles circulaires et carrées sont similaires dès lors que les surfaces des deux éléments sont égales.

Il est aussi intéressant de considérer la boucle comme étant constituée par quatre doublets et parcourue par un courant uniforme en amplitude et en phase.

Le calcul de la composante de champ électrique E_ϕ en champ lointain dans le plan perpendiculaire à la structure (plan yz) ne va dépendre que des doublets 2 et 4 puisque le rayonnement des deux autres va s'annuler dans ce plan de symétrie. Dans ce cas, nous nous retrouvons avec un groupement de deux sources en opposition de phase, dû aux sens opposés des courants. Le calcul a été fait dans le chapitre sur le groupement de sources et donne :

$$E_\phi = -2jE_{\phi 0} \sin \left(\frac{\pi d}{\lambda} \sin \theta \right) \quad [8.1]$$

Avec $E_{\phi 0}$ champ électrique dû à un doublet unique

Si le diamètre $d \ll \lambda$ alors : $E_\phi = -2jE_{\phi 0} \frac{\pi d}{\lambda} \sin \theta$ et comme le champ émis par le doublet est $E_{\phi 0} = \frac{j60\pi Id}{r\lambda}$, il vient, en introduisant l'aire de la boucle, les composantes de champ :

$$\begin{aligned} E_\phi &= \frac{120\pi^2 I}{r} \frac{A}{\lambda^2} \sin \theta \\ H_\theta &= \frac{\pi I}{r} \frac{A}{\lambda^2} \sin \theta \end{aligned} \quad [8.2]$$

Dans le cas général d'une boucle parcourue par un courant uniforme mais dont la taille est quelconque, le champ électrique dépend de la fonction de Bessel de première espèce, d'ordre 1 :

$$E_\phi = \frac{60\pi I}{r} \beta a J_1(\beta a \sin \theta) \quad [8.3]$$

où $\beta = \frac{2\pi}{\lambda}$ et a est le rayon de la boucle supposée circulaire.

Le terme βa représente la circonférence de la boucle exprimée en longueur d'onde, appelée C_λ . Le diagramme de rayonnement est représenté pour trois diamètres différents.

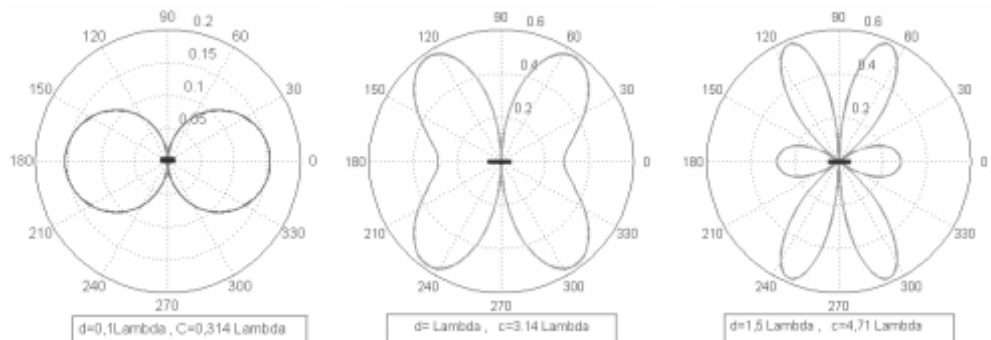


Figure 8.2 – Diagramme de rayonnement pour trois boucles circulaires de diamètre $\lambda/10$; λ ; $1,5\lambda$ dans le plan θ pour $C_\lambda = 0,314$; $3,14$; $4,71$

Nous reconnaissons le diagramme en $\sin \theta$ pour les faibles diamètres, diagramme identique à celui du doublet électrique. Nous pouvons noter que le rayonnement est toujours nul dans la direction perpendiculaire à la boucle.

8.1.2 Résistance de rayonnement de l'antenne boucle à courant uniforme

On commence par calculer la puissance rayonnée en intégrant le vecteur de Poynting dans la sphère entourant l'antenne. On en déduit la résistance de rayonnement en reconnaissant R_r dans la formule $W = R_r I^2$. Pour une boucle de longueur quelconque mais toujours parcourue par un courant uniforme, on trouve que la résistance dépend de la fonction de Bessel de second ordre :

$$R_r = 60\pi^2 C_\lambda \int_0^{2C_\lambda} J_2(x) dx \quad [8.4]$$

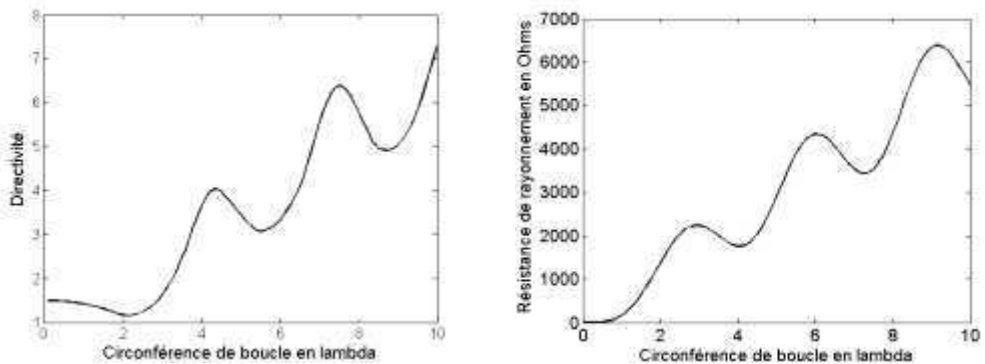


Figure 8.3 – Résistance de rayonnement et directivité de boucle de longueur quelconque

Nous observons les caractéristiques suivantes :

- Elle est très faible pour les boucles à faible diamètre, de nombre de tours n , la formule devient alors :

$$R_r \cong 31\,200 \left(n \frac{A}{\lambda^2} \right)^2 \cong 197 C_\lambda^4 \quad [8.5]$$

ce qui donne $2\,\Omega$ pour une boucle de diamètre $0,1\lambda$.

- Elle augmente très rapidement avec la taille de la boucle, puisqu'en puissance 4 de la circonférence ; sa valeur est de l'ordre de $220\,\Omega$ pour une boucle de diamètre λ .
- Elle varie relativement peu pour des diamètres inférieurs à λ et reste adaptable à une ligne d'impédance caractéristique de 50 ou $75\,\Omega$.

8.1.3 Directivité de l'antenne boucle à courant uniforme

Toujours dans le cas des boucles circulaires parcourues par un courant uniforme, de longueur quelconque, la directivité vaut :

$$D = \frac{2C_\lambda [J_1^2(C_\lambda \sin \theta)]}{\int_0^{2C_\lambda} J_2(x) dx} \quad [8.6]$$

Puisque les fonctions impliquées sont identiques, il est normal de retrouver une forme similaire à celle de la résistance de rayonnement. On remarque que l'on retrouve une directivité de 1,5, correspondant à celle du doublet électrique pour une boucle de faible dimension ($< 0,5\lambda$). Ce résultat était prévisible puisque les propriétés en rayonnement sont similaires.

8.1.4 Considérations sur l'antenne boucle à courant non uniforme

Obtenir l'uniformité du courant sur une boucle dont le diamètre mesure environ λ ou plus n'est pas chose aisée (introduction de déphaseurs dans boucle). En pratique, l'application où l'antenne boucle est réellement utilisée est la réception TV comme antenne d'intérieur. Dans ce cas, sa circonférence C_λ , comptée en longueur d'onde, est de l'ordre de la longueur d'onde. De façon générale, on ne peut plus supposer le courant constant et il est nécessaire de le décrire comme une fonction de l'azimut par une série de Fourier :

$$I(\phi) = I_0 + 2 \sum_{n=1}^m I_n \cos n\phi$$

Par exemple, à la première résonance ($C_\lambda \approx 1$), il est décrit par $I(\phi) = 2I_1 \cos \phi$ et présente donc des ventres pour $\phi = 0, \pi$ et des nœuds pour :

$$\phi = \pi/2, 3\pi/2$$

Dans ce cas, les caractéristiques sont profondément modifiées et il faut tenir compte d'un facteur d'épaisseur :

$$\Omega = 2 \ln (2\pi b/a)$$

avec b rayon moyen de la boucle et a rayon du conducteur de la boucle.

La résistance de rayonnement présente des valeurs croissantes en fonction de C_λ mais avec des résonances ($C_\lambda = 1, 2, 3, \dots$) ou des antirésonances ($C_\lambda = 0,5; 1,5; 2,5, \dots$). Cela permet d'avoir une courbe relativement plate sur une plage de 0,7 à 1,2 (de l'ordre de 100 Ω).

La réactance présente les mêmes résonances et antirésonances. Les valeurs restent compensables ($< 400 \Omega$ pour $\Omega = 10$) par un circuit d'adaptation sur une plage de 0,8 à 1,7 λ .

Ces deux éléments vont permettre l'utilisation de cette boucle dans la bande TV-UHF de 470 à 890 MHz correspondant à 0,9 à 1,7 λ .

La présence de nœuds et de ventres de courant va introduire des dissymétries et des modifications de résultats dans les diagrammes de rayonnement (figure 8.4)

Par exemple, le champ est maximum dans le plan perpendiculaire à la boucle alors qu'il était nul dans le cas du courant uniforme.

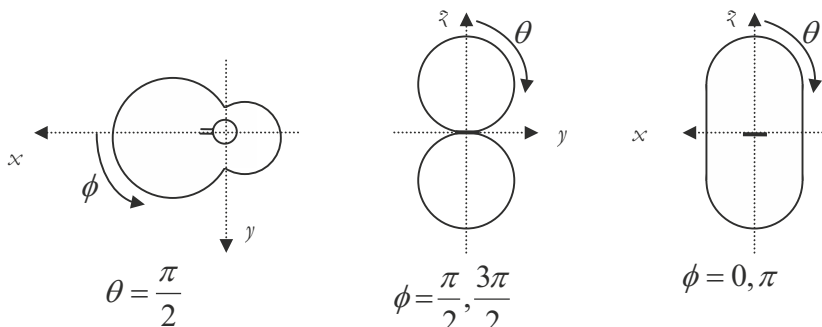


Figure 8.4 – Diagramme de rayonnement dans les plans horizontaux et verticaux de $|E_\phi|$ pour une boucle circulaire de $C_\lambda = 1$.

8.1.5 Efficacité de rayonnement

Lorsque les dimensions de l'antenne boucle sont faibles, il faut tenir compte des pertes ohmiques R_p qui sont du même ordre de grandeur. Ces pertes sont essentiellement dues à la résistance d'effet de peau et donc :

$$R_p = \frac{L}{d} \sqrt{\frac{f\mu}{\pi\sigma}}$$

avec L : circonférence de la boucle

μ : perméabilité du milieu

σ : conductivité du matériau

f : fréquence

L'efficacité de rayonnement devient :

$$k = \frac{1}{1 + \left(\frac{R_p}{R_r}\right)} \quad [8.7]$$

Elle est très faible en général et dépend très fortement de la fréquence de travail.

L'antenne boucle dans laquelle on introduit un barreau de ferrite est appelée antenne ferrite.

L'introduction de ce matériau de perméabilité magnétique μ_r permet d'augmenter la self d'antenne ainsi que la résistance de rayonnement. Il faut considérer la perméabilité effective car il y a passage des lignes de champ magnétique dans l'air. Cette résistance devient, pour une boucle de petite dimension :

$$R_r = 197 \mu_{eff}^2 n^2 C_\lambda^4 \quad [8.8]$$

Cela permet de l'augmenter significativement puisqu'elle augmente en μ_{eff}^2 .

8.2 Antennes à résonateurs

Aujourd'hui, le domaine des télécommunications est en plein essor et l'antenne constitue un élément essentiel dans la chaîne de transmission. Parmi tous les types d'antennes disponibles, on distingue les antennes à résonateurs diélectriques et à bande interdite photonique. Ces deux antennes ont des spécificités bien particulières que nous allons présenter dans ce qui suit.

8.2.1 Antennes à résonateur diélectrique

Les antennes à résonateur diélectrique permettent de compenser un des inconvénients majeurs des antennes plaquées à savoir la bande passante limitée à quelques %, ce qui constitue un handicap pour les communications à hauts débits. De plus, grâce à l'utilisation de matériaux à constante diélectrique élevée ($20 < \epsilon_r < 100$) et à faible tangente de pertes ($\tan \delta < 10^{-4}$), ces antennes présentent des dimensions bien plus petites que l'antenne imprimée résonante demi-onde.

L'antenne à résonateur diélectrique présentée figure 8.5 est constituée d'un résonateur diélectrique (généralement de forme cylindrique ou en anneau) reporté sur un support, qui tient lieu de plan de masse, auquel est associé un dispositif d'alimentation qui assure l'excitation du résonateur. Ce dernier est dépourvu de partie métallique ce qui limite d'autant les pertes ohmiques. On remarque sur la figure que l'excitation a été ici réalisée par couplage d'une ligne micro ruban à travers une fente rectangulaire découpée dans le plan de masse.

Grâce à un facteur de qualité élevé, ces structures ont été dans un premier temps utilisées dans les circuits micro-ondes blindés pour assurer des fonctions de filtrage ou d'oscillateur. Une fois le

résonateur laissé en environnement libre, on constate que le facteur de qualité décroît sensiblement ce qui laisse entrevoir une application en tant qu'antenne puisque la puissance perdue est dorénavant rayonnée.

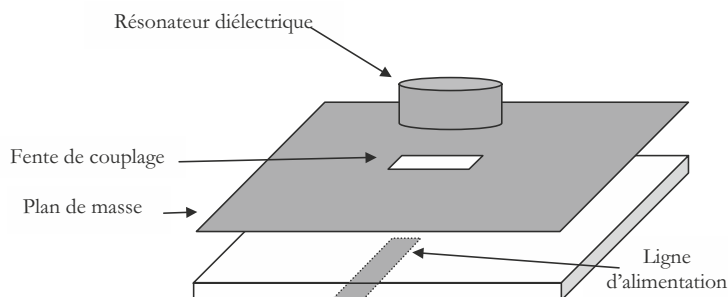


Figure 8.5 – Antenne à résonateur diélectrique.

Précisons que l'efficacité de rayonnement sera d'autant plus élevée que le mode sera convenablement excité. À ce propos, comme dans toute cavité résonante, il existe un très grand nombre de modes qui peuvent être excités. On distingue les modes Transverses Électriques (TE), Transverses Magnétiques (TM) et les modes hybrides. Les antennes à résonateurs diélectriques présentent des diagrammes de rayonnement différents selon le mode excité. Cela nécessite un soin particulier dans la détermination de la position du point d'excitation sous, ou à proximité du résonateur, de façon à n'exciter que le mode désiré en écartant les autres.

Le diagramme de rayonnement d'un résonateur isolé sans plan de masse est proche de celui d'un dipôle électrique ou magnétique selon le mode excité. Cependant, pour des valeurs de permittivité plus faibles (ϵ_r , de 20 à 40) généralement utilisées pour la conception d'antennes, on relève une modification sensible du gain du résonateur. Ce gain est relevé d'environ 3 dB lorsque le résonateur est placé sur un plan de masse à travers un substrat diélectrique.

La plupart des techniques retenues pour alimenter les antennes micro ruban sont compatibles avec l'alimentation des résonateurs diélectriques. La figure 8.6 montre les différentes configurations. La plus usuelle consiste à utiliser une ligne imprimée sur laquelle est déposé le résonateur. L'adaptation est obtenue en ajustant la position de l'élément rayonnant sur la ligne. Particulièrement simple à réaliser, cette solution a pour inconvénient la présence possible d'un rayonnement parasite dû à la ligne et qui peut sensiblement modifier le rayonnement propre du résonateur. Une alternative consiste à alimenter la structure directement par sonde coaxiale. Ici aussi, la position de la sonde déterminera le mode excité et le niveau d'adaptation. Cette solution

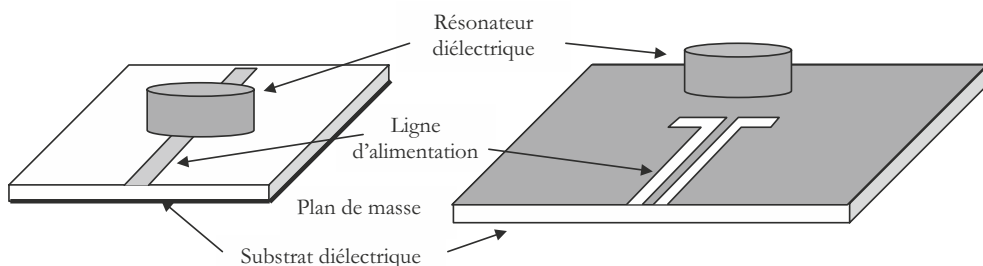


Figure 8.6 – Dispositifs d'alimentation d'antennes à résonateurs diélectriques.

nécessite cependant un perçage du substrat et du résonateur et ne facilite pas l'implantation de l'antenne dans des dispositifs à fort degré d'intégration. L'alimentation par couplage à travers une fente (figure 8.5) est attrayante car elle présente une bonne isolation du résonateur avec la ligne imprimée. En ajustant la résonance de la fente proche de celle du résonateur, il est possible d'augmenter sensiblement la bande passante de la structure. Enfin, l'utilisation d'un guide d'onde coplanaire (CPW) peut être également retenue. Le plan de masse se situe dans ce cas au même niveau que la ligne, ce qui facilite l'intégration des composants discrets et ceci, sans nécessiter de perçage au niveau du substrat. Le couplage est alors obtenu en chargeant la CPW à l'aide d'ouvertures (fentes rectangulaires, circulaires, circuit ouvert...) placées en bout de ligne.

Notons que des variantes au résonateur diélectrique simple ont été proposées dans la littérature. La superposition de résonateurs permet d'obtenir un fonctionnement multifréquences au détriment de l'encombrement tandis qu'en imbriquant des résonateurs entre eux, il est possible d'élargir la bande passante (de 10 % pour un résonateur simple à 40 % pour une structure à deux éléments imbriqués).

8.2.2 Antennes à résonateur à bande interdite électromagnétique (BIE)

Une des principales contraintes dans les antennes concerne l'encombrement. Cela est particulièrement vrai pour les antennes directives à grand gain dont les dimensions (ouverture, épaisseur...) sont bien souvent supérieures à λ (antennes paraboliques, cornet...). Dans le cas particulier des antennes imprimées, l'utilisation d'un substrat de permittivité élevée constitue une solution. Malheureusement, cela se traduit par une diminution des performances de l'antenne. Deux solutions peuvent alors être retenues pour réduire cette dégradation des performances :

- la première consiste, par une technique de micro-usinage, à retirer localement le substrat au dessous et dans l'environnement proche de l'antenne de façon à limiter l'apparition des ondes de surfaces qui altèrent le gain.
- La seconde consiste à utiliser des antennes à Bandes Interdites Électromagnétiques (BIE). On parle ici d'antennes à résonateurs BIE. Il s'agit d'antennes associées à des structures périodiques dont l'application première était de concentrer l'énergie dans une direction privilégiée en augmentant ainsi le gain de la structure rayonnante tout en conservant une épaisseur réduite, principale propriété des antennes planaires.

Les matériaux à bandes interdites électromagnétiques sont des structures périodiques à une, deux ou trois dimensions (figure 8.7). Ils peuvent être constitués d'un empilement de couches diélectriques (cas 1D) ou d'un arrangement périodique volumique de tiges métalliques ou diélectriques (cas 3D) ce qui rend le matériau dans lequel se propage l'onde électromagnétique à la fois anisotrope et dispersif. Cet ensemble crée un filtrage spatial et fréquentiel et provoque l'apparition de bandes interdites. Une modification (ou perturbation) de la périodicité du motif permet cependant, et sous conditions, d'autoriser une bande de fréquences à l'intérieur de la bande interdite. On parle alors de matériau BIE à défaut. La caractérisation d'un matériau BIE est effectuée au moyen du diagramme de dispersion d'où l'on déduit les bandes interdites.

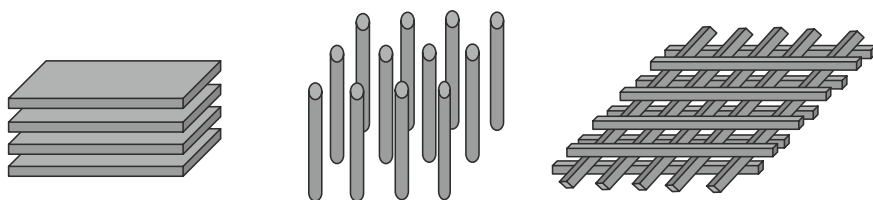


Figure 8.7 – Différents types de matériaux BIE 1D, 2D et 3D.

La conception d'un matériau BIE est liée au choix et à la forme des matériaux qui le composent et à la périodicité de ces motifs. Dans le cas plus général d'une structure tridimensionnelle, les trois axes de périodicités permettent, en théorie, de contrôler la propagation de l'onde électromagnétique dans une direction. Cela ouvre des perspectives intéressantes pour des applications de reconfigurabilité dans les antennes, domaine largement exploité aujourd'hui. Un autre exemple concerne les structures superstrats qui utilisent également les propriétés de périodicité. On retrouve les superstrats au-dessus des antennes sur substrat à haute permittivité pour la conception des antennes miniatures. Une augmentation du gain est obtenue en optimisant la résonance entre les différentes couches.

Appliqués aux antennes imprimées, les matériaux BIE ont été utilisés pour réduire les pertes par ondes de surface, phénomènes bien souvent constatés lorsque l'épaisseur et la permittivité du diélectrique deviennent trop importantes. Une dégradation du gain est alors constatée. Par la suppression des ondes de surfaces, elles permettent également de minimiser les couplages non désirés entre les éléments d'un réseau. La figure 8.8 représente un exemple d'antenne sur substrat haute permittivité ($\epsilon_r = 10,2$) et utilisant un matériau BIE. L'absence de trous métallisés réduit le coût de fabrication puisqu'il s'agit d'une structure imprimée. L'augmentation du gain est ici de 3 dB comparée à la même antenne sans matériau BIE. Il est à noter que la position de l'antenne et l'écartement entre l'antenne et la structure périodique ont une influence sensible sur la fréquence de résonance du patch, sur le niveau d'adaptation et sur les caractéristiques de rayonnement.

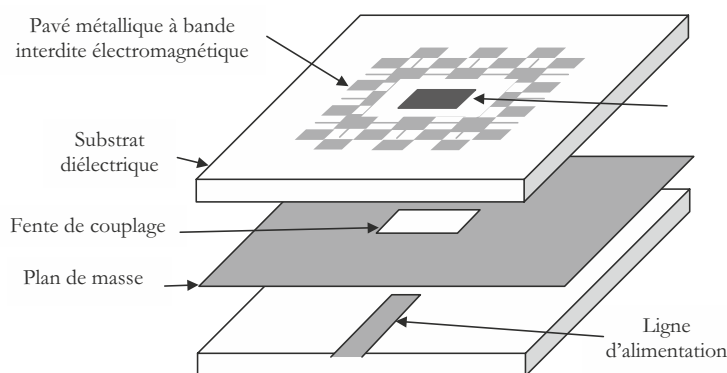


Figure 8.8 – Antenne patch imprimée utilisant un motif à bande interdite électromagnétique.

8.3 Antennes à ondes progressives

Le développement des antennes à long fil supportant des ondes stationnaires ou des ondes progressives s'est fortement ralenti depuis les années soixante. La seule antenne ayant perduré pour des applications de transmission en bande HF *via* l'ionosphère (3-30 MHz) est l'antenne rhombique, appelée aussi antenne losange.

L'antenne à long fil supporte des ondes stationnaires (extrémité en circuit ouvert) ou des ondes progressives (extrémité fermée sur une charge adaptée). Nous signalerons les différences essentielles de leurs caractéristiques.

De façon générale, cette antenne fait partie de la catégorie d'antennes dont les dimensions sont suffisantes pour permettre un fonctionnement en ondes progressives. On trouve deux catégories :

- les antennes terminées par une charge adaptée à l'impédance caractéristique de la ligne pour éviter l'onde retour (antennes en V, losange) ;
- les antennes en circuit ouvert à leur extrémité mais qui sont soit longues (hélice longue) soit épaisses (dipôle linéaire épais).

Dans les deux cas, l'étude sera la même car les répartitions de courant le long des structures seront similaires.

Nous étudierons ici les antennes terminées par leur charge caractéristique. À ce propos et de façon pratique, deux éléments nous éloignent de ces hypothèses de non-réflexion :

- Les réflexions parasites qui peuvent exister, notamment au niveau de l'excitation (jonction ligne antenne).
- La présence de champ dans un environnement non immédiat autour de la ligne fait que seulement une partie du signal est absorbée par cette charge adaptée.

8.3.1 Champ rayonné par l'antenne à long fil unique fonctionnant en ondes progressives

On peut supposer, dans ce cas, que le courant est uniforme sur le fil et peut s'écrire en fonction de l'abscisse spatiale x :

$$i(x) = i_0 e^{-j\beta x}$$

On fait ici l'hypothèse d'une ligne sans pertes, donc seule la phase du courant varie linéairement le long de celle-ci. Cette hypothèse revient à négliger le rayonnement se produisant le long de la ligne, sans erreurs notables.

La structure se comporte comme un ensemble de doublets électriques. Le champ élémentaire rayonné par celui-ci en x dépend du courant qui le traverse :

$$dE = -j \frac{60\pi i(x)}{\lambda r} \sin \theta e^{-j\beta r}$$

Sachant que l'on peut poser $r = r_0 - x \cos \theta$.

Il suffit d'intégrer, afin de calculer le rayonnement du fil dans une direction Δ , formant un angle θ avec l'axe de l'antenne de longueur L . En faisant l'hypothèse que le champ est calculé à grande distance, on obtient :

$$E(\theta) = -j \frac{60 i_0}{\lambda r} \sin \theta e^{-j\beta r_0} \int_0^L e^{-j\beta x(1-\cos \theta)} dx \quad [8.9]$$

Après intégration, on obtient le module du champ :

$$|E(\theta)| = \frac{60 i_0}{r} \frac{\sin \theta}{1 - \cos \theta} \sin \frac{\pi L}{\lambda} (1 - \cos \theta) \quad [8.10]$$

Nous notons les caractéristiques suivantes :

- Le diagramme de rayonnement est de révolution autour du fil.
- Le diagramme s'incline d'un certain angle par rapport à l'axe de l'antenne. Cette inclinaison est fonction de la longueur d'antenne exprimée en longueurs d'onde. Toutefois, cet angle variant relativement peu, cette antenne affiche une directivité peu sensible à la fréquence d'utilisation.
- Comme le champ E est perpendiculaire à la direction considérée, on peut le décomposer en une composante horizontale et une composante verticale.

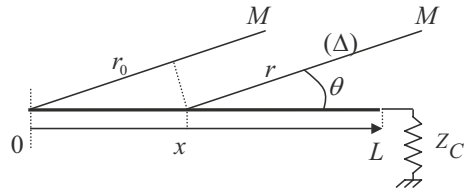


Figure 8.9 – Représentation de l'antenne à long fil supportant une onde progressive

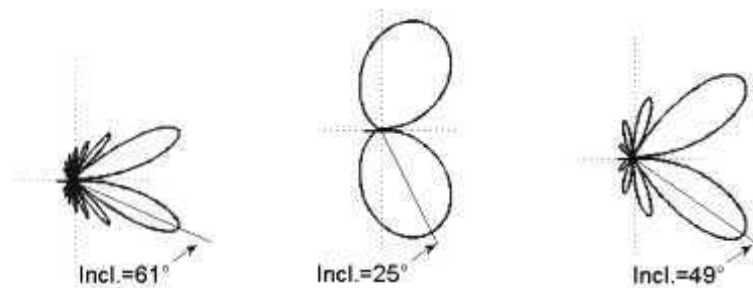


Figure 8.10 – Diagramme de rayonnement pour des longueurs $\lambda/2$, 2λ et 4λ .

- Le champ électrique s'inverse entre deux lobes successifs, phénomène lié à l'inversion des courants dans deux portions successives de fil de longueur égale à la demi-onde.
- Dans chaque lobe symétrique autour de l'axe, le champ rayonné affiche des composantes verticales en opposition de phase, ce qui permettra l'obtention de caractéristiques intéressantes en rayonnement lors du groupement de deux antennes en V.
- De façon plus précise, il faudrait considérer dans le calcul le fait que la vitesse de propagation sur la structure est plus faible que la vitesse de la lumière. Cela a pour conséquence que l'angle d'inclinaison par rapport à la perpendiculaire augmente lorsque v/c diminue (25° pour $L = 0,5\lambda$ avec $v/c = 1$ et 31° pour $v/c = 0,8$). Pour les antennes à fil terminées par leur impédance caractéristique, les vitesses sont très proches mais ce n'est plus le cas pour les antennes de la seconde catégorie (hélice, dipôle épais en circuit ouvert)

Nous pouvons noter certaines différences par rapport à l'étude des mêmes structures fonctionnant en ondes stationnaires :

- le diagramme de rayonnement devient symétrique et n'est plus unidirectionnel comme ici mais les valeurs des angles des minima sont conservées (vrai aussi pour les maxima dès que la longueur dépasse environ 3λ) ;
- il existe un angle d'inclinaison qui éloigne la direction du maximum de $\theta = 90^\circ$;
- l'angle d'ouverture à 3 dB est différent : environ 60° pour l'antenne à ondes progressive, et 78° pour celle fonctionnant en ondes stationnaires.

En pratique, il est très difficile d'utiliser cette antenne isolée dans l'espace. Nous devons donc considérer la proximité du sol pour l'antenne placée horizontalement à une hauteur H . Dans ce cas, il faut introduire le facteur de sol $\sin(\beta H \sin \theta)$ qui conserve seulement la partie supérieure mais en la modifiant légèrement (figure 8.11). Lorsque la hauteur dépasse λ , alors un lobe mineur apparaît entre le lobe principal et la direction de l'antenne.

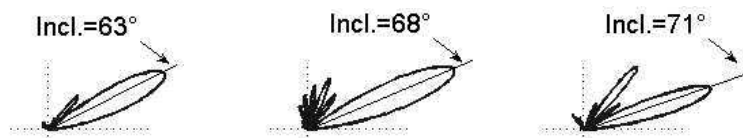


Figure 8.11 – Diagramme de rayonnement de l'antenne à ondes progressives de longueur 4λ au-dessus d'un sol conducteur de hauteur $H = \lambda/2$, $3\lambda/4$ et λ .

8.3.2 L'antenne rhombique ou en losange

Avant d'introduire l'antenne rhombique, il est intéressant de mentionner l'antenne en V pour comprendre le principe retenu. Chaque branche du V sera parcourue par un courant constant qui rayonnera un diagramme élémentaire correspondant à celui d'un fil long. Comme deux lobes consécutifs rayonnent des champs en opposition de phase, si le demi-angle au sommet δ vaut deux fois l'angle d'inclinaison α par rapport à l'axe, on aura une annulation partielle des lobes extérieurs. Dans ce cas, le diagramme possédera un lobe principal unique dont le maximum serait dans l'axe de l'antenne. Cela, bien sûr, à condition d'alimenter les deux fils en opposition. Ce cas de figure est réalisé en alimentant la structure avec une ligne symétrique de type bifilaire (d'impédance caractéristique de 600 à 800 Ω , valeur de la résistance de terminaison). Un avantage pratique de cette antenne, et qui justifie son utilisation, est qu'il n'est plus nécessaire de connecter l'impédance caractéristique à la terre.

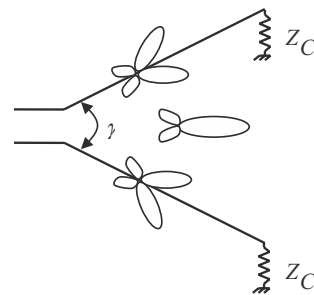


Figure 8.12 – Rayonnement de l'antenne en V

Nous pouvons considérer l'antenne rhombique comme la mise en série de deux antennes en V.

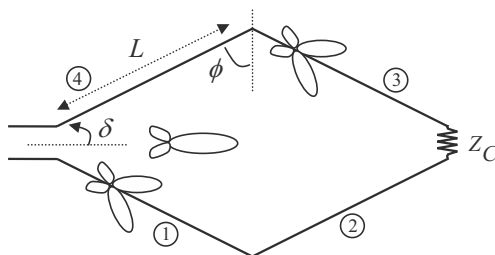


Figure 8.13 – Description géométrique de l'antenne rhombique ou losange

L'antenne horizontale est constituée de quatre branches de longueur L et située au-dessus du sol à une hauteur H . Quand on effectue la sommation des champs rayonnés par chaque doublet élémentaire sur deux fils opposés (brins 1 et 3 par exemple) en tenant compte de leur déphasage respectifs dû à leur différence de distance au point d'alimentation, et au fait qu'ils sont parcourus par des courants inverses, on détermine le champ rayonné dans le plan vertical (en négligeant les impédances mutuelles de couplage) :

$$|E_V| = \frac{2\sqrt{1 - \cos^2 \alpha \cos^2 \phi}}{1 - \sin \phi \cos \alpha} \sin \left(\frac{2\pi H}{\lambda} \sin \alpha \right) \left[\sin^2 \left(\frac{\pi L}{\lambda} \sin \alpha (1 - \sin \phi \cos \alpha) \right) \right] \quad [8.11]$$

avec α : angle d'élévation au-dessus du sol

ϕ : demi-angle obtus du losange

H : hauteur au-dessus du sol

L : longueur d'une branche du losange

Le demi-angle à l'origine δ n'apparaît pas dans cette formule car il est généralement confondu avec l'angle d'inclinaison α par rapport au sol.

Le diagramme est similaire à celui du fil isolé au-dessus d'un sol mais avec un gain double puisqu'il y a renforcement des champs émis par les deux antennes en V dans la direction axiale.

Il existe deux possibilités pour concevoir une antenne, c'est-à-dire déterminer les dimensions géométriques de l'antenne. On peut désirer que la direction du lobe principal corresponde avec

l'angle d'élévation désiré. Plus fréquemment, les dimensions sont déterminées pour privilégier un champ maximum rayonné à l'angle d'élévation désiré.

■ Détermination des dimensions géométriques selon le critère de champ maximum

□ Hauteur au-dessus du sol

Pour obtenir le maximum du champ, il suffit de dériver le champ par rapport à H :

$$\frac{\partial E_V}{\partial H} = 0 \Rightarrow \cos\left(\frac{2\pi H}{\lambda} \sin \alpha\right) = 0 \Rightarrow H = \frac{\lambda}{4 \sin \alpha} \quad [8.12]$$

□ Longueur de chaque branche

Pour obtenir la longueur de chaque branche, il suffit de dériver le champ par rapport à L :

$$\frac{\partial E_V}{\partial L} = 0 \Rightarrow L = \frac{\lambda}{2(1 - \sin \phi \cos \alpha)} = \frac{\lambda}{2 \sin^2 \alpha} \quad [8.13]$$

car on vérifie $\phi = 90^\circ - \alpha$

En fonction du calcul, on peut modifier ces paramètres pour rattraper l'alignement. De façon générale, on a les relations suivantes entre paramètres et performances :

- La longueur peut être augmentée si la hauteur calculée est trop importante.
- La hauteur peut être augmentée si l'on désire baisser l'angle d'élévation.
- L'angle Φ doit être augmenté si on désire réduire la longueur et/ou la hauteur mais, dans ce cas, le losange se ferme et les mutuelles ne peuvent plus être négligées.

À chaque fois que l'on doit s'éloigner du cas optimum, le gain sera réduit.

En ce qui concerne les valeurs de gain d'un losange, le calcul est complexe et on aboutit à des valeurs de l'ordre d'une quinzaine de dB pour un losange de $L = 6\lambda$.

Du point de vue du comportement en fréquence, cette antenne peut permettre un fonctionnement sur une octave de bande, mais il faut alors que l'angle ϕ augmente, ce qui nous limite donc à des angles d'élévation faible.

Elles sont très intéressantes pour la conception d'antennes ultra-large bande car leurs performances sont relativement peu dispersives en fréquence et évitent donc ainsi la déformation des signaux pulsés.

8.4 Antennes à ondes de surface et à ondes de fuite

Une antenne à ondes de surface est une antenne transportant une onde progressive dont l'énergie est concentrée au-dessus de la structure guidante. Par opposition, une antenne à onde de fuite est une antenne transportant aussi une onde progressive mais dont l'énergie contenue dans la structure est soit périodiquement évacuée, soit rayonnée continûment tout au long de l'axe de propagation. Cela permet d'obtenir une direction de rayonnement différente de celle d'un rayonnement en-bout (*end fire*).

En général, ces deux types d'antennes sont à faible épaisseur et trouvent ainsi de nombreuses applications, notamment dans le spatial et l'aéronautique. Elles sont plutôt utilisées en hautes fréquences (au dessus de 500 MHz) mais préférentiellement en hyperfréquences quand les structures impliquées sont de type à guides d'ondes. Typiquement leur largeur de bande est faible (de l'ordre de 10 %) et leur gain est modéré (de l'ordre de 15 dBi). Notons qu'il est difficile de contrôler le diagramme des antennes à ondes de surface alors que, comme pour les réseaux d'antennes, le type à ondes de fuite permet un contrôle plus aisé.

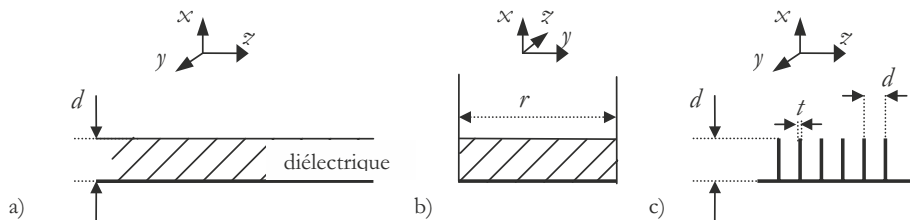


Figure 8.14 – Géométries de structures propageant des ondes de surface selon Oz (a) diélectrique sur métal à plan infini (b) diélectrique sur métal dans guide d'ondes (c) surface métallique corruguée

8.4.1 Équations de propagation d'une onde de surface

Une antenne à onde de surface se présente comme une structure dans laquelle un champ se propage à l'interface de deux milieux de conductivité différente (air-métal) ou de constante diélectrique différente (air-diélectrique).

Le paramètre le plus important pour la conception des antennes à onde de surface est la longueur d'onde dans la direction de propagation. De façon générale, il existe une relation générale permettant de relier les nombres d'ondes définis dans les trois directions entre eux :

$$k_x^2 + k_y^2 + k_z^2 = k^2 \quad [8.14]$$

avec $k = \omega\sqrt{\epsilon_0\mu_0}$: constante de propagation dans le milieu de propagation (air)

Pour une structure définie comme sur la figure 8.14(a), l'onde de surface se propage dans la direction $0z$ avec un champ évanescant dans la direction $0x$. Dans le cas où la propagation ne dépend pas de y (structure infinie dans cette direction), la relation générale se réduit à :

$$-\alpha_x^2 + \beta_z^2 = k^2 \Rightarrow \beta_z^2 = k^2 + \alpha_x^2 \quad [8.15]$$

On vérifie aussi :

$$\frac{\beta_z}{k} = \frac{\lambda}{\lambda_z} = \frac{c}{v_z} \quad [8.16]$$

Il vient l'équation donnant cette constante de propagation selon $0z$, valable pour une onde TM (ne dépend pas de la hauteur de diélectrique) :

$$\frac{\lambda}{\lambda_z} = \frac{c}{v_z} = \sqrt{1 + \left(\frac{\alpha_x \lambda}{2\pi}\right)^2} \quad [8.17]$$

Cette structure propage une onde lente avec un champ électrique de surface légèrement incliné (angle de tilt). Cet angle, lié aux pertes du métal ou la valeur de la permittivité, reste en général très faible (inférieur à 0.1° pour du métal).

On pourrait extraire l'atténuation selon $0x$ à partir de l'équation (4) :

$$\alpha_x = \frac{2\pi}{\lambda} \sqrt{\left(\frac{\lambda}{\lambda_z}\right)^2 - 1} \quad [8.18]$$

Si l'on trace cette atténuation du champ en fonction de l'inverse de λ_z , on s'aperçoit que plus la vitesse de l'onde est lente plus l'énergie est contenue sur une faible distance en s'éloignant de la surface (figure 8.15). Par exemple, 99 % de la puissance est comprise dans une zone de hauteur inférieure à $\lambda/2$.

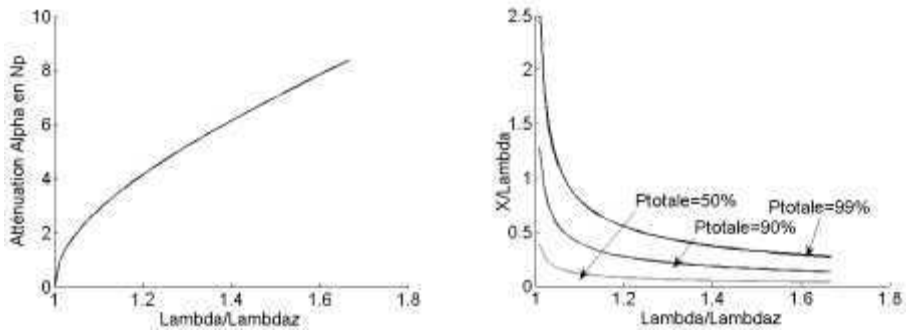


Figure 8.15 – Décroissance de l'onde selon l'axe Ox , perpendiculaire à la direction de propagation.

On aboutit à l'équation donnant les composantes de champs au dessus et sur la surface E_x, H_y, E_z . Par exemple pour E_z , nous aurons :

$$E_z(x, z) = E_z e^{-j\beta_z z} e^{-\alpha_x x} \quad [8.19]$$

Si la structure n'est plus infinie dans la direction Oy , par exemple par l'introduction de deux plans métalliques (figure 8.14(b)), la composante de champ selon y n'est plus constante et il existe une longueur d'onde de coupure liée à la distance entre les plaques (comme dans tout guide d'ondes). Cela introduit une condition supplémentaire pour les équations de continuité des composantes de champ électrique à la surface entre les deux milieux et on obtient pour la longueur d'onde :

$$\frac{\lambda}{\lambda_z} = \frac{c}{v_z} = \sqrt{1 + \left(\frac{\alpha_x \lambda}{2\pi}\right)^2 - \left(\frac{\lambda}{2r}\right)^2} \quad [8.20]$$

avec r : distance entre les plaques métalliques dans la direction Oy

Dans ce cas, comme le second terme est généralement plus faible que le troisième, l'onde se propage plus rapidement que la lumière.

Ces structures sont généralement utilisées pour exciter des discontinuités périodiques (guides d'onde à fentes). Pour l'étude des modes possibles, il faudrait distinguer la propagation des modes TM (qui ne dépendent pas de l'épaisseur du diélectrique) et les modes TE qui dépendent de d .

Un résumé des principales caractéristiques pour les deux types d'antennes est donné dans le tableau 8.1.

Tableau 8.1 – Résumé des principales caractéristiques et différences des antennes à ondes de surface et à ondes de fuite

Caractéristiques	Antennes à ondes de surface	Antennes à ondes de fuite
Méthode d'extraction des champs	Par calcul des champs le long de la structure	Théorie des alignements
Propagation de l'énergie	Au-dessus de la surface, onde lente en général $v < c$	Sous la surface, en général $v > c$
Type de rayonnement	En bout (« end-fire »)	Oblique

8.4.2 Structures d'antennes

■ Antennes à ondes de surface : diélectrique, hélice

Une antenne à ondes de surface va rayonner s'il y a présence de discontinuités. Pour une antenne diélectrique, cela se produira s'il y a une variation de l'épaisseur ou de la constante diélectrique. Pour le mode de propagation particulier dit axial, l'antenne hélice peut propager une onde de surface. Dans ce cas, on peut considérer que l'amplitude du courant est constante sur la structure (sauf au niveau de l'alimentation et de l'extrémité de l'antenne) et que seule la phase varie linéairement. J.-D. Kraus a montré que ce mode existait lorsque la circonférence d'un enroulement était proche ou égale à λ et que cela restait valable sur une octave de bande. Dans ce cas, le champ lointain est celui créé par n spires, chacune étant de longueur λ et déphasées de δ , déphasage d'alimentation de chaque spire :

$$\delta = -\beta \frac{\sqrt{\pi^2 D^2 + S^2}}{p} \quad [8.21]$$

avec $p = \frac{v}{c}$ facteur de vitesse

Pour l'ensemble de la structure (n tours), cela donne un diagramme en :

$$f(\theta) = \frac{\sin \frac{ns}{2}}{\sin \frac{s}{2}} \quad [8.22]$$

avec $s = \beta S \cos \theta + \delta$

Le gain étant donné par :

$$D_i = 15 \left(\frac{C}{\lambda} \right)^2 \frac{nS}{\lambda} \quad [8.23]$$

En mode axial, la polarisation est circulaire droite ou gauche, dépendant du sens d'enroulement.

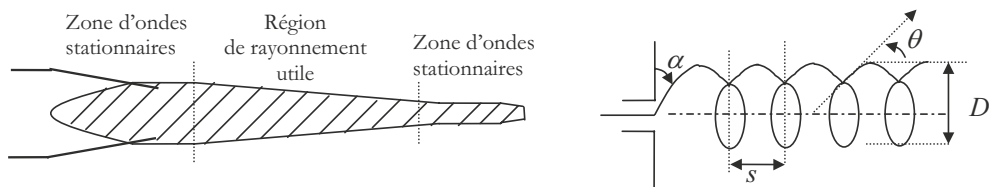


Figure 8.16 – Antenne cigare et hélice à mode axial

■ Antennes à ondes de fuite : guide d'onde fendu

Typiquement, la structure de propagation transporte une onde rapide ($v > c$). Un exemple est un guide d'onde dans lequel des fentes ont été faites sur la partie supérieure ou latérale de façon à couper les lignes de courants. Cela engendre à l'endroit de ces fentes un courant de déplacement, synonyme de rayonnement.

Le contrôle de la puissance rayonnée par chaque fente dépend de son périmètre. On peut définir une conductance par fente, qui permet de faire un modèle électrique du rayonnement. Les fentes sur un plan supérieur peuvent être d'un même côté et sont alors séparées de λ . Il est possible, pour raccourcir la structure, de choisir une distance de séparation de $\lambda/2$ à condition de placer les fentes alternativement de chaque côté pour respecter les conditions de phase d'alimentation.

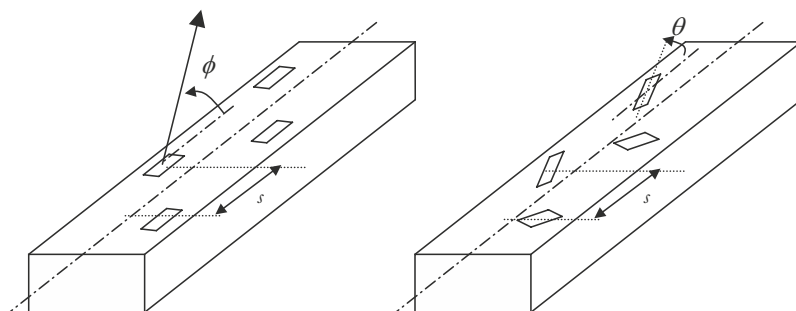


Figure 8.17 – Guides d'ondes fendus rayonnant des ondes de fuite

La relation entre les longueurs d'onde dans l'air, guidée et la distance entre les fentes donne l'angle de rayonnement :

$$\phi = \arccos \left(\frac{\lambda_0}{\lambda_g} + \frac{m}{s/\lambda_0} \right), \quad m : \text{ordre du mode} : 0, \pm \frac{1}{2}, \pm 1, \dots \quad [8.24]$$

On peut remarquer que cet angle dépend de la fréquence de travail due à la dépendance de la longueur d'onde guidée avec la fréquence. Cela signifie qu'il est possible de modifier cet angle en modifiant la fréquence de travail.

