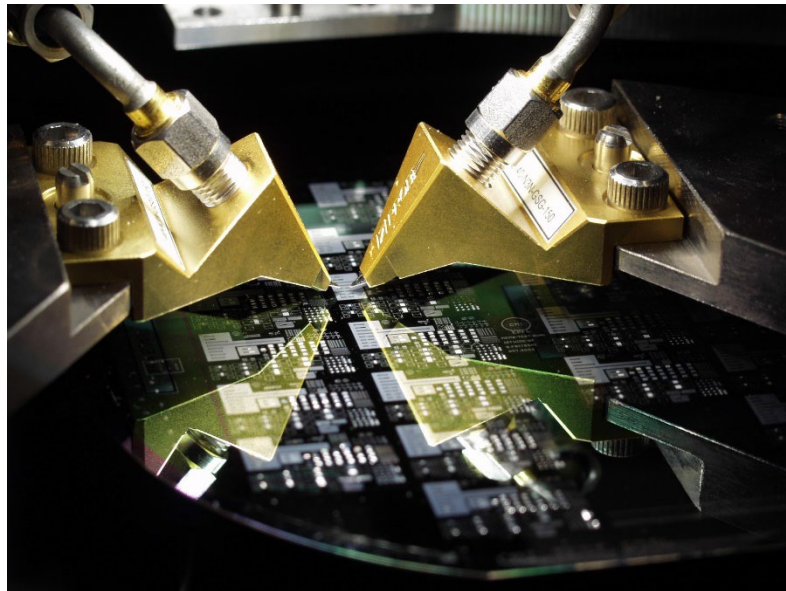


NOTES DE COURS

"Introduction aux microondes et antennes"

SEL, bachelor 5^{ème} semestre



Prof. A. Skrivervik, IEM-MAG

1. Introduction

1.1 Objectif

L'objectif de ce cours est de donner un aperçu des problèmes électromagnétiques et hyperfréquences liés à la transmission guidée et aux circuits hyperfréquences, de manière à dégager une vue d'ensemble du problème. Les thèmes suivants seront abordés :

- la transmission guidée dans les câbles coaxiaux et les guides d'ondes;
- la transmission par circuits imprimés hyperfréquences, notamment les circuits microrubans et les guides d'ondes coplanaires;
- la caractérisation des circuits hyperfréquences et la matrice de répartition;
- les éléments actifs utilisés dans les circuits hyperfréquences, notamment les amplificateurs et les oscillateurs;
- des exemples de systèmes hyperfréquences.

Les problèmes liés à la transmission sans fils, ainsi que la transmission optique, seront vus aux 7ème semestre.

A la fin de ce cours, l'étudiant sera en mesure de comprendre le principe des circuits hyperfréquences utilisés les plus couramment pour les télécommunications hautes fréquences, et sera capable de les interconnecter. Il se sera aussi habitué au vocabulaire spécifique utilisé dans le domaine des télécommunications hyperfréquences.

1.2 Bref historique des hyperfréquences

L'électromagnétisme et les hyperfréquences sont souvent considérées comme des disciplines mûres, car leurs principes fondamentaux furent posés par James Clerk Maxwell dans la seconde partie du XIXème siècle déjà. De plus, ces domaines connurent un formidable essor lors de la seconde guerre mondiale à cause des recherches effectuées à cette époque sur le Radar.

L'époque actuelle est elle aussi très active dans ces disciplines, à cause de l'importance de plus en plus marquée que prennent les télécommunications dans notre société, et plus particulièrement les communications sans fils, qui sont directement liées à l'électromagnétisme et aux hyperfréquences.

1872 : publication de James Clerk Maxwell de "A treatise on Electricity and Magnetism". Cet ouvrage unifie les découvertes antérieures sur l'électricité et le magnétisme et la théorie de l'électromagnétisme, condensée en quatre équations.

1887-1891 : validations des théories de Maxwell par les expériences de Heinrich Hertz.

- 1885-1887 : publication par Oliver Heaviside de commentaires sur les travaux de Maxwell, rendant ceux-ci plus accessibles au commun des scientifiques en simplifiant la notation. Heaviside a par exemple introduit la notation vectorielle dans les équations de Maxwell.
- 1987 : Lord Rayleigh prouve mathématiquement que la propagation d'une onde est possible dans un guide rectangulaire ou circulaire. Aucune vérification expérimentale n'est faite, et le guide d'ondes sombre dans l'oubli.
- 1920 : Découverte du mélangeur, et de la détection par changement de fréquence. Principe du récepteur hétérodyne.
- 1936 : Redécouverte du guide d'ondes par deux hommes de manière indépendante : G.C. Southworth et W.L. Barrow présentent tout les deux un article sur la propagation en guide d'ondes lors d'une conférence en mai.
- 1930 : Premières utilisations du Radar, mais en bande VHF (54-88 MHz).
- 1930 : Détecteur à cristal, qui remplace le détecteur diode à pointe.
- 1937 : Invention par les frères Varian du Klystron, tube pouvant être utilisé comme source ou comme amplificateur hyperfréquences.
- 1940 : Premiers Radar hyperfréquences.
- 1948 : Fondements de la théorie des filtres distribués par Richards.
- 1949 : Première utilisation des ferrites pour la fabrication d'éléments non réciproques (isolateurs, circulateurs).
- 1950 : Début des filtres à cavités multiples, synthétisés selon des caractéristiques de Butterworth ou de Chebychev.
- 1950 : Développement des lignes de transmission hyperfréquence planaire : Tout d'abord le ruban équilibré (stripline), puis la ligne microruban (microstrip line) et le guide coplanaire (co-planar waveguide).
- 1950 : Apparition des amplificateurs TWT (travelling wave tube) et des masers, utilisés comme amplificateurs faible bruit.
- 1960 : Apparition des premiers transistors et circuit intégrés hyperfréquence.
- 1970 : Début des MMIC (monolithic Microwave Integrated Circuits), qui prônent une intégration toujours plus fortes des circuits hyperfréquence.
- 1970 : Apparition des premiers outils de conception assistée par ordinateur, rendus nécessaire par l'intégration toujours plus forte des circuits hyperfréquences.

1.3 Le spectre électromagnétique et son utilisation

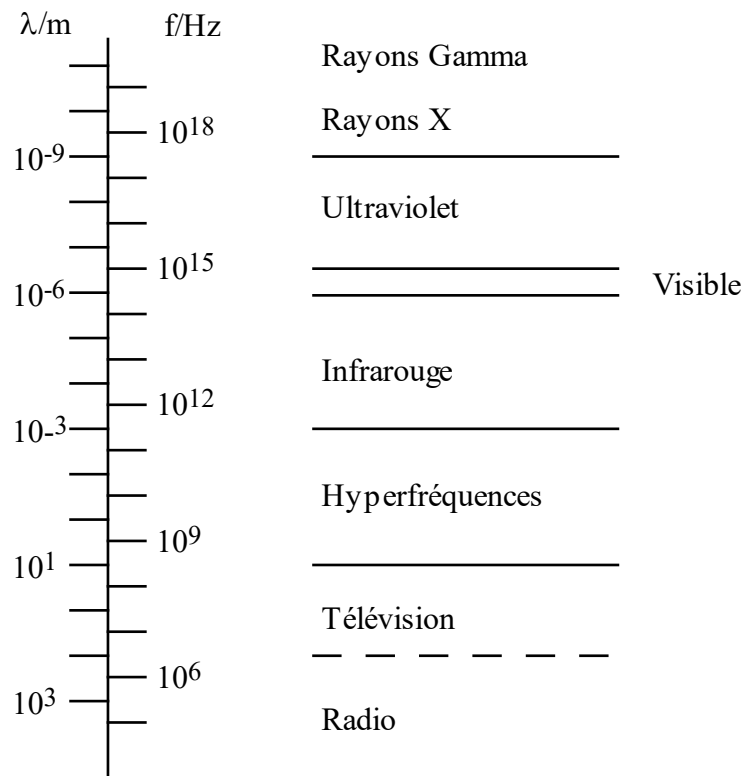


Figure 1.1 : subdivisions du spectre électromagnétique

Bande	Fréquence	Longueur d'ondes	Applications
VLF Very Low Frequency	3 - 30 kHz	100 - 10 km	Navigation, Sonar
LF Low Frequency	30 - 300 kHz	10 - 1 km	Ondes longues radio
MF Medium Frequency	300 kHz - 3 MHz	1 km - 100 m	Goniométrie, radio AM
HF High Frequency	3 - 30 MHz	100 - 10 m	CB, ondes courtes, trafic aérien
VHF Very High Frequency	30 - 300 MHz	10 - 1 m	Radio FM, TV, radar, comm. mobiles
UHF Ultra High Frequency	300 MHz - 3 GHz	1 m - 10 cm	Natel, Satellite, TV, radar, chauffage, comm. mobiles
SHF Supra High Frequency	3 - 30 GHz	10 - 1 cm	Satellite, faisceaux hertziens, radioastronomie
EHF Extremely High Frequency	30 - 300 GHz	10 - 1 mm	Satellite, radar, radioastronomie, comm. militaire

Tableau 1.1 : bandes de fréquences utilisées en transmission sans fils

La plage 300 MHz - 300 GHz est considérée comme le domaine des hyperfréquences (microondes), caractérisée par le fait que les circuits et les appareillages utilisés ont des dimensions comparables à la longueur d'onde.