8.5 Antennes Yagi-Uda

8.5.1 Introduction

C'est l'antenne la plus connue pour la réception des ondes de radiodiffusion TV et FM. Elle est constituée d'un ensemble de dipôles de longueurs différentes dont seulement un est alimenté, les autres brins ayant pour effet d'améliorer les performances en directivité ou de diminuer le rayonnement arrière de façon à obtenir une antenne unidirectionnelle à rayonnement en bout.

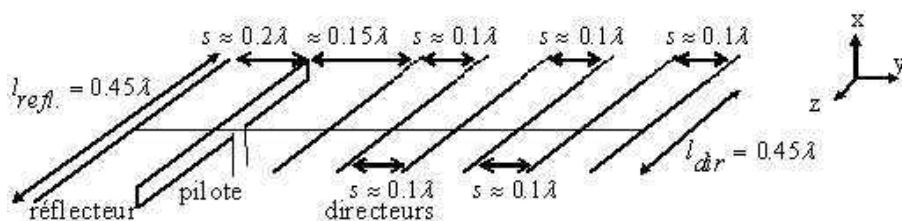


Figure 8.18 – Géométrie de l'antenne Yagi-Uda et ordre de grandeurs des dimensions

8.5.2 Principe de fonctionnement

Il existe deux façons d'étudier cette antenne. Nous pouvons la considérer comme une structure propageant une onde et donc la voir comme une antenne à onde de surface rayonnant en bout (voir section 8.4). Plus classiquement, on déduit ses caractéristiques de l'étude du rayonnement de deux sources d'amplitude égale mais de phase différente.

Un signal d'excitation électrique provenant du câble d'alimentation va entraîner le rayonnement d'un champ électrique par le pilote. Ce champ qui se propage de chaque côté du pilote va créer une fem induite dans le parasite. Cette fem va entraîner l'apparition d'un courant induit

dont l'amplitude et la phase vont dépendre de sa longueur et de la distance de séparation. Nous retrouvons bien le cas de deux sources d'amplitude et de phase données.

En théorie simplifiée, on peut assimiler le rayonnement de l'antenne Yagi à celui de deux dipôles demi-ondes espacés de $\lambda/4$ et alimentés en quadrature (figure 8.19).

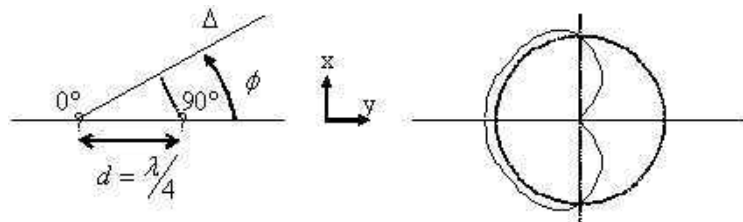


Figure 8.19 – Rayonnement de deux dipôles demi-onde alimentés en quadrature et diagramme de rayonnement associé

La méthode de calcul a été abordée dans le chapitre sur le rayonnement d'un groupement de sources. Dans notre cas, le diagramme de rayonnement dans le plan horizontal est de la forme :

$$E = \cos \left(\frac{\beta d \cos \phi + \delta}{2} \right) \quad [8.25]$$

Le diagramme obtenu est celui d'une cardioïde unidirectionnelle avec un maximum de champ rayonné dans la direction où la phase est retardée et de révolution autour de son axe.

Dans le cas réel de l'antenne Yagi, nous devons considérer deux différences essentielles :

- la distance de séparation entre antennes est plus faible que $\lambda/4$;
- la longueur du brin parasite est différente de la demi-onde.

Il nous faut donc connaître les mutuelles de couplages entre les antennes parallèles alimentées et parasite et il nous faut donc comprendre comment varie le rapport du champ avant sur champ arrière en fonction de la réactance de l'élément parasite.

On définit l'impédance mutuelle comme le rapport de la tension induite, à l'entrée de l'antenne parasite sur le courant circulant dans l'antenne alimentée :

$$Z_{21} = v_{21}/i_1$$

Les tableaux de valeurs des parties réelle et imaginaire de cette mutuelle ont été donnés par de nombreux auteurs (Brown, Pistolkors, Carter...) et reproduit ici dans le cas sans décalage sur l'axe Ox pour deux dipôles demi-onde.

Tableau 8.2 – Parties réelle et imaginaire de l'impédance mutuelle de deux dipôles demi-ondes espacés d'une distance d

$2d/\lambda$	$R_{12} (\Omega)$	$X_{12} (\Omega)$
0	73	42
$\lambda/8$	64	0
$\lambda/4$	41	-28
$3\lambda/8$	12	-37
0.5λ	-12	-30
$5\lambda/8$	-24	-12
1λ	4	17

Dans l'approche proposée par Eyraud *et al.*, le système est modélisé comme un transformateur, cela permet le calcul du courant induit i_2 :

$$\begin{cases} v_1 = Z_{11}i_1 + Z_{12}i_2 \\ 0 = Z_{12}i_1 + Z_{22}i_2 \end{cases} \Rightarrow i_2 = -\frac{Z_{12}}{Z_{22}}i_1 = -i_1 \frac{|Z_{12}|}{|Z_{22}|} e^{j(\gamma_{12}-\gamma_{22})} \quad [8.26]$$

Nous observons les points suivants pour le courant induit i_2 :

- il dépend du courant i_1 ;
- il dépend de la réactance du brin parasite Z_{22} , donc de sa longueur ;
- il dépend de la mutuelle de couplage donc de la géométrie des deux antennes.

Le champ total se calcule alors simplement :

$$E_T = KI_1 \frac{e^{-j\beta r}}{r} + KI_2 \frac{e^{-j\beta r}}{r} e^{-j\beta d \cos \phi \sin \theta} \quad [8.27]$$

qui donne en module :

$$|E_T| = \frac{KI_1}{r} \sqrt{1 + \frac{|Z_{12}|^2}{|Z_{22}|^2} - 2 \frac{|Z_{12}|}{|Z_{22}|} \cos(\gamma_{12} - \gamma_{22} - \beta d \cos \phi \sin \theta)} \quad [8.28]$$

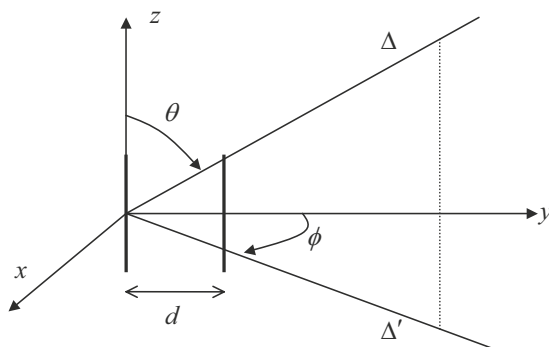


Figure 8.20 – Géométrie de l'antenne à un réflecteur

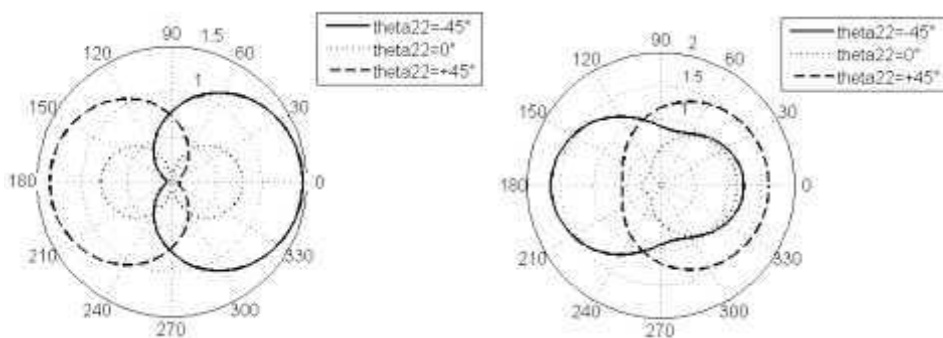


Figure 8.21 – DDR dans le plan horizontal d'un dipôle demi-onde avec son parasite non alimenté pour une distance de $\lambda/8$ et $\lambda/4$. Les cercles d'intensité 1 et 1.5 représentent les DDR du dipôle $\lambda/2$.

Avec l'antenne Yagi-Uda, on observe une propriété fondamentale qui est celle d'une augmentation du gain pour des réseaux à faible espacement entre brins rayonnant par rapport à un réseau de même nombre d'éléments.

8.5.3 Performances

On augmente la directivité en associant un réflecteur et plusieurs directeurs (figure 8.18). En supposant les longueurs des parasites peu différents de la demi-onde, on peut considérer que le potentiel est nul sur l'axe de symétrie Oy . En pratique, on exploite cette caractéristique en soudant les parasites sur cet axe sur une tige métallique qui apporte la rigidité d'ensemble. Cette directivité augmente avec le nombre de directeurs et la nature du réflecteur. Le diagramme de rayonnement affiche généralement un faible rayonnement arrière et latéral.

■ Variation du gain en fonction du nombre de directeurs

Les performances de l'antenne Yagi sont très sensibles aux dimensions. Elle est intrinsèquement bande étroite. On peut élargir la bande d'opération en augmentant la longueur du réflecteur (amélioration en BF) et en diminuant la longueur des directeurs (amélioration en HF). Malheureusement, cela se fait au détriment du gain. La façon la plus efficace et la plus répandue est de remplacer le brin réflecteur par un dièdre, comme nous le verrons par la suite.

Le calcul du gain est lourd dans le cas réel où l'antenne se compose d'un réflecteur, d'un pilote et de plusieurs directeurs. Les impédances mutuelles entre chaque élément permettent de déterminer les courants sur les éléments pilote et parasite.

En pratique, on optimise la longueur du réflecteur (cas d'un dipôle simple) et des directeurs ainsi que les espacements pour un nombre de dipôles et d'encombrement total donné pour un gain désiré (tableau 8.3). L'encombrement total L étant donné par la formule du gain empirique :

$$G = 10 \times L/\lambda \text{ (Valable pour des alignements de longueur inférieure à } 2\lambda \text{)}$$

Tableau 8.3 – Longueurs des parasites pour une yagi d'encombrement total de 0.8λ (*spacing* fixe de 0.2λ) et 1.2λ (*spacing* fixe de 0.25λ).
(Source : P. Viezbicke)

Longueur des parasites	Longueur Yagi	
	0.8λ	1.2λ
Long. réflecteur (l/λ)	0,482	0,482
Long. directeur 1 (l/λ)	0,428	0,428
Long. directeur 2 (l/λ)	0,424	0,420
Long. directeur 3 (l/λ)	0,428	0,420
Long. directeur 4 (l/λ)		0,428
Gain relatif au dipôle ($\lambda/2$)	9,2 dB	10,2 dB

D'un point de vue pratique, nous pouvons noter les caractéristiques suivantes :

- La longueur du réflecteur influence fortement le diagramme de rayonnement et notamment la présence des lobes arrière et latéraux.
- La nature du réflecteur (dipôle simple, triple, dièdre...) influence fortement le gain.
- La longueur du pilote a peu d'influence sur le gain et le diagramme mais plutôt sur l'impédance d'entrée.
- Le gain n'est pas proportionnel à la longueur, il est donc préférable d'utiliser deux Yagi en réseau plutôt qu'une grande Yagi

8.5.4 Amélioration des performances

■ Utilisation de réflecteurs différents

On peut maintenir le rapport champ avant sur champ arrière sur une plus grande bande de fréquence en allongeant ou en remplaçant le dipôle réflecteur par un réflecteur à deux ou trois dipôles ou encore plus efficacement avec un dièdre, intrinsèquement large bande. Dans ce cas, on augmente la tenue en fréquence (on passe de 10 % de bande à l'octave, ce qui devient suffisant pour la bande UHF utilisée en télévision) mais aussi le gain (augmenté de 5 dB avec un dièdre). En effet, si le réflecteur est suffisant en taille alors le rayonnement arrière est redirigé vers l'avant et le gain augmenté.

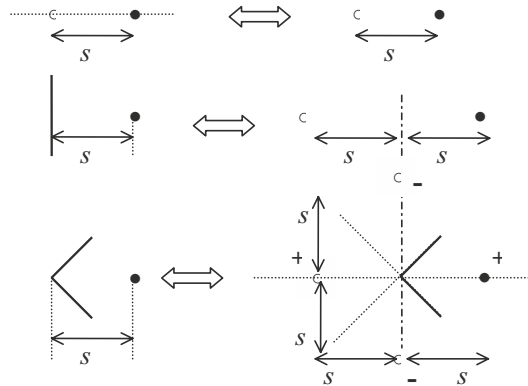


Figure 8.22 – Équivalence de calcul pour l'antenne à réflecteur, à réflecteur plan et réflecteur dièdre

Quand on remplace un dipôle réflecteur par un réflecteur plan, on remplace ce plan par son image. Quand le plan est plié pour former un réflecteur dièdre à 90° (« Active Kraus corner reflector »), on doit considérer dans le calcul 3 images parcourues par des courants identiques au pilote mais dont deux sont en opposition de phase. Dans ce cas, on obtient pour le champ total rayonné :

$$E(\phi) = 2KI_1 \left[\cos\left(\frac{2\pi S}{\lambda} \cos \phi\right) - \cos\left(\frac{2\pi S}{\lambda} \sin \phi\right) \right] \quad [8.29]$$

avec I_1 : amplitude du courant dans chaque élément

K : constante liée à la décroissance du champ en distance

Si on impose une tension V_1 à l'entrée de l'élément pilote, un courant va s'établir qui va dépendre de l'impédance propre du pilote et des mutuelles avec les images :

$$V_1 = Z_{11}I_1 + R_{1p}I_1 + Z_{14}I_1 - 2Z_{12}I_1$$

Si la puissance fournie est W alors le courant aura une amplitude de :

$$I_1 = \sqrt{\frac{W}{R_{11} + R_{1p} + R_{14} - 2R_{12}}}$$

alors que pour le pilote sans réflecteur, le courant serait :

$$I_1 = \sqrt{\frac{W}{R_{11} + R_{1p}}}$$

On obtient donc un gain avec réflecteur par rapport au cas sans réflecteur de :

$$G_{\text{réflecteur}} = \sqrt{\frac{R_{11} + R_{1p}}{R_{11} + R_{1p} + R_{14} - 2R_{12}}} \times \left| \cos\left(\frac{2\pi S}{\lambda} \cos \phi\right) - \cos\left(\frac{2\pi S}{\lambda} \sin \phi\right) \right| \quad [8.30]$$

L'étude serait similaire avec un réflecteur d'angle 60° mais en considérant 6 images. On obtient les caractéristiques suivantes :

- Le gain vaut environ 12 dB pour un angle au sommet de 60° , 10 dB pour 90° et 6 dB pour 180° (réflecteur plan).
- Le gain diminue quand l'espacement augmente.
- Plus l'angle au sommet est élevé, plus l'espacement doit être important.
- Plus l'angle au sommet est élevé, plus la décroissance du gain dû aux pertes ohmiques intervient pour de grands espacements.
- Le diagramme présente un lobe principal sur l'axe Oy pour $s = 0.5\lambda$ puis deux lobes inclinés symétriques à $\pm 30^\circ$ pour $s = \lambda$ puis un lobe principal sur l'axe et deux lobes secondaires (-8 dB) pour $s = 1.5\lambda$.

■ Utilisation du dipôle replié comme pilote

Pour contrecarrer les effets de la baisse de l'impédance d'entrée due à l'utilisation de plusieurs directeurs, il est d'usage de remplacer le dipôle pilote par un dipôle replié. Celui-ci est constitué par deux dipôles demi-onde très fortement couplés et connectés à leurs extrémités. Le fait que les courants soient de même amplitude et en phase permet de connecter les dipôles à leurs extrémités.

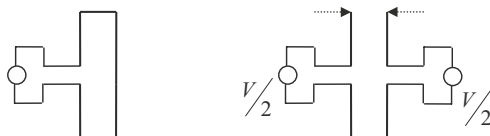


Figure 8.23 – Dipôle replié : représentation et modélisation

Si l'on suppose que chaque dipôle est alimenté par une tension $V/2$, on peut écrire pour le dipôle 1 :

$$\frac{V}{2} = Z_{11}I_1 + Z_{12}I_2$$

avec I_1, I_2 : courants sur les dipôles 1 et 2

Z_{11} : Impédance propre du dipôle 1

Z_{12} : Impédance mutuelle entre les dipôles

Comme les courants sont égaux et que l'impédance mutuelle vaut l'impédance propre (cas des espacements très faibles), l'impédance d'entrée vaudra donc :

$Z_e = \frac{V}{I_1} \cong 4Z_{11}$, ce qui donne environ 280Ω dans le cas d'un dipôle résonant. Cette impédance va diminuer lors de la mise en place des brins directeurs. Nous pourrions alors jouer sur les diamètres respectifs des brins du dipôle replié pour obtenir une adaptation autour de 75Ω .

8.6 Antennes à réflecteurs

8.6.1 Principes d'utilisation et géométrie

De nombreux types d'antennes à réflecteurs sont utilisées selon l'application visée. Par exemple, si l'application nécessite une antenne à gain modéré (12 à 15 dB), il est possible d'utiliser un réflecteur planaire ou dièdre. Si l'application est de type satellite, il est impératif d'utiliser des réflecteurs paraboliques ou cylindro-parabolique qui permettent d'obtenir des gains de l'ordre de 30 dB. Certains points de ce type d'antennes ont été abordés dans le paragraphe 6.2.4.

Nous nous contenterons d'aborder dans ce chapitre les réflecteurs paraboloides.

■ Géométrie de l'antenne parabolique

Pour un paraboloïde de révolution, d'axe Oz , les propriétés suivantes sont vérifiées :

- 1. $FR + RR' = 2f$
- 2. Les angles d'incidence et de réflexion par rapport à la normale n sont égaux.

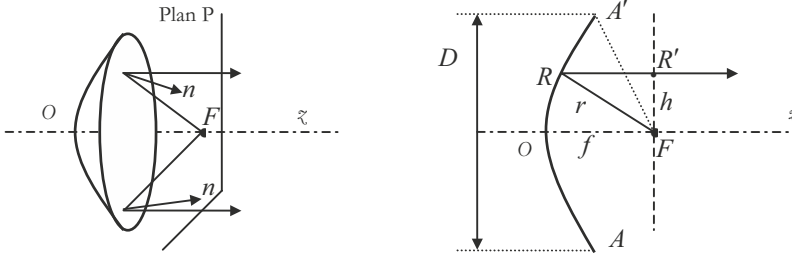


Figure 8.24

Au sens de l'optique géométrique, cela signifie que tous les rayons issus du foyer F vont après réflexion être parallèles à l'axe Oz . De façon équivalente, le plan P passant par F (et tous les plans perpendiculaires à l'axe Oz) sont des surfaces équiphasés dès lors que la source placée au foyer rayonne des ondes sphériques. Nous pouvons aussi dire qu'il existe une onde plane dans ce plan. Le plan AA' est appelé le plan d'ouverture et comme cela a été vu au paragraphe 3.3 les caractéristiques de rayonnement de la parabole seront dépendantes de la répartition des champs sur cette ouverture.

8.6.2 Calcul de la directivité et du gain

Supposons un paraboloïde (figure 8.25), éclairé par une source primaire d'intensité de rayonnement $U_p(\phi)$, le champ électrique sur l'axe Oz aura pour valeur :

$$E = \sqrt{120\pi} \left(\frac{\pi D}{\lambda} \right) \frac{1}{R} \cot \frac{\phi_0}{2} \int_0^{\phi_0} \sqrt{U_p(\phi)} \operatorname{tg} \frac{\phi}{2} d\phi \quad [8.31]$$

- avec ϕ_0 : angle entre l'axe Oz et l'extrémité de la parabole
- R : distance entre l'origine O du système et le point d'observation

Cela permet de calculer la puissance par unité d'angle solide sur l'axe Oz :

$$U = \left(\frac{\pi D}{\lambda} \right)^2 \cot^2 \frac{\phi_0}{2} \left[\int_0^{\phi_0} \sqrt{U_p(\phi)} \operatorname{tg} \frac{\phi}{2} d\phi \right]^2$$

Si l'on écrit que la source primaire possède un gain dans la direction ϕ de :

$$\kappa(\phi) = \frac{4\pi U_p(\phi)}{W_a}$$

avec W_a : puissance totale d'alimentation de la source primaire

L'intensité de rayonnement de l'ensemble devient :

$$U = \frac{W_a}{4\pi} \left(\frac{\pi D}{\lambda} \right)^2 \cot^2 \frac{\phi_0}{2} \left[\int_0^{\phi_0} \sqrt{\kappa(\phi)} \operatorname{tg} \frac{\phi}{2} d\phi \right]^2 \quad [8.32]$$

Comme le gain est donné par $G_0 = \frac{4\pi U}{W_a}$, nous obtenons :

$$G_0 = \left(\frac{\pi D}{\lambda} \right)^2 \cot^2 \frac{\phi_0}{2} \left[\int_0^{\phi_0} \sqrt{\kappa(\phi)} \operatorname{tg} \frac{\phi}{2} d\phi \right]^2 \quad [8.33]$$

Nous reconnaissons :

- un premier terme qui est la directivité maximum d'une ouverture rayonnante $\left(\frac{\pi D}{\lambda} \right)^2$;
- un second terme F_0 que l'on appelle le facteur de gain qui dépend des conditions d'éclairement de l'ouverture et de la géométrie de l'ensemble source primaire et secondaire.

$$F_0 = \cot^2 \frac{\phi_0}{2} \left[\int_0^{\phi_0} \sqrt{\kappa(\phi)} \operatorname{tg} \frac{\phi}{2} d\phi \right]^2 \quad [8.34]$$

Il faut donc choisir correctement la fonction de gain de la source primaire $\kappa(\phi)$ ainsi que les paramètres géométriques ϕ_0 si l'on désire maximiser le gain obtenu pour l'ensemble. Il existe d'ailleurs un maximum de gain pour un certain angle ϕ_0 à éclairement donné. Par exemple, si l'on impose le gain de la source primaire comme une fonction :

$$\kappa(\phi) = \kappa_0 \cos^n(\phi)$$

On obtient alors le gain donné à la figure 8.25.

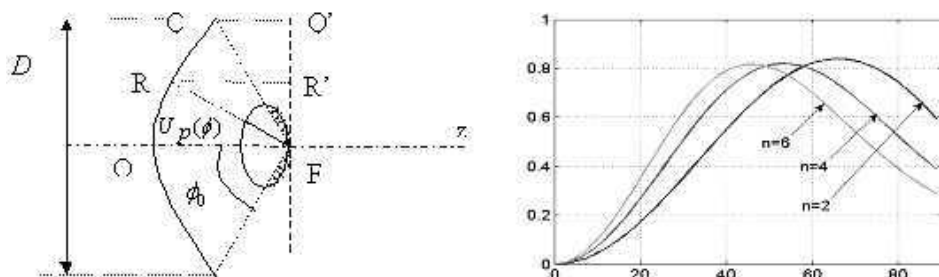


Figure 8.25 – Valeur du facteur de gain d'une parabole en fonction de l'angle ϕ_0 pour différentes formes d'intensité de rayonnement de la source primaire

Il existe un second facteur de réduction du gain qui est le facteur de *spill-over*. Le rayonnement de la source primaire qui n'est pas intercepté par le paraboloïde correspond à de la puissance perdue. Il est donc possible de définir un rendement d'illumination en considérant la puissance

interceptée W_i et la puissance rayonnée par la source primaire W_p :

$$\eta = \frac{W_i}{W_p} = \frac{\int_0^{\phi_0} U_p(\phi) \sin \phi d\phi}{\int_0^\pi U_p(\phi) \sin \phi d\phi} \quad [8.35]$$

Le facteur de gain total de l'antenne sera donc le produit du facteur de gain F_0 par le rendement d'illumination :

$$F_T = F_0 \times \eta \quad [8.36]$$

L'évaluation des intégrales pour une intensité de rayonnement donnée par $U_p(\phi) = U_{p0} \cos^n(\phi)$ aboutit à la figure 8.26. On observe que le rendement atteint d'autant plus rapidement 100 % que la puissance rayonnée par la source primaire sur les bords du paraboloïde est faible, ce qui est le cas pour n grand.

Il est important de réaliser que même si l'onde est en phase en chaque point de l'ouverture lorsque la parabole est éclairée par une source à rayonnement hémisphérique, ce n'est pas le cas de l'amplitude du champ électrique qui est décroissant sur les bords de l'ouverture ce qui justifie donc l'existence du facteur de gain introduit plus haut.

La raison est que, pour le trajet FR+RR', l'atténuation le long de FR est celle d'une onde sphérique en $1/R^2$ tandis que sur RR', celle-ci est quasiment nulle si le faisceau ne diverge pas. Comme pour le trajet FQ+QQ', les distances sont différentes, FQ et QQ' sont différentes de FR et RR', il en résulte une atténuation sur les bords par rapport aux points plus centraux.

Pour obtenir une distribution de champ plus uniforme, il est possible :

- d'augmenter la distance focale en conservant le diamètre, cela revient à augmenter le rapport focal ou rapport F/D ;
- de synthétiser un diagramme de source primaire qui compenserait l'atténuation sur l'ouverture.

Il faut aussi noter que cette perte de gain due à cette illumination atténuée permet d'obtenir un plus faible niveau de lobes secondaires, comme cela a déjà été montré.

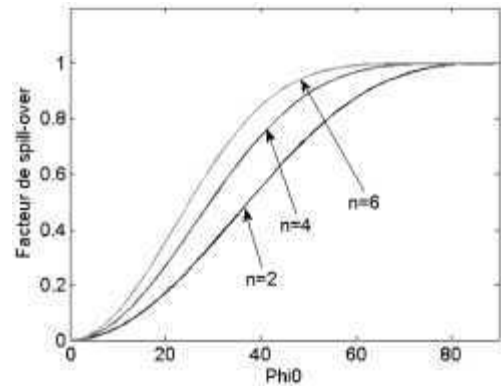


Figure 8.26 – Valeur du facteur de *spill-over* d'une parabole en fonction de l'angle ϕ_0 pour différentes formes d'intensité de rayonnement de la source primaire

8.6.3 Cas des réflecteurs défocalisés

La présence de la source primaire sur le chemin de propagation entraîne deux conséquences néfastes :

- une interaction de l'onde retournée avec l'onde envoyée, donc une modification du TOS de la source primaire ;
- l'apparition d'un effet d'ombre ou de diffraction. Cela a pour conséquence une augmentation des lobes secondaires.

Pour y remédier, on utilise une source décalée ou à offset.

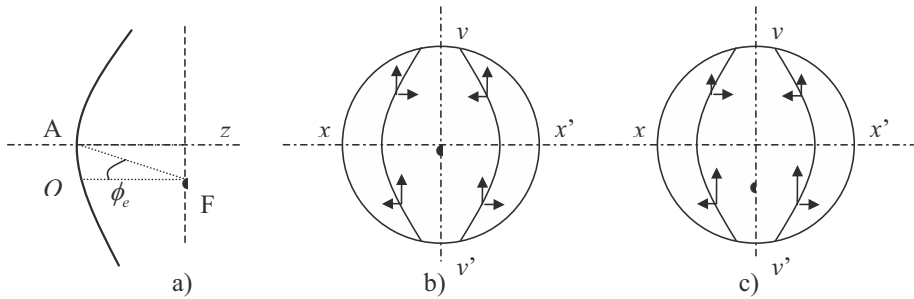


Figure 8.27 – Parabole à source primaire décalée a).
Composantes de polarisation pour une source centrée b) et excentrée c).

Le maximum de rayonnement de la source primaire pointe vers le point A et ne retourne pas vers la source. Le rayon qui retourne vers la source ne correspond plus au maximum de gain primaire. Cette excentration ne peut pas être aussi importante que l'on pourrait le souhaiter car elle entraîne un élargissement du diagramme et donc une diminution du gain. En ce qui concerne la polarisation, dans le cas où le réflecteur est plein, les composantes sont telles que décrites à la figure 8.27b. L'excentration verticale de la source primaire provoque une remontée des lobes de polarisation transversale car il n'y a plus équilibre des composantes verticales dans le plan xx' (figure 8.27c).

La dégradation engendrée sur le gain est fonction du rapport focal et de l'angle d'excentration ϕ_e . Par exemple, elle est insignifiante pour des rapports F/D proches de 1 quel que soit ϕ_e mais peut atteindre 3 dB pour un rapport F/D de 0,5 et une excentration exprimée en nombre de largeurs du lobe principal à 3 dB de 9.

8.6.4 Antenne Cassegrain

On appelle source primaire la source (cornet en général) qui émet les ondes vers le réflecteur secondaire. Dans le cas de l'antenne Cassegrain, on utilise un réflecteur secondaire de forme hyperbolique qui entoure le foyer de la parabole F1.

Le centre de phase de la source primaire est situé au foyer F2 de l'hyperbole. L'idée est que l'onde émise par la source primaire, après réflexion sur l'hyperbole, est sphérique, comme si elle provenait d'une source ponctuelle placée au foyer F1 de la parabole (figure 8.28).

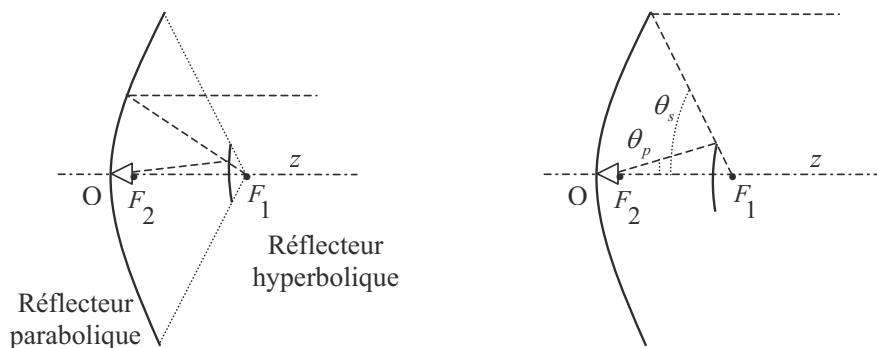


Figure 8.28 – Géométrie de l'antenne Cassegrain
F1 : foyer de la parabole ; F1 et F2 : foyers de l'hyperbole

Si la source primaire éclaire l'hyperbole avec un angle $\max \theta_p$, après réflexion, la parabole est éclairée par la source secondaire avec un angle $\max \theta_s$. Tout se passe donc comme si l'angle d'illumination était augmenté. Cela permet d'avoir des antennes plus compactes qu'un simple paraboloïde.

Le problème de l'effet d'ombre existe toujours, la dimension du réflecteur secondaire doit donc être faible devant la dimension du réflecteur primaire.

Il est aussi possible de s'approcher d'une distribution uniforme de champ dans l'ouverture en modifiant le contour hyperbolique du réflecteur secondaire et même celui du réflecteur primaire. On obtient ainsi une augmentation du facteur de gain mais aussi une augmentation du niveau des premiers lobes secondaires. Comme pour l'antenne défocalisée, la diminution du gain est d'autant plus rapide que la source primaire est excentrée.

■ Diminution de l'effet d'ombre par rotation de polarisation

Dans le cas où l'on travaille avec une polarisation rectiligne (verticale ou horizontale), pour que le réflecteur secondaire ne soit pas vu par l'onde réfléchie par le paraboloïde, il faut faire tourner la polarisation de 90° . On y parvient en utilisant un réseau de fils obliques inclinés à 45° placé à un quart d'onde en avant du réflecteur plein (figure 8.29).

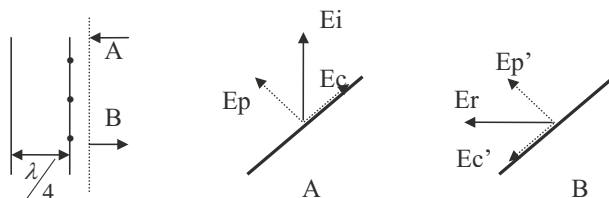


Figure 8.29 – Décomposition des vecteurs en avant du réseau de fils pour l'onde incidente (en A) et l'onde réfléchie (en B)

Le champ incident peut se décomposer en une composante transversale et une composante colinéaire au fil. La composante colinéaire va être déphasée de 180° en raison de la loi de Lenz lors de sa réflexion sur le fil. La composante transversale va traverser le réseau filaire et subir un premier déphasage de 90° pour atteindre le réflecteur plein. Ensuite elle va subir un déphasage de 180° puis à nouveau un déphasage de 90° sur le chemin retour. Elle va donc se retrouver en phase avec la composante transversale de départ. Si l'on effectue une combinaison vectorielle des champs, cela signifie que le vecteur réflecteur est en quadrature spatiale avec le vecteur incident. Le réflecteur hyperbolique qui est constitué d'une nappe de fil de direction perpendiculaire à la polarisation du champ réfléchi (ici verticale) devient transparent pour l'onde.

Il est à noter que le système ne fonctionne que sur une bande restreinte de fréquences puisque cette propriété est liée à la distance entre les fils et le réflecteur exprimée en longueur d'onde.

De façon générale, les antennes Cassegrain affichent plusieurs avantages sur les antennes paraboloïdes :

- un choix des paramètres plus aisé pour l'obtention des caractéristiques ;
- pas d'effet de diffraction de la source primaire ;
- suppression du *spill-over* vers l'arrière de la parabole, qui permet de réduire le lobe arrière et donc de diminuer la température de bruit lorsque l'antenne pointe vers le ciel (application en radioastronomie).

Bibliographie

- BALANIS C.A. — *Antenna theory, analysis and design*, John Wiley & sons, third edition, 2005.
- BROWN G.H. — *Directional antennas*, Proc. IRE, 25, 78-145, janvier 1937.
- COMBES P.F. — *Micro-ondes, tome 2, circuits passifs, propagation, antennes*, Paris, Dunod, 1997.
- EYRAUD, GRANGE, OHANESSIAN — *Théorie et techniques des antennes*, Ed. Vuibert, 1973
- GUILBERT C. — *Pratique des antennes - TV, FM, Réception, Emission*, Paris, Ed. Radio, 1989
- HARPER A.E. — *Rhombic Antenna Design*, Van Nostrand, New York, 1941
- JALIL R., CHEN-TO T. — *A new class of resonant antennas*, IEEE Transactions on Antennas and Propagation, vol. 39, n° 9, septembre 1991.
- JOHNSON R.C., JASIK H. — *Antenna Engineering Handbook*, Éditions McGraw Hill, 2^e édition, 1984.
- KRAUSS J.D. — *A small but effective Flat Top Beam antenn*, Radio, n° 213, 56-58, Mars 1937.
- KRAUSS J.D. — *Antennas*, MacGraw-Hill, 2^e édition, 1988.
- LAPORT E.A. — *Radio Antenna Engineering*, MacGraw-Hill, 1952.
- LO C. T. & LEE S.W. — *Antenna Handbook, Antenna Theory, vol. II.*, New York, Van Nostrand Reinhold, 1993.
- LO T.K. HO C.O, HWANG Y., LAM E.K.W., LEE B. — *Miniature aperture-coupled microstrip antenna of very high permittivity*, Electronics letters, vol. 33, n° 1, janvier 1997.
- LUXEY C., STARAJ R., KOSSIAVAS G., PAPIERNIK A. — *Antennes imprimées, bases et principes*, *Techniques de l'ingénieur*, E 3 310, éditions TI.
- LUXEY C., STARAJ R., KOSSIAVAS G., PAPIERNIK A. — *Techniques et domaines d'applications*, *Techniques de l'ingénieur*, E 3 311, éditions TI.
- NAKANO H., TAGAMI H., YOSHIZAWA A., YAMAUCHI J. — *Shortening ratios of modified dipole antennas*, IEEE Transactions on Antennas and Propagation, vol. 32, n° 4, avril 1984.
- ORFANIDI S. — *Electromagnetic waves and Antennas* — www.ece.rutgers.edu/~orfanidi/ewa
- SIEVENPIPER D., ZHANG L., JIMENEZ BROAS R.F., ALEXOPOULOS N.G., YABLONOVITCH E. — *High-impedance electromagnetic surface with a forbidden frequency band*, IEEE Transactions on Microwave Theory and Techniques, vol. 47, n°11, novembre 1999.
- THOUREL L. — *Les antennes*, Dunod, 1971.

9 • ANTENNES DÉFINIES SELON DES CARACTÉRISTIQUES SYSTÈMES

Les antennes décrites dans cette partie sont classées selon leurs caractéristiques système. On sépare les spécifications de l'antenne selon sa géométrie et la largeur de bande.

9.1 Différentes géométries d'antennes

Les caractéristiques géométriques sont multiples : la forme, le volume occupé, la taille...

9.1.1 Les antennes cornets

Dans ce paragraphe, nous commencerons par les antennes les plus utilisées qui sont les antennes cornets. Elles ont de nombreux avantages, en termes de qualité, de directivité et de puissance. Elles ont une largeur de bande supérieure à celle des dipôles ou des antennes fentes. Elles présentent une bonne adaptation au guide d'onde en s'évasant lentement à partir de celui-ci. Leur rapport d'onde stationnaire (VSWR) est de l'ordre de 1,05 à 1,2.

L'inconvénient majeur de ce type d'antennes est leur poids et leur coût, car elles sont entièrement métalliques. Les antennes satellites sont essentiellement de ce type et, dans ce cas, le poids est un inconvénient majeur qui est contrebalancé par la solidité et le gain de ces antennes. Ce sont aussi les antennes utilisées dans les systèmes radar en raison de la puissance élevée qu'elles supportent. Du fait de leur constitution métallique, elles présentent peu de pertes. Du fait de leur forme, elles sont très directives. Leur directivité peut varier d'une dizaine à une trentaine de décibels, selon leur forme.

Les antennes cornets sont alimentées par un guide métallique, soit rectangulaire, soit circulaire, soit elliptique.

Rappelons quelques caractéristiques des guides (rectangulaire ou circulaire). Soient a_0 et b_0 les longueurs et largeur respectives du guide rectangulaire (figure 9.1). Le champ électrique du mode

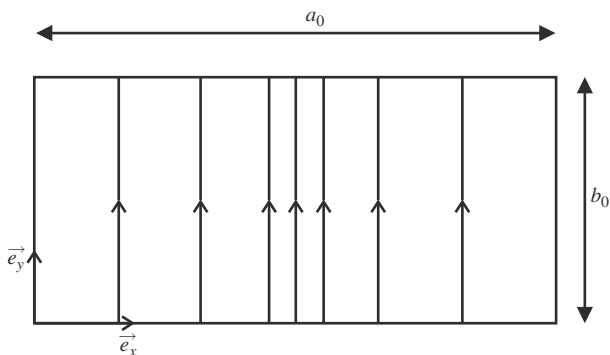


Figure 9.1 – Lignes de champ d'un guide d'onde rectangulaire

fondamental (TE_{10}) est parallèle au petit côté du guide. Il est uniforme dans la direction du petit côté. Il varie selon une loi sinusoïdale dans la direction de la longueur, donnée par :

$$\vec{E} = E_0 \sin\left(\frac{\pi x}{a_0}\right) \vec{e}_y$$

Le mode fondamental d'un guide circulaire est un mode TE_{11} . Ses lignes de champ électrique sont représentées sur la figure 9.2. La dégénérescence de mode du fait de la symétrie circulaire entraîne l'existence d'un mode analogue à celui représenté, dont l'orientation est tournée de 90° .

Le champ électrique du mode TE_{11} est donné en fonction du rayon a_0 du guide :

$$\vec{E} = E_0 \left[\frac{1}{\rho} J_1 \left(1,84 \frac{\rho}{a_0} \right) \sin \varphi \vec{u}_\rho + \frac{1,84}{a_0} J'_1 \left(1,84 \frac{\rho}{a_0} \right) \cos \varphi \vec{u}_\varphi \right]$$

avec $\vec{u}_\rho$ et $\vec{u}_\varphi$ les vecteurs unitaires en coordonnées polaires et a_0 le rayon du guide.

Différents types d'antennes cornets existent (figure 9.3) : les antennes sectorielles plan E (a), les antennes sectorielles plan H (b), les antennes pyramidales (c) et les antennes coniques (d). Dans le cas de l'antenne sectorielle plan E, le cornet conserve la longueur du guide ($a = a_0$) alors que la largeur du guide s'agrandit ($b > b_0$). Pour l'antenne sectorielle plan H, le cornet conserve la largeur du guide ($b = b_0$) alors que $a > a_0$. L'antenne pyramidale présente des dimensions plus grandes que celles du guide.

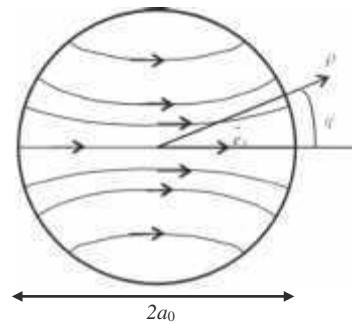


Figure 9.2 – Lignes de champs d'un guide d'ondes circulaire

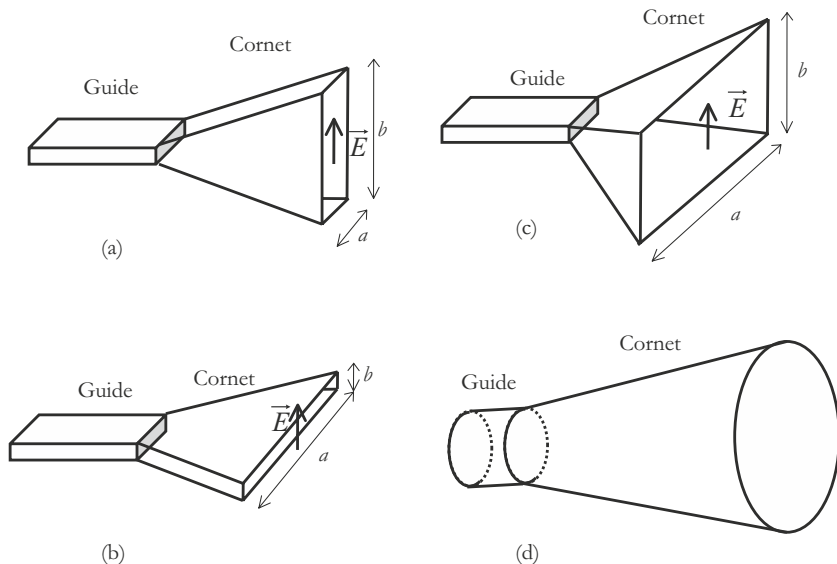


Figure 9.3 – Différents types de cornets

L'onde guidée débouche dans le cornet. Celui-ci s'évase et l'ouverture qui rayonne dans l'air a des dimensions plus grandes que celles du guide. Elle supporte donc, en plus du mode fondamental, des modes d'ordre supérieur. Si le cornet ne s'évase pas trop brutalement on peut considérer avec une bonne approximation que le rayonnement est essentiellement dû au mode fondamental. Les

cornets servent d'adaptation entre le guide et l'espace libre. Du fait de cette transition lente, les antennes cornets ont une largeur de bande plus grande que celle des dipôles ou des antennes à fentes. Dans un premier temps, des règles simples de conception vont être données pour l'utilisation d'antennes cornet. Pour plus de précision, dans un second temps, il faudrait affiner les résultats en utilisant des méthodes numériques.

■ Antenne sectorielle plan E

Le principe de calcul du rayonnement d'une antenne cornet repose sur le fait que le mode fondamental apparaissant à l'embouchure du cornet va créer le champ électromagnétique dans l'espace. Donc à partir de la connaissance du champ sur l'ouverture, on peut calculer le champ dans l'espace. Comme le cornet s'évase parallèlement au champ électrique, l'angle d'ouverture dans le plan E sera nettement plus faible que dans le plan H. La méthode présentée paragraphe 3.7 est utilisée à cette fin. Cependant cette méthode consiste à considérer le champ sur une surface plane. Or dans le cas qui nous intéresse, le champ présente sur ce plan un déphasage variable en fonction de la distance au centre de l'ouverture, calculable par rapport au champ du mode dont l'expression est connue sur la surface équi-phase. La figure 9.4 illustre ce déphasage. Les lignes équi-phase, qui sont les lignes de champ, sont perpendiculaires aux surfaces métalliques qui sont inclinées dans le cas du cornet. Plus l'angle d'ouverture du cornet sera grand, plus l'effet de ce déphasage sera grand. La ligne équi-phase de référence pour le calcul du champ sera donc assimilée à un arc de cercle dans le plan Oyz (plan E).

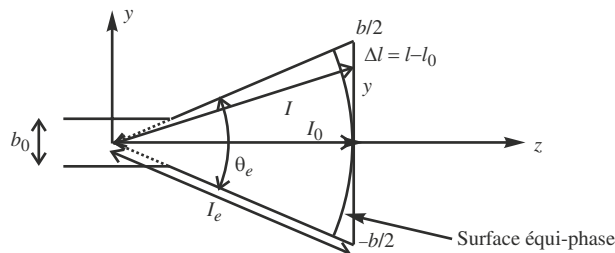


Figure 9.4 – Géométrie du cornet dans le plan E

La méthode de calcul du rayonnement considère le champ sur le plan de l'ouverture comme celui du mode déphasé de φ . Calculons ce déphasage. La longueur l du flanc du cornet est notée l_e pour le cas de l'antenne sectorielle plan E.

La différence de marche est égale à :

$$\Delta l \approx \frac{y^2}{2l_0} \approx \frac{y^2}{2l_e}$$

On en déduit le facteur de déphasage et donc le champ sur l'ouverture :

$$\vec{E} = E_0 \cos\left(\frac{\pi x}{a_0}\right) e^{jk \frac{y^2}{2l_e}} \vec{u}_y$$

Si le déphasage est faible, la directivité de l'antenne sera peu affectée par la non-uniformité de la phase sur l'ouverture. Ce sera le cas si l'angle d'ouverture du cornet est assez faible.

Le déphasage maximal s'obtient au bord du cornet. Il est alors égal à :

$$\Delta\varphi_{\max} = k \frac{b^2}{8l_e}$$

Soit :

$$\Delta\varphi_{\max} = 2\pi \frac{(b/\lambda)^2}{8 (l_e/\lambda)}$$

On admet que, tant que ce déphasage ne dépasse pas $\pi/2$, la directivité n'est pas dégradée. Cette valeur correspond à un optimum de la directivité. En effet, la directivité est d'autant plus grande que l'ouverture diffractante est grande, à condition que les points soient en phase. Lorsqu'on augmente la largeur normalisée de l'ouverture, la directivité croît jusqu'à cette valeur optimale. Tous les points sont alors presque en phase. Si la largeur normalisée continue à croître à partir de la valeur optimale, des sources secondaires viennent s'ajouter en opposition et la directivité diminue. La relation qui suit définit la largeur d'un cornet optimal pour une antenne sectorielle plan E. Elle correspond à un déphasage de $\pi/2$:

$$b_{opt} = \sqrt{2l_e\lambda}$$

La formule de la directivité pour les antennes sectorielles plan E a été calculée par Schelkunoff :

$$D_{0e} = \frac{64al_e}{\pi\lambda b} \left[C^2 \left(\frac{b}{\sqrt{2\lambda l_e}} \right) + S^2 \left(\frac{b}{\sqrt{2\lambda l_e}} \right) \right]$$

avec
$$C(x) = \int_0^x \cos \left(\frac{\pi x^2}{2} \right) dx \quad \text{et} \quad S(x) = \int_0^x \sin \left(\frac{\pi x^2}{2} \right) dx$$

La courbe donnant la directivité rapportée à la longueur normalisée du cornet est représentée, en fonction de la largeur normalisée, sur la figure 9.5. On y retrouve bien le maximum de la directivité. Plus la largeur est grande, plus la directivité optimale est grande. La directivité dépend du paramètre l_e/λ , c'est-à-dire de l'angle avec lequel le cornet s'évase. En effet, pour une valeur de b donnée, pour une valeur plus grande du paramètre l_e/λ , l'angle est plus faible. La transition est plus douce et la surface correspondant à l'ouverture du cornet est plus proche d'une surface équi-phase.

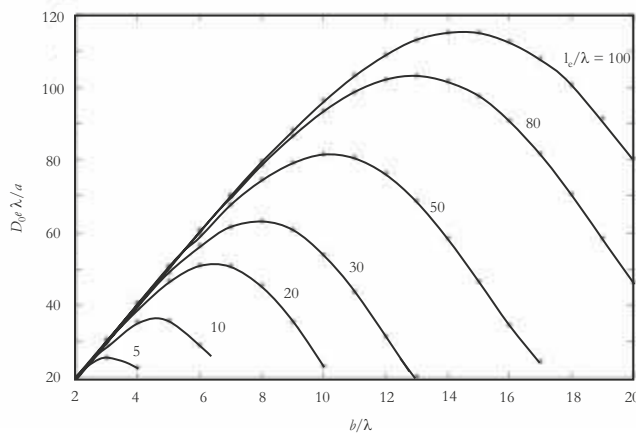


Figure 9.5 – Directivité normalisée de l'antenne sectorielle plan E en fonction de la dimension b normalisée, correspondant à l'élargissement du cornet

La suite de ce paragraphe a pour but de montrer comment calculer le diagramme de rayonnement de l'antenne sectorielle plan E en utilisant la méthode des ouvertures planes du chapitre 3, qui

consiste à calculer la transformée de Fourier du champ électrique sur le plan d'ouverture, soit :

$$\vec{f}_T = E_0 \vec{e}_y \int_{-a/2}^{a/2} \int_{-b/2}^{b/2} e^{j(k_x x + k_y y)} \cos\left(\frac{\pi x}{a}\right) e^{-jk \frac{y^2}{2l_e}} dx dy$$

Cette intégrale est à variables séparables. Calculons chacune des intégrales en x et en y . Selon x :

$$I_x = \int_{-a/2}^{a/2} \cos\left(\frac{\pi x}{a}\right) e^{jk_x x} dx$$

Soit :

$$I_x = \frac{1}{2} \int_{-a/2}^{a/2} \left[e^{j\frac{\pi x}{a}} + e^{-j\frac{\pi x}{a}} \right] e^{jk_x x} dx$$

Après intégration :

$$I_x = \frac{\sin\left(\frac{k_x a}{2} + \frac{\pi}{2}\right)}{k_x + \frac{\pi}{a}} + \frac{\sin\left(\frac{k_x a}{2} - \frac{\pi}{2}\right)}{k_x - \frac{\pi}{a}}$$

Après simplification :

$$I_x = 2 \frac{\pi}{a} \cos\left(\frac{k_x a}{2}\right) \frac{1}{\left(\frac{\pi}{a}\right)^2 - k_x^2}$$

En introduisant la valeur de k_x en fonction des angles ϑ et ϕ :

$$I_x = 2 \frac{\pi}{a} \cos\left(\frac{ka \sin \theta \cos \phi}{2}\right) \frac{1}{\left(\frac{\pi}{a}\right)^2 - (k \sin \theta \cos \phi)^2}$$

L'intégration selon y donne :

$$I_y = \int_{-b/2}^{b/2} e^{jk_y y} e^{jk \frac{y^2}{2l_e}} dy$$

Cette intégrale peut être évaluée numériquement. Après une normalisation faisant intervenir le coefficient s liant la dimension de l'ouverture, la dimension longitudinale du cornet et la longueur d'onde, on obtient :

$$I_y = \int_{-1/2}^{1/2} e^{2\pi j \left(\frac{b \sin \theta \sin \phi}{\lambda} u + 4su^2 \right)} du$$

avec

$$s = \frac{b^2}{8\lambda l_e} \quad \text{et} \quad u = \frac{y}{b}$$

Le paramètre s mesure la différence maximum de trajet entre la surface équi-phase et le plan de l'ouverture, mesurée en longueur d'onde.

La figure 9.6 donne la représentation de la norme de cette intégrale dans le plan E ($\phi = \pi/2$), en normalisant à 1 la valeur maximale pour chaque valeur du paramètre s . C'est en fait la valeur normalisée de la norme du champ électrique dans ce plan, puisque l'intégrale I_x est constante dans ce plan.

Lorsque s est faible, on retrouve pratiquement la répartition de champ en sinus cardinal d'une ouverture plane illuminée par un champ constant. On constate une forte remontée de lobe. Lorsque s augmente, le zéro de champ s'estompe.

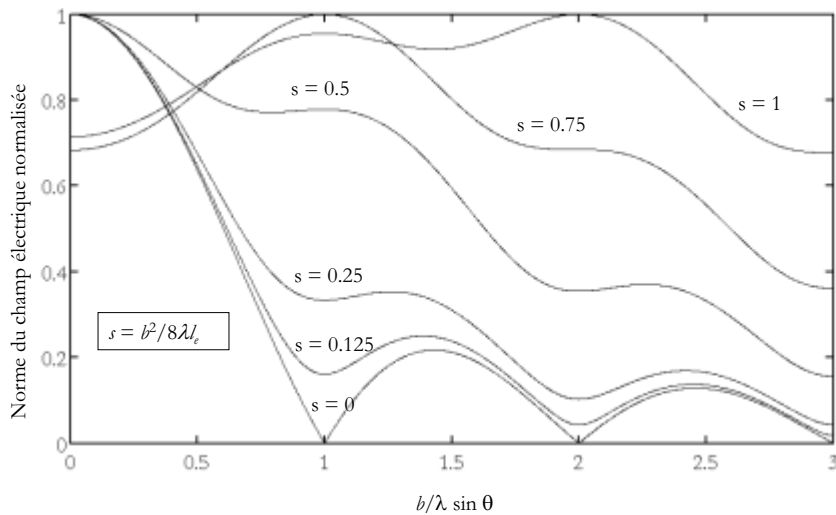


Figure 9.6 – Variation de la norme du champ électrique de l'antenne sectorielle plan E, dans le plan E, pour plusieurs valeurs du paramètre s .

■ Antenne sectorielle plan H

L'étude de l'antenne sectorielle plan H repose sur les mêmes considérations que celles utilisées pour l'antenne sectorielle plan E. Le cornet s'évase, dans ce cas, dans le plan (Ox, Oz) . L'angle correspondant à l'ouverture du cornet dans ce plan est noté θ_h . Comme le cornet s'évase parallèlement au champ magnétique, l'angle d'ouverture dans le plan H est plus faible que dans le plan E. La surface équi-phase de sortie est représentée sur la figure 9.7.

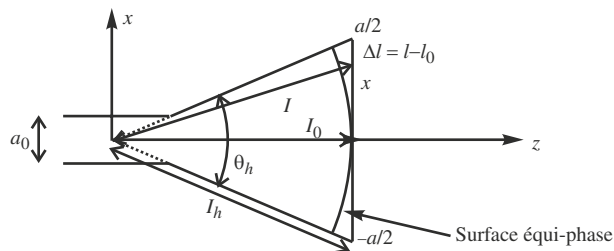


Figure 9.7 – Géométrie du cornet dans le plan H

La directivité D_{0h} dans le plan H a été calculée par Schelkunoff sous la forme de la formule suivante :

$$D_{0h} = \frac{4\pi b l_h}{\lambda a} \{ [C(u) - C(v)]^2 + [S(u) - S(v)]^2 \}$$

avec

$$u = \frac{1}{\sqrt{2}} \left(\frac{\sqrt{\lambda l_h}}{a} + \frac{a}{\sqrt{\lambda l_h}} \right) \quad \text{et} \quad v = \frac{1}{\sqrt{2}} \left(\frac{\sqrt{\lambda l_h}}{a} - \frac{a}{\sqrt{\lambda l_h}} \right)$$

La définition des fonctions C et S a été donnée plus haut.

Le cornet s'évase dans le plan H. La répartition du champ électrique dans l'ouverture est celle de la figure 9.1. Le champ est plus concentré dans la partie centrale de l'ouverture. Cette remarque

permet de comprendre pourquoi l'ouverture optimale est obtenue pour :

$$l_b - l_0 = 3\lambda/8$$

Les effets du déphasage sont moins importants sur les bords puisque le champ y est moins concentré. La dimension optimale est alors donnée par :

$$a \approx \sqrt{3l_b\lambda}$$

On retrouve cette valeur sur la figure 9.8 qui représente la variation de la directivité rapportée à la longueur normalisée en fonction de la longueur normalisée.

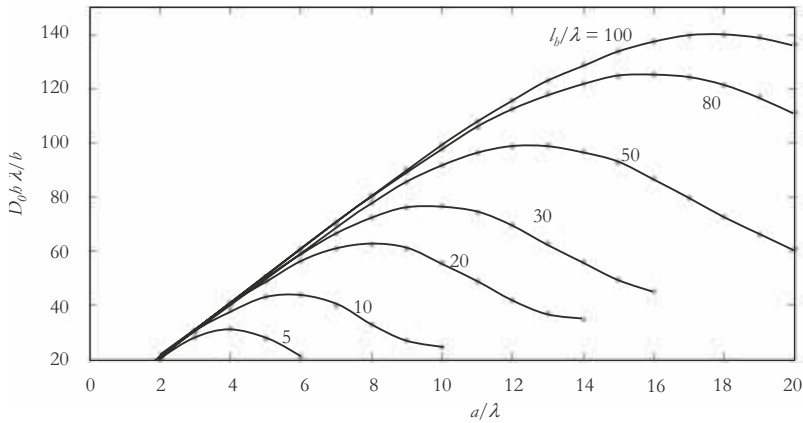


Figure 9.8 – Directivité normalisée de l'antenne sectorielle plan H en fonction de la dimension a normalisée, correspondant à l'élargissement du cornet

On remarque bien que la directivité de l'antenne passe par un maximum qui s'explique de la même façon que dans le cas de l'antenne sectorielle plan E.

On remarque aussi que la directivité de cette antenne est plus grande que celle de l'antenne sectorielle plan E, à paramètres normalisés égaux. Ceci est dû à la meilleure qualité du rayonnement qui est plus concentré vers l'axe de rayonnement, comme on va le voir dans la suite.

La méthode qui vient d'être utilisée est appliquée dans ce qui suit pour le calcul du champ, pour un cornet sectoriel dans le plan H.

La transformée de Fourier du champ sur le plan d'ouverture est donnée par :

$$\vec{f}'_T = E_0 \vec{e}_y \int_{-a/2}^{a/2} \int_{-b/2}^{b/2} e^{j(k_x x + k_y y)} \cos\left(\frac{\pi x}{a}\right) e^{-jk \frac{x^2}{2l_b}} dx dy$$

L'intégration en y est très simplement :

$$I'_y = \int_{-b/2}^{b/2} e^{jk_y y} dy$$

dont l'intégration donne :

$$I'_y = b \frac{\sin\left(\frac{k_y b}{2}\right)}{\frac{k_y b}{2}}$$

Soit encore, en remplaçant k_y par sa valeur en fonction des angles de rayonnement :

$$I'_y = b \frac{\sin\left(\frac{kb \sin \vartheta \sin \varphi}{2}\right)}{\frac{kb \sin \vartheta \sin \varphi}{2}}$$

L'intégration en x donne :

$$I'_x = \int_{-a/2}^{a/2} \cos\left(\frac{\pi x}{a}\right) e^{jk_x x} e^{jk \frac{x^2}{2l_b}} dx$$

Avant d'intégrer numériquement cette intégrale, procédons au changement de variable et à l'introduction du paramètre t :

$$t = \frac{a^2}{8\lambda l_b} \quad \text{et} \quad u = \frac{x}{a}$$

L'intégration numérique s'effectuera sur :

$$I'_x = \int_{-1/2}^{1/2} \cos(\pi u) e^{2\pi j \left(\frac{a \sin \vartheta \cos \varphi}{\lambda} u + 4tu^2 \right)} du$$

La forme de cette intégrale, pour $\phi = 0$, donne la forme du champ dans le plan H. Sa norme représente (figure 9.9) la variation de la norme du champ électrique dans ce plan puisque l'intégrale I_y est constante dans ce plan.

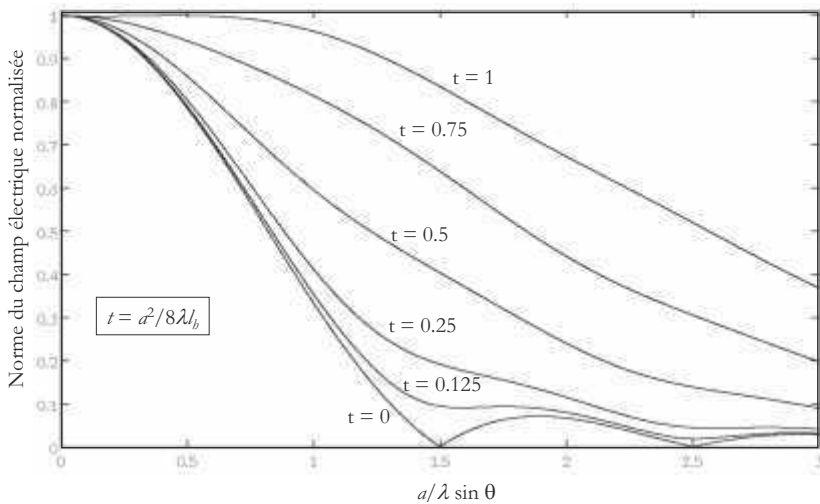


Figure 9.9 – Variation de la norme du champ électrique de l'antenne sectorielle plan H, dans le plan H, pour plusieurs valeurs du paramètre t .

On remarque que le rayonnement de l'antenne sectorielle plan H présente une très faible remontée de lobe. Cela est dû à la forme du champ dans l'ouverture. Le champ diminue dans ce plan, de façon à s'annuler sur les bords. Or, c'est le champ sur les bords qui contribue à la création du rayonnement dans les lobes. Par opposition, pour l'antenne plan E le champ s'annule brutalement et les lobes secondaires sont importants. Par contre le diagramme est plus fin dans le plan E.

■ Antenne pyramidale

L'antenne pyramidale est formée d'un cornet s'évasant dans les deux directions x et y . Elle a alors les propriétés de l'antenne sectorielle plan E, dans le plan E et les propriétés de l'antenne sectorielle plan H, dans le plan H.

Le gain maximal sera obtenu si les deux conditions dans les deux directions sont vérifiées :

$$a \approx \sqrt{3l_h\lambda} \quad \text{et} \quad b_{opt} = \sqrt{2l_e\lambda}$$

Le champ rayonné par l'antenne pyramidale se déduit de son spectre d'onde plane, correspondant à la transformée de Fourier dans l'ouverture qui peut être calculée avec une bonne approximation par :

$$\vec{f}_T = E_0 \vec{e}_y \int_{-a/2}^{a/2} \int_{-b/2}^{b/2} e^{j(k_x x + k_y y)} \cos\left(\frac{\pi x}{a}\right) e^{-jk\left(\frac{y^2}{2l_e} + \frac{x^2}{2l_h}\right)} dx dy$$

La directivité de l'antenne pyramidale se déduit de celles des antennes étudiées précédemment par la formule :

$$D_{0p} = \frac{\pi\lambda^2}{32ab} D_{0e} D_{0h}$$

La surface effective de l'antenne est très proche de 50 % de son aire géométrique.

■ Antenne cornet conique

Les calculs portant sur les ouvertures circulaires sont complexes. Le principe de fonctionnement est cependant le même que celui des ouvertures rectangulaires. On se bornera ici à donner quelques éléments concernant ces antennes.

La condition de maximum de directivité de cette antenne est similaire à celle de l'antenne sectorielle plan H. En appelant l_c la longueur du flanc de l'antenne et d le diamètre de son ouverture, cette condition s'écrit :

$$d \approx \sqrt{3l_c\lambda}$$

La directivité, pour cet optimum, est alors donnée par :

$$D_{0c} = 20 \log\left(\frac{\pi d}{\lambda}\right) - 2,82 \text{ dB}$$

Et sa surface effective est de l'ordre de 50 % de sa surface géométrique.

■ Cornets corrugués

La comparaison de la figure 9.6 à la figure 9.9 montre que le rayonnement dans le plan E est différent du rayonnement dans le plan H, tant du point de vue de l'ouverture de l'antenne que du point de vue de l'importance des lobes secondaires.

Si l'on considère l'antenne sectorielle plan E, le champ E est uniforme dans la direction du champ électrique. Le champ passe donc d'une valeur non nulle au bord du cornet à une valeur nulle à l'extérieur. Cette brusque variation induit une forte remontée de lobe secondaire. Par opposition pour l'antenne sectorielle plan H, le champ électrique varie sous la forme d'un cosinus pour s'annuler lentement sur les bords. Il n'y a alors pratiquement pas de remontée de lobes, par contre l'ouverture normalisée est légèrement plus grande.

Afin d'éviter les remontées de lobes, dans le cas d'une antenne pyramidale, qui nuisent à la directivité de l'antenne, une solution est de transformer le mur électrique perpendiculaire au champ électrique en une paroi dont l'impédance est infinie. Le champ magnétique longitudinal varie alors sous la forme d'un cosinus pour s'annuler sur cette paroi. Le comportement du champ magnétique dans le plan E est alors le même que celui du champ électrique dans le plan H. La symétrie complète est obtenue pour une antenne d'ouverture carrée. Ce type d'antenne peut

rayonner une polarisation circulaire. Du fait du comportement analogue des champs électriques et magnétiques, la polarisation croisée est très faible.

Pour réaliser cette condition, on usine dans la paroi des rainures ou corrugations. Perpendiculaires à la direction de propagation (figure 9.10)

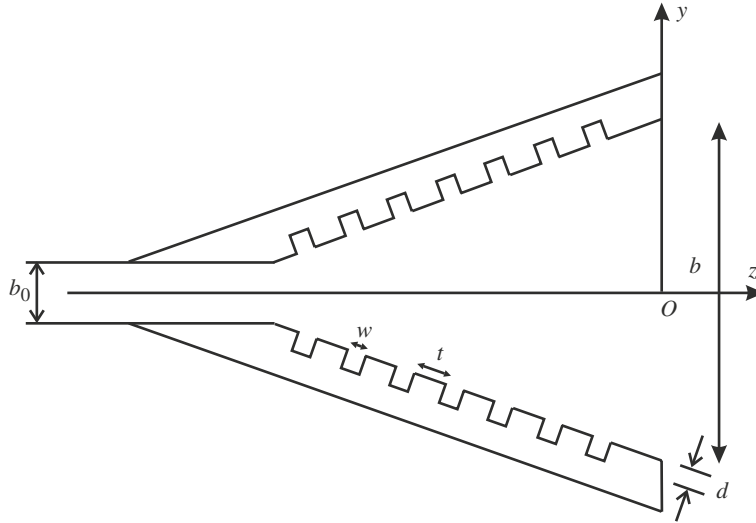


Figure 9.10 – Corruguations dans les parois plan H du cornet

Les crénelures coupent les lignes de champ magnétique (ainsi que les courants surfaciques) dans la direction z . Par contre le champ électrique dans la direction x voit toujours une impédance nulle. On a donc affaire à une surface anisotrope qui présente une impédance nulle dans la direction x et une impédance infinie dans la direction z . Cette propriété est obtenue grâce au comportement périodique de la surface dans la direction z . pour que ce comportement se manifeste, il faut que le nombre de crénelures par longueur d'onde soit important (de l'ordre d'une dizaine).

Afin de réaliser cette surface dans de bonnes conditions, il est nécessaire de considérer t petit par rapport à w . L'impédance ramenée à la surface des corruguations est calculable en considérant que chaque creux constitue un tronçon de ligne court-circuitée au fond de la crénelure. L'impédance moyenne ramenée à la surface est ainsi donnée par :

$$Z_1 = j \frac{w}{w + t} \sqrt{\frac{\mu_0}{\epsilon_0}} \tan(k_0 d) = jX_1$$

Pour que la surface soit vue comme un mur magnétique, il suffit de choisir $d = \lambda/4$.

En fait, pour couper les lignes de courants et diminuer la valeur du champ magnétique, il suffit que l'impédance Z_1 soit capacitive. Ce qui donne comme condition :

$$\frac{\lambda}{4} < d < \frac{\lambda}{2}$$

Cette condition impose une large plage de longueur d'onde pour le fonctionnement qui rend le dispositif large bande. Les largeurs de bandes sont couramment de 60 % et peuvent atteindre 100 %.

Pour comprendre le fonctionnement similaire de ce cornet pour le champ électrique et pour le champ magnétique, considérons la surface réflectrice corruguée, comme horizontale. Pour caractériser la réflexion sur cette surface il faut considérer les deux polarisations de l'onde.

Si le champ électrique est parallèle à cette surface (polarisation horizontale) le coefficient de réflexion est égal à :

$$\rho_E = -1$$

Si le champ magnétique est parallèle à cette surface (polarisation verticale), le coefficient de réflexion est donné par :

$$\rho_H = \frac{\sqrt{\frac{\mu_0}{\epsilon_0}} \cos \theta - jX_1}{\sqrt{\frac{\mu_0}{\epsilon_0}} \cos \theta + jX_1}$$

On constate bien que pour les incidences rasantes ($\theta = \pi/2$) le coefficient de réflexion en H est le même que celui en E. De même si X_1 tend vers l'infini, ce coefficient tend vers -1 . Ceci montre aussi un comportement analogue pour le champ électrique et pour le champ magnétique. Ces conditions entraînent que les champs électrique et magnétique s'annulent sur la surface corruguée donc sur le bord du cornet, réduisant l'effet de la diffraction en concentrant la puissance vers l'axe du cornet. Ces deux polarisations des champs ont des diagrammes de rayonnement très proches et donc la polarisation croisée est très faible.

Les formules suivantes donnent un exemple d'équations d'un mode ayant un champ électrique symétrique en y , qui vérifie les conditions de bord pour un cornet corrugué.

$$\begin{aligned} E_x &= 0 \\ E_y &= -\frac{j\gamma k}{k_y} \cos\left(\frac{\pi}{a}x\right) \cos(k_y y) \\ E_z &= k \cos\left(\frac{\pi}{a}x\right) \sin(k_y y) \\ H_x &= \frac{j\beta_1^2 k}{k_y \omega \mu_0} \cos\left(\frac{\pi}{a}x\right) \cos(k_y y) \\ H_y &= \frac{j\pi k}{a \omega \mu_0} \sin\left(\frac{\pi}{a}x\right) \sin(k_y y) \\ H_z &= -\frac{\gamma \pi k}{a k_y \omega \mu_0} \sin\left(\frac{\pi}{a}x\right) \cos(k_y y) \end{aligned}$$

k_y est la constante de propagation en y .

La constante de propagation dans la direction de l'axe est notée γ .

La relation sur les constantes de propagation donne :

$$\gamma^2 + k_y^2 + \left(\frac{\pi}{a}\right)^2 = \omega^2 \epsilon_0 \mu_0$$

Pour une épaisseur de corrugations nulle, le mode se réduit au mode TE_{10} .

On constate que la composante E_y présente une variation analogue à la composante H_x . Le rôle du champ électrique dans le plan E est le même que celui du champ magnétique dans le plan H.

Pour un cornet conique, les équations mises en jeu sont plus complexes, mais le principe de la surface anisotrope créée par une périodicité des crénelures est le même. Le mode de propagation est un mode hybride HE_{11} qui est un mélange entre un mode TE_{11} et TM_{11} . La surface conique doit présenter une impédance axiale tendant vers l'infini alors que l'impédance azimutale tend vers zéro. Dans ces conditions d'impédance, les deux modes TE_{11} et TM_{11} voient des conditions aux limites identiques. Ils se propagent alors ensemble, à la même vitesse.

9.1.2 Antennes planaires

Les communications hertziennes, les télécommunications spatiales et les radars utilisent le plus souvent des antennes à réflecteur. Ce sont des dispositifs performants qui possèdent un bon rendement, une grande pureté de polarisation et une large bande de fréquences. Dans le cas des applications mobiles, leur poids et leur encombrement deviennent deux inconvénients majeurs. Bien avant d'être appliquée aux antennes micro ruban, dans les années soixante, la technologie dite de circuit imprimée avait été largement mise à contribution notamment dans le domaine de l'électronique. Cette technologie appliquée aux antennes micro ruban (ou antennes patch) présente un certain nombre d'avantages parmi lesquels :

- un faible coût de fabrication,
- elles sont légères et peu encombrantes,
- la possibilité de les imprimer sur des surfaces non-planes dans le cas de substrats souples,
- la possibilité de mise en réseaux pour améliorer la directivité et pour des applications de balayage électronique de l'espace,
- la possibilité de les intégrer dans des appareils électriques nomades,
- la polarisation de l'onde électromagnétique linéaire ou circulaire en ajustant la géométrie et l'excitation de l'élément rayonnant.

Malheureusement, ces antennes présentent également un certain nombre d'inconvénients qui peuvent limiter leur domaine d'applications. On peut noter :

- une bande passante limitée (de 1 à 5 %),
- un faible gain (de l'ordre de 5 dB),
- une pureté de polarisation difficile à obtenir,
- l'excitation possible d'ondes de surface dans le diélectrique,
- des puissances transportées faibles en comparaison aux antennes traditionnelles.

Dans sa version la plus simple, l'antenne micro ruban, ou antenne imprimée, représentée figure 9.11, est constituée d'une plaque de substrat entièrement métallisée d'un côté tandis qu'un film métallique de forme variable (ici carrée) et de dimensions ajustées est déposé sur son autre face. Ce dernier constitue l'élément rayonnant dont les dimensions et les caractéristiques du substrat sur lequel il est déposé fixent, entre autres, la fréquence de résonance. Le plan de masse métallique est suffisamment grand par rapport à l'élément rayonnant de façon à limiter les effets d'ondes de surface qui rayonnent sur les extrémités de la plaque.

En guise de substrat, on trouve des composites à bases de fibres de verre-téflon (polytétrafluoroéthylène $2 < \epsilon_r < 3$, $\tan \delta = 10^{-3}$), du polypropylène ($\epsilon_r = 2,2$, $\tan \delta \approx 3 \cdot 10^{-4}$) mais également des mousses synthétiques contenant de nombreuses poches d'air ($\epsilon_r = 1,03$, $\tan \delta \approx 10^{-3}$).

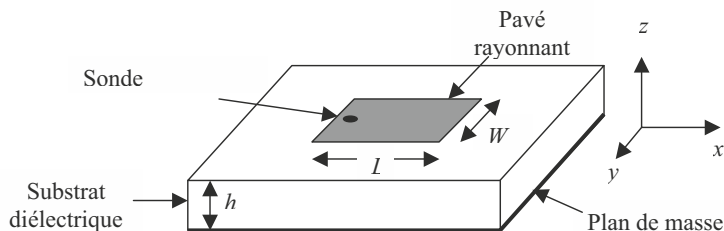


Figure 9.11 – Structure d'une antenne imprimée.

Les dimensions typiques d'une antenne patch sont sa longueur L , sa largeur W et son épaisseur h . D'un point de vue pratique, cette dernière est habituellement fine et bien inférieure à la longueur d'onde de travail ($h < 0,05\lambda_0$), λ_0 représentant la longueur d'onde dans le vide

Une antenne imprimée peut être considérée comme une cavité résonnante ouverte constituée de quatre murs latéraux magnétiques et de deux murs horizontaux électriques. Le rayonnement est provoqué par la fuite du champ aux extrémités entre le patch métallique proprement dit et le plan de masse. Le fonctionnement de l'antenne étant alors illustré à l'aide de deux fentes équivalentes aux bords rayonnants et séparées par la distance L .

Dans sa configuration originale, le comportement de l'antenne est contrôlé à l'aide d'une sonde de courant connectée entre le patch rayonnant et le plan de masse, ce qui va provoquer l'apparition d'un champ électrique à l'intérieur de la cavité. Une condition de résonance qui permet de transférer une puissance maximale à l'antenne consiste à choisir la longueur L légèrement inférieure à la demi-longueur d'onde guidée λ_g dans le diélectrique. Ce fonctionnement correspond à l'excitation du mode fondamental (ou fréquence de résonance fondamentale) qui correspond à la plus faible fréquence excitée.

Dans le cas d'une antenne de forme rectangulaire de dimensions L , W , les fréquences de résonances d'un mode TM_{mn} dans la cavité sont données par la formule suivante :

$$f_{mn} = \frac{c}{2\pi\sqrt{\epsilon_r}} \sqrt{\left(\frac{m}{L}\right)^2 + \left(\frac{n}{W}\right)^2}$$

Où c est la vitesse de la lumière dans le vide et ϵ_r , la permittivité du substrat.

Les conditions aux limites imposées par la nature de la cavité et le choix de la longueur L résonnante vont imposer une répartition des composantes des champs électrique et magnétique à l'intérieur de la cavité. Le champ électrique E_x est maximal et en opposition de phase de part et d'autre des bords rayonnants et sa valeur est nulle au centre de la cavité, ce qui a pour conséquence l'apparition d'un maximum de rayonnement selon la direction normale à l'antenne (figure 9.12). En effet, les champs émis par les deux fentes équivalentes aux bords rayonnants étant, dans ce cas, en phase. Le champ magnétique H_y est nul aux extrémités rayonnantes et, à l'inverse du champ électrique, est maximal au centre de la cavité.

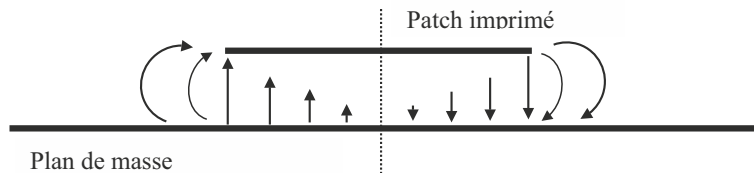


Figure 9.12 – Répartition du champ électrique entre le patch et le plan de masse.

Leurs expressions simplifiées, car ne tenant pas compte en première approximation des effets de bords, s'expriment de la façon suivante :

$$E_z = E_0 \cos \frac{\pi x}{L}$$

$$H_y = H_0 \sin \frac{\pi x}{L}$$

Une représentation simplifiée consiste à considérer l'antenne à l'aide d'un tronçon de ligne de transmission d'impédance caractéristique Z_0 , de longueur L et chargé aux extrémités pour afin de tenir compte du rayonnement électromagnétique dû aux fentes équivalentes. Les expressions de la tension et du courant qui se répartissent le long de cette ligne peuvent s'écrire en première

approximation :

$$V(x) = V_0 \cos \frac{\pi x}{L}$$

$$I(x) = \frac{I_0}{Z_0} \sin \frac{\pi x}{L}$$

Cette écriture offre une représentation simplifiée de la variation de l'impédance d'entrée de l'antenne ($Z_e = V/I$). En effet, si l'on suppose un point d'excitation situé à proximité d'un des bords rayonnants ($x = 0$ ou L) où la tension est maximale et le courant minimum, on obtient alors une impédance d'entrée maximale. À l'inverse, un point d'excitation situé au centre de l'antenne ($x = L/2$) où la tension est cette fois-ci minimale et le courant maximal donnera une impédance minimale proche de zéro. Cette impédance peut, par conséquent, être ajustée par un choix judicieux de la position de la sonde d'excitation ce qui permettra d'adapter l'antenne par rapport à l'impédance interne de la source qui l'alimentera. Les valeurs typiques de l'impédance d'entrée à proximité des bords rayonnants oscillent entre 150 et 300 Ω .

L'impédance d'entrée permet également de déterminer expérimentalement la fréquence de résonance de l'antenne. En effet, si l'on observe les évolutions des parties réelle et imaginaire de l'impédance (figure 9.13), on remarque qu'à la résonance, l'impédance présente un maximum de partie réelle associée à une valeur proche de zéro de la partie imaginaire.

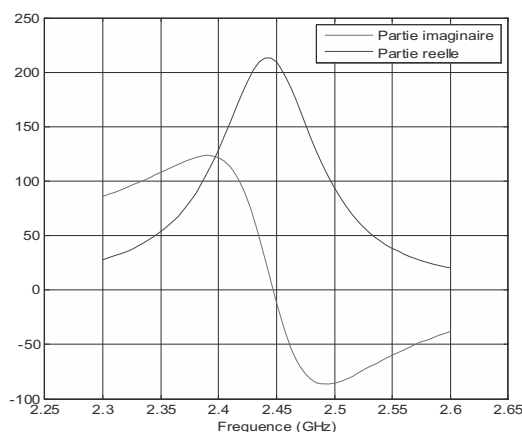


Figure 9.13 – Évolution de l'impédance d'entrée d'une antenne imprimée autour de la résonance.

Les pertes dans l'antenne limitent les performances. L'impédance d'entrée se décompose sous la forme d'une partie réelle Re et d'une partie imaginaire X_e .

La partie réelle se décompose elle-même en deux termes :

$$Re = R_r + R_p$$

où R_r et R_p correspondent respectivement à la résistance de rayonnement et à la résistance de pertes. Cette dernière incluant les pertes dans le conducteur métallique et les pertes dans le diélectrique.

La prise en compte des pertes dans l'antenne permet de donner une interprétation de l'efficacité η (ou rendement) de l'antenne. Il s'agit d'exprimer le rapport de la puissance effectivement rayonnée sur la puissance absorbée par l'antenne qui inclut à la fois la puissance rayonnée mais

également les termes de pertes définies plus haut. On montre que :

$$\eta = \frac{R_r}{R_r + R_p}$$

Un problème récurrent dans la conception des antennes imprimées concerne le choix de la technique d'excitation. L'alimentation par sonde coaxiale est possible mais on préfère souvent utiliser des lignes imprimées qui permettent d'alimenter plusieurs éléments à la fois notamment dans le cas de la mise en réseau des antennes. Nous distinguerons plusieurs types d'alimentations (figure 9.14) dont les principales sont l'excitation par sonde coaxiale, par ligne imprimée (a), par proximité (b) et par couplage à travers une fente dans le plan de masse (c).

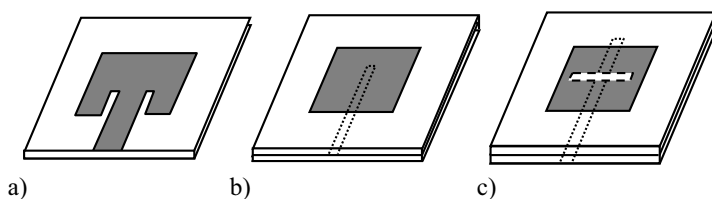


Figure 9.14 – Différentes techniques d'alimentation des antennes imprimées.

L'alimentation par ligne imprimée sur le même plan a pour avantage la simplicité de mise en œuvre. Un seul substrat est ici utilisé et le choix d'une encoche permet d'ajuster l'impédance d'entrée en pénétrant dans l'antenne pour l'adapter à l'impédance de source. L'alimentation par proximité se fait à partir de deux substrats superposés de natures différentes. Le substrat supérieur sera choisi de faible permittivité de façon à favoriser le rayonnement, tandis que le substrat inférieur sera de permittivité élevée de façon à concentrer le champ électromagnétique entre la ligne imprimée et le plan de masse. Enfin, une solution permettant d'isoler la ligne imprimée d'alimentation de l'élément rayonnant consiste à découper une fente dans le plan de masse de façon à coupler la ligne au pavé rayonnant. Cette solution, qui nécessite trois niveaux de métallisation, est attrayante car elle permet d'intégrer des composants actifs sur la ligne imprimée sans nuire au rayonnement de l'antenne compte tenu de la présence du plan de masse entre les deux. Malheureusement, un rayonnement arrière parasite peut apparaître notamment si l'on travaille à une fréquence proche de la résonance de la fente de couplage.

Lorsque les performances de l'antenne en termes d'adaptation, de rayonnement, de polarisation et de gain sont conservées à l'intérieur d'une bande de fréquences, on parle alors de bande passante de l'antenne. S'agissant d'une structure cavité, la bande passante est inversement proportionnelle à la fois au facteur de qualité de l'antenne et à la racine carrée de la constante diélectrique ϵ_r du substrat. La bande passante d'une antenne patch, qui tient compte de l'adaptation exprimée à partir du rapport d'onde stationnaire ($VSWR$), est donnée par de nombreux auteurs par la formule suivante :

$$\frac{\Delta f}{f_0} = \frac{VSWR - 1}{Q \sqrt{VSWR}}$$

où f_0 représente la fréquence de résonance de l'antenne et Q , le facteur de qualité global (incluant le facteur de qualité dû aux pertes par rayonnement, par conduction dans le métal, dans le diélectrique et par ondes de surface).

La bande passante des antennes patch ne dépasse que très rarement 5 %. La démarche consiste donc à réduire ce facteur de qualité pour optimiser la bande passante, ce qui se traduit d'un point de vue pratique par une augmentation de l'épaisseur du substrat. Malheureusement, ceci se traduit par une augmentation du risque d'apparition des ondes de surface.

Les caractéristiques de rayonnement des antennes imprimées planaires dépendent très fortement de la forme géométrique du motif et des caractéristiques du substrat utilisé. Elles sont cependant fondamentales pour connaître la répartition du champ électromagnétique dans l'espace. Nous avons représenté, sur la figure 9.15, un exemple de diagramme de rayonnement obtenu pour une antenne planaire conventionnelle de forme carrée résonnant à la fréquence de 2,45 GHz et déposée sur un substrat de permittivité 2,2 (Duroid RT 5880).

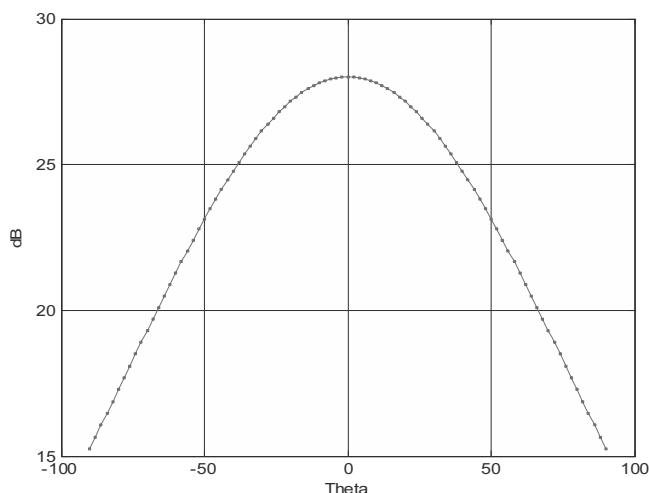


Figure 9.15 – Diagramme de rayonnement d'une antenne patch carrée

On remarque l'absence de rayonnement arrière compte tenu de la présence du plan de masse qui, pour la simulation, présente des dimensions infinies. Ce type d'antenne présente un gain de l'ordre de 4 à 5 dB et une directivité un peu supérieure compte tenu du rendement correct ($\sim 80\%$) obtenu sur ce type de structure. L'amélioration du gain peut être obtenue par une mise en réseau des éléments ou en superposant deux éléments dont les fréquences de résonances sont proches l'une de l'autre.

La connaissance de la polarisation de l'onde électromagnétique émise par une antenne est fondamentale car elle aura des conséquences importantes sur le transfert de puissance entre deux antennes. Naturellement, la polarisation de l'onde électromagnétique émise par une antenne imprimée alimentée par l'un des dispositifs présentés à la figure 9.14 est linéaire. La nature de cette polarisation est étroitement liée à l'orientation des courants à la surface de l'élément rayonnant. L'excitation du mode fondamental TM_{10} ou TM_{01} génère une polarisation linéaire dans le même sens que les courants surfaciques. Le contrôle de la répartition des courants surfaciques se fait en ajustant la position du point d'excitation sur l'élément. Par exemple, une excitation sur la diagonale d'un élément parfaitement carré générera une polarisation linéaire du champ selon cette diagonale. L'orientation oblique du vecteur résultera de l'excitation de deux modes dégénérés de même fréquence mais orientés perpendiculairement l'un par rapport à l'autre.

La polarisation circulaire est obtenue si, en plus de l'orientation perpendiculaire de deux modes dégénérés, il existe un déphasage temporel de 90° entre ces deux modes. Le vecteur résultant de la combinaison de ces deux modes décrit alors un cercle lorsque les ondes se propagent.

On définira la polarisation circulaire droite lorsque un observateur, situé à l'arrière de l'antenne, voit le vecteur résultant tourner dans le sens trigonométrique. Elle sera gauche dans le cas contraire.

Dans le cas particulier des antennes micro ruban, la polarisation circulaire peut être obtenue à partir d'un seul ou de deux points d'excitation. Le choix d'une technique plutôt qu'une autre étant dicté par la bande passante à l'intérieur de laquelle la polarisation circulaire est maintenue. Cette bande passante étant plus élevée dans le cas d'un système à deux points d'excitation.

L'alimentation double est bien adaptée à une structure rayonnante symétrique, dans la mesure où chacune des excitations est associée à un mode de résonance. Dans le cas d'un patch carré, l'excitation de deux modes orthogonaux est obtenue en ajustant la sonde d'alimentation sur chacune des médianes de l'antenne.

Le déphasage temporel de 90° est ajusté au niveau du circuit d'alimentation au moyen de deux méthodes (figure 9.16) :

- ajout d'un tronçon de ligne de longueur $\lambda/4$ (a), la bande passante sera dans ce cas limitée en fréquence,
- utilisation d'un coupleur hybride 3 dB (b) qui permet, outre la séparation de la puissance en deux ondes d'égales énergies, de sélectionner la nature de la polarisation (droite, gauche) en choisissant le point d'excitation, l'accès inutilisé étant chargé sur l'impédance caractéristique de la ligne d'alimentation.

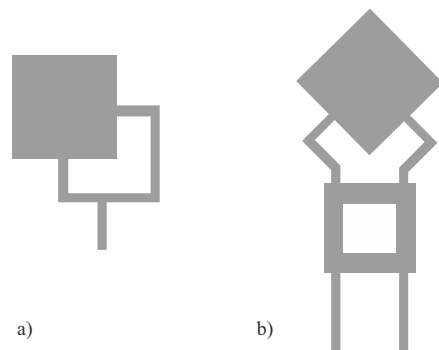


Figure 9.16 – Exemples de deux antennes à polarisation circulaires à double excitation.

L'alimentation à l'aide d'une seule excitation est plus délicate car elle nécessite un ajustement précis de la sonde sur le patch. En effet, cette dernière doit permettre l'excitation de deux modes dégénérés orthogonaux spatialement et en quadrature de phase. Dans la réalité, deux modes dont les fréquences de résonances sont proches l'une de l'autre sont obtenus en introduisant une dissymétrie sur l'élément rayonnant (figure 9.17). L'alimentation de l'antenne se fait alors à la fréquence où le déphasage entre les deux modes est de 90° c'est-à-dire à la fréquence centrale.

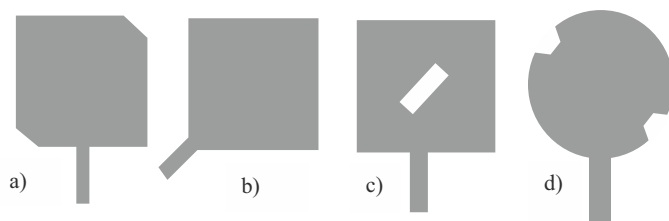


Figure 9.17 – Exemples de dissymétrie permettant d'obtenir une polarisation circulaire à partir d'une excitation unique.

On remarque sur la figure 9.17 que les dissymétries sont créées par une troncature de deux des quatre coins de l'antenne (a), par une légère augmentation de l'une des dimensions de l'élément rayonnant (b) ou par ajout d'un motif (fentes, encoches) sur l'élément rayonnant selon l'une des deux diagonales (c et d). La polarisation circulaire est donc obtenue selon un sens bien défini. Le sens opposé de cette polarisation est possible en ajustant les dissymétries selon la diagonale opposée.

Bien que cette partie présente de façon simplifiée les principales caractéristiques des antennes imprimées, de nombreux ouvrages et communications scientifiques traitent des améliorations

apportées aux caractéristiques des antennes imprimées précédemment citées. L'apport des outils numériques de simulation et la puissance sans cesse croissante des calculateurs a largement contribué au développement des antennes planaires. Les méthodes dites exactes (*full-wave analysis*) ont pris l'ascendant sur les méthodes analytiques approchées. On trouve actuellement des méthodes hybrides ou multi-échelles qui combinent différentes méthodes de façon à améliorer les performances en termes de rapidité de calculs et de précisions sur les résultats.

9.1.3 Antennes miniatures

■ Introduction

Aujourd'hui, l'antenne est au centre des préoccupations dans les applications sans fil telles que la téléphonie mobile, les ordinateurs, le Wi-Fi, le Wi-Max, Bluetooth et le GPS pour ne citer que ceux-là. Les contraintes sur l'intégration se font de plus en plus fortes. L'antenne devient un composant à part entière et doit trouver sa place dans des dispositifs dont les dimensions ne cessent de décroître, diminuant d'autant la place qui lui est réservée dans des proportions parfois critiques. Cette diminution ne devant évidemment en aucun cas réduire les performances radioélectriques de l'élément rayonnant. De plus, l'antenne n'est plus isolée dans son espace qui devient restreint, mais se retrouve à proximité de nombreux composants électroniques présents dans l'appareil. La proximité des éléments métalliques tels que le plan de masse et les lignes d'alimentation se traduit par l'existence de couplages. De plus, la multiplicité des standards de communications actuels (GSM à 880-960 GHz, DCS à 1 710-1 880 GHz et UMTS à 1 920-2 170 GHz, à titre d'exemple pour les télécommunications mobiles) impose de concevoir des antennes de types bi-bande, tri-bande et parfois même quadri-bande. On pourrait conclure en précisant qu'il existe parfois des contraintes d'ordres esthétiques ou ergonomiques qui ne font que compliquer la tâche des concepteurs d'antennes.

La course à l'intégration s'est traduite par l'apparition d'antennes miniatures. L'antenne imprimée sur substrat diélectrique en est un exemple. Dans la recherche d'antennes hypercompactes, les structures rayonnantes électriquement petites (AEP) ont été développées. La caractéristique principale de ces antennes vient de leurs dimensions bien inférieures à la longueur d'onde de travail. On classe les antennes dans cette catégorie lorsque les dimensions ne dépassent pas $\lambda/2\pi$. Certaines structures permettent cependant un fonctionnement multibandes et multipolarisations. Malheureusement, les antennes miniatures présentent des qualités en termes de gain, de bande passante et de rendement qui se dégradent d'autant plus que les dimensions géométriques diminuent. En effet, cette diminution se traduit par une augmentation du coefficient de qualité de la structure et donc par une concentration importante du champ électromagnétique au voisinage de l'antenne augmentant ainsi les pertes par effet Joule et dans le diélectrique.

La plupart des techniques qui permettent de diminuer les dimensions d'une antenne ont été appliquées aux antennes dipôles (plan de court-circuit à mi-hauteur du dipôle, « charge » du brin par un disque capacitif, techniques de repliement...). Elles ont été par la suite efficacement appliquées aux antennes plaquées. Sur celles-ci, en plus des courts-circuits et autres techniques de repliement, on trouve aussi des fentes gravées sur l'élément rayonnant qui ont pour effet de « rallonger » le chemin électrique du courant surfacique, introduisant ainsi de nouvelles résonances à des fréquences inférieures à la fréquence du mode fondamental d'une antenne imprimée demi-onde simple.