Pour comparaison, à la fréquence du réseau (50 Hz) on a une longueur d'onde de 6000 km, tandis que les fréquences optiques visibles, avec des longueurs d'onde de l'ordre de 0.6 micromètres, correspondent à des fréquences de 500 000 GHz !

Les bandes des hyperfréquences sont traditionnellement subdivisées en bandes plus petites. Cette nomenclature provient de l'utilisation de guides d'ondes, et lie une bande de fréquence au guide capable de la propager. Elle reste très utilisée actuellement, malgré le fait que les guides sont de moins en moins utilisés, et est donnée pour information :

Bande	Fréquence
L	1-2 GHz
S	2-4 GHz
C	4-8 GHz
X	8-12 GHz
Ku	12-18 GHz
K	18-26 GHz
Ka	26-40 GHz
Q	40-60 GHz
E	60-90 GHz

Table 1.2 : Nomenclature des bandes hyperfréquences

1.4 Propriétés des hyperfréquences

- Bande passante :

La bande passante absolue d'un système de transmission est directement lié à la fréquence de la porteuse. Les hyperfréquences se situant dans la bande des fréquences élevées du spectre électromagnétique utilisé pour les télécommunications, elles permettent la transmission de débits d'information élevé. La même chose est évidemment vraie pour les fibres optiques, qui se situent à des fréquences encore plus élevées.

- Transparence de la ionosphère :

La ionosphère est un ensemble de couches ionisées qui entourent le globe terrestre à des altitudes situées entre 50 et 10 000 km. La propagation des ondes électromagnétiques à l'intérieur de la ionosphère est similaire à celle d'un guide d'onde. Les fréquences inférieures à quelques dizaines de MHz (fréquence de coupure) sont partiellement ou totalement réfléchies. Les signaux de fréquences supérieures traversent la ionosphère en subissant une distorsion qui décroît avec la fréquence. Les signaux hyperfréquences ne sont donc pratiquement pas affectés par la ionosphère, et sont donc tout indiqués pour des communication avec des satellites.

- Transparence partielle de l'atmosphère :

Les différents gaz qui composent l'atmosphère et les différents corps en suspension n'influencent pratiquement pas les signaux électromagnétiques dont la fréquence est inférieure à 10 GHz.

- Bruit électromagnétique :

La puissance de bruit captée par une antenne pointée vers le ciel est minimale entre 1 et 10 GHz.

- Directivité des antennes :

L'angle d'ouverture du faisceau rayonné par une antenne est proportionnel au quotient de la longueur d'ondes à la plus grande dimension de l'antenne. A dimensions égales, une antenne sera donc plus directive pour une fréquence élevée.

- Interaction avec la matière :

Lorsqu'une onde électromagnétique rencontre un échantillon de matériau, elle est absorbée de façon préférentielle à certaines fréquences. Plus particulièrement, l'eau absorbe fortement toutes les hyperfréquences, propriété qui a permis des applications comme le chauffage par microondes, le traitement thermique de certaines maladies (diathermie) ainsi que la détection et la mesure d'humidité contenue dans les matériaux.

1.5 Applications des hyperfréquences

- Radar
- Détecteurs
- Télécommunication :
 - Faisceaux hertziens (liaisons point à point)
 - Communications spatiales
 - Communications mobiles
- Chauffage à microondes
- Mesure et caractérisation des matériaux
- Radiométrie

1.6 Transmissions hyperfréquences

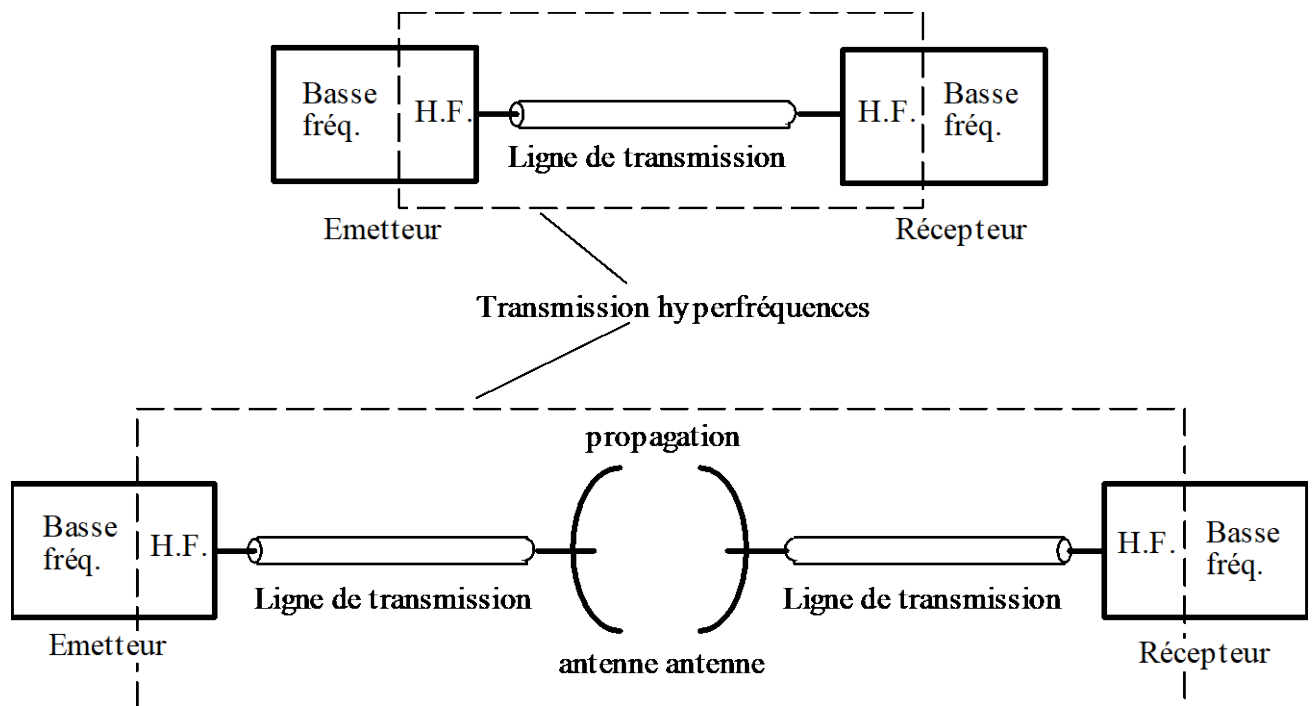


Figure 1.2 : systèmes de transmissions hyperfréquences

Tout système de télécommunication comprend en général des parties fonctionnant à des fréquences différentes. La bande de base (baseband), souvent à relativement basse fréquence, traite de la mise en forme et du codage du signal porteur de l'information. Ce signal module ensuite une porteuse située généralement à une fréquence beaucoup plus élevée (HF, VHF, hyperfréquences ou optique), qui va transporter l'information de l'émetteur au récepteur. Le signal transporté a donc une fréquence beaucoup plus élevée que le signal porteur de l'information. Dans la bande des hyperfréquences, deux modes de propagation du signal sont utilisés : la propagation guidée par des lignes de transmission et la propagation rayonnée dans l'espace libre.

Dans le cadre de ce cours, nous nous concentrerons sur la partie haute fréquence des systèmes de télécommunication dans le cas où cette transmission se fait par hyperfréquence. La propagation guidée sera abordée dans un premier temps, avec les guides d'ondes, câbles coaxiaux et les circuits imprimés hyperfréquences (circuits microruban et guides coplanaires). Puis les éléments spécifiques aux circuits hyperfréquences seront abordés, en mettant le poids sur les applications dans les télécommunications.

La transmission du signal par propagation rayonnée et les communications sans fils seront abordées dans le cours du 7ème semestre.

2. RAPPELS DE THEORIE ELECTROMAGNETIQUE

Dans ce chapitre on partira de la forme générale des équations de Maxwell qui contrôlent tous les phénomènes électromagnétiques. Ces équations seront brièvement commentées et agrémentées d'une présentation succincte des différents acteurs du drame électromagnétique. On aboutira finalement à une version particulière des équations de Maxwell, spécifique aux besoins de ce cours.

Références:

Traité d'Electricité de l'EPFL, vol.III : "Electromagnetisme", chapitre 1

Ramo : "Fields and Waves in Communication Electronics", chapitre 3

2.1 Les équations de Maxwell (version générique)

La forme la plus générale des équations de Maxwell en termes de l'opérateur vectoriel ∇ ("del" ou "nabla", utilisé couramment pour représenter les opérations "rot", "div" et "grad") est la suivante:

$$\begin{aligned}\nabla \times \mathbf{E}(\mathbf{r},t) &= -\partial \mathbf{B}(\mathbf{r},t) / \partial t & \nabla \cdot \mathbf{D}(\mathbf{r},t) &= \rho(\mathbf{r},t) \\ \nabla \times \mathbf{H}(\mathbf{r},t) &= \mathbf{J}(\mathbf{r},t) + \partial \mathbf{D}(\mathbf{r},t) / \partial t & \nabla \cdot \mathbf{B}(\mathbf{r},t) &= 0\end{aligned}$$

Essentiellement, les équations de Maxwell établissent des relations différentielles entre les sources des phénomènes électromagnétiques et leurs effets dans un point de l'espace. Les sources sont les charges et les courants électriques $\mathbf{J}$ et ρ . Leurs effets sont représentés mathématiquement par 4 champs vectoriels $\mathbf{E}$ (champ électrique), $\mathbf{H}$ (champ magnétique), $\mathbf{D}$ (déplacement électrique) et $\mathbf{B}$ (induction magnétique).

Sources et champs doivent en général être considérées comme fonction des coordonnées spatiales (vecteur de position $\mathbf{r}$) et du temps t . Donc $f = f(\mathbf{r},t)$.

En outre, les charges et les courants sont liés par l'équation de continuité :

$$\nabla \cdot \mathbf{J}(\mathbf{r},t) + \partial \rho(\mathbf{r},t) / \partial t = 0$$

En théorie électromagnétique, on travaille le plus souvent en régime sinusoïdal permanent. En effet, même les phénomènes transitoires sont souvent étudiés par décomposition de l'onde temporelle un spectre de fréquences (transformation de Fourier) et par analyse à chaque fréquence, suivie d'une transformation inverse. En régime sinusoïdal de pulsation $\omega = 2\pi f$, on peut donc admettre une dépendance temporelle de type $\cos(\omega t + \phi)$ et écrire

$$A \cos(\omega t + \phi) = \text{Re} [A e^{j\phi} \exp(j\omega t)] = \sqrt{2} \text{Re} \left[\frac{A}{\sqrt{2}} e^{j\phi} \exp(j\omega t) \right] = \sqrt{2} \text{Re} [A_e e^{j\phi} \exp(j\omega t)]$$

A est ici la valeur de crête et $A_e = A/\sqrt{2}$ la valeur efficace.

On remplace alors les vraies grandeurs physiques fonction du temps $f(\mathbf{r},t)$ ($f=\mathbf{E}, \mathbf{D}, \mathbf{H}, \mathbf{B}, \mathbf{J}, \rho$) par des quantités complexes mais indépendantes du temps (phaseurs) $f(\mathbf{r})$ ($f=\mathbf{E}, \mathbf{D}, \mathbf{H}, \mathbf{B}, \mathbf{J}, \rho$), avec la relation:

$$f(\mathbf{r},t) = \sqrt{2} \operatorname{Re} [f(\mathbf{r}) \exp(j\omega t)]$$

Le facteur $\sqrt{2}$ est introduit dans la définition pour que le module du phaseur corresponde à la valeur efficace du signal sinusoïdal.

Dans la plupart de textes anglo-saxons, le facteur $\sqrt{2}$ n'est pas présent dans le définition des phaseurs. La norme du phaseur est alors la valeur de crête et un facteur supplémentaire 1/2 apparaît dans les formules associées aux énergies et aux puissances.

L'introduction de phaseurs permet aussi de remplacer (comme en transformée de Fourier) les dérivées temporelles par des facteurs $j\omega$. Finalement, les équations de Maxwell en termes de vecteurs-phaseurs s'écrivent :

$$\begin{aligned} \nabla \times \mathbf{E}(\mathbf{r}) &= -j\omega \mathbf{B}(\mathbf{r}) & \nabla \cdot \mathbf{D}(\mathbf{r}) &= \rho(\mathbf{r}) \\ \nabla \times \mathbf{H}(\mathbf{r}) &= \mathbf{J}(\mathbf{r}) + j\omega \mathbf{D}(\mathbf{r}) & \nabla \cdot \mathbf{B}(\mathbf{r}) &= 0 \end{aligned}$$

plus l'équation de continuité

$$\nabla \cdot \mathbf{J}(\mathbf{r}) + j\omega \rho(\mathbf{r}) = 0$$

Dans ce cours, on se limitera à l'étude des milieux linéaires. En effet, sauf quelques rares exceptions (substances ferromagnétiques, plasmas) le niveau de puissance et le type des matériaux employés dans la technologie des télécommunications autorise l'utilisation d'un modèle linéaire. D'un point de vue mathématique, un milieu linéaire en régime sinusoïdal permanent est caractérisé par les relations constitutives (traité, vol.III):

$$\mathbf{D} = \epsilon \mathbf{E} \quad ; \quad \mathbf{B} = \mu \mathbf{H}$$

où maintenant ϵ et μ sont deux constantes du milieu (permittivité et perméabilité) qui peuvent en général prendre des valeurs complexes fonction de la fréquence (traité, vol.III):

$$\epsilon = \epsilon' - j \epsilon'' \quad ; \quad \mu = \mu' - j \mu''$$

Une valeur complexe de ϵ implique, quand on revient au phénomène temporel, que les vecteurs $\mathbf{D}(\mathbf{r},t)$ et $\mathbf{E}(\mathbf{r},t)$ oscillent avec la même fréquence ω mais avec un déphasage. Si en plus, $\operatorname{Im}(\epsilon) < 0$, $\mathbf{D}$ est en retard par rapport à $\mathbf{E}$. Le caractère négatif de la partie imaginaire de ϵ est liée à des relations de cause à effet (relations de Kramers-König, analogues à celles de Bode en théorie de circuits, traité IV,7.3.34) et correspond physiquement à l'existence de pertes dans le milieu. Les mêmes considérations sont valables pour le couple $\mathbf{B}, \mathbf{H}$.

Si on est dans l'espace libre (l'air en est une bonne approximation) on a alors les valeurs:
 $\epsilon = \epsilon_0 = 8.854 \cdot 10^{-12}$ [farad/m] et $\mu = \mu_0 = 4\pi \cdot 10^{-7}$ [henry/m]

Dans tout autre milieu linéaire on définit des valeurs relatives complexes:

$$\epsilon_r = \epsilon / \epsilon_0 \text{ et } \mu_r = \mu / \mu_0 .$$

A titre d'exemple, on a pour l'eau à 1 GHz: $\epsilon_r = (80 - j10)$, $\mu_r = 1$.

Il faut finalement remarquer, que la fréquence angulaire d'un phénomène électromagnétique à variation sinusoïdale demeure inaltérée tant que les milieux en jeu sont linéaires.

Les relations constitutives impliquent que seulement deux vecteurs $\mathbf{E}$, $\mathbf{H}$ sont nécessaires pour décrire un phénomène électromagnétique dans un milieu linéaire. On peut maintenant écrire les équations de Maxwell sous la forme:

$$\begin{aligned} \nabla \times \mathbf{E} &= -j\omega\mu\mathbf{H} & \nabla \cdot \mathbf{E} &= \rho/\epsilon \\ \nabla \times \mathbf{H} &= \mathbf{J} + j\omega\epsilon\mathbf{E} & \nabla \cdot \mathbf{H} &= 0 \end{aligned}$$

2.2. Conditions aux limites, énergie et puissance (rappel)

En présence d'une surface séparant deux milieux #1 et #2 de nature différente, les équations de Maxwell doivent être complétées par les conditions aux limites suivantes (fig. 2.1):

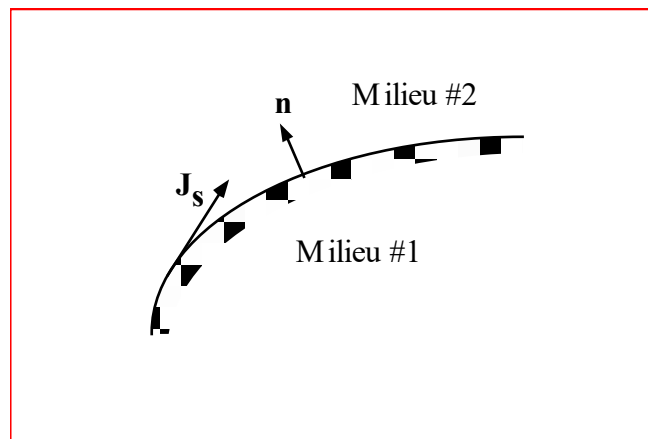


Fig. 2.1 : Conditions aux limites

$$\hat{\mathbf{n}} \times (\mathbf{E}_2 - \mathbf{E}_1) = 0 \quad ; \quad \hat{\mathbf{n}} \times (\mathbf{H}_2 - \mathbf{H}_1) = \mathbf{J}_s$$

où $\hat{\mathbf{n}}$ est le vecteur unitaire normal traversant la surface du milieu #1 vers le #2 et $\mathbf{J}_s$ est un éventuel courant de surface [A/m] pouvant exister dans l'interface entre les milieux.

Les définitions suivantes sont valables en régime sinusoïdal :

$w_e = (1/2) \epsilon \mathbf{E} \cdot \mathbf{E}$ [J/m³] : Valeur moyenne de la densité d'énergie électrique en un point
 $w_m = (1/2) \mu \mathbf{H} \cdot \mathbf{H}$ [J/m³] : Valeur moyenne de la densité d'énergie magnétique en un point.
 $\mathbf{S} = \mathbf{E} \times \mathbf{H}^*$ [W/m²] : Vecteur de Poynting = valeur moyenne du flux de puissance dans un point.

Il s'agit, bien sûr, de moyennes temporelles (voir traité, vol.III).

L'intégration des équations de Maxwell sur un volume v entouré par une surface s donne alors le théorème de Poynting :

$$\int_s ds \mathbf{S} \cdot \hat{\mathbf{n}} + j\omega \int_v dv (w_e + w_m) = - \int_v dv \mathbf{J} \cdot \mathbf{E} \quad , \quad [w]$$

où $\hat{\mathbf{n}}$ est le vecteur unitaire normal à s et dirigé vers l'extérieur de v . Le premier terme de gauche est le flux du vecteur de Poynting, c'est à dire la puissance s'échappant du volume à travers la surface s (rayonnement). Le deuxième terme de gauche correspond à la puissance réactive emmagasinée dans le volume v . La somme de ces deux puissances est égale à celle fournie par les sources de courant.