Nous allons, dans la partie qui suit, donner les principales structures d'antennes miniatures aujourd'hui largement employées dans les terminaux mobiles. Il s'agit d'une liste non exhaustive et le lecteur pourra se reporter à la littérature scientifique particulièrement riche dans ce domaine.

■ Principales structures d'antennes miniatures

Nous ne reprendrons pas dans cette partie les principales caractéristiques du dipôle $\lambda/2$ précédemment présentées et qui constitue l'élément de référence pour l'antenniste. Nous rappellerons simplement que, fonctionnant sur son mode fondamental, son gain est de 2,15 dB et son diagramme de rayonnement, maximum dans l'azimut présente une symétrie de révolution.

Le monopôle $\lambda/4$ constitue une variante au dipôle $\lambda/2$ précédent. Par application du principe des images, on diminue de moitié la hauteur du brin en plaçant un plan métallique parfaitement conducteur en son centre et positionné perpendiculairement à l'axe des brins (figure 9.18). Ce plan permet de représenter la partie manquante du brin (en pointillé) et, de ce fait, ne modifie pas la caractéristique de rayonnement de l'antenne. Cependant, le rayonnement s'effectuant dans le demi-plan, le gain de l'antenne est multiplié par deux soit une augmentation de 3 dB par rapport au gain de l'antenne dipôle, ceci en supposant le plan de masse de dimensions infinies. On retrouve les monopôles dans les antennes pour la réception FM sur les véhicules ou la CB (*Citizen Band*).

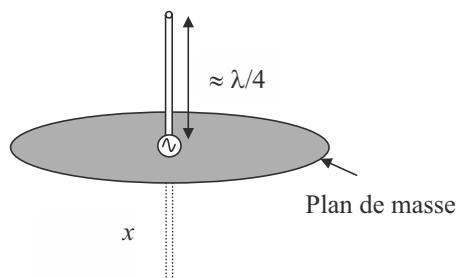


Figure 9.18 – Antenne monopolaire.

Enfin, pour diminuer la longueur du monopôle, on peut disposer une plaque conductrice appelée toit capacitif, généralement de forme carrée ou circulaire, à l'extrémité du brin et parallèlement au plan de masse métallique. Le toit capacitif assurant le complément de longueur électrique manquant. On parle alors d'antenne monopôle quart d'onde chargée. La distance entre le plan de masse et le toit capacitif approprié est de l'ordre de la moitié de celle du brin simple équivalent résonant à la même fréquence. Inversement, pour un monopôle simple, on voit que la mise en place d'un toit capacitif à son extrémité permet à l'ensemble de rayonner à la fréquence moitié de la fréquence pour laquelle rayonne le monopôle simple.

L'antenne fil-plaque est une structure un peu hybride entre le monopôle à toit capacitif et l'antenne imprimée dans la mesure où son rayonnement est de type monopolaire et sa structure, décrite figure 9.19, se présente sous la forme d'une antenne imprimée classique court-circuitée par un fil métallique entre le toit capacitif et la masse. La présence du court-circuit introduit un effet inductif qui, associé à la capacité due au toit, est équivalent à un circuit parallèle résonant. Ceci provoque l'apparition d'un nouveau mode de fonctionnement à une fréquence inférieure à la fréquence du mode fondamental de l'antenne imprimée. Le court-circuit sera placé de sorte qu'il ne modifie pas le comportement du mode fondamental, c'est-à-dire localisé dans une zone où le champ électrique est nul au centre du patch. À la fréquence du mode de résonance parallèle, le gain de la structure est identique à celui d'un monopôle à polarisation rectiligne sur un plan de masse infini soit 5,15 dB dans l'azimut. Typiquement, la résonance parallèle apparaît à une fréquence quatre fois plus petite que celle de l'antenne imprimée. Cela constitue un avantage si l'on souhaite concevoir une antenne de petites dimensions. En effet, un patch chargé fonctionnant à la résonance parallèle verra ses dimensions réduites d'un facteur 4 par rapport à une antenne imprimée simple à la même fréquence.

Une antenne directement déduite du monopôle est l'antenne en F inversée (ou IFA pour *Inverted F Antenna*) (figure 9.20). Il s'agit en fait d'un monopôle replié parallèlement au plan de masse. On réduit ainsi la hauteur de la structure tout en maintenant la longueur résonante. L'effet capacitif qui résulte de la mise en parallèle du brin replié et du plan de masse est compensé par un tronçon de ligne (stub) en court-circuit placé à l'opposé du brin replié. La sonde d'alimentation

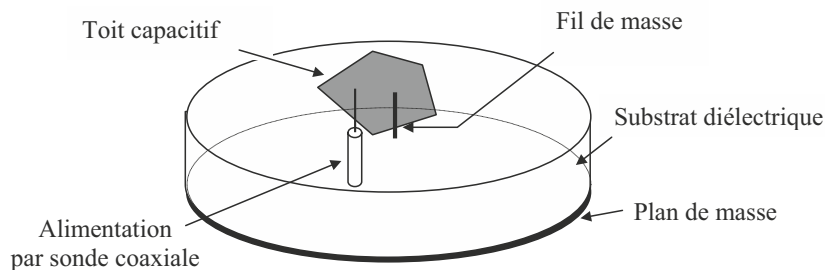


Figure 9.19 – Antenne fil plaque monopolaire.

sera choisie de façon à adapter l'ensemble par rapport à la source. La structure en F inversée présente de sérieux atouts car, en plus de la réduction d'encombrement et de sa relative simplicité d'adaptation, cette antenne rayonne les deux polarisations horizontale et verticale, ce qui est un avantage dans les communications indoor où la propagation par trajets multiples favorise l'apparition d'une composante croisée orthogonale à la polarisation principale émise par la source. L'antenne en F inversée peut également être imprimée sur un substrat diélectrique de façon à la rendre plane. On parle alors d'antenne IFA imprimée à ne pas confondre avec l'antenne PIFA que nous allons présenter dans ce qui suit.

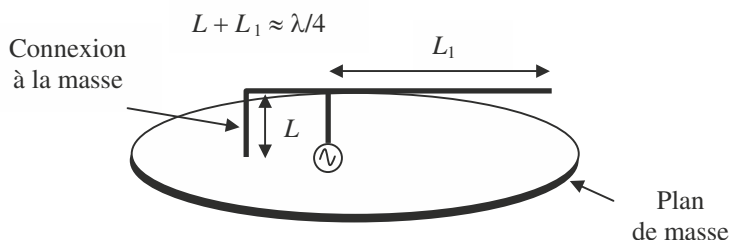


Figure 9.20 – Structure d'une antenne en F inversée.

Le principe de l'antenne F inversée a été développé pour les antennes planaires en F inversées (PIFA pour *Planar inverted F Antenna*). Cette fois-ci, le brin métallique est remplacé par un pavé rayonnant métallique l'assimilant ainsi à une antenne patch conventionnelle mais assortie de quelques particularités (figure 9.21). Le substrat diélectrique est généralement remplacé par de l'air de façon à améliorer les caractéristiques radioélectriques. On utilise un court-circuit qui permet, au même titre que le monopôle, de réduire la dimension résonnante de l'antenne d'un facteur 2. Celui-ci est en effet placé en un point où le champ électrique du mode fondamental est nul. Cependant, sa longueur n'est pas égale à la largeur du pavé rayonnant. Elle est ajustée pour provoquer un effet inductif supplémentaire au niveau du court-circuit. Cet effet a pour conséquence une réduction de la fréquence de fonctionnement et une amélioration de la bande passante de l'antenne qui, sans cet ajout, resterait inférieure à la bande passante d'une antenne patch demi-onde court-circuitée. Le gain n'est que rarement supérieur à 4 dB pour ce type d'antenne. Précisons que la polarisation émise fait apparaître un niveau de polarisation croisée important. On parle alors de polarisation « floue » dans la mesure où les caractéristiques de rayonnement mélangent celles d'un monopôle à celles d'une antenne imprimée conventionnelle. Le principe des antennes PIFA a largement été exploité dans la littérature. Des antennes compactes quadri-bandes ont été développées pour des applications de téléphonie mobile.

Les antennes PIFA sont alors associées à des fentes, des charges capacitives et des patch parasites court-circuités pour obtenir des résonances multiples tout en conservant des dimensions réduites qui permettent l'intégration dans un terminal mobile.

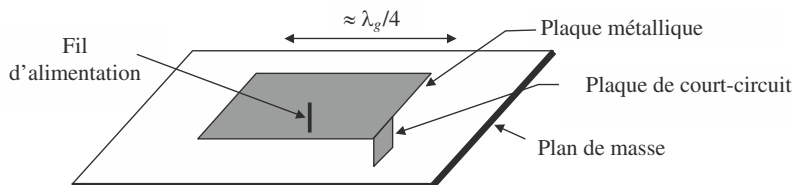


Figure 9.21 – Structure d'une antenne PIFA.

L'antenne à méandres est ici aussi directement déduite du monopôle quart d'onde. Pour diminuer les dimensions de la structure, l'idée consiste à replier le monopôle en plusieurs méandres d'égaux longueurs. La réduction de taille est obtenue en ajustant le nombre de méandres et l'écart entre chacun d'eux (figure 9.22). Notons que nous avons représenté sur la figure l'antenne à méandres originale imprimée sur substrat et sa variante à fentes rayonnantes. Contrairement aux apparences, ce sont généralement les brins les plus courts qui participent au rayonnement de la structure, les courants surfaciques étant en phase, alors qu'ils sont en opposition de phase sur les brins les plus longs. On définit bien souvent cette structure à partir de sa longueur axiale, liée à l'encombrement, et sa longueur équivalente lorsque les méandres sont dépliés. On constate alors que la fréquence de résonance de la structure à méandres est plus élevée que sa version dépliée. Ceci s'explique par les couplages qui existent entre les différents méandres et les effets dus aux coudes pour le repliement des brins. Plus précisément, on estimera les performances de l'antenne en calculant le rapport, ou facteur de réduction, l/L où l représente la longueur axiale de l'antenne à méandres et L la longueur du monopôle qui résonne à la même fréquence. Dans une communication citée en référence, on constate que le facteur de réduction augmente avec le nombre de méandres. De la même façon, le facteur de réduction sera d'autant plus important que le rapport W/a sera élevé. W représentant l'écartement entre méandres et a la section du brin. Un facteur de réduction allant de 0,59 à 0,64 est avancé.

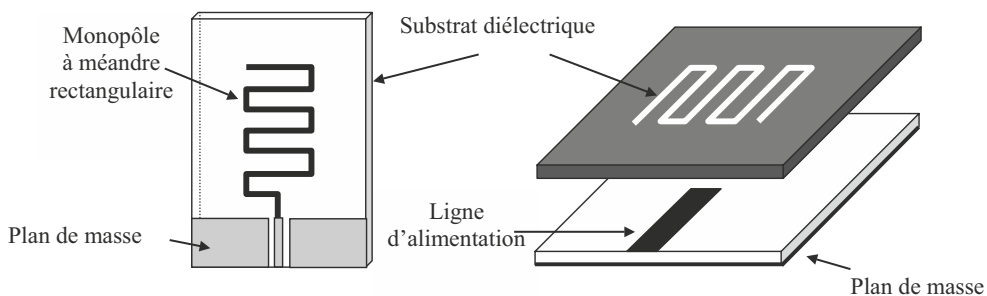


Figure 9.22 – Structures d'antennes monopôles à méandres imprimées (à gauche de type micro ruban, à gauche de type fente).

Le rayonnement d'une antenne à méandres est très proche du monopôle quart d'onde. Le gain, de l'ordre de 1,5 à 2 dB reste cependant inférieur aux 5,15 dB théoriques d'une antenne monopôle. L'antenne en C présente des dimensions bien plus petites (réduction d'un facteur 3) qu'une antenne imprimée simple demi-onde qui fonctionne à la même fréquence. La bande passante est cependant diminuée. Sa conception est établie à partir d'un dipôle imprimé replié sur lequel

on constate qu'à la seconde fréquence de résonance, il existe une symétrie axiale des courants surfaciques sur le dipôle. En ne retenant qu'une moitié du dipôle, on aboutit à l'antenne en C. C'est une structure planaire. Une variante à l'antenne en C est l'antenne double C à éléments superposés (figure 9.23) qui a permis de réduire d'avantage les dimensions de l'antenne en rendant le premier mode de résonance exploitable. L'idée consiste à replier le dipôle imprimé de façon à superposer deux éléments identiques reliés entre eux par un ruban métallique. L'antenne est ici alimentée à l'aide d'une sonde coaxiale sur le pavé inférieur à proximité du ruban de court-circuit. La structure n'est plus planaire mais présente des dimensions de l'ordre de $\lambda/11$, dimensions caractéristiques d'une antenne électriquement petite (AEP). De plus, la bande passante est améliorée en atteignant 1 à 2 %. Des évolutions ont été proposées sur cette antenne telles l'antenne en E, qui permet de sensiblement augmenter la bande passante (26 %) au détriment des dimensions de l'élément rayonnant (on passe de $\lambda/11$ à $\lambda/4$), ou l'antenne en S qui permet un fonctionnement bi-fréquences.

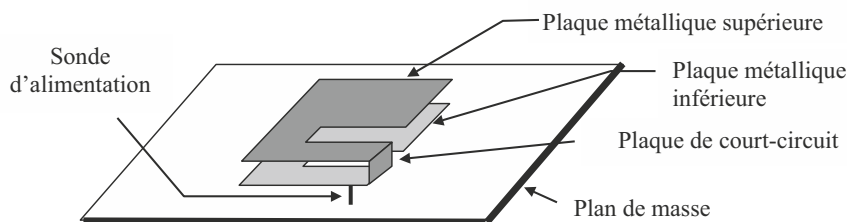


Figure 9.23 – Structure d'une antenne en C à éléments superposés.

Afin de réduire les dimensions des antennes imprimées résonantes demi-onde, un plan de court-circuit a souvent été utilisé. Comme nous l'avons vu au début de ce paragraphe, il s'agit du principe retenu pour le monopôle quart d'onde. Cependant, la conception d'un mur électrique ne va pas sans poser de problèmes. Sur une antenne imprimée rectangulaire simple, des trous métallisés doivent être répartis le long d'une médiane à champ électrique nul. Une alternative à l'utilisation de trous métallisés a été proposée il y a quelques années. Il s'agit de l'antenne en forme de « H » imprimée ou à anneau rectangulaire (figure 9.24). Le principe consiste à partir d'un patch rectangulaire à perturber la répartition surfacique des courants en évitant une partie de la métallisation sur l'élément rayonnant tout en conservant la symétrie de la structure. La symétrie permet de ne pas modifier les caractéristiques de rayonnement. La réduction des dimensions de la structure atteint 25 % pour l'antenne en « H » et 17,5 % pour l'anneau rectangulaire, par rapport à une antenne patch rectangulaire conventionnelle de même fréquence de résonance. Dans leurs configurations originales, l'antenne en « H » présente une largeur de bande plus faible que l'anneau rectangulaire. La largeur de bande de l'antenne patch étant située entre les deux. De nombreuses variantes d'antennes utilisant des encoches plus ou moins larges directement gravées sur l'élément rayonnant ont été élaborées par la suite. Le principe consistant toujours à rallonger le chemin électrique du courant surfacique de façon à diminuer la fréquence du mode fondamental de l'antenne.

Nous venons de voir qu'une réduction de la taille de l'antenne pouvait être obtenue en modifiant la répartition surfacique du courant sur l'élément rayonnant. Une autre solution consiste à utiliser des éléments réactifs (capacité ou inductance) qui permettent également de diminuer les dimensions de l'antenne. Dans la même idée, une alternative consiste à utiliser des matériaux à bande interdite électromagnétique (BIE ou EBG pour *Electromagnetic Bandgap* en anglais). Il s'agit de structures à motifs périodiques constitués de lignes de transmission imprimées à sections variables qui sont assimilées à des capacités et inductances distribuées. Des variantes 3D avec fils de court-circuit entre deux plans métalliques et régulièrement espacés entre eux ont aussi été développées.

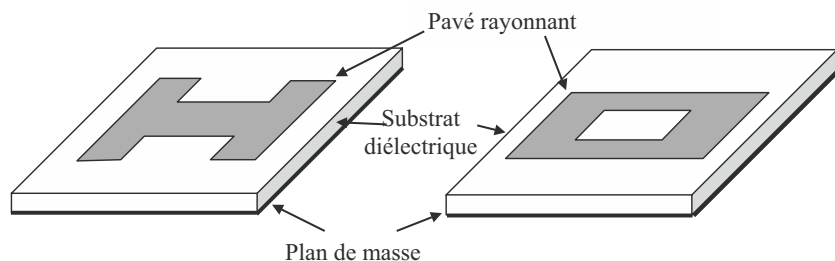


Figure 9.24 – Antennes en « H » et à anneau rectangulaire.

D'une certaine façon, ces structures périodiques peuvent être alors considérées comme des filtres à réjection de fréquence (on parle communément ici de bande interdite). En effet, à l'intérieur de cette bande, la structure présente une surface haute impédance souvent mise à profit pour réduire les ondes de surface qui dégradent le rayonnement principal dans les antennes imprimées. Initialement utilisés en tant que plan de masse, ces dispositifs ont été également développés pour des éléments rayonnants imprimés. La figure 9.25 représente un exemple d'antenne qui utilise un matériau à bande interdite électromagnétique. Les lignes à faibles largeurs représentent les parties inductives et la succession de pavés carrés les parties capacitatives. Le diagramme de dispersion, pour des fréquences situées en dessous de la bande interdite, montrerait la nature à ondes lentes de la structure périodique, ce qui se traduit par une diminution des dimensions physiques de l'antenne. À dimensions égales, il a été montré que la fréquence de résonance du mode fondamental de l'antenne à bande interdite électromagnétique est diminuée de 16 % par rapport à celle du patch plein.

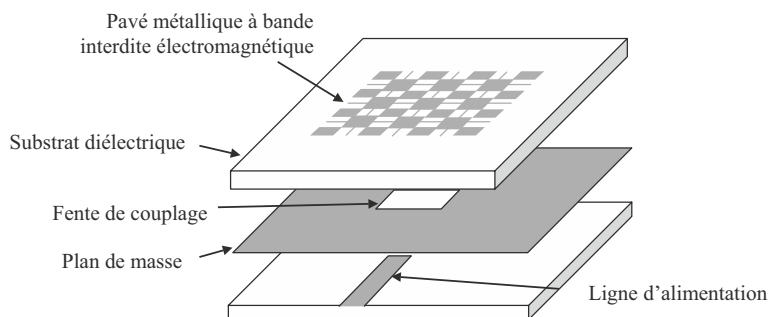


Figure 9.25 – Antenne patch imprimée à bande interdite électromagnétique.

Les dimensions sont étroitement liées à la nature du substrat qui supporte l'antenne et par conséquent à sa permittivité. Pour réduire l'encombrement, il paraît donc judicieux d'augmenter celle-ci. L'utilisation de matériau à permittivité élevée est alors retenue. Le titanate de baryum est un exemple de matériau céramique qui peut être utilisé pour un substrat d'antenne. Sa constante diélectrique s'étend de 38 à 80 et laisse entrevoir des réductions d'échelle importantes. Plus récemment, l'utilisation de l'aluminate de lanthane sous forme cristalline, caractérisé par une permittivité de 23,7 et une faible tangente de pertes (3.10^{-4}) a constitué une alternative originale au substrat diélectrique traditionnel. Malheureusement, le gain d'une antenne imprimée décroît lorsque la permittivité du substrat augmente. Pour corriger cet inconvénient, on superpose à l'antenne un (ou plusieurs) substrat de plus faible permittivité (on parle alors de superstrats) directement au-dessus de l'élément rayonnant (figure 9.26).

Les résonateurs diélectriques utilisent également des matériaux à permittivité élevée. L'absence de métallisation limite les pertes par conduction et par conséquent, augmente l'efficacité de rayonnement. Il est à noter que la plupart des techniques d'alimentation utilisées pour les antennes imprimées peuvent être retenues pour les résonateurs diélectriques. Enfin, ces structures présentent une bande passante relativement élevée (de l'ordre de 10 %). Deux inconvénients majeurs existent cependant :

- l'impossibilité d'intégrer ces antennes sur des surfaces non-planes,
- la difficulté d'usinage du matériau.

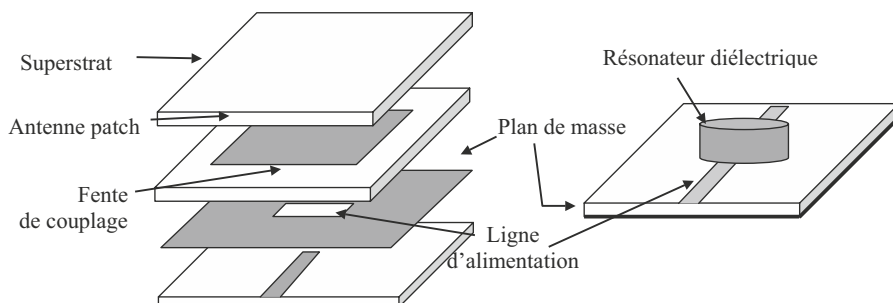


Figure 9.26 – Structures d'antennes utilisant des matériaux à permittivités élevées

9.2 Antennes définies par leur largeur de bande en fréquence

L'utilisation des antennes impose des caractéristiques de largeur de bande en fréquence imposées par le système. On distinguera donc les antennes bande étroite, large bande et ultra-large bande. Certaines antennes multibandes peuvent émettre ou recevoir dans plusieurs bandes de fréquences.

9.2.1 Antennes bande étroite

La largeur de bande a été définie au paragraphe 4.9. Les antennes bande étroite sont caractérisées par leur largeur de bande relative qui est définie par le rapport de la largeur de bande absolue à la fréquence centrale de la bande de fréquences :

$$B_r = \frac{f_2 - f_1}{f_0}$$

Les fréquences extrêmes définissant la bande de fréquences représentent les fréquences au-delà desquelles l'antenne n'a pas de bonnes caractéristiques de rayonnement. Elles correspondent, en général, à des paramètres en réflexion égaux à -10 dB.

Les antennes bande étroite ont des bandes relatives en fréquence qui sont de l'ordre de quelques pour cent à une dizaine de pour cent.

De nombreuses antennes décrites précédemment, reposant sur des phénomènes de résonance, ont une faible bande de fréquences.

9.2.2 Antennes large bande

Pour certaines applications, il est nécessaire de transmettre le signal sur une largeur de bande suffisante. Les antennes ne doivent pas limiter la transmission ou la réception. La notion de bande passante a été définie au paragraphe 4.9.

Une antenne est considérée comme large bande si la fréquence supérieure (f_2) est au moins égale à environ deux fois la fréquence inférieure (f_1). La largeur de bande est alors notée :

$$n : 1, \text{ qui n'est autre que le rapport } f_2 : f_1.$$

L'antenne étant un dispositif de transformation de l'énergie guidée en énergie rayonnée, dont le principe repose sur le phénomène de diffraction, il est bien évident que la largeur de bande d'un tel dispositif est limitée. Nous allons analyser les principes de base qui permettent d'obtenir une grande largeur de bande, et donner quelques exemples d'antennes de ce type, sans prétendre à l'exhaustivité, car les types d'antennes large bande sont nombreux. Certaines antennes utilisent plusieurs principes d'élargissement de bande.

■ Principes de conception d'une antenne large bande

Plusieurs principes concourent à l'élargissement de la bande de fonctionnement.

□ Adaptation

La largeur de bande dépend de l'adaptation de l'antenne. Une antenne réfléchissant très peu le signal sur une grande largeur et rayonnant correctement peut être considérée comme large bande. C'est donc le premier critère à prendre en compte.

□ Transformateur graduel d'énergie

D'un point de vue physique, nous verrons qu'une catégorie d'antennes présentant une large bande est celle pour lesquelles les paramètres géométriques varient lentement, permettant une transformation graduelle du signal en ondes rayonnées, pour différentes longueurs d'onde. Il n'y a alors aucune région de l'antenne présentant des discontinuités susceptibles de créer une zone de diffraction localisée, ou une zone de résonance, dépendant de la fréquence. C'est le cas de l'antenne en V (figure 9.27) dont la partie métallique présente un rapport s/r constant.

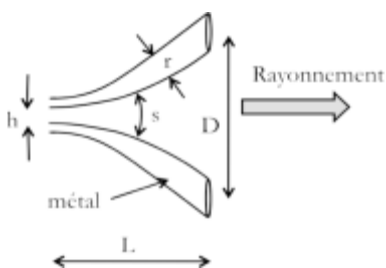


Figure 9.27 – Antenne en V

Le rayonnement correspondant aux petites longueurs d'ondes est émis dans la partie de l'antenne où les parties métalliques sont proches. On admet alors que h est de l'ordre du dixième de longueur d'onde. Les grandes longueurs d'ondes sont émises dans la partie large caractérisée par D , de l'ordre de la demi-longueur d'onde.

$$h \approx \frac{\lambda_{\min}}{10} \quad \text{et} \quad D \approx \frac{\lambda_{\max}}{2}$$

D'où :

$$\frac{\lambda_{\max}}{\lambda_{\min}} = \frac{f_2}{f_1} = \frac{D}{5h}$$

Pour obtenir une antenne large bande, il suffit donc que :

$$D \geq 10h$$

Cette condition est tout à fait réalisable pour certains types de géométrie.

Ce principe est aussi utilisé dans les antennes Vivaldi, formées d'un dispositif de guidage de l'onde, comme une ligne à fente, qui s'évase et rayonne dans l'air.

□ Principe d'auto-complémentarité

Un principe a été proposé par Rumsey, qui peut s'appliquer à certaines antennes pour maximiser la largeur de bande, en considérant la complémentarité de l'air et du métal dans la constitution d'une antenne. Ce principe peut être appliqué aux antennes d'extension infinie, c'est-à-dire celles qui sont suffisamment grandes pour ne pas créer d'ondes réfléchies à leur extrémité. Dans ce cas, et lorsque les parties complémentaires sont identiques, autrement dit qu'elles peuvent se recouvrir par rotation (figure 9.28), l'impédance d'entrée de l'antenne est égale à la moitié de l'impédance du vide. Ce qui est remarquable dans cette observation est l'indépendance de l'impédance en fonction de la fréquence. Ce principe assure donc une grande largeur de bande, qui est bien sûr limitée par la finitude de l'antenne tant du côté de la petite dimension que de la grande dimension.

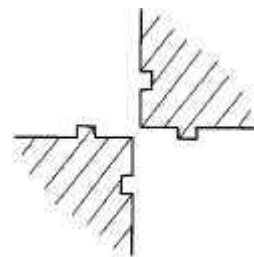


Figure 9.28 – Antenne planeaire auto-complémentaire

□ Application d'un facteur d'échelle en fréquence

Une façon classique de créer une antenne large bande est de créer une forme de l'antenne qui permette de reproduire des phénomènes de rayonnement identiques dans plusieurs bandes de fréquences adjacentes. On aboutit à la conception d'antennes logarithmiques.

□ Couplage

Remarquons aussi que les phénomènes de couplage élargissent la bande passante. Ce principe est utilisé pour obtenir une bande passante plus large pour certaines antennes résonnantes. Les couplages ont alors lieu, soit au niveau de l'excitation, soit au niveau de la forme de l'antenne qui introduit un élément résonnant supplémentaire. Ce phénomène d'élargissement se manifeste dans les réseaux qui ont une bande passante légèrement plus grande que l'antenne élémentaire. Après avoir énoncé quelques principes qui concourent à créer un rayonnement large bande, nous allons décrire quelques antennes qui s'appuient sur un ou plusieurs de ces principes.

■ Antenne à onde progressive

Le premier principe de fonctionnement d'une antenne large bande repose sur la nécessité d'avoir une adaptation large bande en entrée de l'antenne. Lorsque l'antenne présente une réflexion minimale en entrée, la puissance se propageant sur l'antenne est maximale. Cette puissance peut alors être rayonnée. C'est le cas d'une antenne à onde progressive : l'onde se propage et rayonne. Si l'antenne est terminée par une charge adaptée (figure 9.29), il n'y a pas de réflexion à l'extrémité.

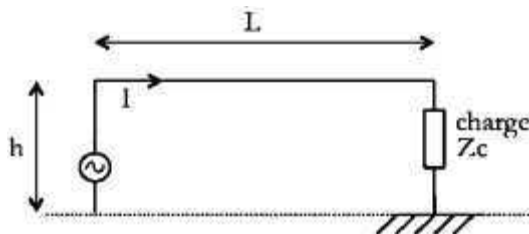


Figure 9.29 – Antenne filaire à onde progressive

Si la charge adaptée est large bande, l'antenne fonctionne aussi sur une large bande. La condition sur la largeur de bande de la charge disparaît si l'antenne est suffisamment longue. En effet, la puissance décroît au cours de la propagation, à cause du rayonnement et devient négligeable,

dans ce cas, à l'extrémité du fil. La valeur de la charge terminale n'a donc que peu d'effet. Les antennes de ce type atteignent une largeur de bande de l'ordre de l'octave. Leur impédance de rayonnement varie entre 200 et 300 Ω . L'inconvénient de ce type d'antenne réside dans la forme du diagramme de rayonnement présentant une inclinaison, par rapport à l'axe du fil, dépendant de la longueur d'onde. Le maximum de rayonnement est obtenu pour l'angle $\theta_{\max}$, tel que :

$$\cos \theta_{\max} = 1 - \frac{0,37}{L/\lambda}$$

De plus, le diagramme de rayonnement est sensible à la proximité du sol. Afin d'obtenir un diagramme mieux maîtrisé, on associe les fils rayonnant selon un losange (voir paragraphe 8.3.2). Remarquons que, si la charge n'est pas bien adaptée et l'antenne assez courte, l'intensité du courant dans le fil prend la forme :

$$I = I_0(e^{-j\beta z} + \Gamma e^{-j\beta z})$$

Le coefficient Γ est lié à la réflexion à l'extrémité. S'il n'est pas nul, un système d'ondes stationnaires s'installe, dû au retour de la puissance vers le générateur. La largeur de bande est ainsi limitée par le phénomène de réflexion. Remarquons que les ondes en retour rayonnent aussi et perturbent le diagramme de rayonnement.

■ Antenne hélicoïdale

L'antenne hélicoïdale est formée de spires enroulées autour d'un cylindre (figure 9.30)

L'antenne hélicoïdale peut être considérée comme une antenne filaire enroulée. Les paramètres importants sont : son diamètre (à rapporter à la longueur d'onde), l'espacement des spires, l'angle des spires et sa longueur. Les deux cas limites sont l'antenne filaire rectiligne et l'antenne boucle.

Deux modes peuvent apparaître :

- le mode normal qui correspond à un rayonnement radial autour de l'axe Oz
- le mode axial qui entraîne un rayonnement maximal selon l'axe Oz .

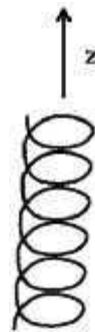


Figure 9.30 – Antenne hélicoïdale

□ Mode normal

Le mode normal apparaît lorsque la taille des spires est plus petite que la longueur d'onde. La projection de l'antenne sur le plan perpendiculaire à l'axe Oz est un cercle, alors que sa projection selon Oz est un segment de la longueur de l'antenne. Il en résulte que le rayonnement est la superposition de celui d'une boucle circulaire perpendiculaire à Oz et d'un fil parallèle à Oz . La fonction caractéristique de rayonnement de ces deux éléments est la même. Cependant leurs polarisations sont perpendiculaires et en quadrature de phase (voir section 3.1.5). Le rayonnement résultant de la superposition de ces deux types de courants est donc radial, de même fonction caractéristique que les dipôles et de polarisation elliptique. On montre que le rayonnement devient circulaire, lorsque l'écart h entre les spires est lié à la circonférence C par la relation :

$$C = \sqrt{2h\lambda}$$

Ce type d'antenne a été très utilisé par une génération de téléphones mobiles, comme un résonateur quart d'onde, placé au-dessus du plan de masse. Ces antennes, fonctionnant sur le mode normal, ne sont pas large bande.

□ Mode axial

Le mode axial apparaît lorsque la circonférence des spires est de l'ordre de grandeur de la longueur d'onde. En effet, considérons deux points opposés de la circonférence. Si celle-ci est égale à la longueur d'onde, un déphasage de π doit exister entre les chemins électriques séparant ces points. Comme les deux points sont placés en opposition sur la boucle, les orientations locales sur la boucle sont opposées, introduisant un déphasage de π . Les points opposés rayonnent donc en phase dans la direction de l'axe Oz .

On montre que ce principe fonctionne bien s'il y a plusieurs spires lorsque la circonférence C reste de l'ordre de la longueur d'onde, en respectant la condition :

$$\frac{3\lambda}{4} \leq C \leq \frac{4\lambda}{3}$$

Cela correspond à une bande relative de fréquences légèrement inférieure à 2.

Le diagramme de rayonnement est obtenu en remarquant que chaque spire peut être considérée comme un élément d'un réseau. Ces antennes se rapprochent du fonctionnement des antennes à ondes progressives. Les antennes fonctionnant dans ce mode permettent d'obtenir une polarisation circulaire et un gain qui peut atteindre environ 15 dB. Elles sont utilisées dans les liaisons satellites.

■ Antennes biconiques

L'antenne biconique est formée de deux cônes symétriques, alimentés par leurs sommets (figure 9.31). On peut la considérer comme la déformation d'un fil épais dont la surface s'incline. En principe, l'antenne biconique est infinie. Le fonctionnement qui va en être présenté se place dans cette hypothèse.

Nous nous plaçons dans le cas d'un mode TEM (Transverse Électrique et Magnétique). Le champ électrique n'a alors qu'une composante selon θ et le champ magnétique, selon φ . L'équation de Maxwell-Ampère, exprimée en coordonnées sphériques, permet de déduire :

$$-\frac{1}{r} \frac{\partial}{\partial r}(rH_\varphi) = j\omega E_\theta \quad \text{et} \quad H_\varphi \sin \theta = cste$$

D'où les expressions du champ électromagnétique :

$$H_\varphi = H_0 \frac{e^{-jkr}}{4\pi r} \frac{1}{\sin \theta} \quad E_\theta = ZH_0 \frac{e^{-jkr}}{4\pi r} \frac{1}{\sin \theta}$$

Par intégration, il est possible d'obtenir le courant et la tension :

$$V(r) = \int_{\theta_0}^{\pi-\theta_0} E_\theta d\theta \quad \text{soit} \quad V(r) = \frac{ZH_0}{2\pi} e^{-jkr} \ln \left(\cot \frac{\theta_0}{2} \right)$$

$$\text{De même } I(r) = \int_0^{2\pi} H_\varphi r \sin \theta d\varphi \quad \text{soit } I(r) = \frac{H_0}{2} e^{-jkr}$$

L'impédance se met donc sous la forme :

$$Z_{ant} = \frac{Z}{\pi} \ln \left(\cot \frac{\theta_0}{2} \right)$$

L'impédance Z du vide est égale à 377Ω . On remarque que Z_{ant} est indépendante de la fréquence. L'antenne biconique infinie est donc large bande.

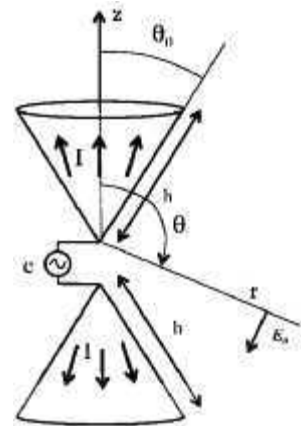


Figure 9.31 – Antenne biconique

Dans la réalité, la structure réalisée a une extension finie. Il se produit donc des réflexions aux extrémités, donnant lieu à un régime d'ondes stationnaires. De ce fait, l'impédance qui est réelle pour la structure infinie devient complexe pour la structure finie. Des modélisations montrent que, lorsque l'angle est faible, la largeur de bande est faible, car l'impédance varie de façon importante avec la fréquence. L'optimum est obtenu lorsque $\theta_0 = 45^\circ$. Cette valeur répond partiellement au principe d'auto-complémentarité, cité au début de ce chapitre.

Il est possible de réaliser des antennes présentant un seul cône au-dessus d'un plan de masse. La théorie des images permet alors de déduire leurs propriétés.

■ Antenne Bow-Tie

L'antenne Bow-Tie présente des ressemblances avec l'antenne bi-conique. Elle est réalisée en structure planaire (figure 9.32).

Cette antenne est large bande pour des longueurs comprises entre 0,3 et 0,8 longueur d'onde, pour des angles supérieurs à 20° . Elle permet d'obtenir facilement une largeur de bande de 2 : 1.

Le rayonnement ne présente pas un gain très important. Il est, au maximum, de 2 à 3 dB supérieur à celui du dipôle $\lambda/2$, selon l'angle θ_0 . Le maximum de rayonnement est obtenu perpendiculairement au plan de l'antenne, de chaque côté du plan.

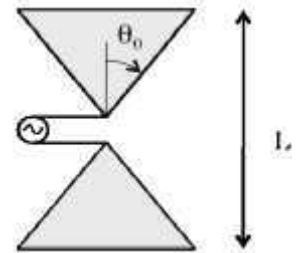


Figure 9.32 – Antenne Bow-Tie

■ Antennes log-périodiques

La conception des antennes logarithmiques repose sur l'idée de reproduire, dans des bandes de fréquences adjacentes, le même phénomène de rayonnement en utilisant un facteur d'échelle en fréquence.

Rappelons le principe de Rumsey qui pose qu'une antenne est de bande infinie, si ses dimensions sont infinies et répond au principe d'auto-complémentarité. Sa forme est donc repérée uniquement par rapport aux angles.

Supposons qu'une antenne large bande soit définie par les dimensions des parties métalliques et diélectriques, repérées par :

$$r = d(\theta, \varphi)$$

Si, cette antenne doit être conçue dans une autre bande de fréquences α fois plus petite que la bande précédente, ses dimensions seront :

$$r' = \alpha d(\theta, \varphi)$$

Il est possible de définir plusieurs bandes de fréquences adjacentes pour couvrir une large bande. Les dimensions de la $n^{\text{ième}}$ bande sont dans le rapport α avec celles de la $(n + 1)^{\text{ième}}$ bande :

$$\frac{r_n}{r_{n+1}} = \alpha$$

Donc le rapport entre la $n^{\text{ième}}$ cellule et la cellule de base est : α^n . L'antenne étant large bande, l'application d'un facteur d'échelle α^n doit redonner les mêmes dimensions quel que soit n . Cela montre que le champ rayonné doit être périodique avec le logarithme de la fréquence.

En fait, les antennes n'étant pas d'extension infinie, des réflexions ont lieu aux extrémités. Afin d'éviter cet inconvénient qui limite la largeur de bande, il est conseillé de concevoir les antennes de façon à ce que le rayonnement soit prépondérant sur la réflexion. Les ondes stationnaires apparaissant sont alors minimisées. Pour cela, des dimensions de l'antenne doivent être au moins de l'ordre de la demi-longueur d'onde correspondant à la fréquence la plus basse.

■ Antenne spirale

L'antenne spirale est une antenne planaire, constituée de zones métalliques délimitées par des spirales (figure 9.33).

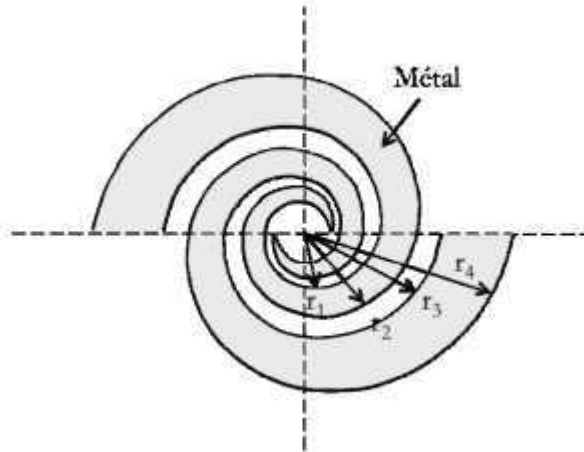


Figure 9.33 – Antenne spirale

L'équation d'une spirale logarithmique est définie par :

$$r = a \exp(b\varphi) \quad [9.1]$$

On en déduit l'équation suivante :

$$b = \frac{1}{r} \frac{dr}{d\varphi} = \cot \beta$$

La figure 9.34 donne la signification géométrique des différents termes.

La construction de l'antenne spirale s'appuie sur le dessin de quatre spirales décalées respectivement de rayons r_1, r_2, r_3, r_4 , selon les équations :

$$\begin{aligned} r_1 &= a \exp(b\varphi) \\ r_2 &= a \exp(b(\varphi - \delta)) \\ r_3 &= a \exp(b(\varphi - \pi)) \\ r_4 &= a \exp(b(\varphi - \pi - \delta)) \end{aligned}$$

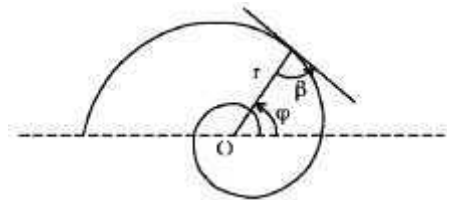


Figure 9.34 – Spirale logarithmique

Comparons le rayon à la longueur d'onde, à partir de l'équation [9.1] :

$$\frac{r}{\lambda} = a \exp \left[b \left(\varphi - \ln \frac{\lambda}{b} \right) \right]$$

Posons :

$$\varphi_0 = \ln \frac{\lambda}{b}$$

Nous en déduisons que le changement de longueur d'onde revient à une rotation d'un angle correspondant au logarithme de la longueur d'onde (ou de la fréquence) rapporté au coefficient b .

Les phénomènes sont donc les mêmes d'une fréquence à l'autre, mais on les retrouve à des endroits qui ont tourné sur la spirale.

En réalité, la spirale est finie. Il est recommandé de fermer les bras de la spirale par une courbe graduelle afin d'éviter les discontinuités abruptes qui génèrent des ondes en retour. Un arc de cercle peu incurvé convient.

Le rayonnement de la spirale est perpendiculaire au plan de celle-ci. La polarisation est circulaire sur l'axe de l'antenne et de sens contraire de chaque côté du plan.

Afin d'avoir une bonne efficacité, il est nécessaire d'alimenter les deux bras de la spirale de façon équilibrée.

La spirale a été présentée ici comme la partie métallique gravée. La structure complémentaire, constituée d'un plan de masse dans lequel a été évidé le métal, appartient aussi à la catégorie des antennes spirales et présente des propriétés analogues.

■ Antenne spirale conique

Cette antenne, représentée sur la figure 9.35, emprunte les propriétés de l'antenne spirale et de l'antenne conique.

□ Antenne log-périodique à éléments rayonnants

Une antenne large bande très utilisée pour les récepteurs en radio astronomie est l'antenne log-périodique formée d'éléments circulaires (figure 9.36).

Les différents éléments ont des longueurs différentes et constituent des résonateurs. Lorsque la fréquence varie, les différents éléments entrent en résonance selon leur longueur. Afin d'assurer une grande bande de fréquences, les rapports entre les rayons des différents éléments doivent respecter les conditions :

$$\frac{r_n}{r_{n+1}} = \alpha \quad \text{et} \quad \frac{\rho_n}{r_n} = \alpha'$$

Le diagramme de rayonnement est maximal de part et d'autre du plan de l'antenne, perpendiculairement à celle-ci.

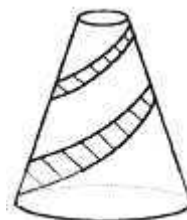


Figure 9.35 – Antenne spirale conique

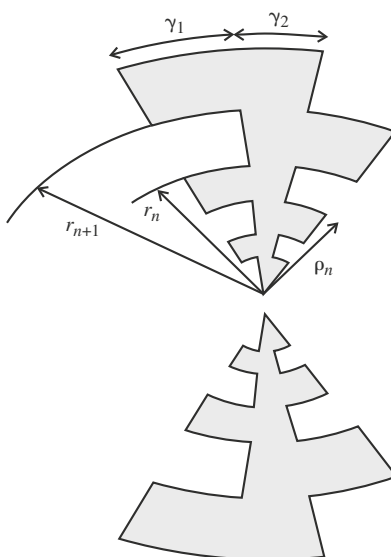


Figure 9.36 – Antenne log périodique à éléments

9.2.3 Antennes Ultra Large Bande (ULB)

Rappelons, dans cette première partie, les principes énoncés dans le paragraphe précédent. Les antennes de type large bande trouvent leurs applications lorsque l'on cherche à couvrir un large spectre de fréquences avec différents objectifs :

- atteindre des débits importants ;
- permettre d'accéder à plusieurs normes de communication avec la même architecture ;
- établir une liaison radar ;
- travailler en mode impulsionnel ;
- utiliser un mode de communication basé sur de l'étalement de spectre...

La bande passante peut être définie de différentes manières en fonction de la grandeur caractéristique prise en compte :

- adaptation de l'antenne ;
- polarisation ;
- gain...

On peut considérer deux types d'antennes large bande, les antennes résonnantes et les antennes à ondes progressives.

Dans le cas des antennes résonnantes, pour obtenir un fonctionnement large bande, on peut soit diminuer la qualité d'une résonance, soit coupler plusieurs résonances entre elles.

L'exemple ci-dessous (figure 9.37) montre le cas d'une antenne fractale pour laquelle la répétition à différentes échelles du même motif, fait apparaître des résonances à plusieurs fréquences.



Figure 9.37 – Cas de l'antenne fractale

Dans le cas des antennes à ouverture progressive, il est possible d'utiliser le concept du passage progressif d'une ligne de transmission à l'espace libre, comme dans le cas de l'antenne de type Vivaldi (figure 9.38).

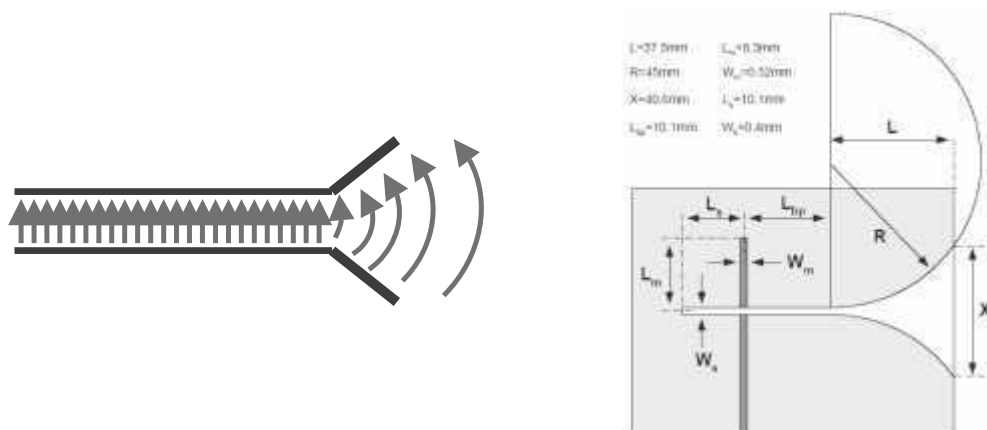


Figure 9.38 – Antenne Vivaldi

Il est également envisageable d'utiliser le principe donné par Rumsey et qui porte sur les antennes indépendantes de la fréquence. Il énonce que si la forme d'une antenne peut être définie par des angles, cette antenne est alors indépendante de la fréquence. Elle est alors confondue avec sa réduction homothétique. L'exemple ci-dessous (figure 9.39) montre un type d'antenne « *Bow-Tie* »

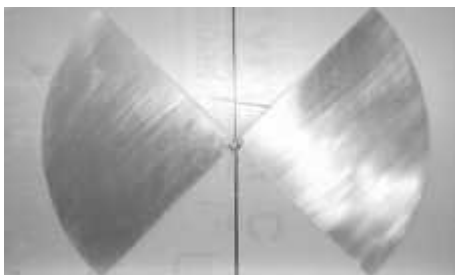


Figure 9.39 – Antenne « *Bow-Tie* »

Une technique assez classique utilisée pour réaliser une antenne large bande repose sur le principe de l'antenne log-périodique où les caractéristiques de l'antenne sont des fonctions périodiques du logarithme de la fréquence. Le principe est transposable en filaire, en planaire ou en volumique (figure 9.40).