2.3 Les courants sources et les courants induits

Dans tout problème d'électromagnétisme on admet l'existence des courants sources $\mathbf{J}_{src}$ qui ne sont pas susceptible d'être modifiés, ni par les champs qu'ils génèrent eux mêmes, ni par aucun autre champ. Ces courants sources génèrent des champs électromagnétiques d'excitation. Si un objet quelconque est placé au voisinage des sources, ces champs d'excitation produisent sur l'objet des courants induits $\mathbf{J}_{ind}$. A leur tour, ces courants induits dans l'objet génèrent des champs diffractés. Les champs totaux sont la somme des champs d'excitation et des champs diffractés, créés par ces deux types de courant.

Il n'y a souvent aucune différence physique entre ces deux courants (dans les deux cas, ce sont essentiellement des électrons en mouvement). On doit toutefois établir une différence conceptuelle, subtile mais essentielle.

* $\mathbf{J}_{src}$ est un courant connu et imposé. Il n'est pas affecté par le champ existant. C'est le courant source donnée du problème.

* $\mathbf{J}_{ind}$ est un courant le plus souvent inconnu qui dépend du champ total.

Dans l'équation de Maxwell

$$\nabla \times \mathbf{H} = j\omega \epsilon \mathbf{E} + \mathbf{J}$$

le courant $\mathbf{J}$ est une valeur totale, et donc fonction des champs . On aimerait mettre en évidence la partie du courant indépendant des champs qui jouera le rôle mathématique de terme inhomogène dans l'équation différentielle. On écrit alors $\mathbf{J} = \mathbf{J}_{\text{src}} + \mathbf{J}_{\text{ind}}$ où $\mathbf{J}_{\text{src}}$ est la partie "source" connue et indépendante du champ, et $\mathbf{J}_{\text{ind}}$ est la partie "induite". Pour des objets composés par des matériaux linéaires, ce courant induit est lié exclusivement au champ électrique total par la loi d'Ohm : $\mathbf{J}_{\text{ind}} = \sigma \mathbf{E}$. On peut alors écrire :

$$j\omega\epsilon \mathbf{E} + \mathbf{J} = j\omega\epsilon \mathbf{E} + \mathbf{J}_{\text{ind}} + \mathbf{J}_{\text{src}} = (j\omega\epsilon + \sigma) \mathbf{E} + \mathbf{J}_{\text{src}} = j\omega\epsilon_T \mathbf{E} + \mathbf{J}_{\text{src}}$$

et le tour de passe-passe est joué car on peut écrire comme nouvelle équation de Maxwell :

$$\nabla \times \mathbf{H} = j\omega\epsilon_T \mathbf{E} + \mathbf{J}_{\text{src}}$$

où on a introduit une permittivité globale

$$\epsilon_T = \epsilon - j \sigma/\omega.$$

De forme analogue, l'équation $\nabla \cdot \mathbf{E} = \rho/\epsilon$ doit être remplacée par $\nabla \cdot \mathbf{E} = \rho_{\text{src}}/\epsilon_T$.

2.4 Les équations de Maxwell (version finale)

Dans ce cours on s'intéresse aux phénomènes de propagation ayant lieu quand le signal électromagnétique a quitté le générateur (les sources) et se propage vers un récepteur. Donc, on considère un milieu sans sources et on écrit les équations de Maxwell sous la forme qui sera toujours utilisée par la suite:

$$\nabla \times \mathbf{E} = -j\omega\mu \mathbf{H} \quad \nabla \cdot \mathbf{E} = 0$$

$$\nabla \times \mathbf{H} = j\omega\epsilon_T \mathbf{E} \quad \nabla \cdot \mathbf{H} = 0$$

Toutefois, pour ne pas alourdir la notation, on dénotera la permittivité totale ϵ au lieu de ϵ_T , mais sa valeur (qui pourrait déjà être complexe au départ si le milieu a des pertes diélectriques) pourra inclure une partie imaginaire négative supplémentaire σ/ω s'il y a des pertes par conduction.

Prenant encore une fois le rotationnel sur les équations à rotationnel et utilisant le calcul vectoriel on montre que dans une région sans sources les champs satisfont des équations d'ondes (ou de Helmholtz)

$$\nabla^2 \mathbf{E} + \omega^2 \mu \epsilon \mathbf{E} = 0 \quad ; \quad \nabla^2 \mathbf{H} + \omega^2 \mu \epsilon \mathbf{H} = 0$$

$$\text{ou en notation compacte : } (\nabla^2 + k^2) \mathbf{E} = 0 \quad ; \quad (\nabla^2 + k^2) \mathbf{H} = 0$$

Pour un matériau et pour une fréquence donnés, la valeur $k = \omega\sqrt{\mu\epsilon}$ est une constante complexe appelée constante de propagation ou nombre d'onde (anglais *wavenumber*). On sait que la solution la plus simple à ces équations dans un milieu infini est l'onde électromagnétique plane transverse dont les champs ont une forme

$$\mathbf{E}(\mathbf{r}) = \mathbf{E}_0 \exp(-jk \hat{\mathbf{n}} \cdot \mathbf{r}) \quad , \quad \text{ou bien} \quad \mathbf{E}(\mathbf{r}, t) = \frac{1}{2} E_0 \cos(\omega t - k \hat{\mathbf{n}} \cdot \mathbf{r})$$

$\hat{\mathbf{n}}$ étant un vecteur unitaire dans la direction de propagation.

La vitesse de propagation d'une telle onde est

$$v = k/\omega = 1/\sqrt{\mu\epsilon} = c_0/\sqrt{\mu_r \epsilon_r}$$

avec c_0 vitesse de la lumière dans le vide, et la longueur d'onde associée vaut simplement

$$\lambda = 2\pi/k = v/f$$

3. LA PROPAGATION ÉLECTROMAGNÉTIQUE GUIDÉE

Dans le chapitre 2, aucune hypothèse n'a été faite sur la géométrie du milieu supportant les phénomènes électromagnétiques. De ce fait, la forme finale obtenue pour les équations de Maxwell peut être appliquée avec le même bonheur aux problèmes de rayonnement dans un milieu ouvert, de propagation guidée par une structure à symétrie de translation et d'oscillations dans une région bornée. On va maintenant spécialiser ces équations au problème de la propagation guidée, sujet de ce cours.

Références:

Traité d'Electricité de l'EPFL, vol.III : "Electromagnetisme", section 2.4 (séparation de variables)

Ramo : "Fields and Waves in Communication Electronics", chapitre 8

3.1 Généralités

3.1.1 Coordonnées et composantes

Nous admettrons que toutes les géométries étudiées possèdent une symétrie de translation le long de l'axe des coordonnées z . De ce fait, la coordonnée "axiale" ou "longitudinale" z va jouer un rôle spécifique et devra être différenciée des autres deux coordonnées "transverses", qui peuvent être (x,y) mais aussi (ρ,ϕ) ou tout autre système de coordonnées dans le plan. On introduira la notation générique $\mathbf{t} = (t_1, t_2)$ pour ces coordonnées transverses.

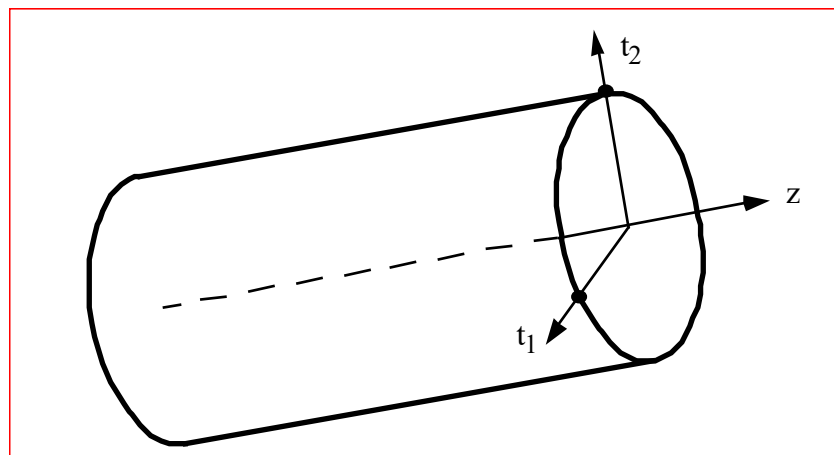


Fig. 3.1 Structure de guidage d'ondes

Ainsi, on écrira tout vecteur $\mathbf{v}$ comme: $\mathbf{v} = v_z \hat{\mathbf{z}} + \mathbf{v}_t$, mettant en évidence la composante selon z et groupant les deux autres composantes transverses dans un vecteur transverse. En particulier pour l'opérateur vectoriel ∇ , on écrit :

$$\nabla = (\partial/\partial z)\hat{\mathbf{z}} + \nabla_t, \text{ et pour le laplacien : } \nabla^2 = \partial^2/\partial z^2 + \nabla_t^2$$

On introduit la technique de séparation de variables (traité, vol.III) et toute fonction f des trois coordonnées sera considérée comme un produit

$$f(t_1, t_2, z) = T(t_1, t_2) Z(z)$$

Chacune des 6 composantes scalaires du champ électromagnétique ($\mathbf{E}$ et $\mathbf{H}$) satisfait une équation d'onde :

$$(\nabla^2 + k^2)f = 0 \quad ; \quad f = E_{t1}, E_{t2}, E_z, H_{t1}, H_{t2}, H_z$$

La technique de séparation de variables transforme cette équation et donne :

$$\frac{\nabla_t^2 T}{T} + k^2 = -\frac{1}{Z} \frac{d^2 Z}{dz^2} = -\gamma^2$$

où γ est la constante complexe de séparation associée au processus de séparation.

3.1.2 Dépendance longitudinale : exposant de propagation

L'équation gouvernant la partie longitudinale $Z(z)$ a une solution analytique simple. On trouve que la dépendance des phénomènes électromagnétiques de propagation guidée avec la coordonnée longitudinale z est toujours du type :

$$Z(z) = A \exp(\gamma z) + B \exp(-\gamma z)$$

exactement comme pour les tensions et courants dans une ligne de transmission.

On peut alors parler d'onde incidente $\exp(-\gamma z)$ et d'onde réfléchie $\exp(+\gamma z)$. Le sens physique de la constante γ est maintenant clair : il s'agit de l'exposant longitudinal de propagation. En général $\gamma = \alpha + j\beta$. La partie réelle α est l'affaiblissement linéique ([neper/m] ou [dB/m]). La partie imaginaire β est le déphasage linéique (mesuré en [rad/m]).

Pour simplifier les expressions, on va considérer seulement des ondes "incidentes" se propageant selon l'axe z positif. Alors, avec $B=0$, on peut intégrer l'amplitude de l'onde incidente dans les fonctions transverses T et écrire en forme vectorielle compacte :

$$\begin{aligned} \mathbf{E}(t_1, t_2, z) &= \mathbf{e}(t_1, t_2) \exp(-\gamma z) \\ \mathbf{H}(t_1, t_2, z) &= \mathbf{h}(t_1, t_2) \exp(-\gamma z) \end{aligned}$$

où les parties transverses $\mathbf{e}, \mathbf{h}$ dépendent de la section droite de la structure de guidage.

Quand on voudra trouver des expressions pour la partie réfléchie, il suffira de changer formellement dans les formules γ par $-\gamma$.

3.1.3 Dépendance transverse : valeurs et vecteurs propres

Les parties transverses $\mathbf{e}$, $\mathbf{h}$ sont solution des équations aux valeurs propres :

$$\begin{aligned}(\nabla_t^2 + k^2 + \gamma^2) \mathbf{e} &= (\nabla_t^2 + k_c^2) \mathbf{e} = 0 \\(\nabla_t^2 + k^2 + \gamma^2) \mathbf{h} &= (\nabla_t^2 + k_c^2) \mathbf{h} = 0\end{aligned}$$

Les valeurs propres admissibles k_c ($k_c^2 = k^2 + \gamma^2$) sont déterminées par les conditions aux limites associées à la section droite du guide d'ondes. Pour chaque valeur propre, on trouve des vecteurs propres $\mathbf{e}$, $\mathbf{h}$, qui sont appelées couramment “modes” de la structure guidée.

Aussi, à chaque valeur propre correspond une valeur de l'exposant de propagation :

$$\gamma = \sqrt{k_c^2 - k^2} = \sqrt{k_c^2 - \omega^2 \mu \epsilon}$$

Par conséquent, la relation entre les valeurs propres et la fréquence déterminera la nature de γ (réelle, imaginaire ou complexe) et donc les caractéristiques de la propagation.

3.1.4 Composantes longitudinales et transverses

Les vecteurs transverses $\mathbf{e}$, $\mathbf{h}$ contiennent en général six composantes scalaires. On doit donc en principe résoudre 6 équations scalaires aux valeurs propres, chacune avec ses conditions aux limites particulières.

Néanmoins, il est évident que ces six composantes ne sont pas indépendantes mais liées par les quatre équations de Maxwell. Il devrait donc être possible de se concentrer sur deux composantes et d'exprimer les quatre restantes en fonction des ces deux.

Le choix le plus logique consiste à considérer les composantes longitudinales e_z , h_z comme fonctions de base et tâcher d'obtenir les relations donnant les quatre composantes transverses en fonction des deux axiales. D'une certaine façon, e_z , h_z peuvent être considérés comme “potentiels” desquels dérivent les autres composantes.

Les relations souhaitées s'obtiennent directement en introduisant dans les équations de Maxwell en rotationnel des expressions pour les champs :

$$\mathbf{E} = (e_z \hat{\mathbf{z}} + \mathbf{e}_t) \exp(-\gamma z) \quad ; \quad \mathbf{H} = (h_z \hat{\mathbf{z}} + \mathbf{h}_t) \exp(-\gamma z) \quad ; \quad \nabla = (\partial/\partial z) \hat{\mathbf{z}} + \nabla_t$$

où la décomposition “axial-transverse” est mise en évidence. On trouve aisément

$$\begin{aligned}\nabla_t \times \mathbf{e}_t &= -j\omega\mu (\mathbf{h}_z \hat{\mathbf{z}}) \quad ; \quad \nabla_t \times (\mathbf{e}_z \hat{\mathbf{z}}) - \gamma \hat{\mathbf{z}} \times \mathbf{e}_t = -j\omega\mu \mathbf{h}_t \\ \nabla_t \times \mathbf{h}_t &= +j\omega\varepsilon (\mathbf{e}_z \hat{\mathbf{z}}) \quad ; \quad \nabla_t \times (\mathbf{h}_z \hat{\mathbf{z}}) - \gamma \hat{\mathbf{z}} \times \mathbf{h}_t = +j\omega\varepsilon \mathbf{e}_t\end{aligned}$$

et finalement :

$$\begin{aligned}k_c^2 \mathbf{e}_t &= -\gamma \nabla_t \mathbf{e}_z + j\omega\mu \hat{\mathbf{z}} \times \nabla_t \mathbf{h}_z \\ k_c^2 \mathbf{h}_t &= -\gamma \nabla_t \mathbf{h}_z - j\omega\varepsilon \hat{\mathbf{z}} \times \nabla_t \mathbf{e}_z\end{aligned}$$

Pour une onde réfléchie, se propageant selon l'axe z négatif avec un exposant $\exp(+\gamma z)$, il suffira de changer formellement $+\gamma$ par $-\gamma$ dans toutes les équations précédentes.

3.1.5 Résumé : procédure de calcul

En général, l'étude d'une structure guidée comportera donc les étapes suivantes:

a) résoudre dans la section droite de la structure les équations aux valeurs propres

$$(\nabla_t^2 + k_c^2) \mathbf{e}_z = 0 \quad ; \quad (\nabla_t^2 + k_c^2) \mathbf{h}_z = 0$$

avec les conditions aux limites pertinentes. Trouver en particulier les valeurs propres k_c admissibles et les modes $\mathbf{e}_z, \mathbf{h}_z$ associés.

b) Calculer l'exposant de propagation $\gamma = \alpha + j\beta = \sqrt{k_c^2 - k^2}$. Pour une onde incidente, prendre le signe de la racine donnant $\text{Im}(\gamma) > 0$.

c) Calculer les composantes transverses avec

$$\begin{aligned}k_c^2 \mathbf{e}_t &= -\gamma \nabla_t \mathbf{e}_z + j\omega\mu \hat{\mathbf{z}} \times \nabla_t \mathbf{h}_z \\ k_c^2 \mathbf{h}_t &= -\gamma \nabla_t \mathbf{h}_z - j\omega\varepsilon \hat{\mathbf{z}} \times \nabla_t \mathbf{e}_z\end{aligned}$$

d) Construire les champs de l'onde incidente comme :

$$\begin{aligned}\mathbf{E}(t_1, t_2, z) &= (\mathbf{e}_t + \mathbf{e}_z \hat{\mathbf{z}}) \exp(-\gamma z) \\ \mathbf{H}(t_1, t_2, z) &= (\mathbf{h}_t + \mathbf{h}_z \hat{\mathbf{z}}) \exp(-\gamma z)\end{aligned}$$

e) Refaire les calculs pour une onde réfléchie. Il suffit de choisir l'autre branche dans la racine donnant γ , ce qui revient formellement à changer $+\gamma$ par $-\gamma$ dans les formules.

f) Trouver les amplitudes des ondes incidente et réfléchie utilisant la théorie des lignes de transmission. Pour ceci, il faut connaître les conditions “en bout de ligne”. Puis combiner les parties incidentes et réfléchies pour obtenir les champs totaux.

g) Si l’on souhaite les expressions temporelles sont facilement récupérables. Par exemple, le champ E_z d’une onde incidente d’amplitude complexe $A = |A| \exp(j\varphi_A)$ est donné par :

$$E_z(t_1, t_2, z, t) = \sqrt{2} |A| e_z(t_1, t_2) \exp(-\alpha z) \cos(\omega t - \beta z + \varphi_A)$$

3.2 Les modes de propagation

Les solutions possibles de l’équation de Helmholtz relatives à la dépendance transverse des champs sont appelées modes de la structure de guidage. Chaque mode représente une configuration du champ électromagnétique sous laquelle un signal peut se propager. Il est donc de la plus haute importance pour évaluer la qualité d’un canal de transmission de connaître les modes pouvant exister pour une combinaison quelconque géométrie-fréquence -paramètres électriques des matériaux.

Références:

“Fields and Waves in Communication Electronics”, chapitre 8

3.2.1 Classification

Les solutions des équations de Helmholtz

$$(\nabla_t^2 + k_c^2) e_z = 0 \quad ; \quad (\nabla_t^2 + k_c^2) h_z = 0$$

se classifient de la manière suivante:

	Modes TEM	Modes TM ou E	Modes TE ou H	Modes hybrides
e_z	0	$\neq 0$	0	$\neq 0$
h_z	0	0	$\neq 0$	$\neq 0$

Tous ces types de mode peuvent exister dans la nature. En général, une structure de guidage avec une géométrie donnée peut supporter plusieurs familles de modes. Ainsi on verra que la propagation dans une ligne coaxiale peut se faire sous forme de mode TEM, TM ou TE tandis que les fibres optiques supportent des modes hybrides.