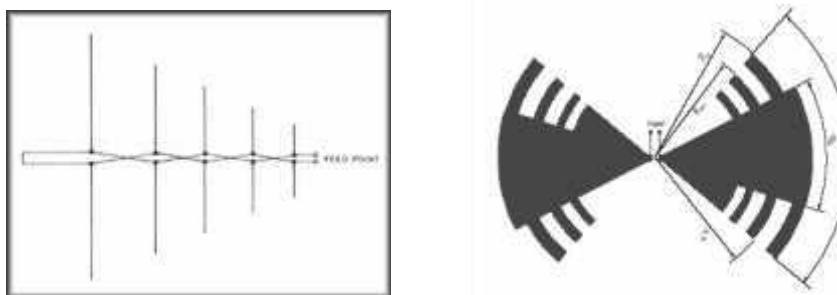


Figure 9.40 – Antennes log-périodiques

Une dernière possibilité dont la théorie a été énoncée précédemment, repose sur le principe de la spirale logarithmique planaire ou en volume (figure 9.41).



Figure 9.41 – Antenne spirale logarithmique

■ Intérêt des transmissions ULB

Les transmissions ULB (UWB en anglais) étaient, il y a encore une vingtaine d'années, réservées aux communications militaires. Les technologies ULB ont très fortement évolué durant les dernières années, lorsqu'en février 2002, la FCC (*Federal Communications Commission*) a mis en place une régulation permettant l'utilisation de ces technologies pour des applications grand public. En libérant aux États-Unis une bande continue de 7,5 GHz, le FCC a permis l'ouverture à de nouvelles applications pour les transmissions à très haut débit. Par contre, la puissance émise étant limitée à 0,5 mW pour toute la bande FCC, les applications ne pourront être que de proximité. Une autre approche orientée vers le bas débit et destinée aux réseaux de capteurs par exemple, a été envisagée. Elle utilise alors un mode de transmission de type impulsionnel.

□ La normalisation

L'ULB a d'abord eu une première définition donnée par Taylor. Il s'agissait de systèmes qui transmettent et reçoivent des ondes dont la largeur de bande relative (LB) est supérieure ou égale à 0,25 avec :

$$LB = \frac{f_h - f_l}{f_c} \quad f_c = \frac{f_h + f_l}{2}$$

Cette première définition a été modifiée et remplacée par une nouvelle proposée par la FCC. Selon cette nouvelle définition, un signal ULB est un signal dont la bande passante à -10 dB dépasse à tout moment 500 MHz et 20 % de la fréquence centrale.

La principale bande destinée à l'ULB se situe entre 3,1 et 10,6 GHz. Cette bande représente environ 7 GHz et pourrait donc être divisée en 14 sous-bandes de 500 MHz. Un système utilisant toute la bande, ou un ensemble de sous-bandes (voire une sous-bande), sera considéré comme un système ULB. Mais il doit bien sûr respecter les normes en vigueur dans le pays considéré.

Le principe de base des systèmes ULB est de pouvoir cohabiter dans des bandes de fréquences déjà utilisées par d'autres systèmes de communications. Il permet donc de ne pas passer par un mécanisme d'allocations de licences ou de se trouver confiner dans des bandes de fréquences dites sans licence (bandes ISM 2,4 et 5,2 GHz par exemple). Par contre, les systèmes ULB ne doivent pas brouiller les systèmes existants, d'où l'importance de l'aspect réglementaire.

□ La réglementation nord américaine

En Amérique du Nord, la réglementation de l'ULB a été mise en place par la FCC avec la publication du rapport « *First Report and Order* » en février 2002.

L'émission des signaux ULB pour les communications est autorisée sans licence pour des applications « *indoor* » et pour des liaisons mobiles point à point en « *outdoor* ».

Les appareils « *indoor* » doivent être conçus pour ne fonctionner qu'en ce milieu et ne doivent pas être dirigés intentionnellement vers l'extérieur.

D'autre part, les systèmes « *outdoor* » autorisés sont les appareils portables ne reposant pas sur une infrastructure fixe.

En ce qui concerne les systèmes de communication, la puissance moyenne d'émission des signaux est limitée par les masques représentés à la figure 9.42.

La limite de $-41,3$ dBm/MHz correspond à une mesure de champ électromagnétique de $500 \mu\text{V/m}$ dans toute bande de 1 MHz, à 3 mètres de l'antenne d'émission.

Cette limite a été établie par la Partie 15 des textes de la FCC en référence à la puissance isotrope rayonnée équivalente (PIRE) des émissions non intentionnelles.

La puissance pic est également limitée dans le rapport de la FCC. Elle est mesurée autour de la fréquence pour laquelle le rayonnement est maximum et se calcule à partir de :

$$P_{pic}^{lim}(RBW) = 20 \log_{10} \left(\frac{RBW}{50} \right)$$

où RBW est la bande de résolution de la mesure exprimée en MHz.

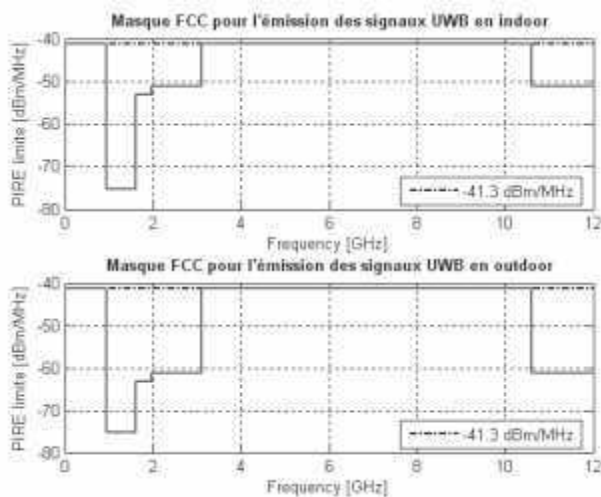


Figure 9.42 – Masques FCC

Pour $RBW = 50$ MHz, la puissance pic ne doit pas dépasser 0 dBm, soit 1 mW.

□ La réglementation européenne

En mars 2006, a été publiée la décision finale de l'ECC, sur les conditions d'utilisation de la technologie ULB dans les bandes inférieures à 10,6 GHz.

La décision s'applique aux technologies ULB de largeur de bande supérieure à 500 MHz dans la bande au-dessous de 10,6 GHz qui sont exemptées de licence et opèrent selon le principe de non-interférence et non-protection.

Le masque actuel pour les systèmes ULB en Europe est donné en figure 9.43.

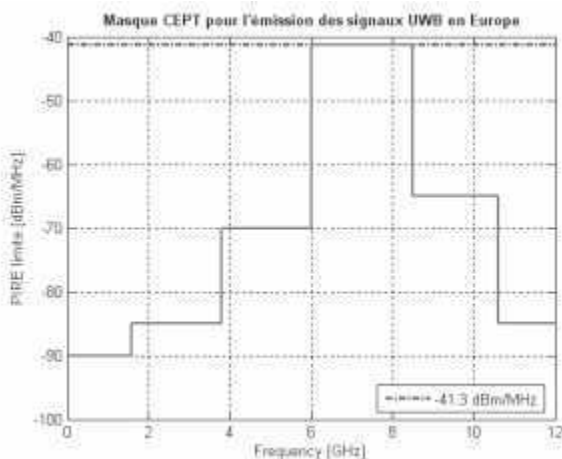


Figure 9.43 – Masque pour l'Europe

□ Les différentes techniques de transmission

Plusieurs approches ont été diversement étudiées et peuvent être réparties en :

- ULB mono bande :
 - approche impulsionnelle : IR UWB (*Impulse Radio – Ultra Wide Band*)
- ULB multibandes :
 - approche MB-OFDM (*Multi-bandes OFDM*)
 - approche MB-OOK (*Multi-bandes On-Off-Keying*)

Dans le cas de l'approche mono bande, le signal correspond à un train d'impulsions de type large bande modulée en amplitude ou en position (figure 9.44).

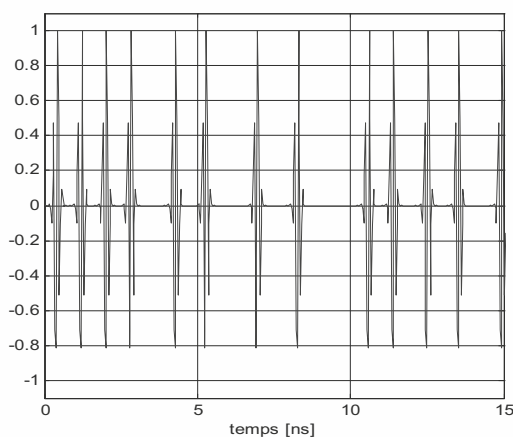


Figure 9.44 – Mode de transmission de type IR – UWB

En ce qui concerne l'approche multibandes, deux techniques ont été développées. La première (MB-OFDM) est utilisée une bande de fréquences divisée en plusieurs sous-porteuses (figure 9.45). Chaque sous-porteuse est modulée puis transmise.

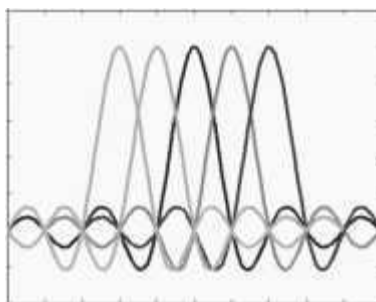


Figure 9.45 – Mode de transmission de type OFDM

L'approche MB-OOK est de type impulsionnel, mais le signal est divisé en plusieurs sous-bandes (figure 9.46).

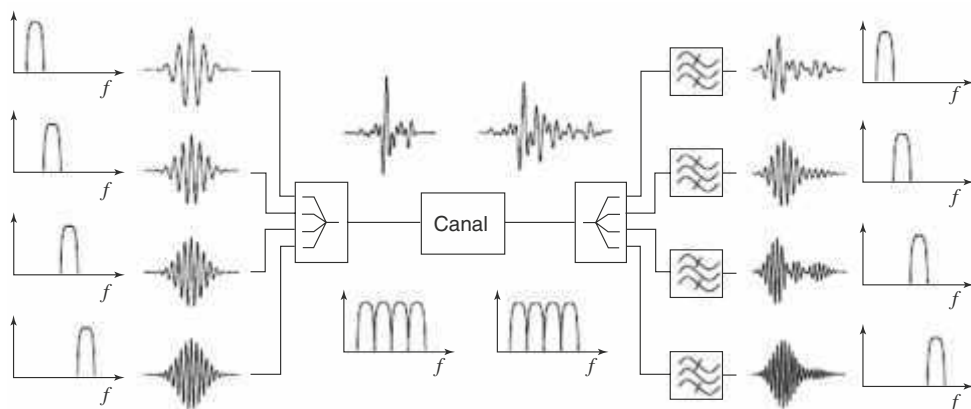


Figure 9.46 – Mode de transmission de type MB – OOK

■ Transmissions ULB en mode impulsionnel

□ Comparaison d'une transmission ULB impulsionnelle par rapport à une transmission en bande étroite

Le principe d'une transmission ULB impulsionnelle repose sur l'émission d'un signal fortement limité dans le temps, de type impulsion. En transmission en bande étroite, l'émission peut se faire de façon continue. De ce fait, l'analyse fréquentielle d'un signal ULB présente une large occupation spectrale en comparaison d'un signal bande étroite (figure 9.47).

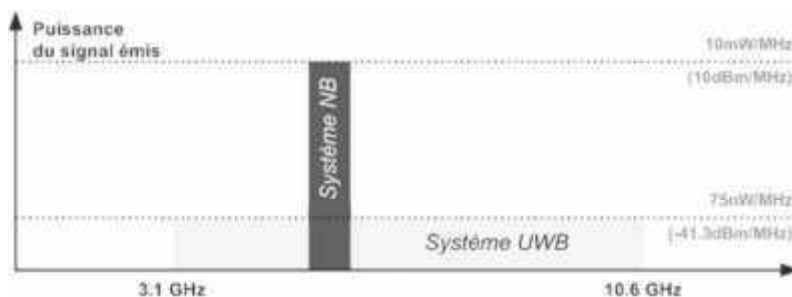


Figure 9.47 – Transmission ULB versus bande étroite

On peut donc remarquer que, dans le cas des applications ULB, l'augmentation de la largeur de bande du signal se fait au détriment de la densité spectrale de puissance (DSP) afin de garantir la coexistence avec les systèmes existants.

■ Capacité du canal

En communications numériques, une relation fondamentale est celle de Shannon qui donne la capacité C d'un canal à bruit blanc additif gaussien et à bande limitée.

$$C = B \log_2 (1 + SNR)$$

SNR = rapport signal sur bruit.

Cette formule est réductrice car elle s'applique au cas le plus favorable en termes de propagation et d'interférence entre symboles. Mais elle montre que la capacité du canal croît de manière

logarithmique avec la puissance transmise alors qu'elle est proportionnelle à la bande passante du signal.

Lorsque l'on considère des canaux multitrajets, cette caractéristique liée à la capacité du canal est conservée, même s'il n'est plus possible d'extraire la valeur numérique de la capacité de cette formule.

Par conséquent, la diminution de la puissance émise au profit d'une largeur de bande plus grande s'avère être intéressante si l'on veut gagner du débit.

■ Problématique de caractérisation et de conception des antennes en ULB

□ Les particularités du canal ULB

En raison de la largeur de bande des systèmes ULB, il n'est pas possible d'utiliser l'ensemble des théories mises en place pour les systèmes à bande étroite.

Dans un système à bande étroite, tous les trajets arrivent dans un intervalle de temps inférieur à la résolution du récepteur tandis que dans le cas de l'ULB, la résolution du récepteur est inférieure à l'étalement du canal. Dans l'analyse d'un système ULB, il est important de prendre en compte l'effet du canal.

C'est un canal difficile avec une forte atténuation surtout en milieu indoor, avec une décroissance exponentielle de la puissance, avec de nombreux multitrajets, et un étalement temporel important (200 ns max). Cela nécessitera donc, dans les systèmes ULB de type impulsif, de conserver un intervalle de garde.

Le modèle de canal généralement retenu dans le domaine des communications ULB en intérieur est celui de Saleh et Valenzuela. Ce canal est modélisé dans le cadre de l'IEEE802.15.3a.

Les graphiques suivants montrent la réponse du canal dans le cas d'une liaison à hauts débits, courte portée, « indoor ».

Quatre cas sont représentés :

- CM1 : visibilité directe (LOS), entre 0 et 4 mètres (figure 9.48 a)
- CM2 : absence de trajet direct (NLOS), entre 0 et 4 mètres (figure 9.48 b)
- CM3 : absence de trajet direct (NLOS), entre 4 et 10 mètres (figure 9.48 c)
- CM4 : configuration NLOS difficile avec un nombre et une densité des trajets très importants (figure 9.48 d)

L'étalement du canal pour le meilleur cas en visibilité directe (CM1) est d'environ 40 ns, tandis que dans un environnement sans visibilité, il est proche de 80 ns (CM3) et il atteint presque 180 ns pour la CM4.

L'information sur l'étalement du canal est très importante pour le dimensionnement de l'architecture et dans antennes ULB, il détermine la période de répétition des impulsions.

□ Les antennes utilisables en ULB

Comme dans tous les systèmes de communications, l'antenne est un élément clef, mais encore plus dans le cas des communications ULB.

Dans le cas de ce type d'antenne, il sera nécessaire de trouver un compromis largeur de bande/rendement/intégration/coût.

Fonctionnant en large bande, leur comportement est plus difficile à analyser.

De plus, la fonction antenne, ne peut plus être considérée indépendamment du reste de l'architecture. L'antenne est un élément permettant le transfert d'énergie et le rayonnement. Le transfert d'énergie est lié à l'adaptation et au rendement.

Les distorsions du signal proviennent de l'un des phénomènes suivants :

- désadaptation en fréquence ;
- distorsion d'amplitude ;
- distorsion de phase.

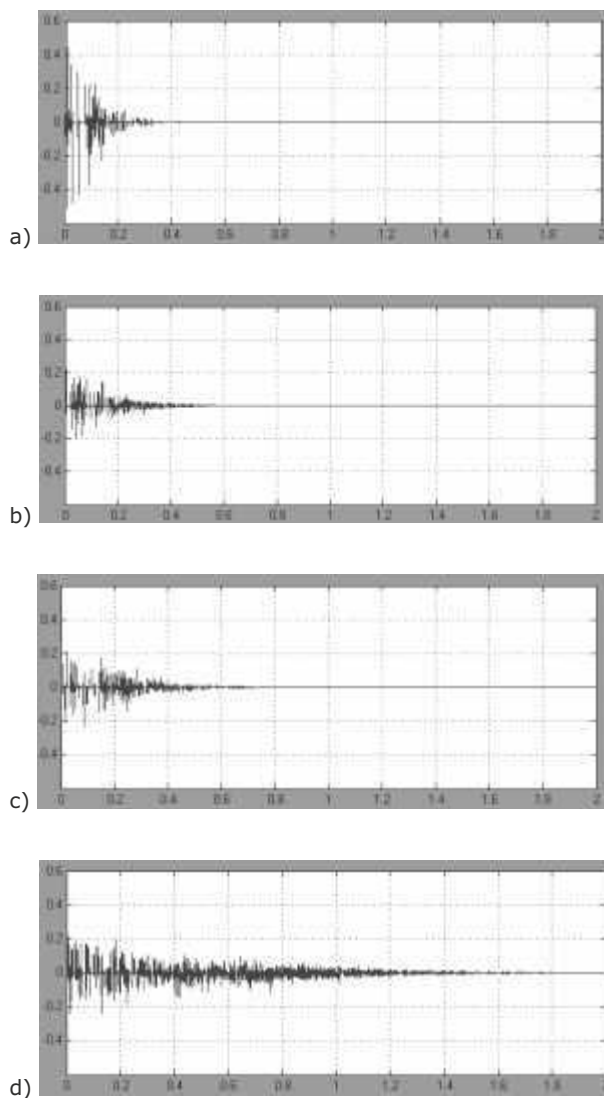


Figure 9.48 – Modèles des canaux ULB

Les antennes peuvent être classées en trois grandes familles :

- les antennes à bande étroite
- les antennes large bande
- les antennes indépendantes de la fréquence.

☐ **Prise en compte de l'effet de l'antenne**

Il est communément admis qu'étant donné que les impulsions ULB se caractérisent par une très grande largeur de bande, les antennes d'émission et de réception se comportent comme des filtres, ce qui produit une déformation de l'impulsion et des spectres des signaux.

Les graphiques suivants montrent l'effet de l'antenne sur une impulsion, dans le cas d'une application ULB fonctionnant dans la bande 3,1-5,1 GHz.

L'impulsion utilisée est donnée en figure 9.49. Elle est caractérisée par $\tau = 5$ ns.

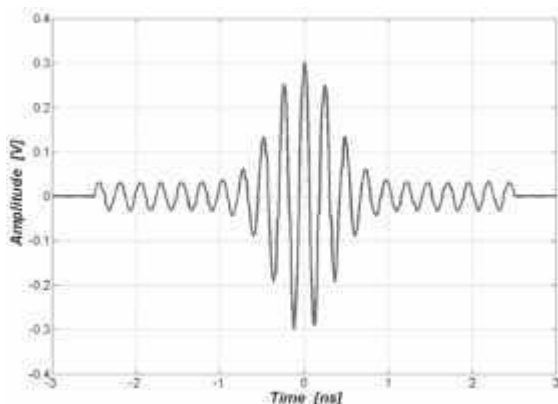


Figure 9.49 – Forme de l'impulsion ULB

Nous considérons deux types d'antennes :

- une antenne dipôle $\lambda/2$, de type bande étroite, de longueur 31,4 mm et caractérisée par un gain maximum de 2,7 dBi. (figure 9.50)



Figure 9.50 – Antenne dipôle

- une antenne « Diamond », de type large bande, de base et hauteur identiques (31,4 mm) et caractérisée par un gain de 2,8 dBi. (figure 9.51)

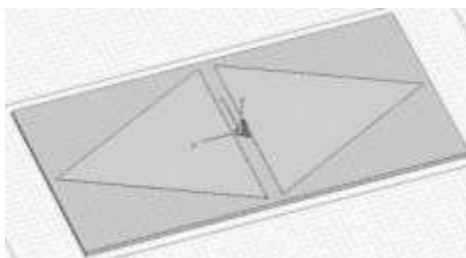


Figure 9.51 – Antenne Diamond

Les simulations ont été effectuées sous CST Microwave.

Afin d'analyser les phénomènes sur un cas simple, nous présentons les caractéristiques pour une impulsion gaussienne en entrée de l'antenne.

Les résultats en coefficient de réflexion et champ rayonné dans le domaine temporel, sont donnés en figures 9.52 et 9.53

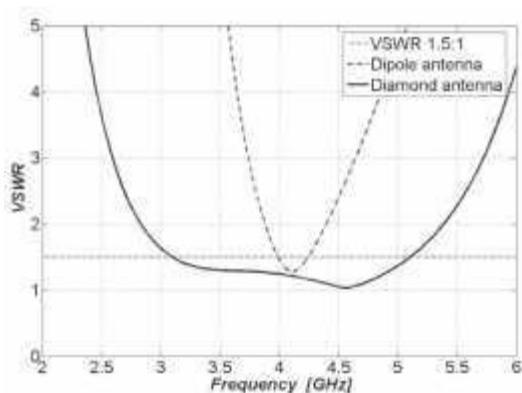


Figure 9.52 – Coefficients de réflexion des deux antennes

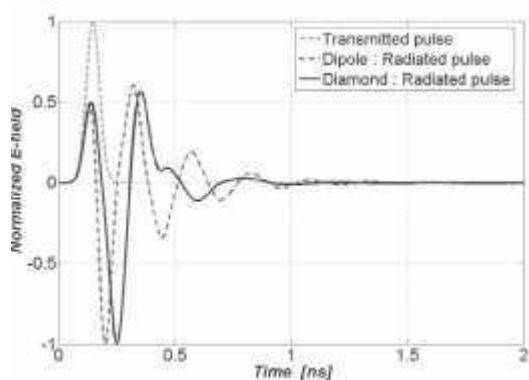


Figure 9.53 – Étalement temporel des deux antennes

Considérons maintenant l'impulsion choisie pour une communication ULB respectant un masque de fréquences. La réponse temporelle et la FFT du champ rayonné sont représentées en figures 9.54 et 9.55.

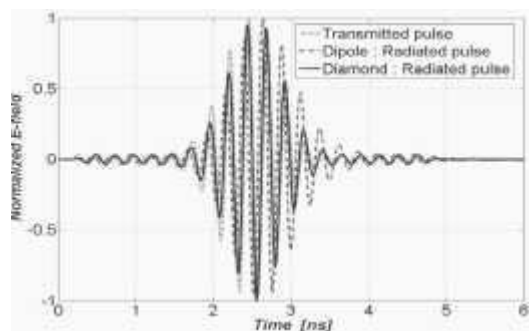


Figure 9.54 – Réponse temporelle des antennes dans le cas d'une impulsion réelle

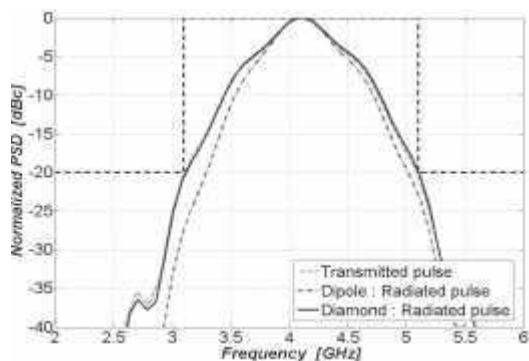


Figure 9.55 – FFT du champ rayonné pour les deux antennes

On constate ici que, dans la bande et avec l'impulsion choisie, les variations de l'enveloppe du signal sont suffisamment lentes pour être « suivies » par le dipôle.

Il y a peu d'étalement temporel dû à l'antenne.

La conclusion importante, est qu'il est indispensable d'analyser la largeur de bande, la forme de l'impulsion, pour déterminer les contraintes sur l'antenne.

□ Exemples d'antennes ULB

Dans les systèmes bande étroite, il est fréquent de retrouver des antennes résonantes, accordées sur la fréquence centrale d'un ensemble de canaux afin par exemple d'y maximiser l'efficacité.

Il est dès lors supposé que l'antenne a un gain et un rendement constants dans la bande considérée, à l'inverse d'un système ULB dont les grandeurs caractéristiques de l'antenne telles que le gain ou la position des centres de phase varient en fonction de la fréquence. Ces variations ont comme conséquence la distorsion de l'impulsion transmise.

En effet, ces variations s'interprètent comme une fonction de filtrage appliquée en sortie d'antenne. La variation du gain correspond à une amplitude non constante dans la bande passante du filtre tandis que la variation des centres de phases est assimilée à un retard de groupe non constant au niveau de ce même filtre. L'impulsion est donc filtrée : ses allures temporelle et fréquentielle sont modifiées.

L'approche la plus répandue pour réaliser une antenne ULB est basée autour d'une antenne dipôle. Afin d'augmenter la largeur de bande d'un dipôle, il est possible d'élargir ce dipôle au niveau de son alimentation pour former une antenne « Diamond » ou alors au niveau de ses extrémités pour obtenir une antenne Bow-Tie. D'autres variantes planaires peuvent être dérivées d'un dipôle à l'image des antennes elliptique et demi-elliptique comme l'illustre la figure 9.56. De manière générale, ces antennes présentent en entrée une impédance différentielle de 100 Ohms.

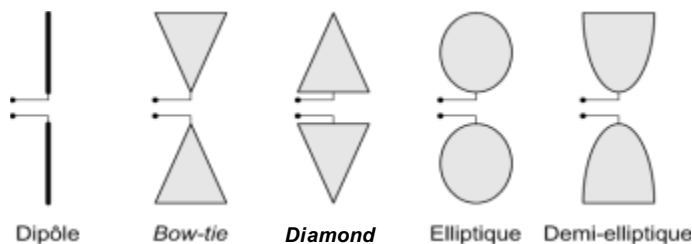


Figure 9.56 – Exemples d'antennes planaires dérivées du dipôle

Bibliographie

First report and order, ET Docket No. 98-153 Rapport, Federal Communication Commission, 2002.

AÏSSAT H., CIRIO L., GRZESKOWIAK M., LAHEURTE J.M., PICON O. — *Reconfigurable circularly polarized antenna for short-range communication systems*, IEEE Transactions on Microwave Theory and Techniques, vol. 54, n°6, pp. 2856-2863, June 2006.

BALANIS C.A. — *Antenna theory, analysis and design*, John Wiley & sons, third edition, 2005.

BARRET T.W. — *History of UltraWideBand (UWB) radar & communications : pioneers and inventors*, Progress in Electromagnetics Symposium 2000, Cambridge, MA, juillet 2000.

BRUNSON L.K. *et al.* — *Assessment of Compatibility Between Ultra Wideband Devices and Selected Federal Systems*, U.S. Department of Commerce, NTIA special publication 01-43, janvier 2001.

CEPT/ECC — *ECC Decision of 24 March 2006 on the harmonized conditions for devices using Ultra-Wideband (UWB) technology in bands below 10.6GHz*, Doc. ECC/DEC/(06)04, 2006.

CEPT/ECC — *Final CEPT Report in response to the second EC Mandate to CEPT to harmonize radio spectrum use for Ultra-wideband Systems in the European Union*, RSCOM05-73, 2005.

COMBES P.F. — *Micro-ondes, tome 2, circuits passifs, propagation, antennes*, Paris, Dunod, 1997.

DORE J.-B., UGUEN B., PAQUELET S. et MALLEGOL S. — *UWB Non-coherent high data rates transceiver Architecture and implementation*, IWCT'05, Oulu, 2005.

ELLIOTT R.S. — *Antenna Theory and Design*, New Jersey, IEEE Press John Wiley & Sons, 2003.

European Commission — Radio Spectrum Committee — *Mandate to CEPT to harmonize radio spectrum use for Ultra-Wideband Systems in the European Union*, Doc. RSCOM04-08 EN, 2004.

JALIL R., CHEN-TO T. — *A new class of resonant antennas*, IEEE Transactions on Antennas and Propagation, vol. 39, n° 9, septembre 1991.

JOHNSON R.C. — *Antenna Engineering Handbook, Third Edition*, New-York, McGraw-Hill, 1996.

KRAUS J.D., MARHEFKA R.J. — *Antennas for all applications*, McGraw-Hill Higher Education, 2002.

KRAUSS J.D. — *Antennas*, MacGraw-Hill, 2^e édition, 1988.

LO C. T., LEE S.W. — *Antenna Handbook, Antenna Theory, vol. II.*, New York, Van Nostrand Reinhold, 1993.

LO T.K. HO C.O., HWANG Y., LAM E.K.W., LEE B. — *Miniature aperture-coupled microstrip antenna of very high permittivity*, Electronics letters, vol. 33, n° 1, janvier 1997.

LUXEY C., STARAJ R., KOSSIAVAS G., PAPIERNIK A. — *Antennes imprimées, bases et principes, Techniques de l'ingénieur*, E 3 310, éditions TI.

LUXEY C., STARAJ R., KOSSIAVAS G., PAPIERNIK A. — *Techniques et domaines d'applications, Techniques de l'ingénieur*, E 3 311, éditions TI.

MALLEGOL S., COUPEZ J.P., PERSON C., LESPAGNOL T., PAQUELET S., BISIAUX A. — *Microwave (De) Multiplexer for Ultra-Wideband (UWB) Non-Coherent High Data Rates Transceiver*, Microwave Conference 2006, Manchester, pp. 1825-1828.

MARCHALAND D., VILLEGAS M., BAUDOIN G., TINELLA C., BELOT D. — *Novel Pulse Generator Architecture dedicated to Low Data rate UWB Systems*, European Microwave Week 2005, Proc. Conf. ECWT/EuMC, pp. 229-232, 1687-1, 3 – 7 octobre 2005 CNIT La Défense, Paris.

MARCHALAND D., VILLEGAS M., BAUDOIN G., TINELLA C., BELOT D. — *System concepts dedicated to UWB transmitter*, European Microwave Association Letters, vol. 2, juin 2006, pp. 116-121.

MARSDEN K., LEE H., HA D.S., ET LEE H. — *Low Power CMOS Reprogrammable Pulse Generator for UWB Systems*, IEEE Conference on Ultra Wideband Systems and Technologies, pp. 443-447, 2003.

NAKANO H., TAGAMI H., YOSHIKAWA A., YAMAUCHI J. — *Shortening ratios of modified dipole antennas*, IEEE Transactions on Antennas and Propagation, vol. 32, n° 4, avril 1984.

PAQUELET S., AUBERT L.-M. — *An energy adaptive demodulation for high data rates with impulse radio*", IEEE Radio and Wireless Conf. RAWCON, Atlanta, 2004.

PAQUELET S., AUBERT L.-M. et UGUEN B. — *An impulse radio asynchronous transceiver for high data rates*, IEEE Joint UWBST&IWUWBS2004 Conf., Kyoto, 2004.

PEZZIN M., KEIGNART J., DANIELE N., DE RIVAZ S., DENIS B., MORCHE D., ROUZET P., CATENOZ R., RINALDI N. — *Ultra Wideband : the radio link of the future ?* Annals of Telecommunications, vol. 58, n°3-4, 2003.

PIOCH S., LAHEURTE J.M. — *Size reduction of microstrip antennas by means of periodic metallic patterns*, Electronics Letters, vol. 39, n° 13, juin 2003.

RUDGE A.W., MILNE K., OLVER A.D., KNIGHT P. — *The Handbook of Antenna Design, Vol 1.*, London, UK, Peter Peregrinus Ltd., 1982.

SCHANTZ H.G., FULLERTON L — *"The diamond dipole : a Gaussian impulse antenna"* Antennas and Propagation Society International Symposium, Juillet 2001.

SCHANTZ H.G. — *Introduction to ultra-wideband antennas*, IEEE Conference on Ultra Wideband Systems and Technologies, pp. 1-9, 2003.

SCHANTZ Hans G. — *Measurement of UWB Antenna Efficiency*, The IEEE Semiannual Vehicular Technology Conference VTC2001 Spring.

SHENG H., ORLIK P., HAIMMOVICH A.M., CIMINI L.J., ZHANG Jr et J. — *On the Spectral and Power Requirements for Ultra-Wideband Transmission*, IEEE International Conference on Communications, vol. 1, pp. 738-742, 2003.

SIEVENPIPER D., ZHANG L., JIMENEZ BROAS R.F., ALEXOPOULOS N.G., YABLONOVITCH E. — *High-impedance electromagnetic surface with a forbidden frequency band*, IEEE Transactions on Microwave Theory and Techniques, vol. 47, n°11, novembre 1999.

SIWIAK K., WITHINGTON P., PHELAN S. — *Ultra-wide band radio : the emergence of an important new technology*, Vehicular Technology Conference, 2001. VTC 2001 Spring. IEEE VTS 53rd, Volume : 2, 6-9 mai 2001.

STUTZMANN W.L., GARY A.T. — *Antenna Theory and Design*, New Jersey, A. John Wiley & Sons, Inc., 1998.

SUAREZ M., VILLEGAS M., BAUDOIN G. — *Non uniform band pass filter bank for an UWB MB-OOK transceiver architecture*, ICWMC 2007.

ULABY F.T., MOORE R.K., FUNG A.K. — *Microwave Remote Sensing, Active and Passive, Vol. I*, Norwood, Artech House, 1981.

WEISENHORN M. et HIRT W., *Impact of the FCC Average- and Peak Power Constraints on the Power of UWB Radio Signals*, PULSERS Project, ANNEX 1 for D3b4b, septembre 2004.

WIN M.Z., SCHOLTZ R.A. — *Impulse Radio : How it works*, IEEE Communications Letters, vol. 2, pp. 36–38, 1998.

WITHINGTON P. II., FULLERTON L. — *An impulse radio communications system*, in Proc. of the International Conference on Ultra-Wide Band, Short Pulse Electromagnetics, Brooklyn NY, USA, pp. 113-120, octobre 1992.

10.1 Introduction

L'outil numérique est largement utilisé dans la conception des antennes. Il facilite, par exemple, l'optimisation des dimensions géométriques ou de la forme de l'élément rayonnant de façon à maximiser les performances en termes de rayonnement, d'adaptation ou de qualité de polarisation. Cependant, dans bien des cas, dans un souci de simplification et pour minimiser les durées de simulations, un certain nombre d'approximations sont souvent effectuées. Cela se traduit par des décalages plus ou moins importants dans les résultats simulés obtenus ce dont l'utilisateur ne percevra pas toujours les effets et les conséquences.

Pour quantifier les performances réelles des antennes et les comparer aux résultats simulés, la mesure dans un environnement adapté s'avère incontournable. Non pas qu'il faille systématiquement considérer les résultats expérimentaux comme référence — en effet les sources d'erreurs sont multiples le long d'une chaîne d'acquisition — mais ils permettent de tenir compte des caractéristiques géométriques et électromagnétiques réelles des systèmes mesurés.

Dans ce qui suit, nous allons détailler les principales méthodes qui permettent de caractériser expérimentalement les antennes par leur diagramme de rayonnement, leur gain et leur polarisation. Notons que le théorème de réciprocité nous signifie que, quelle que soit la position de l'antenne sous test (émettrice ou réceptrice pour les besoins de la mesure), les résultats en termes de rayonnement seront identiques.

10.2 Rappels sur les différentes zones de rayonnement

Nous allons rappeler les trois zones de rayonnement d'une antenne. Pour cela, nous allons considérer une antenne circulaire de diamètre D suffisamment grand par rapport à la longueur d'onde. L'antenne est excitée par une onde de fréquence fixe.

En s'éloignant de l'antenne, on constate qu'il existe trois zones qui sont :

- la zone de Rayleigh ou zone proche, qui s'étend de l'ouverture jusqu'à une distance de $D^2/2\lambda$, le champ qui se propage à la structure d'une onde plane et reste concentré dans la direction normale à l'ouverture,
- la zone de Fresnel qui s'étend jusqu'à $2D^2/\lambda$ où le faisceau commence à diverger et finira par se transformer d'une onde plane en une onde sphérique,
- la zone de Fraunhofer ou zone lointaine, à partir de laquelle la propagation de l'onde s'effectue par ondes sphériques uniquement. On constate qu'à partir de là, la mesure du diagramme de rayonnement n'est pas modifiée quelle que soit la distance. De ce fait, la mesure des caractéristiques de rayonnement devra être effectuée à l'intérieur de cette zone. La densité de puissance décroît alors en $1/R^2$, où R représente la distance entre le centre de phase de l'antenne, où sont centrées les ondes sphériques, et le point d'observation.

10.3 Diagramme de rayonnement et directivité

La connaissance du diagramme de rayonnement est une des principales caractéristiques des antennes car il caractérise la distribution spatiale du champ électromagnétique rayonné. Une représentation graphique est associée à cette distribution, généralement normalisée à 0 dB. Bien qu'il soit possible de mesurer le champ et d'en déduire la fonction caractéristique de rayonnement, c'est bien souvent à partir de la mesure de la puissance qu'est construit le diagramme de rayonnement. Par définition, une représentation tridimensionnelle illustre la répartition spatiale de la puissance dans toutes les directions d'observation c'est-à-dire pour toutes les valeurs de θ et ϕ sachant que, pour chaque position, les deux composantes orthogonales E_θ et E_ϕ sont mesurées (figure 10.1).

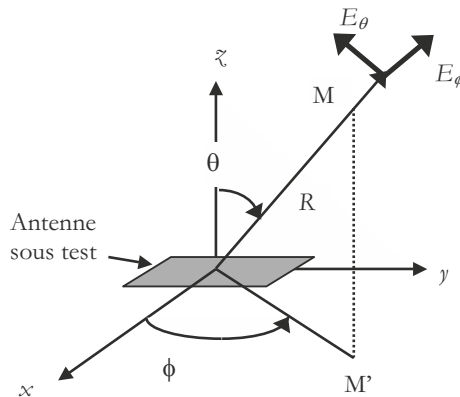


Figure 10.1 – Système de coordonnées utilisées pour la mesure du diagramme de rayonnement.

D'un point de vue pratique, la mesure du diagramme de rayonnement est bien souvent effectuée dans deux plans principaux orthogonaux qui sont les plans $\phi = 0^\circ$ et 90° . On parle aussi de plan E et plan H . La polarisation du champ émis par l'antenne n'étant *a priori* pas connue, il faudra, en plus de la mesure du rapport axial qui donne une information sur la qualité de la polarisation émise, mesurer les quatre diagrammes qui sont :

- $E_\theta(\phi = 0^\circ, \theta)$: diagramme de la composante du champ électrique E_θ dans le plan $\phi = 0^\circ$,
- $E_\theta(\phi = 90^\circ, \theta)$: diagramme de la composante du champ électrique E_θ dans le plan $\phi = 90^\circ$,
- $E_\phi(\phi = 0^\circ, \theta)$: diagramme de la composante du champ électrique E_ϕ dans le plan $\phi = 0^\circ$,
- $E_\phi(\phi = 90^\circ, \theta)$: diagramme de la composante du champ électrique E_ϕ dans le plan $\phi = 90^\circ$.

Notons enfin que la mesure la plus représentative du comportement de l'antenne devrait être effectuée en toute logique lorsque celle-ci est installée dans son environnement naturel. Bien souvent, les mesures des caractéristiques de rayonnement seront effectuées à l'intérieur d'un espace fermé et isolé des perturbations extérieures qui peuvent entacher les résultats. De plus, nous avons vu que la distance minimale à observer entre la source et l'antenne sous test, pour se situer en zone de champ lointain, est intimement liée à la fréquence de fonctionnement et aux dimensions géométriques de l'antenne. On appréhende aisément la contrainte que représente cette distance, notamment lors de mesures en basses fréquences.

Nous allons présenter dans ce qui suit les principales techniques de mesures du rayonnement électromagnétique d'antennes effectuées à l'intérieur d'une chambre anéchoïque, et les différentes variantes associées qui permettent de reproduire la zone de champ lointain tout en s'affranchissant de la contrainte sur la distance entre source et antennes sous test.

10.3.1 Chambre anéchoïque

La plupart des mesures d'antennes s'effectuent dans une chambre anéchoïque (figure 10.2). Il s'agit d'un espace fermé entièrement tapissé d'absorbants sous lesquels est disposé un feuillard métallique d'aluminium pour se rapprocher des caractéristiques d'une cage de Faraday et ainsi, se prémunir des agressions électromagnétiques extérieures. Afin de limiter les réflexions parasites dues à la source sur les parois de la chambre, l'utilisation de mousses absorbantes en polyuréthane chargées de particules de carbone est souvent retenue pour absorber les ondes incidentes aux parois. Ces mousses, caractérisées par leur réflectivité, peuvent prendre des formes variées, cependant la forme pyramidale est très souvent utilisée. En effet, l'idée consiste, grâce à cette forme particulière, à passer de façon graduelle, d'une impédance d'onde au sommet de la pyramide égale à 377Ω à celle d'un milieu dissipatif à la base de la pyramide. Notons que les performances des absorbants seront d'autant meilleures que la fréquence est élevée. La hauteur des pyramides est proportionnelle à la longueur d'onde de travail et, par conséquent, inversement proportionnelle à la fréquence.

L'utilisation de chambres de formes adaptées (pyramidale du côté de la source puis carrée autour de l'antenne sous test) peut constituer une alternative intéressante. De même, l'utilisation d'absorbants plus courts (un peu moins performants) peut être envisagée dans les zones où le niveau de champ est faible (à l'arrière de l'émetteur par exemple). Précisons que les performances des absorbants se dégradent sensiblement lorsque l'on s'écarte de la normale à la paroi. Ceci est d'autant plus vrai pour les absorbants pyramidaux dont les performances sont optimales en incidence normale. Pour atténuer ces effets, l'utilisation d'absorbant de forme arrondie à l'extrémité peut être retenue. Ici aussi, la forme de la chambre peut constituer une alternative pour diminuer les réflexions parasites. Enfin, les absorbants d'utilisation courante dans les chambres se présentent sous la forme de panneaux carrés de 61 cm de côté et supportent une densité de puissance incidente de l'ordre de $0,3 \text{ W/cm}^2$. Un absorbant de forme pyramidale de hauteur 12 cm présente une atténuation de -40 dB à 10 GHz pour une incidence normale. Cependant, une perte d'atténuation de 5 dB pour un angle d'incidence de 35° est courant pour ce type d'absorbant.

Notons enfin que les absorbants plans et souples peuvent être utilisés pour masquer les mâts ou les plateaux tournants qui supportent les antennes. Leurs performances en réflectivité sont limitées autour de -10 dB .

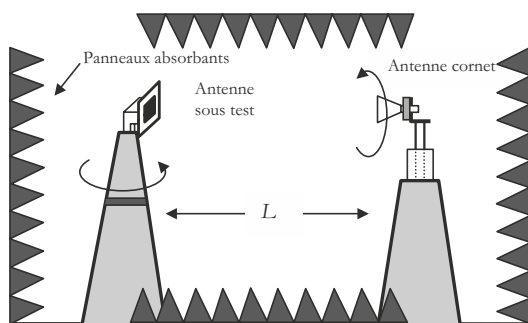


Figure 10.2 – Chambre anéchoïque pour la mesure d'antennes

À l'intérieur de la chambre, une antenne source de référence (bien souvent une antenne cornet directive de façon à limiter au mieux les réflexions sur les parois) est alignée avec l'antenne sous test qui sera judicieusement placée dans un espace où la mesure peut être effectuée sans perturbations avec un minimum de réflexions. Nous définirons un peu plus loin cet espace,

fondamental, communément appelé zone calme ou tranquille. Le principe de réciprocité permet de choisir indifféremment la source (soit l'antenne de référence, soit l'antenne sous test).

Les positionneurs (ou plateaux tournants) employés doivent permettre les mesures en élévation et en azimut. Deux axes de rotation sont donc *a priori* nécessaires au niveau de l'antenne sous test en réception. La vitesse de rotation (généralement pas à pas) sera choisie suffisamment basse de façon à éviter les distorsions sur les diagrammes notamment sur les antennes massives où des effets d'inertie peuvent apparaître. D'un point de vue pratique, une solution consiste à n'utiliser que deux positionneurs à un seul axe de rotation chacun. En effet, le premier positionneur autorise le mouvement en rotation de l'antenne sous test selon son axe vertical tandis que le second permet à l'antenne d'émission (ici le cornet) de tourner selon son axe horizontal (figure 10.2). Partant d'une position de référence pour laquelle la polarisation des deux antennes est la même, il suffit, pour mesurer la composante principale $E_\theta(\theta, \phi = 0^\circ \text{ ou } 90^\circ)$, de maintenir fixe l'antenne cornet à polarisation verticale et de faire pivoter l'antenne sous test sur son axe vertical de façon à relever le niveau de puissance reçue en fonction de l'angle d'élévation θ . La composante croisée $E_\phi(\theta, \phi = 0^\circ \text{ ou } 90^\circ)$ sera obtenue en effectuant la même mesure mais en ayant au préalable fait pivoter le cornet à 90° . La détermination des autres diagrammes s'effectue en pivotant l'antenne sous test à 90° par rapport à sa position d'origine. La multitude de relevés imposent une automatisation de la procédure de mesure. L'ensemble est presque systématiquement asservi par un dispositif électronique et informatique qui permet de contrôler le fonctionnement des positionneurs mais également d'assurer l'acquisition, le stockage et éventuellement le traitement et la visualisation des données recueillies.

Pour des mesures à des fréquences inférieures aux ondes VHF où les effets dus aux réflexions sur le sol sont difficiles à éviter, on utilise des chambres qui sont dépourvues d'absorbants au sol. Dans ce cas, le sol est constitué d'une surface plane parfaitement réfléchissante de très faible rugosité et l'on utilise la réflexion pour créer des interférences constructives au niveau de l'antenne sous test dans la zone tranquille. On ajustera les hauteurs respectives des antennes d'émission et de réception de façon à faire coïncider l'antenne sous test avec le premier lobe d'interférence entre les rayons direct et réfléchi. Ainsi, pour une hauteur H_e de la source par rapport au sol et une distance L entre source et antenne sous test fixées, la hauteur H_r entre récepteur et sol devra satisfaire la relation suivante :

$$H_r = \frac{\lambda L}{4 H_e}$$

où λ représente la longueur d'onde de travail.

Cette approche est également utilisée lorsque les mesures sont réalisées en espace libre sur des antennes de grandes dimensions ou lorsque la distance L est trop importante pour des mesures en chambre. Dans cet espace dépourvu d'absorbants, on s'assurera de l'absence d'obstacles dans le champ de mesure de façon à ne considérer que la réflexion au sol. Précisons qu'aux longueurs d'ondes millimétriques, le sol n'a que peu d'influence sur les mesures.