On décrit par la suite les caractéristiques principales de chaque type de mode.

3.2.2 Modes TEM

Un mode TEM (Transverse ElectroMagnétique) est caractérisé par l'absence des composantes du champ dans le sens de la propagation ($e_z = h_z = 0$). On est donc en présence des champs purement transverses comme dans le cas de l'onde plane dans un milieu infini. Dans toute structure guidée, les champs transverses sont donnés par (§3.1.4)

$$\begin{aligned} k_c^2 \mathbf{e}_t &= -\gamma \nabla_t e_z + j\omega\mu \hat{\mathbf{z}} \times \nabla_t h_z \\ k_c^2 \mathbf{h}_t &= -\gamma \nabla_t h_z - j\omega\varepsilon \hat{\mathbf{z}} \times \nabla_t e_z \end{aligned}$$

On voit aisément que si e_z et h_z sont nuls, les champs transverses seront aussi en principe nuls et on arrive à une solution triviale. La seule possibilité pour obtenir des composantes transverses non nulles est de forcer la valeur $k_c=0$, ce qui implique: $\gamma = jk = j\omega\sqrt{(\mu\varepsilon)}$, valeur identique à celle associée à une onde plane de la même fréquence.

Les champs transverses ne peuvent alors être calculés avec les relations ci-dessus, mais les équations de Maxwell s'écrivent alors :

$$\begin{aligned} \nabla_t \times \mathbf{e}_t &= 0 & ; & & -\gamma \hat{\mathbf{z}} \times \mathbf{e}_t &= j\omega\mu \mathbf{h}_t \\ \nabla_t \times \mathbf{h}_t &= 0 & ; & & -\gamma \hat{\mathbf{z}} \times \mathbf{h}_t &= +j\omega\varepsilon \mathbf{e}_t \end{aligned}$$

Donc, les champs transverses ont un rotationnel nul et dérivent des potentiels. Par exemple

$$\mathbf{e}_t = -\nabla_t V \quad \text{et} \quad \mathbf{E}_t = (-\nabla_t V) \exp(-\gamma z) \quad \text{avec} \quad \nabla_t^2 V = 0$$

Il suffit donc de résoudre l'équation de Laplace dans la section de la structure en étude avec les conditions aux limites appropriées, comme en électrostatique.

En particulier, on sait qu'à l'intérieur d'un conducteur creux le champ électrostatique est nul. Il est donc impossible d'obtenir un mode de propagation TEM dans des structures du type guide d'onde métallique, et il faut en général deux conducteurs isolés pour assurer l'existence d'un mode TEM.

Les équations de Maxwell montrent aussi que dans un mode TEM le champ magnétique transverse est lié au champ électrique par la relation :

$$\mathbf{h}_t = (\gamma / j\omega\mu) \hat{\mathbf{z}} \times \mathbf{e}_t = (j\omega\varepsilon / \gamma) \hat{\mathbf{z}} \times \mathbf{e}_t$$

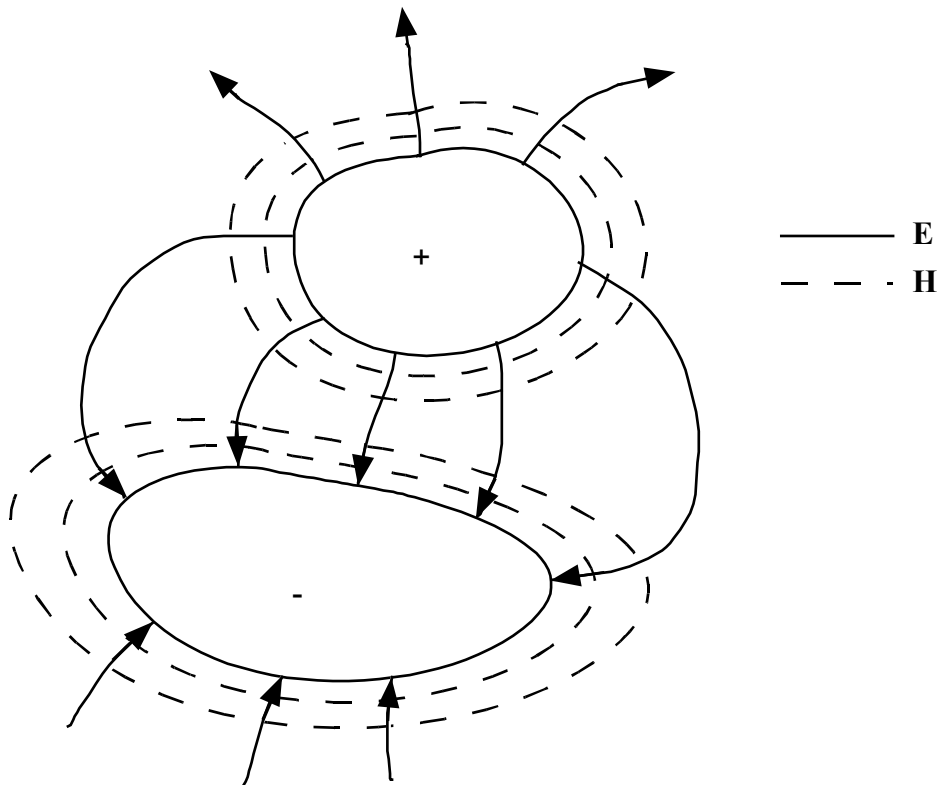
Le rapport d'amplitudes $|\mathbf{e}_t|/|\mathbf{h}_t|$ doit avoir les dimensions d'une impédance et il est appelé impédance modale ou d'onde, Z_{mod} . On trouve pour les modes TEM:

$$Z_{\text{mod(TEM)}} = j\omega\mu/\gamma = \gamma/j\omega\epsilon = \sqrt{(\mu/\epsilon)}.$$

Donc, l'impédance d'onde d'un mode TEM est égale à l'impédance du milieu supportant la propagation, comme dans le cas d'une onde plane.

Courant, tension et impédance caractéristique

Nous avons vu dans ce qui précède que les champs électromagnétiques ont un comportement statique dans le cas d'un modes TEM. Ceci implique qu'un courant et une tension peuvent être définis de manière univoque le long d'une ligne de transmission supportant un mode TEM. Considérons par exemple la ligne de transmission bifilaire représentée en coupe dans la figure ci-dessous :



Ligne TEM bifilaire

La différence de potentiel entre le conducteur + et le conducteur - est donnée par l'intégrale du champ électrique sur entre ces deux conducteurs. Vu le caractère statique du champ électrique (il dérive d'un potentiel), cette intégrale ne dépend pas du chemin d'intégration choisi entre les deux conducteurs, et définit donc une tension unique dans le plan de coupe considéré :

$$V = \int_+ \vec{E} \cdot d\mathbf{l}$$

On peut donc définir une onde de tension le long de la ligne de transmission, de la même manière qu'une onde de champ électrique :

$$V(z) = V^+ e^{-j\gamma z} + V^- e^{j\gamma z}$$

De même, le courant total circulant sur le conducteur + peut être obtenu par la loi d'Ampère, qui pour un mode TEM s'écrit :

$$I = \oint_{C^+} \mathbf{H} \cdot d\mathbf{l}$$

Où C^+ est n'importe quel chemin fermé entourant le conducteur + sans contenir le conducteur -. Le long de la ligne de transmission, le courant s'écrit donc comme :

$$I(z) = I^+ e^{-j\gamma z} - I^- e^{j\gamma z}$$

Les rapports $\frac{V^+}{I^+}$ et $\frac{V^-}{I^-}$ sont constants le long de la ligne, et ont la dimension d'une impédance. On nomme ce rapport l'impédance caractéristique de la ligne de transmission. Cette impédance, qui est aussi égale à la racine carrée du rapport entre l'inductance linéique de la ligne par sa capacité linéique dépend donc essentiellement du milieu de propagation et de la géométrie de la ligne.

3.2.3 Modes TM

Un mode TM (Transverse Magnétique) est caractérisé par une composante axiale du champ magnétique nulle. Le champ magnétique est donc transverse, mais en revanche le champ électrique possède une composante e_z dans la direction de propagation, solution de

$$(\nabla_t^2 + k_c^2) e_z = 0$$

avec les conditions aux limites appropriées (par exemple, si la structure en étude comporte des conducteurs parfaits alors $e_z=0$ sur leurs surfaces).

Les champs transverses peuvent être obtenus comme:

$$k_c^2 \mathbf{e}_t = -\gamma \nabla_t e_z \quad ; \quad \mathbf{h}_t = \frac{j\omega\epsilon}{\gamma} \hat{\mathbf{z}} \times \mathbf{e}_t$$

On remarque que dans les modes TM, l'exposant de propagation vaut $\gamma = \sqrt{k_c^2 - \omega^2\mu\epsilon}$ et l'impédance modale Z_{mod} est donnée par :

$$Z_{\text{mod}} = \gamma / j\omega\epsilon = (\sqrt{\omega^2\mu\epsilon - k_c^2}) / \omega\epsilon$$

Ces deux paramètres dépendent donc de la géométrie à travers la valeur propre k_c .