10.3.2 Base compacte

L'obtention d'un champ parfaitement planaire, ayant les caractéristiques d'une onde plane dans la zone tranquille, est obtenue en éloignant suffisamment les deux antennes qui se font face. Selon la fréquence utilisée et l'ouverture de l'antenne, cette distance pouvait devenir contraignante. C'est la raison pour laquelle ont été développées les chambres anéchoïques compactes dans lesquelles on utilise des réflecteurs pour générer un champ électromagnétique planaire de bonne qualité dont la répartition transverse est constante en amplitude et en phase. Ceci est réalisé dans un volume suffisamment grand de façon à optimiser l'espace de mesure pour l'antenne même lorsque celle-ci est en mouvement. Les paraboles sont souvent utilisées mais nous verrons plus loin que l'utilisation de lentilles diélectriques constitue une intéressante alternative.

La figure 10.3 représente le principe d'une chambre compacte réalisée à l'aide d'un réflecteur unique.

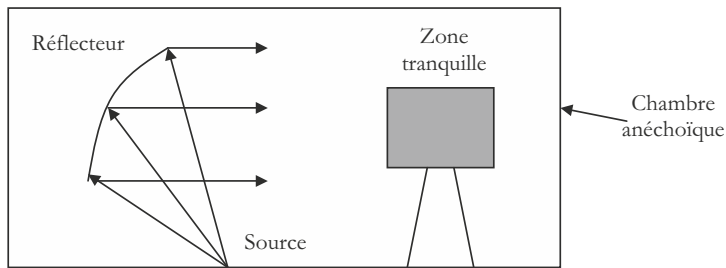


Figure 10.3 – Représentation d'une chambre anéchoïque compacte à réflecteur unique

On y trouve la source qui peut être constituée d'une ou plusieurs antennes de façon à couvrir plusieurs bandes de fréquences. L'onde issue de la source se propage par ondes sphériques pour être transformée *via* le réflecteur en onde localement plane. La qualité du champ à l'intérieur de la zone tranquille est quantifiée à travers les grandeurs que sont l'ondulation et l'apodisation qui traduisent la variation de l'amplitude et de la phase autour d'une valeur médiane. Nous reparlerons de ces paramètres un peu plus loin dans ce chapitre. L'espace qui constitue la zone calme (ou tranquille) sera d'autant plus grand que les dimensions du réflecteur sont importantes et la qualité de ce dernier (état de surface, rugosité...) influera sensiblement sur la nature de l'onde dans la zone calme.

De plus, du fait de la présence d'une source excentrée, il est nécessaire de réduire l'illumination sur les bords du réflecteur de façon à limiter la détérioration du champ dans la zone calme et l'augmentation de la composante croisée. Diverses solutions sont alors proposées :

- l'utilisation de réflecteurs conformés à bords roulés ou a serrations (ajout d'une structure dentée sur les bords du réflecteur) plutôt que des réflecteurs cylindriques, plus faciles à fabriquer mais présentant de moins bonnes performances au niveau de la composante croisée. De cette façon, le champ est diffracté sur les bords vers les absorbants au détriment cependant d'une perte de puissance vers la zone calme ;
- l'utilisation d'un système à double (ou à triple) réflecteurs car bien que plus complexe à mettre en œuvre, ils permettent de limiter la remontée du niveau de la composante croisée du champ électromagnétique qui doit rester inférieure à -40 dB pour des mesures fiables ;
- l'utilisation de lentilles diélectriques pour collimater le faisceau et transformer une onde sphérique en onde plane permet de s'affranchir du réflecteur et simplifie la conception du dispositif (figure 10.4).

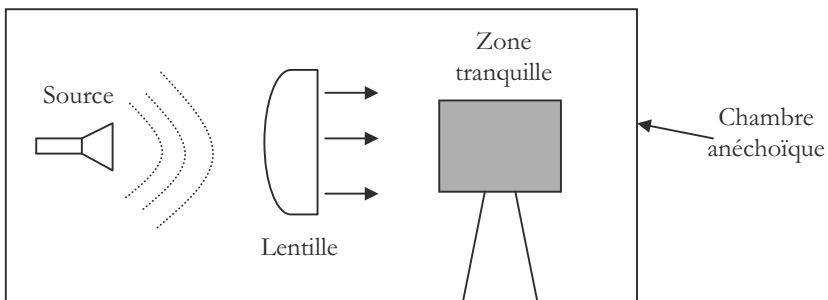


Figure 10.4 – Principe d'une chambre anéchoïque compacte à lentille diélectrique

Les matériaux diélectriques couramment utilisés sont le téflon, le polyéthylène ou la réxolite. Ils présentent des constantes diélectriques faibles ce qui limite les réflexions parasites à la surface de la lentille. Il est possible aussi de tailler les bords de la lentille de la même façon que les réflecteurs pour atténuer les effets du champ diffracté mais là aussi, l'état de surface du collimateur est fondamental et le coût de fabrication augmente sensiblement lorsque l'on souhaite diminuer la rugosité de l'élément.

10.3.3 Mesure en champ proche

Si la distance entre antennes d'émission et de réception est trop faible, il est possible de réaliser les mesures en champ proche, c'est-à-dire à proximité de l'antenne sous test. Cette approche s'appuie sur le fait que, le champ lointain est déduit du champ proche mesuré dans l'ouverture de l'antenne par Transformée de Fourier. En général, l'antenne sous test est immobile et parfaitement ajustée par rapport à la sonde qui se déplace par balayage successif le long de l'élément rayonnant en décrivant un plan plus large que l'ouverture géométrique de l'antenne. Le pas d'échantillonnage est généralement inférieur à $\lambda_0/2$ de façon à éviter le repliement du spectre. La technique de mesure en champ proche planaire présente cependant certaines contraintes car la mesure du rayonnement arrière, des lobes secondaires ou des effets dus à la diffraction sur les bords d'une antenne ne peuvent s'effectuer sans modifier la position de l'antenne. L'influence de la sonde doit être prise en compte dans la mesure du diagramme. Particulièrement, sa pureté de polarisation est un élément fondamental. Il s'agit d'une procédure bien adaptée à la mesure des antennes à ouverture (cornet, parabole...) où le rayonnement est principalement concentré dans une direction. Pour tenir compte du rayonnement dans toutes les directions, des mesures cylindriques et sphériques peuvent être effectuées. Dans ce dernier cas, le plus utilisé, la sonde est cette fois-ci immobile et l'antenne sous test se déplace par rotation de façon à décrire les plans en azimut et en élévation. Cette approche, plus rigoureuse, nécessite cependant un traitement mathématique plus complexe et un nombre de points de mesures considérablement plus important.

10.4 Conditions sur la mesure du diagramme de rayonnement

Nous avons vu plus haut qu'une distance L de $2D^2/\lambda$ doit être observée pour considérer une propagation par onde sphérique. À titre d'exemple, si l'on considère une antenne caractérisée par une ouverture de 0,5 m et une fréquence de fonctionnement de 3 GHz, la condition impose d'éloigner les deux éléments rayonnants de 5 mètres au moins. Si la fréquence est de 10 GHz, la distance passe à 16,6 mètres. Il devient alors difficile d'effectuer les mesures dans un espace « standard » de la dimension d'une pièce. La mesure en espace libre constitue une alternative.

$L = 2D^2/\lambda$ représente également la distance minimale qui permet de quantifier l'erreur de phase au niveau de l'antenne sous test (figure 10.5). Pour cet écartement et en considérant $\delta = \lambda/16$, l'erreur de phase est de $22,5^\circ$. L'augmentation de l'écartement L va réduire l'erreur de phase. On constate alors que la mesure d'antennes de grande ouverture à des fréquences élevées impose des contraintes fortes sur l'écartement L .

La condition dite de Fraunhofer est une première étape dans l'élaboration du dispositif de mesures. À l'intérieur d'une chambre, on définit la zone tranquille (ou zone calme) comme étant la zone où les mesures peuvent être effectuées avec un minimum d'erreurs dues aux réflexions parasites. Dans cet espace localisé, le champ a localement la structure d'une onde plane et les perturbations dues aux réflexions parasites sur les absorbants ont une influence limitée. Deux grandeurs vont s'appliquer : l'une à l'amplitude et l'autre à la phase. Pour l'amplitude, on introduira les termes d'apodisation et d'ondulation. Pour la phase, on ne parlera que d'ondulation. Sur

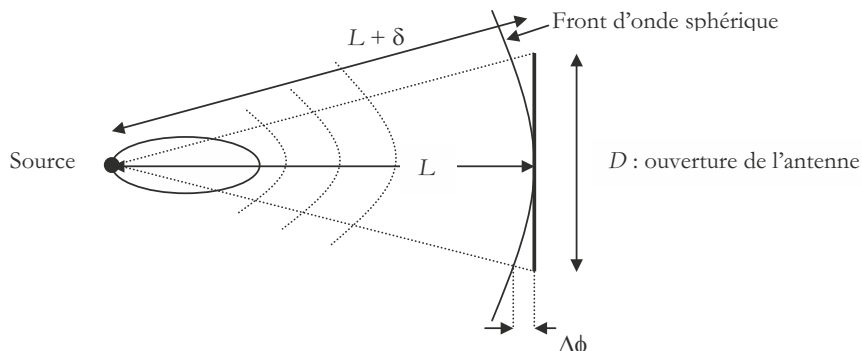


Figure 10.5 – Illustration de l'erreur de phase observée au niveau de l'antenne sous test.

la figure 10.6, on constate que l'ondulation traduit une variation autour d'une valeur moyenne, l'apodisation représente la variation de la valeur moyenne tolérée pour la définition de la zone tranquille. Les valeurs couramment utilisées pour l'ondulation sont de $\pm 0,5$ dB pour l'amplitude et $\pm 5^\circ$ pour la phase. À l'intérieur de cette zone, on peut considérer le champ comme étant quasi-uniforme.

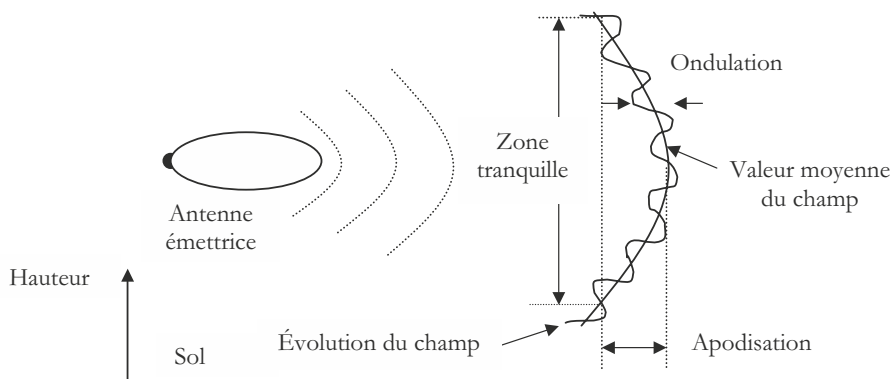


Figure 10.6 – Définition de la zone tranquille.

Malheureusement, il existe d'autres sources d'erreurs dont il faudra également avoir connaissance. En particulier, en basses fréquences où les distances entre antennes doivent être importantes, il existe un risque de couplage avec le champ proche réactif.

La nécessité d'un parfait alignement entre émetteur et récepteur, notamment lorsqu'on effectue des relevés de composantes croisées, impose une contrainte forte sur les positions relatives des deux éléments.

La prise en compte de l'atténuation atmosphérique peut être nécessaire à des fréquences supérieures à 15 GHz où la modification de l'indice de réfraction peut provoquer des variations d'amplitudes dans les mesures.

La qualité et le blindage des câbles de liaison sont importants dans la mesure où la présence d'un courant à la périphérie d'un conducteur provoquera un rayonnement parasite.

Enfin, une mauvaise adaptation d'impédance entre les instruments de mesure et les antennes peut se traduire par des erreurs dans la détermination des grandeurs caractéristiques telles que le gain ou la directivité.

10.5 Gain

Le gain d'une antenne est l'un des paramètres fondamentaux qui caractérisent un élément rayonnant. Par définition, le gain dans une direction est le rapport de la densité de puissance rayonnée par l'antenne dans cette direction sur la puissance rayonnée par la source isotrope équivalente. Au même titre que le diagramme de rayonnement, le gain peut être mesuré à l'intérieur d'une chambre anéchoïque, dont nous avons parlé un peu plus haut, pour des fréquences généralement supérieures à 1 GHz. En deçà, les contraintes sur l'espacement entre antennes imposent d'effectuer ces mesures plutôt en espace libre et les ouvertures des antennes source de références deviennent trop importantes pour négliger les réflexions au sol.

La mesure du gain nécessite souvent l'emploi d'antennes de référence étalon. Les deux antennes les plus utilisées sont l'antenne dipôle résonnante $\lambda/2$ et l'antenne cornet, toutes deux à polarisation linéaire. La première présente l'avantage d'une grande pureté de polarisation mais n'est que peu directive, ce qui accentuera les réflexions parasites. La seconde est plus directive mais peut présenter un niveau de polarisation croisée plus élevé. Les gains sont respectivement de 2,1 dB pour le dipôle et s'étendent de 15 à 25 dB pour le cornet, selon l'ouverture géométrique.

Il existe deux méthodes pour la mesure du gain en champ lointain : la mesure absolue du gain et la mesure par comparaison.

10.5.1 Mesure absolue du gain

Cette procédure s'appuie sur la formule de Friss. Elle donne la puissance reçue lors de la transmission d'une onde en espace libre *via* deux antennes de gains distincts G_e (en émission) et G_r (en réception), espacées d'une distance R . Littéralement, la puissance reçue par le récepteur s'exprime de la façon suivante :

$$P_r = P_e (G_e G_r) \left(\frac{\lambda}{4 \pi R} \right)^2$$

P_e et P_r représentent respectivement les puissances émises et reçue (en watt), λ étant la longueur d'onde en espace libre. Cette formulation a été établie en supposant l'antenne de réception placée en champ lointain, les antennes parfaitement adaptées et alignées en polarisation. Si les antennes d'émission et de réception sont identiques, alors le gain recherché s'exprime par :

$$G_{e|\text{dB}} = G_{r|\text{dB}} = \frac{1}{2} \left(10 \text{ Log } \frac{P_r}{P_e} + 10 \text{ Log } \left(\frac{\lambda}{4 \pi R} \right)^2 \right)$$

où P_e , P_r , R et λ sont les grandeurs mesurées.

Cette première approche nécessite l'usage de deux antennes identiques. Dans l'hypothèse où l'on ne dispose pas de deux antennes identiques, ce qui est une situation courante, il est nécessaire dans ce cas de disposer de trois antennes distinctes et par conséquent de déterminer le gain de l'antenne par trois mesures associant l'ensemble des trois antennes de gain G_1 , G_2 et G_3 :

$$G_{i|\text{dB}} + G_{j|\text{dB}} = 20 \text{ Log } \left(\frac{4 \pi R}{\lambda} \right) + 10 \text{ Log } \left(\frac{P_{ri}}{P_{ej}} \right) \quad \text{avec } i, j = 1, \dots, 3 \quad \text{et } i \neq j$$

Trois mesures sont réalisées par permutation, de façon à obtenir un système de trois équations à trois inconnues et déterminer ainsi le gain de chaque antenne.

10.5.2 Mesure du gain par comparaison

Une autre méthode couramment employée utilise une seule antenne de gain parfaitement connu. La détermination du gain est obtenue par deux mesures successives.

La première étape consiste à mesurer la puissance reçue à la réception avec l'antenne sous test. La seconde est effectuée en permutant l'antenne sous test avec le cornet de gain connu. On obtient ainsi deux puissances P_{test} et $P_{\text{réf}}$, l'environnement de mesure n'étant pas modifié entre les deux configurations. On en déduit le gain de l'antenne sous test en appliquant la formule ci-dessous :

$$G_{\text{dB}} = G_{\text{réf|dB}} + 10 \text{ Log} \left(\frac{P_{\text{test}}}{P_{\text{réf}}} \right)$$

L'inconvénient d'une telle méthode est l'erreur qui peut apparaître lorsque le gain de l'antenne sous test est proche de celui de l'antenne de référence. Dans ce cas, l'influence des réflexions parasites devient sensible.

Dans le cas d'une antenne à mesurer dont la polarisation est circulaire, on déduit le gain de l'antenne de deux façons :

- Soit l'on dispose de deux antennes de référence dont l'une est à polarisation circulaire droite et l'autre gauche. Le gain s'exprime à partir de la mesure de chacune des composantes.
- Soit l'on ne dispose que d'une antenne à polarisation linéaire. On utilise le fait qu'une onde polarisée circulairement peut être décomposée en deux ondes polarisées linéairement chacune. Le gain est alors obtenu à partir de la mesure des deux gains distincts. L'un (G_{testH}), lorsque l'antenne de référence est positionnée de façon à émettre une onde polarisée horizontalement, et l'autre (G_{testV}), lorsqu'elle émet une onde à polarisation verticale. Le gain global s'exprime alors par :

$$G_{\text{dB}} = G_{\text{testH|dB}} + G_{\text{testV|dB}}$$

On appréhende mieux dans cette configuration l'intérêt qu'il y a à utiliser des antennes de référence à très faible niveau de polarisation croisée de façon à ne pas trop perturber la mesure de la composante principale.

10.5.3 Mesure du gain à partir du champ proche

Le calcul du gain à partir de la mesure du champ proche est obtenu à l'aide de la relation de Bracewell. Elle s'exprime de la façon suivante :

$$G = \frac{4 \pi A}{\lambda^2} \frac{1}{\frac{1}{A} \iint_A \left[\frac{E(x,y)}{E_{\text{moy}}} \right] \left[\frac{E(x,y)}{E_{\text{moy}}} \right]^* dx dy}$$

où * représente le complexe conjugué, $E(x,y)$ le champ électrique mesuré dans l'ouverture (en V/m), A l'ouverture géométrique sur laquelle est prélevé le champ (en m²) et E_{moy} , le champ moyen dans l'ouverture défini de la façon suivante :

$$E_{\text{moy}} = \frac{1}{A} \iint_A E(x,y) dx dy$$

10.5.4 Mesure de la directivité

La directivité D d'une antenne peut être déduite de la détermination du diagramme de rayonnement par la relation suivante :

$$D = \frac{4 \pi}{\iint_{\Omega} f(\theta, \phi) \sin \theta \, d\theta \, d\phi}$$

où $f(\theta, \phi)$ représente le diagramme de rayonnement normalisé et a pour expression :

$$f(\theta, \phi) = \frac{(E_{\theta}^2 + E_{\phi}^2)}{(E_{\theta}^2 + E_{\phi}^2)_{MAX}}$$

La directivité de l'antenne est donc déterminée en intégrant sur la sphère le diagramme de rayonnement en puissance normalisé. Il ne prend donc pas en compte les pertes et l'éventuelle désadaptation de l'antenne. Rappelons que le rapport entre gain et directivité représente l'efficacité η (ou rendement) de l'antenne dont nous reparlerons un peu plus loin dans ce chapitre.

Dans le cas des antennes directives en présence d'un lobe principal prépondérant, on peut approcher la directivité à partir de la mesure de la largeur du lobe principal de la façon suivante :

$$D = \frac{41253}{\Delta_{\theta^{\circ}} \Delta_{\phi^{\circ}}}$$

Dans cette expression $\Delta_{\theta^{\circ}}$ $\Delta_{\phi^{\circ}}$ représentent respectivement les angles d'ouverture (en degrés) à mi-puissance dans deux plans orthogonaux.

10.6 Polarisation

La polarisation correspond à l'orientation d'un vecteur de champ électrique dans le plan orthogonal par rapport à la direction de la propagation. Si le vecteur de champ électrique est toujours orienté dans la même direction, l'onde est polarisée linéairement. Si le vecteur de champ électrique tourne autour de la direction de la propagation, l'onde est polarisée circulairement ou plus généralement elliptiquement.

La connaissance du diagramme de rayonnement n'est pas suffisante pour quantifier la nature de l'onde émise par l'antenne. Pour caractériser l'onde émise, on détermine le rapport axial (RA), le sens de rotation (droite ou gauche) et l'angle d'inclinaison de l'ellipse τ . On trouvera dans la littérature citée en référence l'ensemble des formulations qui permettent de caractériser l'ellipse de polarisation. Précisons ici que la nature de la polarisation n'est pas uniforme dans toute la sphère qui entoure l'antenne. On peut la déterminer dans la direction normale à l'antenne mais, dans certaines applications, sa connaissance est utile lorsque l'on s'écarte de la normale et l'on peut définir un angle d'ouverture à l'intérieur duquel la qualité de la polarisation pour laquelle a été conçue l'antenne reste correcte.

Le rapport axial est défini comme étant le rapport du grand axe sur le petit axe de l'ellipse de polarisation (figure 10.7). Il donne une indication fondamentale sur la polarisation de l'onde qui se propage. Pour une polarisation circulaire de bonne qualité, on fixera un seuil de rapport axial à 2 dB par exemple. On peut déterminer expérimentalement ce paramètre de deux façons :

- en utilisant une antenne à polarisation linéaire,
- à l'aide de deux antennes à polarisation circulaire (droite et gauche).

Dans le cas d'une antenne à polarisation linéaire, l'antenne sous test est fixe et l'on fait pivoter l'antenne de référence autour de son axe horizontal. On décrit ainsi l'ellipse de polarisation d'où

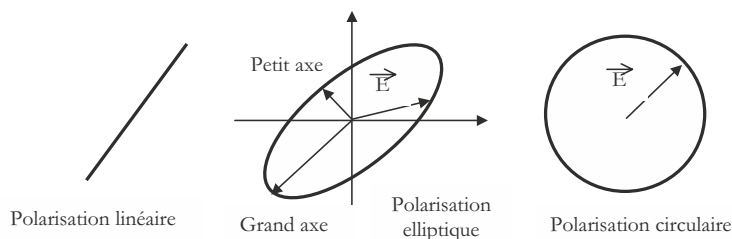


Figure 10.7 – Différents états de polarisation d'une onde électromagnétique.

l'on déduit le rapport axial expérimental. L'angle d'inclinaison τ est directement déduit du tracé résultant de la mesure. Ici aussi, l'antenne de référence à polarisation linéaire (bien souvent un cornet) doit présenter un niveau de composante croisée le plus faible possible afin de limiter les erreurs sur la mesure du rapport axial. On constate que cette approche, fiable et simple, peut nécessiter un temps de mesure long si l'on souhaite caractériser l'antenne sur la totalité de la sphère.

Pour minimiser le temps de mesure, on évite la rotation complète de l'antenne de référence. On ne recherche dans ce cas, que la valeur maximale du champ (qui correspond au grand axe de l'ellipse) et la valeur minimale (petit axe) à 90° du grand axe.

De la même façon, connaissant le centre de rotation de l'antenne, on peut reconstituer mathématiquement l'ellipse de polarisation à partir de trois mesures de champs qui correspondent à trois orientations distinctes de l'antenne de référence.

Dans la seconde configuration, l'antenne sous test est toujours fixe, et l'on utilise deux antennes de référence à polarisation circulaire droite et gauche positionnées relativement proche l'une de l'autre. La procédure consiste à mesurer successivement les niveaux de champ (E_{droite} et E_{gauche}) reçus sur ces deux antennes et à déterminer le rapport axial de la façon suivante :

$$RA = \frac{E_{\text{droite}} + E_{\text{gauche}}}{E_{\text{droite}} - E_{\text{gauche}}}$$

Cette méthode est moins gourmande en durée de mesure. De plus, selon le signe du rapport axial, on peut déduire le sens de la polarisation de l'onde. En effet, une valeur positive du RA correspondra à une onde polarisée elliptiquement à droite tandis que, dans le cas contraire, il s'agira d'une onde polarisée elliptiquement à gauche.

Enfin, précisons que les antennes de référence qui présentent une très bonne qualité de polarisation circulaire sont les antennes filaires hélicoïdales. Celles-ci sont d'autant plus performantes (rapport axial proche de l'unité) que le nombre de tours qui les constituent est important.

10.7 Impédance

L'impédance d'entrée constitue également une autre grandeur fondamentale de l'antenne. Celle-ci varie avec la fréquence. La mesure se ramène à un problème de ligne de transmission et la détermination du coefficient de réflexion S_{11} permet de déduire l'impédance d'entrée de l'antenne. Celle-ci sera mesurée en chambre anéchoïque ou en espace libre en l'absence d'obstacles. Cependant il peut être là aussi intéressant d'effectuer la mesure lorsque l'antenne est dans son environnement. La détermination expérimentale du coefficient de réflexion peut être effectuée à l'aide d'un dispositif à ligne fendue à l'intérieur de laquelle se déplace une sonde de mesure. La procédure consiste à localiser avec précision sur cette ligne, en déplaçant la sonde, la position et l'amplitude des ventres et des nœuds de tension de façon à en déduire à la fois le module et

la phase du coefficient de réflexion à partir de la mesure du rapport d'onde stationnaire (ROS). Cette mesure étant répétée pour chaque fréquence qui constitue la bande d'étude. Cependant, un appareil spécifique est aujourd'hui très largement employé pour cette mesure, il s'agit de l'analyseur vectoriel de réseau (VNA). Il fournit en module et phase, les valeurs des coefficients qui constituent la matrice de répartition $[S]$ d'un quadripôle dans des bandes de fréquences s'étendant de 45 MHz à 120 GHz. Sans entrer dans le détail du fonctionnement complexe d'un analyseur, nous indiquerons simplement qu'il est constitué d'amplificateurs, de mélangeurs, de détecteurs, d'oscillateurs, de diviseurs, d'isolateurs et de coupleurs pour n'en citer que les principaux éléments. Un coupleur est un composant passif à quatre accès qui permet de séparer les ondes incidentes et réfléchies sur une ligne. Il constitue donc un élément capital pour la détermination des paramètres de la matrice $[S]$. Cependant, un coupleur n'est pas idéal et un certain nombre d'erreurs lui sont imputables compte tenu de ses imperfections sur sa directivité, son couplage et notamment son isolation, trois paramètres qui le caractérisent. La figure 10.8 illustre le principe d'un réflectomètre à deux coupleurs utilisé dans l'analyseur pour la détermination vectorielle du coefficient de réflexion S_{11} d'une antenne.

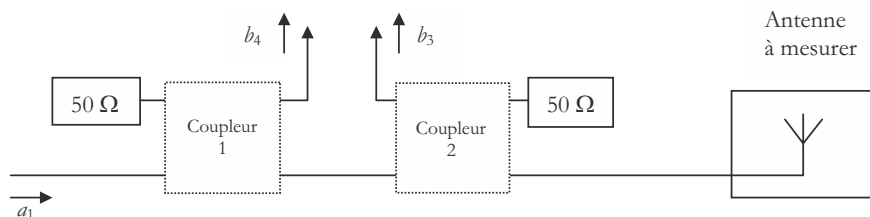


Figure 10.8 – Principe du réflectomètre à deux coupleurs

Dans cette illustration où l'on suppose que les lignes d'accès sont d'impédances caractéristiques $50\ \Omega$ (deux accès ont des charges adaptées à la ligne), on montre que le coefficient de réflexion s'exprime par :

$$S_{11} = \frac{1}{\gamma} \frac{b_3 \beta - b_4 \delta}{b_4 \beta - b_3 \delta}$$

où γ et β sont les paramètres intrinsèques au coupleur et δ représente un défaut d'isolation du composant. À cela, il faudra ajouter les erreurs dues aux autres composants mais également à la désadaptation, à la réponse en fréquence qui caractérise les différences d'amplitude et de phase entre les ports d'accès de l'analyseur.

Ces sources d'erreurs, si elles sont préalablement connues avec précision, peuvent être judicieusement compensées par un calcul élaboré.

Cette opération est réalisée en effectuant sur l'analyseur une calibration, une fois la bande de fréquences d'étude définie. Cette calibration a pour but de déterminer l'ensemble des termes d'erreur (au nombre de 12 dans un analyseur vectoriel à deux ports) dus aux imperfections des réflectomètres. Elle est effectuée à l'aide de charges étalons parfaitement calibrées que sont la charge adaptée, le circuit ouvert et le court-circuit. La totalité formant un ensemble (kit de calibration) généralement fourni avec l'appareil et proposé suivant différentes connectiques utilisées dans la gamme des hyperfréquences. Il est à noter que cette procédure doit être effectuée avec le plus grand soin car une manipulation hasardeuse associée à des étalons de piètres qualités ou non entretenus introduiront nécessairement des erreurs sur les coefficients de la matrice $[S]$, ce dont l'opérateur ne pourra s'apercevoir lors de la mesure.

10.8 Efficacité

La détermination de l'efficacité constitue l'une des étapes fondamentales pour quantifier les performances d'une antenne. Ce paramètre indique la quantité de puissance qui sera réellement rayonnée par l'antenne par rapport à la puissance fournie à cette même antenne. Le rapport du gain de l'antenne sur la directivité pour une direction donnée donne également l'efficacité dans cette direction. Selon la définition choisie pour le gain, on peut inclure dans les pertes de l'antenne à la fois les pertes dans les parties métallique et diélectrique mais également les pertes par désadaptation entre l'antenne et sa ligne d'alimentation. On comprend bien qu'ici, il est bien souvent indispensable de prendre en compte l'environnement réel de l'antenne (circuit d'adaptation, radôme ou la proximité d'un usager dans le cas d'antennes pour mobiles) car celui-ci aura un impact direct sur l'efficacité.

Les principales méthodes pour déterminer l'efficacité d'une antenne sont au nombre de trois. On distingue :

- la méthode de directivité/gain,
- la méthode radiométrique,
- la méthode de Wheeler cap.

La méthode de directivité/gain comme son nom l'indique nécessite la connaissance du gain G et de la directivité D de l'antenne. L'efficacité est déduite de ces mesures par la relation :

$$\eta_{\text{ray}} = \frac{G}{D}$$

Si dans le cas d'antennes directives on peut déduire la directivité de façon approchée à partir des angles d'ouvertures à -3 dB, dans le cas général, cette mesure nécessite la détermination de la puissance rayonnée par l'antenne en intégrant les composantes tangentielles du champ électrique autour de la sphère entourant l'antenne. On appréhende aisément la complexité de cette mesure tridimensionnelle en chambre anéchoïque. Les incertitudes de mesures seront d'autant plus importantes que la directivité de l'antenne sous test sera faible compte tenu des réflexions parasites possibles sur les parois de la chambre et les effets dus aux interférences entre le câble d'alimentation et l'antenne.

La méthode radiométrique consiste à évaluer le rendement de l'antenne sous test à partir de mesures successives de la puissance de bruit captée par cette antenne et une antenne de référence à faibles pertes de caractéristiques connues lorsque celles-ci sont placées dans des environnements caractérisés par des températures de bruit parfaitement identifiées. Les environnements communément choisis sont la chambre anéchoïque ($T = 290$ K) et le ciel ($T = 10-50$ K). Il s'agit cependant d'une procédure délicate à mettre en œuvre car elle nécessite l'utilisation de matériel spécifique (radiomètre à faible bruit en réception, réflecteur pour limiter l'influence des émissions venues du sol pour la mesure de l'antenne sous test vers le ciel...).

La méthode de Wheeler cap s'inspire de la définition de l'efficacité de rayonnement η_{ray} exprimée à partir des résistances de pertes R_{pertes} et de rayonnement R_{ray} . Elle a pour expression :

$$\eta_{\text{ray}} = \frac{R_{\text{ray}}}{R_{\text{pertes}} + R_{\text{ray}}}$$

La connaissance de ces deux résistances permet donc de déterminer l'efficacité de rayonnement. Malheureusement, une mesure en espace libre de l'impédance d'entrée de l'antenne sur l'analyseur de réseau vectoriel ne permet pas d'isoler de façon simple ces deux termes. J. Wheeler a proposé d'isoler la résistance de pertes de cette mesure en plaçant l'antenne sous test à l'intérieur d'une cavité fermée parfaitement hermétique et de dimensions ajustées. À l'origine, une sphère blindée de rayon $\lambda/2\pi$ permettrait d'annuler la résistance de rayonnement et d'en déduire la résistance

de pertes. Des cavités de formes rectangulaires ou cylindriques ont été par la suite proposées. Les dimensions seront choisies de façon à repousser les modes de résonances propres de la cavité à des fréquences qui ne perturberont pas la mesure. La difficulté de cette méthode réside dans la forme et la parfaite isolation de la cavité blindée qui doit être choisie de façon à ne pas modifier la distribution surfacique du courant sur l'antenne qui pourrait altérer la valeur de R_{pertes} . Ceci est d'autant plus vrai lorsque l'antenne possède un plan de masse et n'est recouverte que d'une demi-cavité. Celle-ci devra par conséquent présenter un contact parfait avec le plan de masse de façon à éviter les fuites électromagnétiques. Bien adaptée à la caractérisation des antennes électriquement petites quasi omnidirectionnelles (dont les dimensions sont inférieures à la longueur d'onde), la méthode de Wheeler cap nécessitera l'application de différentes méthodes de post-traitement qui permettront de déduire l'efficacité de l'antenne des mesures de façon à tenir compte de ces paramètres d'ajustement.

Nous avons présenté dans ce chapitre les principales grandeurs qui peuvent être mesurées sur les antennes. Les méthodes présentées sont conventionnelles mais fiables. Des variantes peuvent être mises en œuvre pour améliorer la précision des résultats ou diminuer le temps de mesure. À ce titre, le lecteur pourra se rapporter à la bibliographie ci-après pour une description plus détaillée de ces méthodes.

Bibliographie

- BALANIS C.A. — *Antenna theory, analysis and design*, John Wiley & Sons, 3^e édition, 2005.
- CHANG D.C., YANG C.C., YANG S.Y. — *Dual-reflector system with a spherical main reflector and shaped subreflector for compact range*, IEEE Proceedings, Microwaves, Antennas and Propagation, vol. 144, n° 2, 1997.
- EVANS G.E. — *Antenna measurement techniques*, Artech House, Boston, 1990.
- HIRASAWA K. AND HANEISHI M. editors, — *Analysis, design, and measurement of small and low-profile antennas*, Artech House, 1992.
- HIRVONEN T., ALA-LAURINAHO J., TUOVINEN J., RÄISÄNEN A.V. — *A compact antenna test range based on a hologram*, IEEE Transactions on Antennas and Propagation, vol. 45, n° 8, 1997.
- IEEE Standard Test and Procedures, ANSI/IEEE Std 149-1979.
- KRAUS J.D. — *Antennas*, McGraw-Hill, 1988.
- KRAUS J.D., MARHEFKA R.J. — *Antennas for all applications*, McGraw-Hill Higher Education, 2002.
- KUMMER W.H., GILLESPIE E.S. — *Antenna measurements-1978*, Proceedings of the IEEE, vol. 66, n° 4, 1978.
- LO Y.T., LEE S.W. — *Antenna Handbook*, Van Nostrand Reinhold publisher, New-York, 1993.
- OLVER A.D., SALEEB A.A. — *Lens-type compact antenna range*, Electronics Letters, vol. 15, n° 14, juillet 1979.
- PARINI C.G., PRIOR C.J. — *Radiation pattern measurements of electrical large antennas using a compact antenna test range at 180 GHz*, Electronics Letters, vol. 24, n° 25, décembre 1988.
- POTIER P., SAMSON J. — *Salles anéchoïques électromagnétiques*, Techniques de l'ingénieur, E 6 227, éditions TI.
- SLATER D. — *Near-field antenna measurements*, Artech House, Norwood, 1991.

11 • MODÉLISATION NUMÉRIQUE DU RAYONNEMENT DES ANTENNES

11.1 Introduction aux méthodes de modélisation numérique

Les principes de fonctionnement des antennes viennent d'être présentés dans les chapitres précédents. Toutes les caractéristiques des antennes découlent des équations de Maxwell qui ont été développées grâce à des méthodes analytiques. Ces méthodes ont comme seules limites la complexité mathématique sur laquelle le concepteur s'arrête à partir d'un certain degré de difficulté.

Les méthodes numériques prennent alors le relais des méthodes analytiques. Bien entendu, les résultats en sont approchés. Ce dernier point constitue une question fondamentale qui se pose lors de l'utilisation de méthodes numériques : comment celles-ci convergent-elles vers la solution ?

Les méthodes numériques nécessitent une discrétisation des variables et des grandeurs afin de les estimer avec l'outil informatique. En effet, l'ordinateur ne peut effectuer les opérations que pas à pas, mais en très grand nombre. Ces opérations, en nombre fini, constituent les étapes de la discrétisation. Le résultat du calcul discrétisé est forcément approché. Le rôle du numéricien est alors de donner un résultat mais aussi de préciser l'incertitude sur celui-ci. Il a la même démarche qu'un expérimentateur. On parle d'ailleurs souvent d'expérience numérique.

Nous verrons, dans les différentes méthodes présentées, comment s'effectue la discrétisation. Celle-ci peut intervenir très tôt au cours du développement, comme dans la méthode des différences finies ou, au contraire en fin de calcul, comme dans la méthode des moments.

La méthode la plus adaptée pour la détermination du rayonnement des antennes est la méthode des moments. C'est la méthode la plus naturelle qui s'appuie sur l'étude des fonctions de Green. Les développements analytiques sont poussés le plus loin possible selon une méthode intégrale. En dernier recours, la méthode des moments est utilisée pour exprimer la solution sous forme matricielle. Elle donne des résultats fréquentiels.

Une autre méthode fréquentielle utilisée, lorsque la complexité des structures est trop grande, est la méthode des éléments finis. Elle s'appuie sur un maillage de l'espace 3D et sur une formulation variationnelle des équations de Maxwell. Elle est, de ce fait, très robuste. C'est une méthode générale, utilisée dans tous les domaines de la physique. Elle permet de traiter de structures très complexes, mais demande une puissance de calcul très importante. Notons qu'elle fait l'objet d'importants développements mathématiques depuis de nombreuses années.

La troisième méthode, aussi très utilisée, est la méthode des différences finies. Elle s'appuie sur un maillage en trois dimensions. Elle présente plusieurs avantages qui sont, d'une part sa simplicité de mise en œuvre et d'autre part, la place mémoire réduite qu'elle requiert. De nombreux développements ont porté sur une variante de cette méthode dans le domaine temporel : la FDTD (*Finite Difference Time Domain*). Les résultats, obtenus dans le domaine temporel, peuvent s'exprimer dans le domaine fréquentiel par une transformée de Fourier en une seule opération. Il est alors facile de déterminer les caractéristiques des antennes en fonction de leur fréquence sans avoir à effectuer une simulation par point de fréquence. C'est un avantage important dans l'étude des

antennes, en particulier pour certaines antennes à résonance aiguë. On ne risque pas alors de manquer une résonance en raison d'un balayage en fréquence trop lâche.

De nombreux simulateurs électromagnétiques ont été développés par des équipes spécialisées. Certains sont commercialisés et constituent une aide précieuse pour la conception d'antennes. Pour les utiliser correctement, mieux vaut avoir une bonne connaissance des méthodes sur lesquelles ils reposent afin d'éviter des interprétations erronées de certains résultats. Il faut se dire qu'un simulateur donne presque toujours un résultat qui est à valider impérativement. Chaque simulateur a ses limites qu'il vaut mieux connaître avant un investissement coûteux qui n'aboutira pas forcément aux résultats escomptés.

Les trois méthodes présentées ici ne sont pas les seules. Elles constituent une base de méthodes numériques suffisamment variée permettant de décrire certains modes de raisonnements utilisés par les numériciens. Signalons aussi des méthodes ayant des points communs avec ces méthodes : la FIT (*Finite Integration Technique*), la méthode TLM (*Transmission Line Matrix*), la BEM (*Boundary Element Method*) qui sont des méthodes très répandues en électromagnétisme.

Chaque méthode présente des avantages mais aussi des inconvénients. La tendance est, à l'heure actuelle, d'hybrider les méthodes entre elles afin de ne garder que les avantages de l'une et de compenser ses inconvénients en associant une autre méthode. Cela donne lieu à de nombreuses recherches dans ce domaine.

11.2 Méthode intégrale associée à la méthode des moments

Nous présenterons, dans ce paragraphe, l'utilisation de la méthode intégrale associée à la méthode des moments. L'association de ces deux méthodes est souvent appelée d'une façon raccourcie : méthode des moments. C'est avant tout le principe de la méthode intégrale qui est utilisée. La méthode des moments apparaît en fin de calcul comme outil de résolution.

Le développement de la méthode s'appuiera sur l'exemple d'une antenne planaire, de type micro ruban. La généralisation à d'autres types de technologies (micro fente, coplanaire...) s'effectue sans difficulté. Cette méthode permet de calculer le champ électromagnétique rayonné par la structure rayonnante considérée à partir d'une estimation des courants sur la surface. C'est pourquoi on la place dans la catégorie des méthodes 2,5D.

11.2.1 Principe de la méthode

La méthode intégrale repose sur le principe d'équivalence (paragraphe 3.2). Une onde électromagnétique excitatrice (onde plane ou onde guidée) interagit avec une surface constituée de différents matériaux qui diffractent l'onde dans toutes les directions. La géométrie et la constitution de ces matériaux diffractant déterminent la forme du rayonnement.

Nous appliquerons la méthode au cas d'antennes planaires. C'est dans ce cas la forme de la surface métallique diffractant qui impose la forme du diagramme de rayonnement de l'antenne. Soit $\vec{E}^e, \vec{H}^e$ le champ incident ou excitateur et $\vec{E}^d, \vec{H}^d$ le champ diffracté.

Le champ excitateur induit sur la surface des courants électriques de surface $\vec{J}_s$. Les charges associées sont notées : ρ_s . Les charges sont liées à la densité de courant par l'équation de conservation locale de la charge [2.5] qui permet de ne considérer finalement que les sources sous forme de courants. On supposera que le métal est infiniment fin.

Le champ diffracté est solution des équations de Maxwell.

Dans cette méthode, le champ électromagnétique est calculé à partir du potentiel vecteur $\vec{A}$, dont la définition est rappelée au paragraphe 2.2.

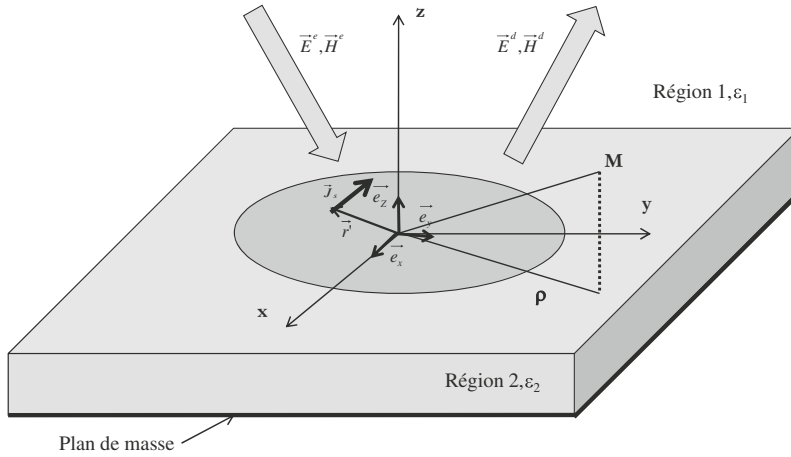


Figure 11.1 – Structure de type micro ruban rayonnante

La méthode repose sur les propriétés des fonctions de Green. Une forme du tenseur de Green associé à la structure pour le potentiel vecteur, permet de relier le potentiel vecteur au courant source du champ électromagnétique, selon l'expression :

$$\vec{A}(\vec{r}) = \iiint \overline{\overline{G}}_A(\vec{r}, \vec{r}') \vec{J}(\vec{r}') d\vec{r}' \quad [11.1]$$

Rappelons les notations classiques utilisées dans cet ouvrage : le point d'observation M est repéré par $\vec{r}$ et les sources sont repérées par $\vec{r}'$.

$\overline{\overline{G}}_A(\vec{r}, \vec{r}')$ est le tenseur (ou dyade) de Green relatif au potentiel vecteur, qui tient compte à la fois du point d'observation, du point source créant le potentiel et de la nature de la structure diffractante. Ce tenseur est de dimension (3×3) . Il caractérise la structure considérée. Les fonctions de Green qui caractérisent le vide ont déjà été introduites (chapitre 3). Chaque structure a un tenseur de Green approprié. On calcule très classiquement, par exemple, le tenseur de Green de structures multicouches. Dans la suite, nous considérerons une antenne planaire de type micro ruban, constituée d'une seule couche de substrat de permittivité ϵ_r et d'un plan de masse. Un motif métallique est gravé sur le substrat. Celui-ci engendre un phénomène de diffraction. La perméabilité magnétique considérée est celle du vide. Pour présenter la méthode, le métal est considéré comme parfait et sans épaisseur.

Le calcul de l'expression [11.1] est celui d'un produit de convolution. Le tenseur de Green représente l'effet, au point M, de la source située dans un environnement donné. La grandeur importante dans ce calcul est la distance de la source au point d'observation caractérisée par le vecteur $\vec{r} - \vec{r}'$. Le sens physique du produit de convolution apparaît bien comme le résultat intégré sur toutes les sources de l'effet en M de l'action des sources élémentaires.

Dans le cas d'une structure rayonnante de type micro ruban, les sources de courant, placées sur la surface métallique d'épaisseur infiniment fine sont notées : $\vec{J}_s(\vec{r}_s)$.

Pour bien spécifier que les sources n'existent qu'en surface, on remplace $\vec{r}'$ par $\vec{r}_s$. L'intégrale définissant le potentiel vecteur se réduit alors à l'intégrale de surface :

$$\vec{A}(\vec{r}) = \iint \overline{\overline{G}}_A(\vec{r}, \vec{r}_s) \vec{J}_s(\vec{r}_s) d\vec{s}$$

Le calcul du potentiel vecteur repose sur le même principe :

$$V(\vec{r}) = \iint G_V(\vec{r}, \vec{r}_s) \rho_s(\vec{r}_s) d\vec{s}$$

On se limitera au calcul du potentiel vecteur, sachant que le potentiel scalaire s'en déduit en utilisant la jauge de Lorentz [2.20].

Le tenseur $\overline{\overline{G}}_A(\vec{r}, \vec{r}_s)$ possède neuf composantes. Il se projette selon l'expression :

$$\begin{aligned} \overline{\overline{G}}_A = & G_A^{xx} \vec{e}_x \vec{e}_x + G_A^{xy} \vec{e}_x \vec{e}_y + G_A^{xz} \vec{e}_x \vec{e}_z + \\ & G_A^{yx} \vec{e}_y \vec{e}_x + G_A^{yy} \vec{e}_y \vec{e}_y + G_A^{yz} \vec{e}_y \vec{e}_z + G_A^{zx} \vec{e}_z \vec{e}_x + G_A^{zy} \vec{e}_z \vec{e}_y + G_A^{zz} \vec{e}_z \vec{e}_z \end{aligned}$$

Cette notation tensorielle permet d'effectuer les développements vectoriels de façon naturelle en multipliant successivement les vecteurs à partir de la droite. Ainsi lorsque le tenseur s'applique au vecteur courant, on trouve neuf termes.

Prenons l'un d'entre eux pour expliquer le développement tensoriel, en effectuant la multiplication à partir de la droite :

$$\vec{e}_y \vec{e}_z \left(\vec{J} \right) = \vec{e}_y J_z$$

La nature du problème impose un certain nombre de simplifications. Ainsi, un courant orienté selon l'axe Ox ne crée pas de composante du potentiel vecteur dans la direction Oy et réciproquement :

$$G_A^{xy} = G_A^{yx} = 0$$

Comme le courant n'a pas de composante en z :

$$G_A^{xz} = G_A^{yz} = G_A^{zx} = G_A^{zy} = 0$$

Dans le cas étudié, le tenseur $\overline{\overline{G}}_A$ se réduit donc à quatre termes :

$$\overline{\overline{G}}_A = G_A^{xx} \vec{e}_x \vec{e}_x + G_A^{yy} \vec{e}_y \vec{e}_y + G_A^{zz} \vec{e}_z \vec{e}_z + G_A^{zy} \vec{e}_z \vec{e}_y$$

Exprimons le produit tensoriel de la dyade de Green par le vecteur courant :

$$\overline{\overline{G}}_A \cdot \vec{J} = J_x G_A^{xx} \vec{e}_x + J_y G_A^{yy} \vec{e}_y + (G_A^{zz} J_z + G_A^{zy} J_y) \vec{e}_z$$

Les termes de cette expression vont être calculés pour la structure considérée. Il faudra donc utiliser toutes les conditions aux limites liées à la structure puisque ce sont elles qui imposent la forme de la solution des équations différentielles.

11.2.2 Conditions aux limites

Chaque équation de Maxwell impose des conditions aux limites à vérifier par les champs, donc aussi par les potentiels.

■ Champ électrique sur les conducteurs

Le champ électrique tangentiel doit être nul sur les surfaces conductrices :

$$\vec{E}_T = \vec{0}$$

Or le champ électrique sur la surface s'exprime selon l'expression [2.19] dans le domaine harmonique :

$$\vec{E} = -j\omega \vec{A}(\vec{r}_s) - \overrightarrow{\text{grad}}V(\vec{r}_s) \quad [11.2]$$

On doit donc vérifier :

$$\vec{e}_z \wedge \left(j\omega \vec{A}(\vec{r}_s) + \overrightarrow{\text{grad}}V(\vec{r}_s) \right) = \vec{0} \quad [11.3]$$

Cette expression conduit à vérifier les deux équations suivantes sur les potentiels :

$$\vec{e}_z \wedge \vec{A}(\vec{r}_s) = \vec{0}$$

et

$$V(\vec{r}_s) = 0$$

On utilise donc par la suite les relations :

$$A_x(\vec{r}_s) = 0 \quad \text{et} \quad A_y(\vec{r}_s) = 0$$

La jauge de Lorenz s'écrit sous la forme :

$$\text{div} \vec{A} + \frac{1}{c^2} \frac{\partial V}{\partial t} = 0 \quad [11.4]$$

En utilisant cette expression, le potentiel scalaire étant nul sur le métal, la condition suivante doit s'appliquer sur les surfaces métalliques :

$$\frac{\partial \left(\vec{e}_z \cdot \vec{A}(\vec{r}_s) \right)}{\partial z} = 0 \quad [11.5]$$

■ Conditions générales à l'interface en $z = 0$

Les potentiels sont continus sur l'interface en $z = 0$. Cela conduit aux conditions :

$$\vec{A}(\vec{r}_1) = \vec{A}(\vec{r}_2) \quad [11.6]$$

et

$$V(\vec{r}_1) = V(\vec{r}_2) \quad [11.7]$$

Dans ces expressions, $\vec{r}_1$ et $\vec{r}_2$ repèrent deux points très proches de l'interface et situés de part et d'autre de celle-ci, en $z = 0$. Les potentiels sont représentés de façon générique. Dans la suite, il faudra distinguer les potentiels dans la région 1 (au-dessus de l'interface air/diélectrique) et le potentiel dans la région 2 (en dessous de l'interface) qui seront notés : $\vec{A}_1(\vec{r}_s)$, $\vec{A}_2(\vec{r}_s)$ et $V_1(\vec{r}_s)$, $V_2(\vec{r}_s)$. On vérifie donc :

$$\vec{A}_1(\vec{r}_s) = \vec{A}_2(\vec{r}_s) \quad [11.8]$$

et

$$V_1(\vec{r}_s) = V_2(\vec{r}_s) \quad [11.9]$$

■ Discontinuité du champ magnétique tangentiel sur les conducteurs

La discontinuité au niveau des conducteurs se traduit par la relation en champ magnétique :

$$\vec{H}_{T1} - \vec{H}_{T2} = \vec{J}_s \wedge \vec{e}_z$$

Rappelons que de façon générique :

$$\vec{H} = \frac{1}{\mu_0} \text{rot} \vec{A} \text{ et donc } \vec{H}_T = \frac{1}{\mu_0} \left\{ \begin{array}{l} \frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \\ \frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \end{array} \right\}$$

Les potentiels étant égaux sur l'interface, leur dérivée selon les directions tangentielles x et y sont égales et n'apparaissent donc pas dans la différence entre les composantes tangentielles du champ magnétique. Seules restent les dérivées par rapport à z . La relation de discontinuité sur le champ magnétique s'exprime par :

$$\frac{1}{\mu_0} \vec{e}_z \wedge \left(\frac{\partial \vec{A}_1(\vec{r}_s)}{\partial z} - \frac{\partial \vec{A}_2(\vec{r}_s)}{\partial z} \right) = -\vec{e}_z \wedge \vec{J}_s(\vec{r}_s) \quad [11.10]$$

■ Continuité de la composante normale du vecteur déplacement

Considérant qu'il n'y a pas de charge à l'interface, la composante normale du vecteur déplacement est continue. Ceci s'exprime par :

$$\vec{e}_z \cdot (\vec{D}_2 - \vec{D}_1) = 0$$

Soit :

$$\vec{e}_z \cdot (\epsilon_2 \vec{E}_2 - \epsilon_1 \vec{E}_1) = 0$$

La définition du champ électrique en fonction des potentiels [11.2], implique :

$$j\omega(\epsilon_1 - \epsilon_2)A_z(\vec{r}_s) + \epsilon_1 \frac{\partial V_1(\vec{r}_s)}{\partial z} - \epsilon_2 \frac{\partial V_2(\vec{r}_s)}{\partial z} = 0 \quad [11.11]$$

Dans cette expression, le terme A_z représente indifféremment la valeur de la composante en z du potentiel dans le milieu 1 ou dans le milieu 2, puisqu'ils sont continus en $\vec{r}_s$.

■ Condition de rayonnement de Sommerfeld

Les conditions précédentes s'appliquent à distance finie, aux interfaces entre les matériaux. Pour terminer le calcul, il faut ajouter une condition aux limites à l'infini. Cette condition est la condition de Sommerfeld qui impose une équation différentielle lorsque r tend vers l'infini. On a déjà vu que le champ à grande distance s'exprime par le produit de fonctions à variables séparées, d'une part selon la distance r et d'autre part selon deux directions de l'espace (θ et ϕ).

La condition sur la fonction scalaire Ψ traduisant les variations du champ selon la distance s'exprime par la condition de Sommerfeld :

$$\lim_{r \rightarrow \infty} \left(r \left(\frac{\partial \Psi}{\partial r} + jk\Psi \right) \right) = 0 \quad [11.12]$$

11.2.3 Calcul du tenseur de Green

Les différents termes du tenseur de Green sont des fonctions de l'espace qui doivent vérifier les conditions énoncées précédemment. Comme ces fonctions dépendent à la fois du point où est localisée la source et du point d'observation M, on établit d'abord leurs expressions en considérant la source située à l'origine. Une translation du point source permet ensuite de trouver les expressions générales.

On se limitera, dans un premier temps, à considérer la structure constituée de couches parallèles de matériaux infinis :

- le substrat ;
- le plan de masse.

La composante en z , perpendiculaire aux plans des couches, joue donc un rôle important. Selon le plan des couches et compte tenu du fait que les matériaux sont infinis dans les directions x et y , la symétrie radiale s'impose (figure 11.1). On choisit donc de repérer les points par la longueur de leur projection sur le plan : ρ , et par leur hauteur : z et par l'angle θ . Il en va de même pour les vecteurs considérés.

■ Solutions types de l'équation de Helmholtz

Chaque terme du tenseur de Green est représenté par une fonction scalaire qu'on suppose à variables séparables :

$$\psi(\rho, \phi, z) = f(\rho) g(\phi) h(z)$$

Cette fonction est solution de l'équation de Helmholtz qui s'écrit, en coordonnées cylindriques :

$$\frac{1}{\rho} \frac{\partial}{\partial \rho} \left(\rho \frac{\partial \psi}{\partial \rho} \right) + \frac{1}{\rho^2} \frac{\partial^2 \psi}{\partial \phi^2} + \frac{\partial^2 \psi}{\partial z^2} + k^2 \psi = 0$$

La constante de propagation a pour valeur :

$$k = \sqrt{\epsilon_0 \epsilon_r \mu_0 \omega}$$

Tenant compte de l'hypothèse sur la fonction $\psi(\rho, \phi, z)$, les trois équations suivantes doivent être vérifiées :

$$\frac{1}{\rho} \frac{\partial}{\partial \rho} \left(\rho \frac{\partial f}{\partial \rho} \right) + \left[(k_\rho \rho)^2 - n^2 \right] f = 0 \quad [11.13]$$

$$\frac{\partial^2 g}{\partial \phi^2} + n^2 g = 0 \quad [11.14]$$

$$\frac{\partial^2 h}{\partial z^2} + k_z^2 h = 0 \quad [11.15]$$

n est un entier et la décomposition du vecteur de propagation impose :

$$k^2 = k_z^2 + k_\rho^2$$

Le type de solution de l'équation [11.13] est soit une fonction de Bessel de première ou de seconde espèce, soit une fonction de Hankel de première ou de deuxième espèce de la variable $(k_\rho \rho)$. Seule cette dernière solution s'annule lorsque r tend vers l'infini avec k_ρ complexe. Cette propriété lui permet de vérifier la condition de rayonnement de Sommerfeld. C'est donc la fonction qu'on choisira pour représenter la variation radiale des termes du tenseur de Green. Cette fonction élémentaire a pour variable le produit de la distance radiale sur le substrat par la projection k_ρ du vecteur de propagation sur le plan du substrat :

$$H_n^{(2)}(k_\rho \rho)$$

La variation globale en ρ est une somme continue de toutes les fonctions élémentaires solutions de l'équation [11.13] :

$$f(\rho) = \int_0^\infty f(k_\rho) H_n^{(2)}(k_\rho \rho) dk_\rho \quad [11.16]$$

La solution de l'équation [11.14] concernant g se met sous forme, soit d'une combinaison de fonctions trigonométriques, soit d'une combinaison de fonctions exponentielles. La géométrie

du dispositif rayonnant impose une symétrie par rapport à l'axe Oz qui entraîne une solution constante en ϕ , par conséquent n est nul.

La solution de l'équation [11.15] concernant h se met sous forme d'une combinaison de fonctions exponentielles de type :

$$e^{\pm jk_z z}$$

avec $\text{Re}(k_z) > 0$ et $\text{Im}(k_z) < 0$

■ Application au calcul des termes du tenseur de Green

Les solutions de l'équation de Helmholtz sont valables dans chaque région (figure 11.1). La région 1 correspond à l'air et la région 2 au diélectrique.. Nous ne présenterons ici que les calculs concernant une des composantes du tenseur de Green afin d'exposer le principe de la méthode. Les autres termes sont déduits de la même façon. Nous allons exprimer le terme correspondant à la composante en xx dans chacune des régions.

Il existe un vecteur de propagation dans chaque région. Cependant ces vecteurs vérifient l'égalité de leur composante tangentielle k_p . Cette propriété résulte de l'indépendance des phases des ondes dans les deux milieux en fonction de x et de y . C'est cette propriété qui permet de déduire les lois de Descartes de la réfraction. Seules diffèrent les composantes longitudinales : k_{z1} et k_{z2} qui sont telles que :

$$k_1^2 = k_p^2 + k_{z1}^2$$

et

$$k_2^2 = k_p^2 + k_{z2}^2$$

Les vecteurs de propagation s'expriment dans chacune des régions par :

$$k_1^2 = \epsilon_0 \epsilon_{r1} \mu_0 \omega^2$$

$$k_2^2 = \epsilon_0 \epsilon_{r2} \mu_0 \omega^2$$

Dans la suite, la région 1 représentant l'air, sa permittivité relative sera égale à un :

$$\epsilon_{r1} = 1$$

D'après ce qui vient d'être dit, les termes en xx du tenseur de Green s'écrivent dans chaque région sous la forme :

$$G_{A1}^{xx} = \int_0^\infty H_0^{(2)}(k_p \rho) f_1(k_p) dk_p e^{-jk_{z1} z} \quad [11.17]$$

$$G_{A2}^{xx} = \int_0^\infty H_0^{(2)}(k_p \rho) f_2(k_p) dk_p (C_1 e^{-jk_{z2} z} + C_2 e^{+jk_{z2} z}) \quad [11.18]$$

Selon la relation [11.3], sa valeur est nulle sur le plan de masse. Cela impose :

$$C_1 e^{jk_{z2} h} + C_2 e^{-jk_{z2} h} = 0 \quad [11.19]$$

La continuité de la composante tangentielle du potentiel en $z = 0$ entraîne :

$$f_1(k_p \rho) = f_2(k_p \rho) (C_1 + C_2) \quad [11.20]$$

La discontinuité du champ magnétique entraîne une relation faisant intervenir le courant de surface. En coordonnées cylindriques, le courant surfacique élémentaire selon l'axe Ox s'exprime sous la forme :

$$\vec{J}_s = \frac{\delta(\rho)}{2\pi\rho} \vec{e}_x$$

Or la distribution de Dirac s'exprime selon la fonction de Hankel de deuxième espèce :

$$\frac{\delta(\rho)}{2\pi\rho} = \frac{1}{4\pi} \int_0^\infty H_0^{(2)}(k_\rho\rho) k_\rho dk_\rho$$

La relation [11.10] devient alors :

$$\frac{1}{\mu_0} \left(\frac{\partial A_{1x}}{\partial z} - \frac{\partial A_{2x}}{\partial z} \right) = -\frac{1}{4\pi} \int_0^\infty H_0^{(2)}(k_\rho\rho) k_\rho dk_\rho$$

Ceci impose une relation sur l'intégrant qui conduit à :

$$jk_{z1}f_1(k_\rho\rho) + jk_{z2}f_2(k_\rho\rho) (C_2 - C_1) = \frac{\mu_0}{4\pi} k_\rho \quad [11.21]$$

Les relations [11.19], [11.20] et [11.21] conduisent à :

$$f_1(k_\rho\rho) = \frac{\mu_0}{4\pi} k_\rho \frac{1}{jk_{z1} + jk_{z2} \coth(jk_{z2}h)}$$

Le calcul de la fonction analogue dans le milieu 2 donne :

$$f_2(k_\rho\rho) = \frac{f_1(k_\rho\rho)}{C_1 (1 - e^{2jk_{z2}h})}$$

Le résultat final obtenu est donc :

$$G_{A1}^{\text{xx}} = \frac{\mu_0}{4\pi} \int_0^\infty H_0^{(2)}(k_\rho\rho) \frac{k_\rho}{DTE} dk_\rho e^{-jk_{z1}z}$$

$$G_{A2}^{\text{xx}} = \frac{\mu_0}{4\pi} \int_0^\infty H_0^{(2)}(k_\rho\rho) \frac{k_\rho}{DTE} dk_\rho \frac{\sinh[jk_{z1}(z+h)]}{\sinh(jk_{z1}h)}$$

Avec :

$$DTE = jk_{z1} + jk_{z2} \coth(jk_{z2}h)$$

Le calcul des autres termes du tenseur de Green s'effectue selon le même principe présenté en détail dans l'article de J. R. Mosig.

11.2.4 Équation électrique intégrale

Dans ce paragraphe, nous montrons comment les conditions sur le champ électrique, au niveau des conducteurs, permettent d'imposer la forme du champ diffracté. Ces conditions vont aboutir à une équation intégrale.

Les potentiels vecteur et scalaire en chaque point de l'espace sont calculés à partir des expressions du tenseur de Green. Nous allons montrer qu'il est alors possible de déterminer les champs électriques et magnétiques diffractés.

Le champ total est formé de la superposition du champ excitateur $\vec{E}^e(\vec{r}_s)$ et du champ diffracté $\vec{E}^d(\vec{r}_s)$. Ce champ total doit vérifier les conditions aux limites imposées par la structure diffractant. En particulier, si la structure est formée d'un certain nombre de conducteurs repérés par l'indice i , la relation suivante doit être vérifiée si les conducteurs sont parfaits :

$$\vec{e}_z \wedge \left[\vec{E}^d(\vec{r}_s) + \vec{E}^e(\vec{r}_s) \right] = 0 \quad [11.22]$$

Si les conducteurs sont de bons conducteurs, ils présentent une impédance de surface non nulle, donnée par la relation :

$$Z_{si} = (1 + j) \sqrt{\frac{\mu_0 \pi f}{\sigma_i}}$$

Cette relation fait intervenir la conductivité du conducteur i et la fréquence. L'impédance de surface est liée à l'effet de peau.

La relation [11.22] sur les conducteurs s'écrit alors :

$$\vec{e}_z \wedge \left[\vec{E}^d(\vec{r}_s) + \vec{E}^e(\vec{r}_s) \right] = \vec{e}_z \wedge \sum_i Z_{si} \vec{J}_{si}(\vec{r}_s)$$

Le champ diffracté est remplacé par son expression en fonction des potentiels :

$$\vec{e}_z \wedge \left[-j\omega \vec{A}(\vec{r}_s) - \overrightarrow{\text{grad}} V(\vec{r}_s) + \vec{E}^e(\vec{r}_s) \right] = \vec{e}_z \wedge \sum_i Z_{si} \vec{J}_{si}(\vec{r}_s)$$

Les potentiels sont remplacés par leurs expressions en fonction des sources diffractant : les densités de courant et de charges surfaciques :

$$\begin{aligned} \vec{e}_z \wedge \left[j\omega \iint \overrightarrow{G}_A(\vec{r}, \vec{r}_s) \vec{J}_{si}(\vec{r}_s) d\vec{s} + \overrightarrow{\text{grad}} \iint G_V(\vec{r}, \vec{r}_s) \rho_{si}(\vec{r}_s) d\vec{s} + \sum_i Z_{si} \vec{J}_{si}(\vec{r}_s) \right] \\ = \vec{e}_z \wedge \vec{E}^e(\vec{r}_s) \quad [11.23] \end{aligned}$$

Cette dernière expression fait apparaître, au second membre, le champ exciteur qui impose la forme de la solution. Dans le premier membre figurent les inconnues.

Par ailleurs, la densité de charges est liée à la densité de courant par la relation :

$$\rho_s = \frac{1}{j\omega} \text{div} \left(\vec{J}_{si}(\vec{r}_s) \right)$$

Ceci permet de ne faire apparaître dans l'équation [11.23] que la densité surfacique de courant comme seule inconnue du problème.

L'étape suivante consiste à résoudre cette équation par la méthode des moments.

11.2.5 Méthode des moments

La méthode des moments est une méthode de décomposition d'une grandeur inconnue sur une base de fonctions. Les coefficients de la projection sur la base deviennent les inconnues du problème. Cette méthode conduit à une résolution matricielle, comme nous allons le voir sur l'application à la résolution de l'équation intégrale qui a été établie dans le paragraphe précédent. L'équation intégrale [11.23] se met sous une forme symbolique en représentant toutes les opérations du membre de gauche sous la forme d'un opérateur linéaire L :

$$L(\vec{J}_s) = \vec{e}_z \wedge \vec{E}^e(\vec{r}_s) \quad [11.24]$$

L'inversion de cet opérateur conduit à la solution en termes de densité de courant surfacique :

$$\vec{J}_s = L^{-1} \left(\vec{e}_z \wedge \vec{E}^e(\vec{r}_s) \right)$$

La méthode des moments permettra de conduire à la valeur approchée de la densité de courant surfacique. Entrons dans le détail des principes de calcul.

■ Discrétisation

La méthode consiste à mailler les conducteurs afin de résoudre l'équation [11.24]. Le maillage le plus couramment employé est formé de quadrilatères (figure 11.2)

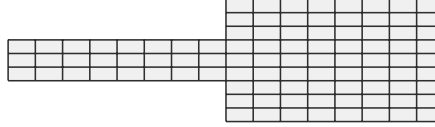


Figure 11.2 – Maillage de la surface conductrice en quadrilatères

Développons le courant surfacique sur une base de vecteurs $\vec{j}_k$:

$$\vec{J}_s = \sum_k a_k \vec{j}_k \quad [11.25]$$

Définissons un produit interne entre vecteurs sous la forme :

$$\langle \vec{U}, \vec{V} \rangle = \iint \vec{U} \cdot \vec{V} dS \quad [11.26]$$

Introduisons un type particulier de vecteurs appelés fonctions test $\vec{u}_i$. Ces fonctions permettent de construire des produits internes définis selon [11.26] et mettant en jeu le courant de surface. La multiplication de l'équation [11.24] par une fonction test conduit à :

$$\langle \vec{u}_i, L(\vec{J}_s) \rangle = \langle \vec{u}_i, \vec{e}_z \wedge \vec{E}^e \rangle$$

Utilisons le développement du courant de surface sur les fonctions de base $\vec{j}_k$ et la linéarité de L , on obtient successivement :

$$\langle \vec{u}_i, L(\vec{J}_s) \rangle = \left\langle \vec{u}_i, L\left(\sum_k a_k \vec{j}_k\right) \right\rangle = \left\langle \vec{u}_i, \sum_k a_k L(\vec{j}_k) \right\rangle \quad [11.27]$$

Soit encore :

$$\langle \vec{u}_i, L(\vec{J}_s) \rangle = \sum_k a_k \langle \vec{u}_i, L(\vec{j}_k) \rangle$$

Posons :

$$L_{ik} = \langle \vec{u}_i, L(\vec{j}_k) \rangle \quad \text{et} \quad \langle \vec{u}_i, L(\vec{J}_s) \rangle = b_i$$

Avec ces notations, l'équation [11.27] équivaut à :

$$\sum_k L_{ik} a_k = b_i$$

En utilisant un nombre de fonctions test égal au nombre de fonctions de base, on obtient un système linéaire se mettant sous la forme :

$$LA = B \quad [11.28]$$

Le second membre B est connu. Les termes de la matrice L sont calculables puisqu'ils sont égaux au produit interne entre un vecteur de test et le vecteur résultant de l'application de l'opérateur L sur une fonction de base. Seuls les coefficients a_k sont à déterminer par inversion du système [11.28]. Après inversion, le vecteur courant surfacique peut être calculé par l'application de [11.25].

■ Choix des fonctions de base et des fonctions test

Le choix des fonctions de base est crucial dans l'application des principes d'une méthode numérique. En effet, plus la forme des fonctions de base est proche de la solution cherchée, plus la méthode a des chances de donner des résultats précis.

Les fonctions de bases utilisées pour la méthode intégrale appliquée aux antennes planaires sont des fonctions triangulaires dans la direction de propagation.

Pour comprendre l'utilité de ces fonctions, prenons un exemple à une dimension. Les fonctions de bases sont définies par :

$$t_n(x - x_n) = \begin{cases} 1 - \frac{|x - x_n|}{a} & \text{si } |x - x_n| < a \\ 0, & \text{ailleurs} \end{cases}$$

Avec $x_n = na$

Donc la décomposition d'une fonction selon cette base se met sous la forme :

$$f(x) = \sum_{n=-\infty}^{+\infty} a_n t_n(x - x_n)$$

Un exemple simple utilisant quatre fonctions de base est représenté sur la figure 11.3 :

$$f(x) = \sum_{n=0}^4 a_n t_n(x - x_n), \text{ en prenant :}$$

$$\text{pour } a = 1 : \quad a_0 = 1 \quad a_1 = 2.5 \quad a_2 = 2 \quad a_3 = 1.5 \quad a_4 = 1$$

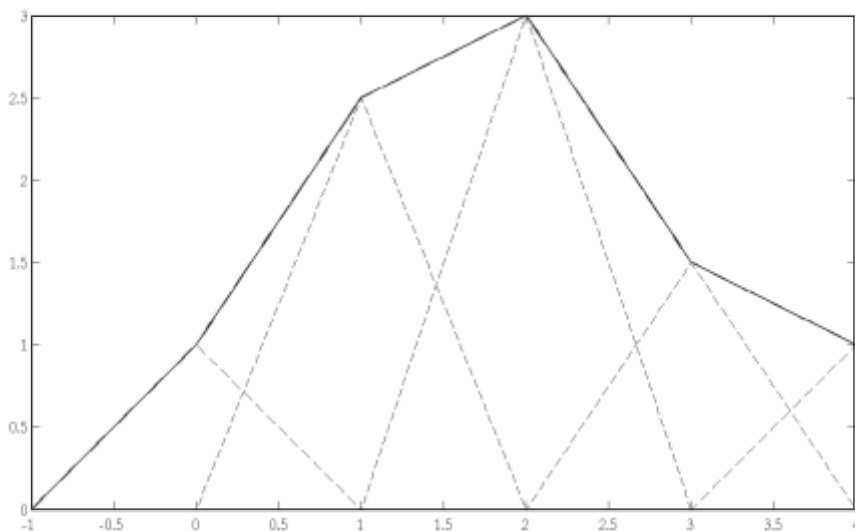


Figure 11.3 – Représentation de $f(x)$ comme somme de fonctions triangulaires

La description d'une fonction sera d'autant plus précise que la valeur du paramètre a sera petite. Montrons, sur l'exemple d'une ligne micro ruban, comment se fait le choix des fonctions de base. Nous étendrons ensuite ce choix au cas d'une antenne planaire en technologie micro ruban.

Les densités surfaciques de courant sont essentiellement orientées dans la direction de l'axe Ox de la ligne. Les courants se développent donc selon des fonctions de base vectorielles orientées selon Ox . En choisissant des fonctions triangulaires selon Ox et constantes selon Oy , les fonctions de base s'expriment par :

$$\vec{j}_k = t_k(x - x_k) \vec{e}_x$$

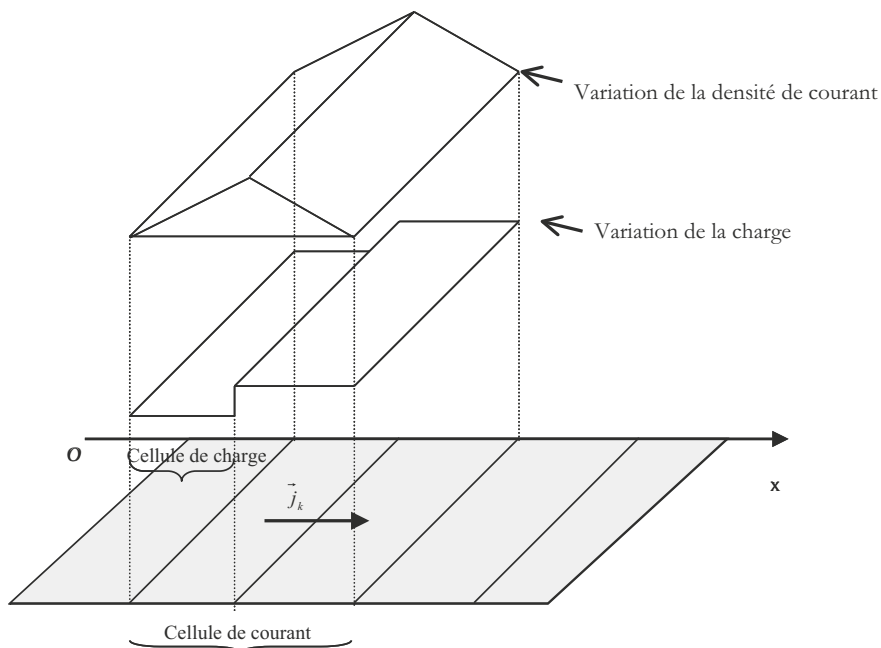


Figure 11.4 – Représentation des fonctions de base du courant selon les cellules de courant.

L'équation de continuité de la charge relie la densité de courant à la densité de charges par l'équation :

$$\text{div} \vec{j}_s + j\omega\rho_s = 0$$

On distingue alors deux types de cellules : les cellules de courant de largeur égale à $2a$, selon l'axe Ox , et les cellules de charges de largeur égale à a . Le choix de fonctions de base triangulaires pour la densité de courant entraîne que chaque cellule de charges supporte une densité uniforme de charges.

Ce principe du choix des cellules de courant s'étend facilement à deux dimensions. Il suffit de prendre les fonctions analogues en x et en y , et d'en prendre le produit. Il est alors facile de tapisser la surface de l'antenne de cellules à deux dimensions analogues aux cellules de charges et de courant présentées.

■ Avantages et inconvénient de la méthode intégrale

La méthode intégrale présente de nombreux avantages car c'est à la fois une méthode reposant sur le calcul analytique et utilisant un maillage en surface. Elle permet d'obtenir le champ en tout point de l'espace à partir d'une description surfacique. Cela permet de la placer dans les méthodes 2,5D.

Le maillage introduit une souplesse d'utilisation en fonction de la puissance de calcul dont on dispose et de la précision attendue. Plus le maillage est fin, plus les résultats sont précis. Signalons un inconvénient de la méthode qui réside dans le fait que les calculs des fonctions de Green ont été faits en supposant un plan de masse infini. Dans certains cas d'antennes de faible taille, la finitude du plan de masse entraîne des résultats légèrement différents.

11.3 La méthode des éléments finis

Les efforts de développement de la méthode des éléments finis ont été nombreux dans tous les domaines de la physique. Ils ont conduit à des théories mathématiques poussées. Nous donnerons ici quelques notions très simples permettant de comprendre les processus de simulation utilisant cette méthode.

La méthode des éléments finis repose sur une interpolation par élément. Ceux-ci sont définis comme des sous-ensembles de l'ensemble de définition d'une grandeur physique représentée par une fonction. En trois dimensions, l'espace est décomposé en éléments qui sont la plupart du temps tétraédriques. Les éléments parallélépipédiques sont moins utilisés. La grandeur physique est interpolée sur chacun de ces éléments grâce à des fonctions adaptées.

11.3.1 Principes de base de la méthode des éléments finis

Nous allons exposer le principe de la méthode à une dimension, en considérant une grandeur scalaire (le potentiel par exemple).

Soit une fonction $f(x)$ dont on connaît la valeur en certains points $(x_1, x_2, \dots, x_N)$ de l'ensemble total de définition. Cette fonction est approchée par une fonction $f_a(x)$ qui prend les valeurs de la fonction f aux points $(x_1, x_2, \dots, x_N)$.

On définit l'erreur $e(x)$ de cette approximation comme :

$$e(x) = f(x) - f_a(x)$$

L'erreur d'approximation est nulle pour l'ensemble des points $(x_1, x_2, \dots, x_N)$.

$f_a(x)$ peut être recherchée comme une fonction polynomiale :

$$f_a(x) = \alpha_0 + \alpha_1 x + \alpha_2 x^2 + \dots + \alpha_n x^n$$

La valeur de la fonction f_a est imposée pour l'ensemble des points $(x_1, x_2, \dots, x_N)$. On obtient donc N équations par rapport aux paramètres $(\alpha_0, \alpha_1, \dots, \alpha_n)$. Ces paramètres sont calculables si l'ordre des polynômes est $n = N - 1$.

En transformant l'équation précédente, on obtient :

$$f_a(x) = b_0 P_0(x) + b_1 P_1(x) + b_2 P_2(x) + \dots + b_n P_n(x)$$

$P_i(x)$ représentant un polynôme d'ordre i , où i est au maximum égal à n .

Cette écriture de la fonction d'approximation est très générale. En prenant comme paramètres les valeurs $y_i = f(x_i)$, de la fonction aux points $(x_1, x_2, \dots, x_N)$, la fonction $f_a(x)$ s'écrit :

$$f_a(x) = y_1 a_1(x) + y_2 a_2(x) + \dots + y_n a_n(x)$$

Cette forme de f_a s'appelle l'approximation nodale de la fonction f . Les fonctions $a_i(x)$ doivent vérifier :

$$a_i(x_j) = \delta_{ij}$$

où δ_{ij} est le symbole de Kronecker :

$$\delta_{ij} = 0 \text{ si } i \neq j$$

$$\delta_{ij} = 1 \text{ si } i = j$$

x_i sont les nœuds de l'interpolation, y_i sont les variables nodales et a_i les fonctions d'interpolation. La méthode des éléments finis consiste à restreindre l'intervalle de définition de la fonction à des intervalles plus petits, appelés éléments, sur lesquels il est plus facile d'appliquer la définition de la fonction d'approximation. Ces éléments, dans le cas à une dimension, sont des segments. Sur chaque élément, on applique la fonction d'interpolation telle qu'elle vient d'être présentée. Cette forme a l'avantage d'assurer la continuité de la fonction sur tout l'intervalle de définition.

Si l'on définit un segment compris entre deux nœuds voisins, l'interpolation est tout simplement linéaire sur l'élément. On parle d'élément fini du premier ordre. Si l'on définit un segment contenant trois nœuds, l'interpolation est quadratique. On parle d'éléments finis du second ordre. L'approximation nodale ne fait intervenir que des variables nodales situées sur l'élément et sur sa frontière.

11.3.2 Principe de base de la méthode des éléments finis à trois dimensions

Le principe de la méthode précédente se généralise à deux ou trois dimensions.

Soit une fonction scalaire, dépendant de l'espace $f(x, y, z)$ et $f_a(x, y, z)$ sa fonction approchée. L'erreur d'approximation est nulle aux points $M_1, M_2, \dots, M_N$.

Les coordonnées du point M_i sont (x_i, y_i, z_i) .

D'après la définition de l'approximation nodale, $f_a(x, y, z)$ s'écrit :

$$f_a(x, y, z) = y_1 a_1(x, y, z) + y_2 a_2(x, y, z) + \dots + y_n a_n(x, y, z)$$

où $y = f(x_i, y_i, z_i)_{\cdot i}$

Les propriétés des fonctions d'interpolation sont les mêmes qu'à une dimension :

$$a_i(x_i, y_i, z_i) = \delta_{ij}$$

L'étape suivante consiste à définir les éléments, c'est-à-dire les volumes sur lesquels on restreint la définition de l'approximation nodale. Les éléments les plus usuels sont les tétraèdres. On maille ainsi l'espace à l'aide de tétraèdres qui ne présentent pas de recouvrement et dont la réunion constitue l'ensemble global de définition de la fonction. Chaque sommet d'un tétraèdre est repéré par un entier (figure 11.5). La fonction d'interpolation, définie sur chaque élément, fait intervenir les nœuds d'interpolation. Lorsque les sommets du tétraèdre constituent les nœuds d'interpolation (cas de l'ordre 1), la fonction d'interpolation se met sous la forme :

$$f_a(x, y, z) = y_i a_i(x, y, z) + y_j a_j(x, y, z) + y_k a_k(x, y, z) + y_l a_l(x, y, z)$$

Cette opération constitue l'étape de discrétisation. En effet, la fonction continue $f_a(x, y, z)$ est exprimée en fonction des valeurs discrétisées de la fonction prises aux nœuds qui sont des constantes. Les fonctions d'interpolations assurent les propriétés de continuité.

À l'ordre 1, les fonctions d'interpolation se mettent sous la forme :

$$a_i(x, y, z) = \frac{1}{6V_e} \left\{ \begin{aligned} &(x - x_j) [(y_k - y_j)(z_l - z_j) - (z_k - z_j)(y_l - y_j)] \\ &- (y - y_j) [(x_k - x_j)(z_l - z_j) - (z_k - z_j)(x_l - x_j)] \\ &+ (z - z_j) [(x_k - x_j)(y_l - y_j) - (y_k - y_j)(x_l - x_j)] \end{aligned} \right\}$$

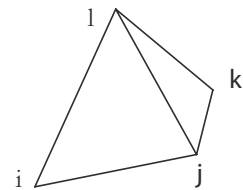


Figure 11.5 – Élément de base tétraédrique servant à la construction d'un maillage en trois dimensions

V^e est le volume de l'élément considéré. Les expressions des quatre fonctions d'interpolation linéaire se déduisent de celle-ci par permutation circulaire.

Lorsque les fonctions sont quadratiques (ordre 2), le nombre de nœuds à prendre en compte est 10. Les nœuds sont constitués des sommets des tétraèdres et des points milieux des arêtes.

11.3.3 Application de la méthode aux problèmes d'électromagnétisme

La méthode consiste à résoudre une équation différentielle traduisant le problème physique, associée à des conditions aux limites. Une équation différentielle présente généralement des familles de solutions. Les conditions aux limites permettent de choisir une des solutions. Prenons l'exemple simple d'un problème concernant une grandeur scalaire (le potentiel V par exemple). L'équation régissant la forme du potentiel peut être, par exemple, l'équation de Poisson :

$$\Delta V + \frac{\rho}{\epsilon_0} = 0 \quad \text{dans le volume de calcul}$$

$$V = V_1 \quad \text{sur la frontière du volume}$$

De façon générale, le problème consiste à résoudre un système de la forme :

$$L(V) + f = 0 \quad \text{dans le volume de calcul}$$

$$\Gamma(V) \quad \text{imposé sur la frontière du volume}$$

$L(V)$ est un opérateur différentiel donc linéaire. f est une fonction connue à l'intérieur du volume de calcul. $\Gamma(V)$ est une fonction imposée sur la surface frontière du volume. Les fonctions f et Γ sont appelées fonctions des sollicitations.

Les solutions en V de l'équation différentielle vérifient donc l'équation intégrale :

$$I(V) = \iiint \Phi [L(V) + f(V)] d\tau = 0, \forall \Phi$$

L'intégration a lieu sur l'ensemble du volume de calcul. On utilise souvent comme fonction Φ , qui joue le rôle d'une fonction test, la différentielle du potentiel δV , de façon à faire apparaître pour $I(V)$ la différentielle d'une forme bilinéaire. La méthode est dite alors de type Galerkin :

$$I(V) = \iiint \delta V [L(V) + f(V)] dv = 0$$

Dans certains cas (la majorité en électromagnétisme), dits conservatifs, $I(V)$ s'exprime comme la différentielle d'une fonction $U(V)$:

$$I(V) = \delta U(V)$$

La solution est donc obtenue lorsque $U(V)$ est stationnaire. La condition mathématique pour obtenir la solution s'écrit alors :

$$\delta U(V) = 0$$

$U(V)$ est appelée la **fonctionnelle** du problème. On reconnaît, dans cette méthode, le principe utilisé en physique pour chercher la solution d'un problème (par exemple, le principe de moindre action, du minimum d'énergie...). La solution est obtenue pour une valeur qui rend extrémale une grandeur physique de type intégrale (énergie, trajet optique...). La solution est alors obtenue avec une grande précision puisque la dérivée première de la fonctionnelle est nulle. Cette méthode qui consiste à calculer l'extremum d'une fonctionnelle pour obtenir la solution du problème s'appelle méthode **variationnelle**.

■ Exemple de calcul de potentiel

Nous allons résoudre l'équation de Poisson dans un volume V_1 avec cette méthode :

$$\Delta V + \frac{\rho}{\epsilon_0} = 0$$

Les conditions aux limites de type Dirichlet sur S_D et de type Neumann sur S_N s'expriment par :

$$\left\{ \begin{array}{ll} V = V_D & \text{sur } S_D \\ \frac{\partial V}{\partial n} = f_N & \text{sur } S_N \end{array} \right\}$$

Le problème sous forme intégrale est équivalent à :

$$I(V) = \iiint \delta V \left(\frac{\partial^2 V}{\partial x^2} + \frac{\partial^2 V}{\partial y^2} + \frac{\partial^2 V}{\partial z^2} + \frac{\rho}{\epsilon_0} \right) dv = 0$$

Une intégration par parties conduit à :

$$I(V) = - \iiint \left(\frac{\partial \delta V}{\partial x} \frac{\partial V}{\partial x} + \frac{\partial \delta V}{\partial y} \frac{\partial V}{\partial y} + \frac{\partial \delta V}{\partial z} \frac{\partial V}{\partial z} - \delta V \frac{\rho}{\epsilon_0} \right) dv + \iint \delta V \frac{\partial V}{\partial n} ds$$

La première partie de cette expression est une intégration dans tout le volume de calcul. Le second terme est une intégrale à deux dimensions sur la surface limitant le volume de calcul.

δV est la différentielle du potentiel. Elle est donc nulle sur les surfaces S_D pour lesquelles le potentiel est imposé. L'intégration à deux dimensions se réduit donc à celle qui s'effectue sur la surface S_N .

Cette expression de l'intégrale obtenue après une intégration par parties permet de faire apparaître des ordres de dérivation moindres. Ainsi, il est possible de ne considérer que des fonctions dérivables une fois. Cette transformation conduit à ce qui est appelé la formulation faible du problème.

On constate que :

$$I(V) = \delta U(V)$$

si :

$$U(V) = - \iiint \left[\frac{1}{2} \left(\frac{\partial V}{\partial x} \right)^2 + \frac{1}{2} \left(\frac{\partial V}{\partial y} \right)^2 + \frac{1}{2} \left(\frac{\partial V}{\partial z} \right)^2 - \frac{\rho}{\epsilon_0} V \right] dv + \iint V f_N ds$$

Cette expression est liée à l'énergie électrostatique totale contenue dans le volume.

Le problème continu vient d'être présenté : pour trouver la solution du problème, il suffit de trouver l'extremum de la fonctionnelle. La fonction potentiel est inconnue, sauf en des points particuliers (sur les frontières du volume, par exemple) On utilise donc une approximation par éléments finis de la fonction potentiel.

Le volume étant maillé en éléments finis, notés (e), on calcule $U(V)$ sur chaque élément, noté $U^e(V)$.

$$U(V) = \sum_e U^e(V)$$

Le calcul de $U^e(V)$ s'effectue à l'aide de l'approximation nodale présentée plus haut. Si l'élément e est constitué des nœuds (i, j, k, l), le potentiel V s'écrit :

$$V(x, y, z) = V_i a_i(x, y, z) + V_j a_j(x, y, z) + V_k a_k(x, y, z) + V_l a_l(x, y, z)$$

Soit :

$$V(x, y, z) = \sum V_i a_i(x, y, z)$$

De même :

$$\delta V = \sum \delta V_{ja_j}(x, y, z)$$

Pour présenter la méthode, nous utiliserons la formulation forte du problème. Le calcul s'effectue de la même façon avec la formulation faible, mais présente un nombre plus grand de termes.

La différentielle de la fonctionnelle s'écrit :

$$\delta U^e = \iiint \left(\sum \delta V_{ja_j}(x, y, z) \right) \left[L \left(\sum V_{ia_i}(x, y, z) \right) + f \right] dv$$

Soit en utilisant la linéarité de l'opérateur L :

$$\delta U^e = \sum \delta V_j \left[\iiint a_j(x, y, z) \sum L(a_i(x, y, z)) V_i dv + \iiint a_j(x, y, z) f dv \right]$$

En introduisant le vecteur colonne formé des quatre potentiels, V_i , aux nœuds d'interpolation sous la forme : $\{V_i\}$ et son vecteur transposé : $\{V_i\}^t$, on obtient :

$$\delta U^e = \{\delta V_j\}^t [(K^e) \{V_i\} - \{f^e\}]$$

$\{f^e\}$ est le vecteur élémentaire des sollicitations.

La différentielle totale dans le volume de calcul est obtenue en sommant les différentielles calculées sur chacun des éléments.

$$\delta U = \sum_e \delta U^e$$

En tenant compte de tous les points du volume de calcul, on obtient une expression qui se met sous la forme :

$$\delta U = \{\delta V_m\}^t [(K) \{V_n\} - \{F\}]$$

Le vecteur $\{V_n\}$ est un vecteur colonne qui regroupe tous les nœuds d'interpolation du volume de calcul. Le vecteur $\{F\}$ est le vecteur global des sollicitations dont les éléments sont non nuls à l'endroit où il existe une densité volumique de charge.

(K) est la matrice globale. L'opération qui consiste à calculer tous les termes de la matrice globale à partir des termes des matrices élémentaires s'appelle l'assemblage.

Cette expression doit être vérifiée quelles que soient les variations du potentiel. Donc :

$$[(K) \{V_n\} - \{F\}] = 0$$

Ceci conduit donc à la résolution d'un système linéaire à N équations, si N représente le nombre total de nœuds du système.

11.3.4 Problème d'électromagnétisme dans le cas général

Dans ce paragraphe, nous allons montrer comment transformer la méthode présentée précédemment sur une grandeur scalaire, en vue de traiter de problèmes vectoriels.

■ Formulation du problème vectoriel en trois dimensions

Le cas le plus général qui nous intéresse ici est celui du calcul du champ électromagnétique dans un volume de calcul donné. On montre qu'il suffit de calculer soit le champ électrique, soit le champ magnétique.

L'équation différentielle vérifiée par l'un ou l'autre des champs est l'équation de Helmholtz :

$$\Delta \vec{E} + k^2 \vec{E} = 0$$

$$\Delta \vec{H} + k^2 \vec{H} = 0$$

le nombre d'onde k dans le vide étant tel que :

$$k^2 = \epsilon_0 \mu_0 \omega^2$$

Des conditions aux limites sont imposées sur les frontières du volume de calcul. Différents types de conditions peuvent être imposés :

La condition de mur électrique s'écrit :

$$\vec{E} \wedge \vec{n} = 0 \text{ ou } \vec{H} \cdot \vec{n} = 0$$

La condition de mur magnétique, utilisée dans certains cas de symétrie, s'écrit :

$$\vec{H} \wedge \vec{n} = 0 \text{ ou } \vec{E} \cdot \vec{n} = 0$$

Lorsque le métal n'est pas parfait, on place une condition d'impédance de surface.

Lors de simulations de milieux ouverts, une condition d'adaptation simulant l'espace libre est nécessaire.

Le problème physique est posé à partir de l'équation différentielle et des conditions aux limites. L'étape suivante consiste à construire une fonctionnelle. Selon le choix de la variable (champ électrique, ou champ magnétique), l'une des deux fonctionnelles peut être construite :

$$F(\vec{E}) = \iiint \left((\vec{\text{rot}} \vec{E})^2 - k^2 \epsilon_r \mu_r \vec{E} \cdot \vec{E} \right) dv$$

$$F(\vec{H}) = \iiint \left((\vec{\text{rot}} \vec{H})^2 - k^2 \epsilon_r \mu_r \vec{H} \cdot \vec{H} \right) dv$$

En pratique, pour le sujet qui nous concerne, les milieux sont non magnétiques et $\mu_r = 1$. La valeur du champ pour laquelle la fonctionnelle atteint un extremum donne la solution. Pour trouver le champ, on le décompose sur chacun des éléments en utilisant les fonctions d'approximation. Malheureusement les fonctions scalaires d'interpolation, présentées plus haut pour une grandeur scalaire telle que le potentiel, ne conviennent pas pour décrire le champ électrique ou le champ magnétique, car elles conduisent, dans certains cas, à des solutions non physiques aberrantes, nommées solutions parasites. De nombreuses explications sont données à ce phénomène. Sans entrer dans les développements mathématiques, on peut dire simplement que les fonctions d'interpolation présentées plus haut ne permettent pas de rendre compte de propriétés vectorielles. En particulier, un vecteur tel que le champ électrique est à divergence nulle. Or il est impossible de tenir compte de cette contrainte en utilisant simplement les fonctions d'approximation scalaires, à moins d'utiliser un terme de pénalisation qui alourdit beaucoup la méthode. Une méthode plus élégante consiste à utiliser des fonctions d'approximation vectorielles, encore appelées éléments de Whitney.

Avec des fonctions d'approximations scalaires on écrirait :

$$\vec{E}(x, y, z) = \sum_i \vec{E}_i a_i(x, y, z)$$

$\vec{E}_i$ représente les valeurs du champ aux nœuds d'interpolation.