3.2.4 Modes TE

Un mode TE (Transverse Électrique) est caractérisé par une composante axiale du champ électrique nulle. Le champ électrique est donc transverse, mais en revanche le champ magnétique possède une composante h_z dans la direction de propagation solution de

$$(\nabla_t^2 + k_c^2) h_z = 0$$

avec les conditions aux limites appropriées (par exemple, si la structure en étude comporte des conducteurs parfaits alors $\partial h_z / \partial n = 0$ sur leurs surfaces).

Les champs transverses peuvent être obtenus comme :

$$k_c^2 \mathbf{h}_t = -\gamma \nabla_t h_z \quad ; \quad \mathbf{e}_t = \frac{j\omega\mu}{\gamma} \mathbf{h}_t \times \hat{\mathbf{z}}$$

On remarque que dans les modes TE, l'exposant de propagation vaut $\gamma = \sqrt{k_c^2 - \omega^2\mu\epsilon}$ et l'impédance modale Z_{mod} est donnée par :

$$Z_{\text{mod}} = j\omega\mu / \gamma = \omega\mu / \sqrt{\omega^2\mu\epsilon - k_c^2}$$

Ces deux paramètres dépendent donc de la géométrie à travers la valeur propre k_c .

3.2.5 Tableau récapitulatif

Voici sous forme compacte l'ensemble de champs et de paramètres caractérisant chacune des trois familles de modes de propagation. La racine carrée dans la valeur de l'exposant de propagation γ doit être toujours choisie de façon à avoir $\arg(\gamma) \in [0 ; \pi/2]$. Ceci correspond à un mode se propageant selon l'axe z positif (mode "incident") avec un éventuel affaiblissement. Le tableau ci-dessous n'est donc valable que pour une onde incidente. Pour obtenir les équations associées à une onde réfléchie, il suffit de changer systématiquement γ par $-\gamma$ dans le tableau.

	Modes TEM	Modes TM	Modes TE
Equation de base	$\nabla_t^2 V = 0$	$(\nabla_t^2 + k_c^2) e_z = 0$	$(\nabla_t^2 + k_c^2) h_z = 0$
γ	$\sqrt{-\omega^2\mu\epsilon} = j\omega\sqrt{\mu\epsilon}$	$\sqrt{k_c^2 - \omega^2\mu\epsilon}$	$\sqrt{k_c^2 - \omega^2\mu\epsilon}$
E_z	0	$e_z \exp(-\gamma z)$	0
H_z	0	0	$h_z \exp(-\gamma z)$
$\mathbf{E}_t$	$-\nabla_t V \exp(-\gamma z)$	$-(\gamma / k_c^2) \nabla_t E_z$	$(j\omega\mu / \gamma) (\mathbf{H}_t \times \hat{\mathbf{z}})$
$\mathbf{H}_t$	$(\gamma / j\omega\mu) (\hat{\mathbf{z}} \times \mathbf{E}_t)$	$(j\omega\epsilon / \gamma) (\hat{\mathbf{z}} \times \mathbf{E}_t)$	$-(\gamma / k_c^2) \nabla_t H_z$
$Z_{\text{mod}} = \mathbf{e}_t / \mathbf{h}_t $	$\sqrt{\mu/\epsilon}$	$\sqrt{\omega^2\mu\epsilon - k_c^2} / \omega\epsilon$	$\omega\mu / \sqrt{\omega^2\mu\epsilon - k_c^2}$

3.3. Le guide à plaques parallèles

Références:

“Fields and Waves in Communication Electronics”, section 8.3

Un des systèmes de transmission électromagnétique le plus simple est le guide à plaque parallèles, constitué par deux plans métalliques parallèles séparés par une distance a (la “hauteur” du guide (figure 3.2)).

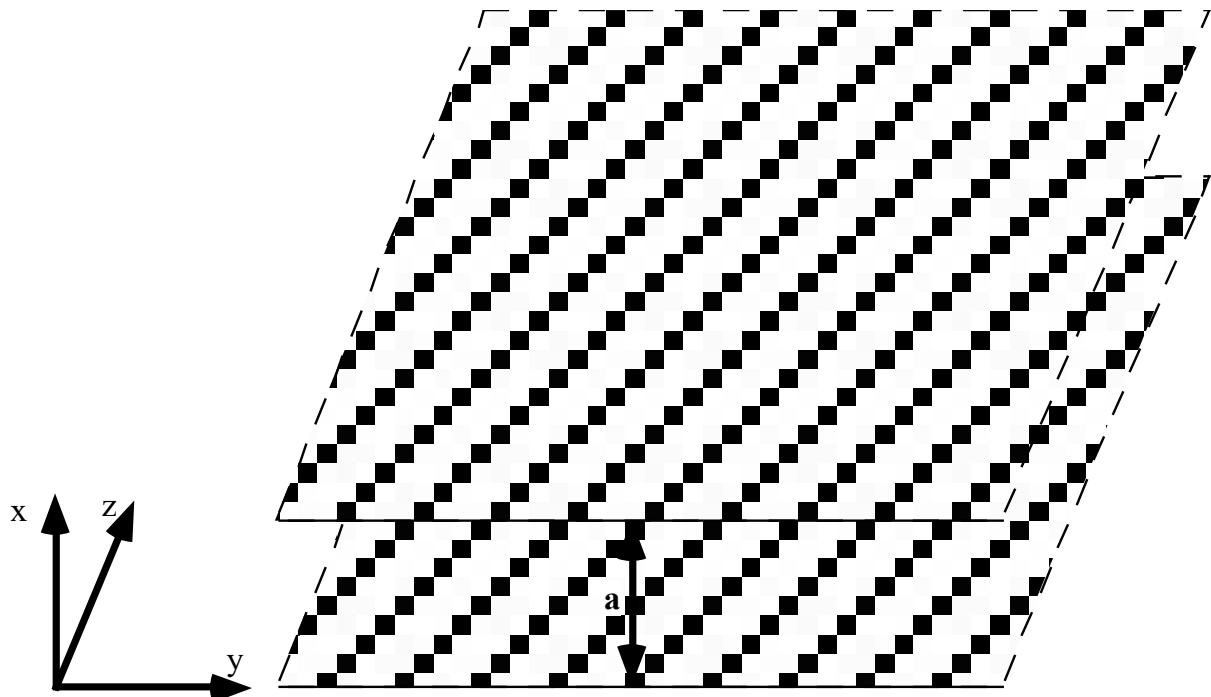


Figure 3.2 : guide à plaques parallèles

Pour l’analyse, on considère les plans métalliques infinis et on les place, respectivement, en $x=0$ et en $x=a$. La propagation s’effectue comme d’habitude selon z , et l’espace entre les plans est rempli d’un diélectrique de paramètres ϵ, μ .

Puisqu’on admet que la structure est infinie selon la coordonnée transverse y , les champs sont indépendantes de cette coordonnée. La résolution des équations résumées dans la section §3.2.5 est alors très simple. On trouve les résultats suivants :

	Modes TEM	Modes TM	Modes TE
k_c	0	$m\pi / a$	$m\pi / a$
ω_c	0	$m\pi / (a\sqrt{\mu\epsilon})$	$m\pi / (a\sqrt{\mu\epsilon})$
γ	$\sqrt{-\omega^2\mu\epsilon} = j\omega\sqrt{\mu\epsilon}$	$j\omega\sqrt{\mu\epsilon} \sqrt{1 - (\omega_c / \omega)^2}$	$j\omega\sqrt{\mu\epsilon} \sqrt{1 - (\omega_c / \omega)^2}$
E_z	0	$E_0 \sin \frac{m\pi x}{a} e^{-\gamma z}$	0
H_z	0	0	$H_0 \cos \frac{m\pi x}{a} e^{-\gamma z}$
Et	$E_0 \hat{x} e^{-\gamma z}$	$-\frac{\gamma a}{m\pi} E_0 \cos \frac{m\pi x}{a} \hat{x} e^{-\gamma z}$	$-\frac{j\omega\mu a}{m\pi} H_0 \sin \frac{m\pi x}{a} \hat{y} e^{-\gamma z}$
Ht	$\frac{E_0}{\sqrt{\mu/\epsilon}} \hat{y} e^{-\gamma z}$	$-\frac{j\omega\epsilon a}{m\pi} E_0 \cos \frac{m\pi x}{a} \hat{y} e^{-\gamma z}$	$\frac{\gamma a}{m\pi} H_0 \sin \frac{m\pi x}{a} \hat{x} e^{-\gamma z}$
$\mathbf{J}_s(x=0)$	$\frac{E_0}{\sqrt{\mu/\epsilon}} \hat{z} e^{-\gamma z}$	$-\frac{j\omega\epsilon a}{m\pi} E_0 \hat{z} e^{-\gamma z}$	$H_0 \hat{y} e^{-\gamma z}$
$\mathbf{J}_s(x=a)$	$-\mathbf{J}_s(x=0)$	$-(-1)^m \mathbf{J}_s(x=0)$	$-(-1)^m \mathbf{J}_s(x=0)$

3.4. Le Guide rectangulaire

La géométrie du guide rectangulaire est illustrée à la figure 3.3. Il consiste en un tube de section rectangulaire $a \times b$, que l'on suppose infini dans la direction z .