Avec les éléments de Whitney, on écrit :

$$\vec{E}(x, y, z) = \sum_s e_s \vec{w}_s(x, y, z)$$

e_s représente les valeurs de la circulation du champ sur une arête et $\vec{w}_s(x, y, z)$ les éléments de Whitney qui sont définis par :

$$\vec{w}_s(x, y, z) = a_j(x, y, z) \vec{\nabla} a_i(x, y, z) - a_i(x, y, z) \vec{\nabla} a_j(x, y, z)$$

s est un entier qui caractérise une arête située entre les nœuds i et j . Il y a donc six fonctions d'approximation par tétraèdre pour des éléments finis à l'ordre 1. On les appelle éléments finis d'arête d'ordre 1.

Ce type d'éléments inséré dans l'expression de la fonctionnelle permet d'obtenir, après différentiation, un système linéaire dont l'inversion donne les circulations des champs sur les arêtes.

■ Formulation du problème pour les antennes

Les éléments finis sont très utilisés pour la simulation d'antennes. Le principe de leur utilisation est celui qui est décrit dans le paragraphe précédent. Cependant, la difficulté de la modélisation du rayonnement d'antennes est la définition du volume de calcul. En effet, on cherche à déterminer le rayonnement en champ lointain et, dans ce cas, le volume de calcul est trop grand. Ce problème est contourné en utilisant une forme du principe d'équivalence. On définit un volume V_C , appelé volume de calcul, entourant la source à une distance raisonnable (figure 11.6). Le volume de calcul est limité par la surface S .

La méthode des éléments finis est appliquée dans le volume de calcul en imposant des conditions de rayonnement sur la surface S . Ces dernières permettent théoriquement de simuler l'espace libre. Les champs sont complètement déterminés par éléments finis dans le volume de calcul et en particulier sur la surface S . D'après le principe d'équivalence, la connaissance du champ tangentiel sur la surface S permet de calculer le champ dans tout l'espace. Cette étape est appelée transformation champ proche champ lointain. Lorsque le champ lointain est calculé, il est facile de déterminer les caractéristiques de rayonnement de l'antenne. La précision attendue de la méthode réside dans un choix approprié de la distance de la surface S à l'antenne. Des études de convergence doivent être faites sur ce paramètre afin d'estimer la précision.

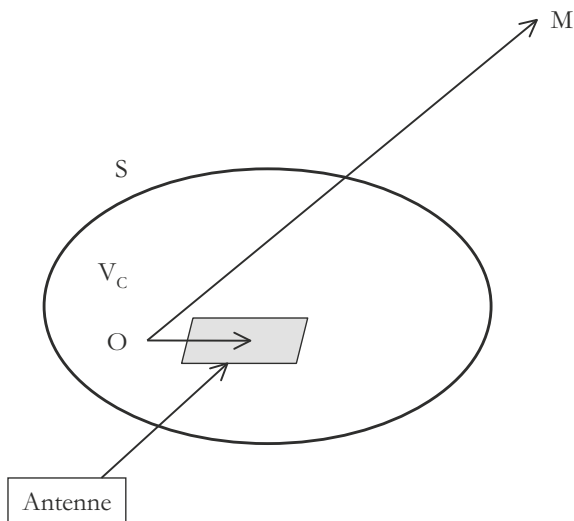


Figure 11.6 – Volume de calcul

Différentes conditions de rayonnement existent. Nous n'entrerons pas dans la description détaillée des opérateurs décrivant une condition d'espace libre, sachant que cela fait l'objet de nombreux travaux de recherche en analyse numérique. Les principes utilisés sont basés, soit sur une condition d'impédance sur la surface S , soit sur la représentation du champ en une décomposition de modes qui sont absorbés à un ordre plus ou moins élevé, soit sous forme d'une condition qui annule les ondes sur la surface S . Remarquons qu'aucune condition n'est

absolument parfaite. La performance d'une méthode se mesure au taux de réflexion des ondes sur la surface S : plus ce taux est faible, meilleurs seront les résultats pour la simulation d'antennes. Le problème du choix de la condition à imposer sur S pour limiter le volume de calcul est un problème général rencontré dans toutes les méthodes à maillage (méthode des différences finies par exemple).

■ Avantages de la méthode des éléments finis

Les simulateurs électromagnétiques commerciaux, basés sur la méthode des éléments finis, utilisent des éléments d'arêtes. Ils comportent, en plus du noyau de calcul concernant les éléments finis, de puissants logiciels graphiques, permettant des représentations en trois dimensions des dispositifs et des post-traitements pour le calcul des grandeurs dérivées, comme la puissance, l'impédance.

Certains logiciels travaillent en utilisant les champs, d'autres en utilisant les potentiels électromagnétiques. Tous sont, en tout cas, une aide précieuse à la simulation d'antennes car ils permettent des calculs de structures complexes rayonnantes en trois dimensions.

La méthode des éléments finis est une méthode très générale, permettant de prendre en compte des géométries très complexes. Elle est robuste. Pour obtenir la précision désirée, il suffit d'affiner le maillage aux endroits où les champs présentent des gradients importants. Ceci est fait dans les logiciels du commerce qui utilisent un maillage adaptatif. Cette méthode exige des ordinateurs possédant une mémoire importante et la résolution de problèmes non triviaux requiert des temps de calcul importants. Le résultat d'une simulation donne un point de fréquence. Pour obtenir une réponse dans toute une bande de fréquence, il est nécessaire de cumuler plusieurs simulations.

Toutes ces remarques font de la méthode des éléments finis une méthode puissante et précise dont le seul inconvénient est de nécessiter des moyens informatiques importants.

11.4 La méthode des différences finies dans le domaine temporel (FDTD)

La méthode des différences finies repose sur l'utilisation d'un maillage pour exprimer la discrétisation des équations de Maxwell. Le volume de calcul, défini par l'ensemble des mailles, est limité. Pour calculer les caractéristiques des antennes qui rayonnent à grande distance, une condition simulant l'espace libre doit être imposée à la frontière du maillage. Ce point est délicat et fait l'objet de nombreuses recherches. Nous donnerons quelques éléments sur ces conditions de rayonnements.

La méthode des différences finies repose sur une approximation des dérivées. On utilise pour cela la définition de la dérivée centrée, définie ci-dessous. Soit une fonction f à une dimension et f' sa dérivée :

$$f'(x) = \lim_{h \rightarrow 0} \left[\frac{f(x + h/2) - f(x - h/2)}{h} \right]$$

Son approximation devient, pour h petit :

$$f'(x) \approx \frac{f(x + h/2) - f(x - h/2)}{h} \quad [11.29]$$

La méthode des différences finies peut être appliquée dans le domaine fréquentiel ou dans le domaine temporel. Dans le domaine fréquentiel, elle résulte de la discrétisation de l'équation de Helmholtz et conduit à une simulation par point de fréquence.

La méthode temporelle, FDTD (*Finite Difference in Time Domain*) lui est actuellement préférée. Elle permet le suivi d'un signal électromagnétique au cours du temps dans un dispositif en trois

dimensions. Le calcul de la réponse fréquentielle du signal s'effectue par une transformation de Fourier. Sa facilité de mise en œuvre et sa souplesse d'adaptation en ont fait son succès.

11.4.1 Maillage du volume de calcul

La méthode consiste à mailler l'espace en parallélépipèdes dont les arêtes sont orientées selon les trois axes (Ox , Oy , Oz). Les six composantes du champ électromagnétique sont définies dans une maille élémentaire introduite par Yee définie sur la figure 11.7.

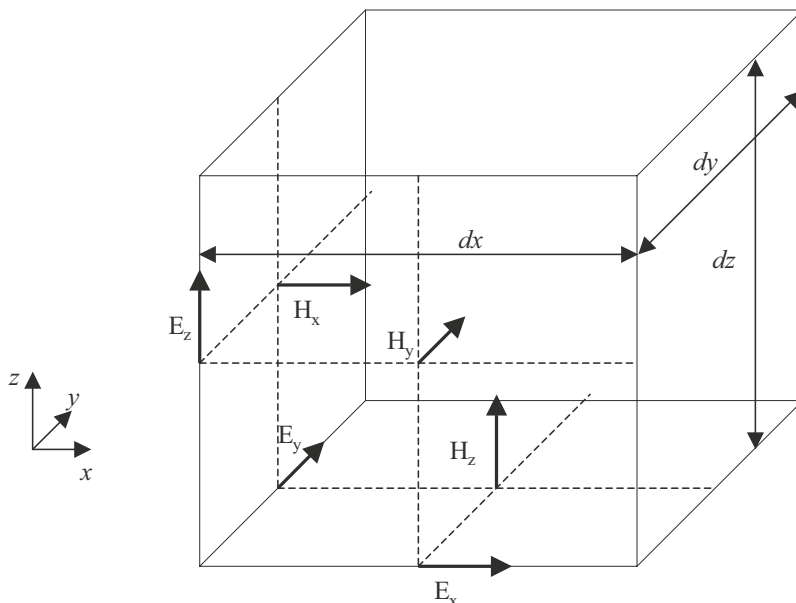


Figure 11.7 – Maille de Yee.

On constate sur cette figure que les composantes du champ électrique sont portées par les arêtes de la maille, alors que les composantes du champ magnétique sont situées au centre des faces et perpendiculairement à celles-ci. Cette disposition permet d'imposer facilement les conditions de passage à la traversée de surfaces séparant les diélectriques. Il est nécessaire de placer les faces tangentiellement aux interfaces entre deux matériaux.

11.4.2 Discrétisation

La méthode des différences finies repose sur une approximation des dérivées dans l'espace et dans le temps. On utilise pour cela la définition de la dérivée centrée, définie ci-dessus [11.29].

La FDTD traite les équations de Maxwell temporelles suivantes :

$$\begin{aligned}\vec{\text{rot}}(\vec{E}) &= -\frac{\partial \vec{B}}{\partial t} \\ \vec{\text{rot}}(\vec{H}) &= \vec{j} + \frac{\partial \vec{D}}{\partial t}\end{aligned}$$

Pour exposer la méthode, les constituants diélectriques sont supposés linéaires, homogènes, isotropes. Dans ce cas, le vecteur déplacement dans le matériau se met alors sous la forme :

$$\vec{D} = \epsilon \vec{E}$$

Les milieux considérés sont non magnétiques, ceci conduit à l'équation :

$$\vec{B} = \mu_0 \vec{H}$$

Compte tenu des remarques précédentes, les deux équations à traiter sont donc :

$$\vec{\text{rot}}(\vec{E}) = -\mu_0 \frac{\partial \vec{H}}{\partial t} \quad [11.30]$$

$$\vec{\text{rot}}(\vec{H}) = \vec{j} + \epsilon \frac{\partial \vec{E}}{\partial t} \quad [11.31]$$

On constate que ce sont des équations couplées dans l'espace et le temps, faisant intervenir le champ électrique $\vec{E}$ et le champ magnétique $\vec{H}$. Dans l'équation [11.31], on ne considérera, pour exposer la méthode, que l'équation en absence de courants. Ceux-ci peuvent être introduits ensuite en utilisant la loi d'Ohm locale, liant le courant au champ électrique.

Les dérivées des grandeurs apparaissant dans les équations de Maxwell sont remplacées par leurs approximations selon la définition de la dérivée centrée soit par rapport à l'espace, soit par rapport au temps.

Rappelons l'expression du rotationnel :

$$\vec{\text{rot}}(\vec{E}) = \begin{pmatrix} \frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} \\ \frac{\partial E_x}{\partial z} - \frac{\partial E_z}{\partial x} \\ \frac{\partial E_y}{\partial x} - \frac{\partial E_x}{\partial y} \end{pmatrix}$$

La première composante de l'équation [11.30] s'exprime selon :

$$\frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} = -\mu_0 \frac{\partial H_x}{\partial t} \quad [11.32]$$

Nous allons décrire en détail le processus de discrétisation sur cette équation. Cette équation doit être vérifiée à tout instant et en tout point de l'espace.

■ Discrétisation temporelle

Considérons l'équation [11.32] écrite à l'instant t . La dérivée de la composante en x du champ magnétique par rapport au temps, calculée au temps t , est exprimée selon la définition [11.29] de la dérivée centrée. Elle fait intervenir le champ magnétique au temps $t - dt/2$ et au temps $t + dt/2$, le temps étant découpé en intervalles dt :

$$\frac{\partial H_x}{\partial t}(t) \approx \frac{H_x(t + dt/2) - H_x(t - dt/2)}{dt}$$

Pour calculer le rotationnel du champ électrique au temps t , il est donc nécessaire de connaître le champ magnétique aux instants demi-entiers.

Le temps t est décomposé en pas de temps dt :

$$t = n \cdot dt$$

Le rotationnel du champ électrique est calculé à l'instant t , puisqu'il est égal à la dérivée temporelle du champ magnétique. Le champ électrique doit donc être déterminé à l'instant t , alors que le champ magnétique doit être déterminé aux instants décalés de $dt/2$ par rapport à t . Les grandeurs sont quantifiées temporellement et on utilisera par la suite les notations :

$$\vec{E}^n \quad \text{et} \quad \vec{H}^{n+1/2}$$

■ Discrétisation spatiale

Le rotationnel du champ électrique est évalué sur le même principe. Chacune des coordonnées du rotationnel [11.30] est exprimée, selon l'approximation de la dérivée centrée.

Prenons la première composante qui fait intervenir des dérivées des composantes du champ électrique par rapport à y et à z . Calculons tout d'abord la dérivée de E_z par rapport à y . Selon la maille de Yee, on constate que la composante du champ électrique selon z se situe sur l'arête de la maille, en y , alors que le champ magnétique se trouve au centre de la face correspondante, en $y + dy/2$. La dérivée $\frac{\partial E_z}{\partial y}$ doit donc être évaluée en $y + dy/2$ puisqu'elle est liée à la composante en x du champ magnétique. Pour calculer la dérivée centrée en $y + dy/2$ il faut donc aussi faire intervenir la composante E_z en $y + dy$ qui appartient à la maille suivante :

$$\frac{\partial E_z}{\partial y}(y + dy/2) \approx \frac{E_z(y + dy) - E_z(y)}{dy}$$

Les composantes interviennent dans le calcul selon le schéma suivant :

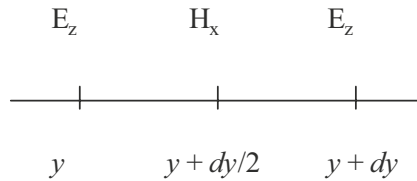


Figure 11.8 – Schéma de discrétisation spatiale

Le calcul s'effectue de la même façon pour la dérivée selon z de la composante du champ électrique selon y . Il faut avoir recours à la composante dans la cellule placée en $z + dz$.

Lors de la discrétisation spatiale, les cellules sont repérées par des entiers : i, j, k , selon les trois directions. Les distances sont aussi quantifiées selon les relations suivantes :

$$x = i \cdot dx, \quad y = j \cdot dy, \quad z = k \cdot dz.$$

La projection selon x du rotationnel du champ électrique, au temps $t = n \cdot dt$ s'écrit donc, après discrétisation spatiale :

$$\frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} \approx \frac{E_z^n(i, j+1, k) - E_z^n(i, j, k)}{dy} - \frac{E_y^n(i, j, k+1) - E_y^n(i, j, k)}{dz}$$

L'expression continue du rotationnel est donc remplacée par une expression formée de grandeurs discrètes qui sont les seules à pouvoir être traitées par ordinateur.

La forme discrétisée en espace et en temps de l'équation [11.30] se décompose en :

$$\begin{aligned}
 H_x^{n+1/2}(i, j, k) &= H_x^{n-1/2}(i, j, k) \\
 &\quad - \frac{dt}{\mu_0} \left[\frac{E_z^n(i, j+1, k) - E_z^n(i, j, k)}{dy} - \frac{E_y^n(i, j, k+1) - E_y^n(i, j, k)}{dz} \right] \\
 H_y^{n+1/2}(i, j, k) &= H_y^{n-1/2}(i, j, k) \\
 &\quad - \frac{dt}{\mu_0} \left[\frac{E_x^n(i, j, k+1) - E_x^n(i, j, k)}{dz} - \frac{E_z^n(i+1, j, k) - E_z^n(i, j, k)}{dx} \right] \\
 H_z^{n+1/2}(i, j, k) &= H_z^{n-1/2}(i, j, k) \\
 &\quad - \frac{dt}{\mu_0} \left[\frac{E_y^n(i+1, j, k) - E_y^n(i, j, k)}{dx} - \frac{E_x^n(i, j+1, k) - E_x^n(i, j, k)}{dy} \right]
 \end{aligned}$$

Il existe un autre groupe d'expressions discrétisées correspondant à l'équation [11.31] :

$$\begin{aligned}
 E_x^{n+1}(i, j, k) &= E_x^n(i, j, k) \\
 &\quad + \frac{dt}{\epsilon} \left[\frac{H_z^{n+1/2}(i, j, k) - H_z^{n+1/2}(i, j-1, k)}{dy} - \frac{H_y^{n+1/2}(i, j, k) - H_y^{n+1/2}(i, j, k-1)}{dz} \right] \\
 E_y^{n+1}(i, j, k) &= E_y^n(i, j, k) \\
 &\quad + \frac{dt}{\epsilon} \left[\frac{H_x^{n+1/2}(i, j, k) - H_x^{n+1/2}(i, j, k-1)}{dz} - \frac{H_z^{n+1/2}(i, j, k) - H_z^{n+1/2}(i-1, j, k)}{dx} \right] \\
 E_z^{n+1}(i, j, k) &= E_z^n(i, j, k) \\
 &\quad + \frac{dt}{\epsilon} \left[\frac{H_y^{n+1/2}(i, j, k) - H_y^{n+1/2}(i-1, j, k)}{dx} - \frac{H_x^{n+1/2}(i, j, k) - H_x^{n+1/2}(i, j-1, k)}{dy} \right]
 \end{aligned}$$

On constate que les indices repérant les cellules ne sont pas pris dans le même ordre pour les deux groupes d'équations. Ceci est dû au fait que dans ce dernier groupe d'équations, le champ électrique est calculé en fonction du rotationnel du champ magnétique. Or les composantes du champ électrique d'une cellule sont comprises entre les composantes du champ magnétique de la même cellule et les composantes du champ magnétique de la cellule précédente.

Le principe de la méthode repose sur le calcul, par les formules précédentes, du champ électrique aux instants entiers. Le champ magnétique est alors calculable à l'instant demi-entier suivant puisque toutes les composantes sont connues aux instants précédents. Le calcul se poursuit par pas de temps de $dt/2$. La méthode FDTD est une méthode itérative explicite.

11.4.3 Excitation

Il reste alors à initialiser le calcul. Pour cela, on impose une excitation sous forme d'un champ électrique correctement choisi selon le type de ligne.

Ainsi pour une ligne micro ruban, par exemple, un champ électrique uniforme est imposé juste sous la ligne, entre la masse et la ligne, dans le plan d'excitation (figure 11.9).

Pour une ligne coplanaire, deux champs électriques opposés sont imposés dans chaque fente dans le plan d'excitation. Cela constitue le mode fondamental de propagation d'une ligne coplanaire. Si cette ligne présente des dissymétries, le mode fente apparaît et il perturbe, en général, le bon fonctionnement de l'antenne.

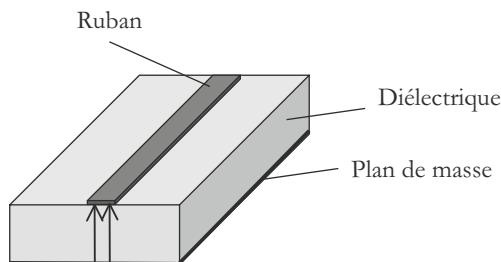


Figure 11.9 – Excitation d'une ligne micro ruban

On choisit, pour l'excitation, des champs électriques proches des modes fondamentaux des lignes. Même si le champ imposé n'est pas rigoureusement celui du mode fondamental, car sa description n'en est pas très précise (effets de bord non pris en compte), le calcul itératif permet de rétablir la forme réelle du champ le long de la ligne. Cette description spatiale de l'excitation ne suffit pas car elle ne porte pas d'information temporelle. C'est pourquoi on multiplie cette excitation spatiale par une fonction $u(t)$ variant avec le temps.

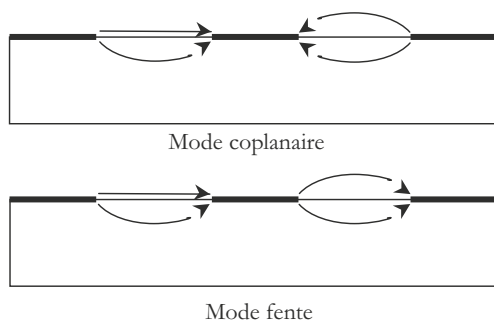


Figure 11.10 – Mode coplanaire et mode fente d'une ligne coplanaire

Si l'excitation est imposée uniquement à l'instant initial, puis supprimée, elle porte toutes les fréquences puisqu'elle correspond à la transformée de Fourier d'un signal sous forme d'une impulsion de Dirac.

Les lignes ne sont en fait utilisées que dans une certaine bande de fréquences, par exemple du continu à une fréquence maximale $f_{\max}$. On choisit alors une excitation temporelle de forme gaussienne :

$$u(y, t) = u_0 e^{-a^2(t-t_0)^2}$$

Étant donné que la transformée d'une gaussienne est une gaussienne, le spectre en fréquence a la forme d'une gaussienne centrée sur la fréquence nulle. La fréquence maximale considérée est en rapport avec le coefficient a . On choisit $f_{\max} = a/2$. Cela permet d'obtenir la réponse fréquentielle dans une grande partie du spectre jusqu'à la fréquence maximale.

L'avantage de ce type d'excitation est d'obtenir un signal limité dans le temps. La transformée de Fourier en est facilitée car l'intégrale numérique à effectuer est limitée.

Lorsque les lignes ne laissent passer qu'une bande de fréquences, un signal sinusoïdal modulé par une gaussienne est mieux adapté. Son spectre est une gaussienne centrée sur la fréquence de la sinusoïde dont la largeur est celle de la transformée de Fourier de la gaussienne. On choisit la fréquence de la sinusoïde en milieu de bande.

L'excitation est imposée au début du calcul, puis par itérations successives, le champ électromagnétique est calculé dans tout le volume pour chaque pas de temps. Les caractéristiques fréquentielles des dispositifs sont obtenues par transformée de Fourier

11.4.4 Conditions particulières à la traversée de deux diélectriques

La présence d'une interface diélectrique impose des conditions particulières, déduites des équations de Maxwell. L'équation :

$$\vec{\text{rot}}(\vec{H}) = \epsilon \frac{\partial \vec{E}}{\partial t}$$

requiert un traitement spécial puisque la permittivité varie brutalement à la traversée du diélectrique. Il suffit de décomposer cette équation dans les deux milieux et d'écrire séparément :

$$\begin{aligned}\vec{\text{rot}}(\vec{H}_1) &= \epsilon_1 \frac{\partial \vec{E}_1}{\partial t} \\ \vec{\text{rot}}(\vec{H}_2) &= \epsilon_2 \frac{\partial \vec{E}_2}{\partial t}\end{aligned}$$

On montre que ces équations sont satisfaites en imposant, dans la couche intermédiaire, une valeur de la permittivité égale à la moyenne des permittivités des matériaux situés de part et d'autre.

11.4.5 Conditions aux limites particulières

■ Plan métallique

Sur un plan métallique parfait, le champ électrique tangentiel est nul. Il suffit alors d'imposer cette condition sur les composantes du champ se situant sur le plan.

■ Conditions d'adaptation

Dans certaines zones, il est nécessaire d'imposer des conditions d'adaptation. Ces conditions sont introduites pour simuler la propagation d'une onde vers l'infini. On pourra ainsi modéliser des lignes infinies ou des milieux ouverts.

Les conditions d'adaptation peuvent être très complexes. L'objet de ce paragraphe n'est pas de toutes les décrire. Nous présentons ici la plus simple, la condition de Mur de premier ordre, qui donne une bonne approximation de l'adaptation. Cette condition s'applique sur un plan et rend compte du fait que l'onde se propage dans un sens avec une vitesse de phase v_ϕ . Dans le cas d'une onde partant dans la direction x , sans réflexion, cette équation de propagation s'écrit :

$$\frac{\partial \vec{E}}{\partial x} - \frac{1}{v_\phi} \frac{\partial \vec{E}}{\partial t} = 0$$

Elle se décompose, dans le plan perpendiculaire, en $x = 0$, où s'applique la condition d'adaptation, selon les deux équations portant sur les composantes, en y et z , du champ électrique, tangentielles au plan, selon :

$$E_{y,z}^{n+1}(0, j, k) = E_{y,z}^n(1, j, k) + \frac{v_\phi dt - dx}{v_\phi dt + dx} [E_{y,z}^{n+1}(1, j, k) - E_{y,z}^n(0, j, k)] \quad [11.33]$$

Cette condition n'est satisfaisante que pour des structures non dispersives, pour lesquelles la vitesse de phase ne dépend pas de la fréquence. On peut cependant l'utiliser dans une bande de fréquences où la vitesse de phase présente une faible variation.

■ Conditions de rayonnement

Comme nous l'avons déjà dit, le volume de calcul étant fini, il faut, pour simuler un rayonnement dans l'espace, placer à la frontière de celui-ci une condition dite de rayonnement. Cette condition traduit le fait que l'onde part du volume, sans réflexion à la frontière.

La condition d'adaptation vue dans le paragraphe précédent pourrait être utilisée. C'est une condition approximative, dite du premier ordre. L'onde se propageant dans le vide, la vitesse de phase est remplacée par la vitesse de la lumière. Elle donne de bons résultats pour un rayonnement arrivant majoritairement perpendiculairement au plan servant de frontière. Dans de nombreux cas, le rayonnement est à symétrie sphérique et cette approximation est grossière, sauf si la frontière est placée très loin. Cela entraîne alors un volume de calcul important. De plus, pour les ondes n'arrivant pas perpendiculairement au plan frontière, la condition ne traduit pas une bonne adaptation, puisque l'angle d'incidence doit intervenir dans la condition.

C'est pourquoi, d'autres conditions de rayonnement ont été proposées. Ces conditions traduisent le fait que le milieu entourant le volume de calcul absorbe complètement le rayonnement sortant. On les appelle des conditions d'absorption ou en anglais : *Absorbing Boundary Conditions* (ABC). Il existe de nombreuses variantes de ces conditions.

À l'heure actuelle, il est admis que la méthode PML (*Perfectly Matched Layer*) proposée par J. P. Bérenger est l'une des plus efficace. Elle consiste à considérer qu'une onde sortant du volume de calcul rencontre un milieu non physique dont les propriétés imposent à l'onde de se propager à la vitesse c de la lumière. Nous ne présenterons pas ici en détail cette méthode qui sort du cadre de l'ouvrage, mais nous en donnerons quelques principes.

Le matériau artificiel entourant le volume de calcul présente une conductivité électrique et une conductivité magnétique. Il est d'autre part anisotrope. Cela entraîne des valeurs différentes des conductivités selon les directions. Afin de vérifier que la vitesse de propagation de l'onde est égale à c , il est nécessaire que le milieu vérifie la condition suivante entre les différentes composantes des conductivités électrique σ et magnétique σ^* :

$$\frac{\sigma_{x,y,z}}{\epsilon_0} = \frac{\sigma_{x,y,z}^*}{\mu_0}$$

Chaque composante des champs électrique et magnétique est décomposée en tenant compte des directions transverses. Ainsi la composante en x du champ électrique est décomposée en :

$$E_{xy} \text{ et } E_{xz}$$

La discrétisation des équations de Maxwell conduit alors à douze équations qui sont résolues par la méthode itérative proposée plus haut.

Il est nécessaire de prévoir dans le calcul quelques couches (de 5 à 10) de ce matériau absorbant afin de modéliser un milieu infini. Cette modification à la frontière du volume entraîne cependant une réflexion, même si elle est très faible. Afin d'atténuer cette réflexion un artifice est utilisé qui consiste à augmenter progressivement la conductivité des différentes couches de façon à passer d'une couche dont les propriétés sont très proches de celles du vide à celle d'un matériau très absorbant.

11.4.6 Calcul de la constante de propagation

Supposons qu'une onde se propage dans la direction x dans une structure de guidage. Le champ électrique évolue dans la direction de propagation selon l'exponentielle de la phase Φ :

$$\Phi = \omega t - \beta x$$

Il suffit d'utiliser les champs en x et en $x + l$, pour une fréquence donnée, pour pouvoir calculer la constante de propagation β car la relation suivante est vérifiée :

$$\vec{E}(x + l, y, z) = \vec{E}(x, y, z) e^{-j\beta l}$$

On en déduit :

$$\beta = \frac{1}{jl} \ln \left[\frac{\vec{E}(x, y, z)}{\vec{E}(x + l, y, z)} \right]$$

Les valeurs des champs électriques aux points situés en x et $x + l$, à une fréquence donnée, sont obtenues en effectuant une transformée de Fourier sur le signal temporel résultant de la FDTD. Au lieu d'utiliser le champ en un point, il est courant de calculer la circulation du champ entre la masse et le ruban métallique. On revient alors à la notion de différence de potentiel. Lorsque l'approximation en potentiels ne suffit pas, on calcule la puissance en utilisant l'expression du vecteur de Poynting.

11.4.7 Calcul des coefficients des paramètres S

La méthode pour calculer les paramètres S de la matrice de répartition ressemble beaucoup à celle qui vient d'être présentée. Il suffit de calculer, à une fréquence donnée, les potentiels obtenus par circulation du champ électrique sur chaque plan de référence de la matrice de répartition.

La connaissance de la matrice de répartition permet de remonter aux paramètres de l'antenne, tels que son coefficient de réflexion et donc son impédance d'entrée.

11.4.8 Diagramme de rayonnement

La détermination du diagramme de rayonnement de l'antenne simulée par FDTD, s'effectue en utilisant une transformation champ proche — champ lointain. Cette transformation consiste à considérer le volume de calcul comme une boîte, appelée boîte de Huyghens, sur les parois de laquelle on connaît les champs par FDTD. On applique alors le principe d'équivalence pour déterminer les courants équivalents sur cette frontière. Le rayonnement des courants équivalents détermine le champ à grande distance. Cette dernière partie du calcul utilise les expressions présentées dans le chapitre 3 concernant le rayonnement des courants.

11.4.9 Conclusions sur la méthode

La méthode des différences finies temporelle est bien adaptée à la simulation d'antennes puisqu'elle permet à la fois d'en connaître les caractéristiques du point de vue du circuit (adaptation, impédance d'entrée) et du point de vue du rayonnement.

C'est une méthode facile à mettre en œuvre, qui permet d'obtenir rapidement des résultats dans une bande de fréquences. Elle permet d'obtenir des résultats précis à condition d'utiliser une mémoire informatique importante et de disposer de temps de calcul. Le maillage, étant de type parallélépipédique, ne permet pas de générer toutes les formes géométriques avec précision. Il est mieux adapté aux géométries présentant des angles droits que des courbes. Plus le maillage est fin, plus la précision attendue est grande.

Dans ce bref exposé de la méthode, nous n'avons décrit la méthode que sur des matériaux sans pertes. Il est possible de prendre en compte les pertes en complexifiant légèrement l'algorithme. La force de la FDTD réside dans sa simplicité de mise en œuvre et dans le fait que les résultats sont obtenus en fonction du temps. Il suffit alors d'une transformée de Fourier pour déterminer les paramètres fréquentiels.

Bibliographie

BOSSAVIT A. — *Électromagnétisme, en vue de la modélisation*, Paris, Springer Verlag France, 1993.

DHATT G, TOUZOT G. — *Une présentation de la méthode des éléments finis*, Paris, Maloine et Québec, Les Presses de l'université Laval, 1981.

MATTHEW S. — *Numerical techniques in electromagnetics, 2nd edition*, CRC Press LLC, Floride, USA, 2000

MOSIG J. R. — *Comparison of Quasi-static and Exact Electromagnetic Fields from a Horizontal Electric Dipole Above a Lossy Dielectric Backed by an Imperfect Ground Plane*. IEEE Transactions on microwave theory and techniques, vol. 34, n° 4, avril 1986, pp. 379-387.

TAFLOVE A. — *Computational Electrodynamics : The Finite-Difference Time-Domain Method*, Boston London, Artech House Publishers, 1995.

WONG M.F, PICON O., FOUAD HANNA V. — *A Finite Element Method Based on Whitney Forms to solve Maxwell Equations in the Time Domain*, IEEE Transactions on Magnetics, vol.31, n° 3, may 1995, pp. 1618-1621.

A

- absorbant, 325
- adaptation, 93
- aire d'absorption, 83
- alimentation, 94
 - de réseaux, VI, 121
 - par couplage, 291
 - par fente, 96
 - par ligne imprimée, 291
 - par proximité, 291
- altimètre, 178
- analyseur vectoriel de réseau, 334
- ANFR, 127
- angle
 - solide, 76
 - solide de l'antenne, 76
- antenne
 - à balayage, 112
 - à diversité, 214
 - à méandres, 297
 - à onde de fuite, 260
 - à onde de surface, VI, 260
 - à ondes progressives, 256, 302
 - à réflecteur, 270
 - à résonateur, 253
 - à résonateur BIE, 255
 - à résonateur diélectrique, 253
 - active, 8
 - bande étroite, 300
 - bande passante, 291
 - biconique, 304
 - boucle, 6, 48, 249
 - Bow-Tie, 305, 318
 - Cassegrain, 7, 154, 273
 - cornet, 6, 277
 - cornet conique, 285
 - d'émission, 5
 - de réception, 5
 - demi-onde, 109, 111
 - Diamond, 318
 - dipolaire, 5
 - double C, 298
 - en bandes décadiques, 138
 - en C, 297
 - en E, 298
 - en H, 298
 - en S, 298
 - en V, 259
 - filaire, 44
 - formée, 157
 - fractale, 308
 - grégorienne, 157
 - hélicoïdale, 303
 - hybride, 9
 - IFA, 295
 - intégrée, 9
 - intelligente, 8
 - isotrope, 77
 - large bande, 99, 300
 - log périodique, 305
 - miniature, 294, 295
 - monopolaire, 295
 - multifaisceaux, 115
 - panneaux, 144, 164
 - patch, 8
 - pertes, 335
 - PIFA, 296, 297
 - planaire, 288
 - plaquée, 8
 - pyramidale, 285
 - reconfigurable, 115
 - réflecteur parabolique, 7
 - rhombique ou losange, 259
 - sectorielle, 143
 - sectorielle plan E, 279
 - sectorielle plan H, 282
 - spirale, 306

spirale conique, 307
 Vivaldi, 308
 Yagi-Uda, 264
 apodisation, 327
 approximation nodale, 350
 assemblage, 354
 axe de rayonnement, 74

B

bande
 de cohérence, 221, 222
 EHF, 129
 HF, 129
 ISM, 128, 129
 LF, 128
 MF, 128
 passante, 97
 SHF, 129
 UHF, 129
 VHF, 129
 VLF, 128
 bilan de liaison, 99
 BLAST, 221
 DBLAST, 244
 VBLAST, 244
 blocage de l'ouverture, 152
 boîte de Huyghens, 365
 Bracewell
 relation de, 331
 brillance, 187

C

cône d'ouverture, 81
 canal
 capacité, 313
 MIMO, 227
 capacité
 canal MIMO, 221, 237
 canal MIMO déterministe, 238
 d'un canal, 237
 de coupure, 239, 240
 ergodique, 239
 cellule, 141
 centre de phase, 104
 chambre anéchoïque, 325

champ
 de Neumann, 22
 diffracté, 51
 électromoteur, 22
 lointain, 69–71
 proche (zone de $\sim$), 69
 rayonné (ouverture rectangulaire), 67
 tangential, 65
 très proche, 69
 charges images, 34
 codage spatio-temporel, 220, 221, 235, 240, 244, 245
 code
 d'Alamouti, 235, 240
 spatio-temporels linéaires, 242
 coefficient
 d'obliquité, 70
 de réflexion, 133
 combinaison de signaux, 218
 combineurs, 122
 condition
 d'absorption, 364
 d'adaptation, 363
 de rayonnement, 356, 363
 de rayonnement de Sommerfeld, 342
 PML, 364
 conservation
 de l'énergie, 29
 coplanaire, 255
 corps noir, 187, 190
 corrugations, 286
 couplage, 302
 couverture radar, 171

D

DBLAST, 244
 débordement (*spill-over*), 152
 décalage angulaire, 87
 densité surfacique de puissance, 29, 73
 directivité, 78
 déphasage linéique, 113
 déphaseurs, 121
 à ferrites, 122
 à lignes, 122
 en parallèle, 121
 en série, 121
 diagramme de rayonnement, 42, 74
 calcul, 105

diagramme diversité de, 216
 diffraction, 9, 51, 213
 ouverture rectangulaire, 57
 diffusion de Mie, 178
 diffusomètre, 179
 directivité, 44, 78
 d'une antenne, 77
 discrétisation, 347
 distance intercapteurs, 198
 distribution
 et contrôle optique, 124
 spatiale d'intensité, 200
 diversité, VI, 213, 231
 d'antennes en émission, 234
 d'antennes en réception, 233
 d'espace, 233
 de trajets, 232
 fréquence, 217
 ordre, 231
 temporelle, 232
 diviseurs, 122
 doublet électrique, 41
 dualité grandeurs électriques et magnétiques,
 54, 55
 dyade, 339

E

écartométrie, 172
 effet de sol, 138
 effet Doppler, 220, 221, 223
 effet Pockels, 124
 efficacité, 92, 335
 dans le lobe principal, 77
 de l'antenne, 82
 de rayonnement, 80
 éléments
 de Whitney, 355
 parasites, 216
 ellipse de polarisation, 87
 émissivité, 190
 environnement
 indoor, 214
 épaisseur de peau, 19
 étalement
 Doppler, 221, 223, 225
 temporel, 221, 222, 225

F

facteur
 d'adaptation, 89
 d'antenne, 83
 d'obliquité, 68
 de mérite, 158
 de réduction, 297
 fading, 228
 de Rayleigh, 228, 229
 de Rice, 228
 ergodique, 228
 non ergodique, 228
 par bloc, 228
 fauchée, 185
 filtrage spatial, 208
 filtre spatial, 201
 fonction
 de filtrage, 318
 de Green du vide, 26
 objectif, 119
 fonction caractéristique
 de rayonnement, 74
 de rayonnement normalisée, 74
 fonctionnelle, 352
 formation
 de faisceau, 117
 de voies, 202
 formulation vectorielle, 61
 formule de Friis, 100
 fréquence
 diversité, 217
 Doppler, 183
 Fresnel
 ellipsoïde, 135
 formules, 133

G

gain, 82, 330
 d'une antenne, 82
 de multiplexage, 245
 en débit, 220
 en diversité, 216, 218, 220, 245
 linéaire, 100
 Green
 fonctions, 339
 tenseur, 342, 344
grid-dip, 139

H-I

hauts débits, 213
 identité de Green, 24, 25
 impédance, 333
 d'entrée, 92
 d'onde, 21
 mutuelle, 92
 propre, 92
 intégrale de Fresnel-Kirchhoff, 68
 intensité, 77

J-K

jauge de Lorentz, 23
 kit de calibration, 334

L

largeur de bande, 97
 lentilles, 115
 ligne micro ruban, 95
 lignes à retard, 125
 lobe
 principal, 76
 secondaire, 76, 81
 loi
 d'illumination, 80
 de Planck, 187
 Rayleigh Jeans, 189
 LOS (*Line Of Sight*), 227
 luminance, 188

M

macrocellules, 141
 matériau
 à pertes, 32
 anisotrope, 16
 BIE, 256
 BIE ou BEG, 298
 homogène, 16
 isotrope, 16
 linéaire, 16
 non linéaire, 16
 parfait, 16
Maximum Ratio Combining, 233
 méthode
 à haute résolution, 209
 autorégressive, 209
 de Capon, 206

de directivité/gain, 335
 de Pisarenko, 212
 de Weeler cap, 335
 des éléments finis, 337, 350
 des différences finies, 337, 357
 des moments, 337, 338, 346
 intégrale, 338, 349
 MUSIC, 212
 radiométrique, 335
 variationnelle, 352

micro ruban, 253
 MIMO, 220
 modulateur
 de Mach-Zender, 124
 direct, 124
 externe, 124
 monopôle, 295
 multibonds, 138
 multiplexage spatial, 220, 221, 236
 multitrajets, 314
 mur
 électrique, 355
 magnétique, 355

N-O

NLOS (*Non Line Of Sight*), 227
 octave, 99
 onde
 entrante, 27, 63
 évanescence, 64, 69
 plane, 21, 73
 propagée, 64
 sortante, 27, 63
 stationnaire, 93
 ondulation, 327
 ordre de diversité, 231
 ouverture
 angulaire, 76
 circulaire, 59
 diffractante, 51
 plane, 61
 rectangulaire, 56

P

parabole, 149
 paramètres de Stokes, 90

- patch
 - circulaire, 95
 - presque carré, 94
 - rectangulaire, 94
 - perméabilité
 - du vide, 12
 - magnétique, 16
 - relative, 16
 - permittivité, 16
 - du vide, 12
 - relative, 16
 - phase stationnaire, 66
 - polarisation
 - circulaire, 88
 - circulaire droite, gauche, 292
 - co-polarisation, 91
 - croisée, 91
 - diversité de, 216
 - elliptique, 85
 - H et V, 85
 - rectiligne, 84, 88
 - sens, 87
 - TE, 134
 - TE, TM, 85
 - TM, 133
 - potentiel
 - avancé, 27
 - électromagnétique, V, 21
 - retardé, 24, 27
 - scalaire, 22
 - vecteur, 21
 - pouvoir de séparation, 158
 - Power Imbalance*, 214
 - principe
 - d'équivalence, 53
 - d'Huyghens, 9, 51
 - probabilité de coupure, 240
 - problème équivalent, 35, 52
 - champ électrique, 53
 - champ magnétique, 52
 - problème des trajets multiples, 134
 - puissance
 - de bruit, 100
 - isotrope rayonnée équivalente, 83
 - PIRE, 83, 157
- R**
- RADAR, 167
 - radar, 114, 176
 - à ouverture latérale, 184
 - à synthèse d'ouverture, 179, 186
 - bi-statique, 180
 - de poursuite, 168
 - de recherche, 168
 - Doppler, 182
 - imageur, 179, 185
 - interférométrique, 180
 - météorologique, 178
 - modulé continûment en fréquence, 183
 - mono-statique, 180
 - pulsé, 181
 - radioamateurs, 135
 - radioastronomie, 194
 - radiométrie, 176
 - hyperfréquence, 186
 - radiomètre, 180
 - rapport
 - axial, 88
 - réactance, 92
 - récepteurs, 243
 - linéaires, 243
 - MMSE, MMSE-DFE, MMSE-SIC, 243
 - non-linéaires, 243
 - SIC, 243
 - ZF Zero Forcing, 243
 - réciprocité, 5, 32, 34
 - recombinaison
 - des signaux en réception, 218
 - reconfigurabilité, 114
 - réflecteur, 115, 268
 - réflexion minimale, 93
 - réseau, 103
 - à fentes, 7
 - cellulaire, 141
 - cosécantés, 118
 - de communications, 140
 - facteur ($\sim$ de), 104
 - hertzien, 140
 - microcellulaire, 141
 - mobiles, 140
 - picocellulaire, 141
 - reconfigurable, 103
 - synthèse, 118

résistance
 d'antenne, 92
 d'entrée, 92
 de pertes, 290
 de rayonnement, 44, 46, 92, 290

S

satellites, 144
 sélectivité en fréquence, 221, 222, 228, 229, 236, 238
 sensibilité, 158
 signature, 176
 sondeurs ionosphériques, 179
 source, 12
 décalée, 153, 272
 spectre, 127, 176
 d'ondes planes
 propriétés, 64
 radioélectrique, 128
 sphère de Poincaré, 89
 SRD, 159
 superstrat, 256, 299
 surface
 effective, 79
 rayonnante, 10
 système de type Oméga, 129

T

Té magique, 122
 télédétection passive, 186
 température
 apparente, 191
 d'antenne, 191
 de brillance, 190
 de bruit, 100
 temps
 de cohérence, 221, 223
 diversité, 218
 terme d'obliquité, 71
 théorème
 d'Ampère, 15
 d'Ostrogradski, 13, 14
 d'unicité, 30
 de Gauss, 14

de Lorentz, 34
 de Nyquist, 190
 de réciprocité, 33
 de Stokes, 15
 de Van Cittert-Zernicke, 200

théorie
 des images, 34
 des rayons, 149
 tilt, 144
 toit capacitif, 295
 trajet
 multiple, 130, 213
 optimal, 130
 transducteur, 5
 transformateur imparfait, 82
 transmission ULB (UWB), 310
 trou de serrure, 229

U-V

UIT, 127
Uncorrelated scattering, 225
 unicité, 52
 variation
 angulaire, 74
 linéique de phase, 108
 radiale, 74
 VBLAST, 244
 vecteur
 de Poynting, V, 27
 directionnel, 198
via hole, 94
 vide, 11

W-Y-Z

WSS *Wide Sense Stationary*, 225
 Yee maille de, 358
 zone
 calme ou tranquille, 326
 de couverture, 157
 de rayonnement, 323
 de visibilité, 107
 Fraunhofer, 70
 Fresnel, 69
 Rayleigh, 69

Odile Picon et coll.

Préface de Maurice Bellanger

LES ANTENNES

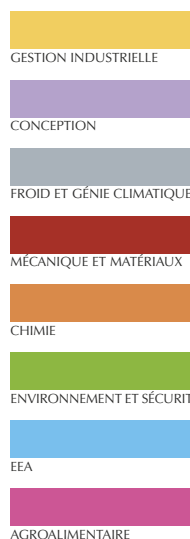
Théorie, conception et applications

La course à l'innovation concernant les systèmes de communication entraîne des études poussées dans le domaine des antennes. Dans ce contexte, les **méthodes de conception, de mesures et de modélisation** constituent une aide considérable.

Cet ouvrage est conçu comme un outil de travail à l'attention des ingénieurs et des techniciens confrontés à des choix concernant le type d'antennes à concevoir. Il propose successivement :

- Des **bases théoriques** sur le fonctionnement des antennes.
- Un **classement** des nombreux types d'antennes existants (antennes résonnantes, à ondes de fuite, Yagi-Uda...) avec leurs caractéristiques et leurs utilisations ainsi que de nombreux **aspects très innovants** tels que la mise en œuvre de systèmes complexes multi-antennes (traitement d'antenne, diversité, MIMO).
- Des **éléments de conception applicables** à chaque type, complétés par des **principes de mesure et de simulation numérique**.

L'ensemble des chapitres s'organise en partant du principe général du rayonnement des antennes pour aller vers leurs diversités. Cet ouvrage peut être lu, soit de façon théorique comme base aux calculs de rayonnement, soit de façon pratique dans le but de choisir un type d'antennes et de poursuivre jusqu'à sa conception. Il contient les éléments indispensables pour des études en recherche et développement dans le domaine des antennes.



ODILE PICON,
LAURENT CIRIO,
CHRISTIAN RIPOLL,
GENEVIÈVE BAUDOIN,
JEAN-FRANÇOIS
BERCHER,
MARTINE VILLEGAS.

Cet ouvrage est le résultat du travail d'une équipe d'auteurs, spécialistes en antennes, électromagnétisme, systèmes électroniques et traitement du signal, tous enseignants chercheurs à l'université de Paris-Est. Les auteurs ont une large expérience de l'enseignement des méthodes exposées et de leur mise en œuvre. Ils enseignent pour la plupart dans le master *Systèmes de communications hautes fréquences de l'UPE*. Leur expérience est issue des travaux de recherche menés en partie en collaboration avec des partenaires du milieu industriel.

L'USINE NOUVELLE