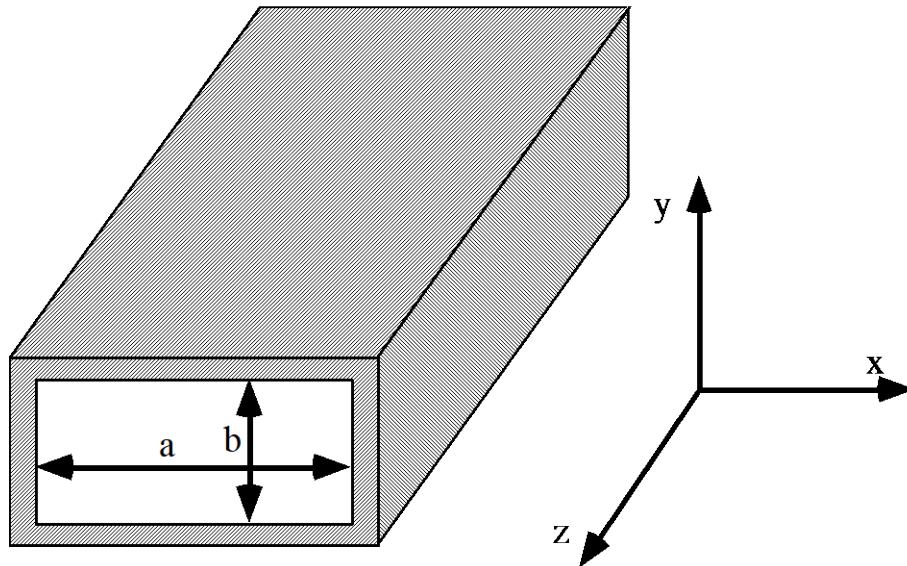


Figure 3.3 : Guide d'ondes rectangulaire

Le guide est formé de quatre parois conductrices situées en $x=0$, $x=a$, $y=0$ et $y=b$. La propagation de l'onde se fait selon la direction z . Comme ce type de guide n'est formé que d'un seul conducteur, il ne peut pas transmettre de modes TEM : en effet, à l'intérieur d'un guide creux formé d'un seul conducteur, aucun champ statique ne peut exister, et nous avons vu que la dépendance transverse des champs d'un mode TEM est identique au champ statique (cf. §3.2.2)

3.4.1 Modes TM

Les modes transverses magnétiques ont une composante longitudinale du champ électrique non nulle ($E_z \neq 0$) alors que la composante longitudinale du champ magnétique est nulle ($H_z = 0$). Nous résolvons donc l'équation d'ondes pour e_z :

$$\left(\nabla_t^2 + k_c^2\right)e_z = \frac{\partial^2 e_z}{\partial x^2} + \frac{\partial^2 e_z}{\partial y^2} + k_c^2 e_z = 0$$

Cette équation peut être résolue à l'aide de la technique de la séparation des variables. On suppose que

$$e_z = e_{zx}(x)e_{zy}(y)$$

donc que e_z peut être écrit comme le produit d'une fonction dépendant de x uniquement avec une fonction dépendant de y uniquement. L'équation d'ondes peut alors s'écrire sous la forme

$$\frac{d^2 e_{zx}}{dx^2} \frac{1}{e_{zx}} + \frac{d^2 e_{zy}}{dy^2} \frac{1}{e_{zy}} = -k_c^2$$

Cette relation devant être valable pour toutes les valeurs de x et de y , il est nécessaire que chacun des termes de la somme du membre de gauche soit égal à une constante :

$$\frac{d^2 e_{zx}}{dx^2} \frac{1}{e_{zx}} = -k_x^2, \quad \frac{d^2 e_{zy}}{dy^2} \frac{1}{e_{zy}} = -k_y^2, \quad k_x^2 + k_y^2 = k_c^2$$

Les solutions à ces deux équations différentielles sont données par

$$e_{zx} = A \cos k_x x + B \sin k_x x$$

$$e_{zy} = C \cos k_y y + D \sin k_y y$$

$$k_x^2 + k_y^2 = k_c^2$$

et donc

$$e_z = (A \cos k_x x + B \sin k_x x)(C \cos k_y y + D \sin k_y y)$$

$$k_x^2 + k_y^2 = k_c^2$$

Les six constantes A, B, C, D, k_x et k_y sont déterminées à l'aide des conditions aux limites : A doit être nul pour satisfaire la condition d'un champ tangentiel nul en $x=0$, C doit être nul pour satisfaire la condition d'un champ nul en $y=0$. Pour obtenir des champs tangentiels nuls en $x=a$ et $y=b$, nous avons deux solutions : soit B ou D est nul, et la solution est triviale. Soit

$$k_x = \frac{m\pi}{a} \quad \text{et} \quad k_y = \frac{n\pi}{b} \quad mn \neq 0$$

Nous obtenons donc finalement

$$\begin{aligned}
e_z &= E_0 \sin \frac{m\pi}{a} x \sin \frac{n\pi}{b} y \\
E_z &= E_0 \sin \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z} \\
\gamma &= \sqrt{k_c^2 - \omega^2 \epsilon \mu}
\end{aligned}$$

On remarque que ces modes ne peuvent se propager que pour un γ imaginaire, donc pour

$$\omega > \omega_c = \frac{1}{\sqrt{\epsilon \mu}} \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}$$

En dessous de cette pulsation, γ est réel et l'onde est affaiblie dans le guide.

3.4.2 Modes TE

Les modes transverses électriques ont une composante longitudinale du champ magnétique non nulle ($H_z \neq 0$) alors que la composante longitudinale du champ électrique est nulle ($E_z = 0$).

Nous résolvons donc l'équation d'ondes pour H_z :

$$(\nabla_t^2 + k_c^2)h_z = \frac{\partial^2 h_z}{\partial x^2} + \frac{\partial^2 h_z}{\partial y^2} + k_c^2 h_z = 0$$

Cette équation peut à nouveau être résolue à l'aide de la technique de la séparation des variables, et on obtient :

$$\begin{aligned}
h_z &= (A \cos k_x x + B \sin k_x x)(C \cos k_y y + D \sin k_y y) \\
k_x^2 + k_y^2 &= k_c^2
\end{aligned}$$

Le calcul des constantes est un peu moins direct dans ce cas. On commence par dériver e_x et e_y à partir de h_z :

$$\begin{aligned}
e_x &= -\frac{j\omega\mu}{k_c^2} \frac{\partial h_z}{\partial y} = -\frac{j\omega\mu k_y}{k_c^2} (A \cos k_x x + B \sin k_x x)(-C \sin k_y y + D \cos k_y y) \\
e_y &= \frac{j\omega\mu}{k_c^2} \frac{\partial h_z}{\partial x} = \frac{j\omega\mu k_x}{k_c^2} (-A \sin k_x x + B \cos k_x x)(-C \sin k_y y + D \cos k_y y)
\end{aligned}$$

Deux constantes, B et D doivent être nulles pour que e_x s'annule en $y=0$ et e_y en $x=0$. De même, e_x doit s'annuler en $y=b$ et e_y en $x=a$. La seule solution non triviale est alors donnée par :

$$k_x = \frac{m\pi}{a} \quad , \quad k_y = \frac{n\pi}{b}$$

Nous obtenons donc finalement :

$$\begin{aligned} h_z &= H_0 \cos \frac{m\pi}{a} x \cos \frac{n\pi}{b} y \\ H_z &= H_0 \cos \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z} \\ \gamma &= \sqrt{k_c^2 - \omega^2 \varepsilon \mu} \end{aligned}$$

On remarque que ces modes ne peuvent se propager que pour un γ imaginaire, donc pour

$$\omega > \omega_c = \frac{1}{\sqrt{\varepsilon \mu}} \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}$$

En dessous de cette pulsation, γ est réel et l'onde est affaiblie dans le guide.

3.4.3 Tableau récapitulatif

	Modes TM	Modes TE
k_c	$\sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}, mn \neq 0$	$\sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}, m+n \neq 0$
ω_c	$\frac{1}{\sqrt{\mu\varepsilon}} \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}$	$\frac{1}{\sqrt{\mu\varepsilon}} \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2}$
α	$k_c \sqrt{1 - \left(\frac{\omega}{\omega_c}\right)^2}, \omega < \omega_c$	$k_c \sqrt{1 - \left(\frac{\omega}{\omega_c}\right)^2}, \omega < \omega_c$
β	$k \sqrt{1 - \left(\frac{\omega_c}{\omega}\right)^2}, \omega > \omega_c$	$k \sqrt{1 - \left(\frac{\omega_c}{\omega}\right)^2}, \omega > \omega_c$
Ez	$E_0 \sin \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z}$	0
H _z	0	$H_0 \cos \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z}$
E _x	$-\frac{\gamma m \pi}{a k_c^2} E_0 \cos \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z}$	$\frac{j \omega \varepsilon n \pi}{b k_c^2} H_0 \cos \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z}$
E _y	$-\frac{\gamma n \pi}{b k_c^2} E_0 \sin \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z}$	$-\frac{j \omega \varepsilon m \pi}{a k_c^2} H_0 \sin \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z}$
H _x	$\frac{j \omega \varepsilon n \pi}{b k_c^2} E_0 \sin \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z}$	$\frac{\gamma m \pi}{a k_c^2} H_0 \sin \frac{m\pi}{a} x \cos \frac{n\pi}{b} y e^{-\gamma z}$
H _y	$-\frac{j \omega \varepsilon m \pi}{a k_c^2} E_0 \cos \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z}$	$\frac{\gamma n \pi}{b k_c^2} H_0 \cos \frac{m\pi}{a} x \sin \frac{n\pi}{b} y e^{-\gamma z}$
Z	$\sqrt{\frac{\mu}{\varepsilon}} \sqrt{1 - \left(\frac{\omega_c}{\omega}\right)^2}$	$\frac{\sqrt{\frac{\mu}{\varepsilon}}}{\sqrt{1 - \left(\frac{\omega_c}{\omega}\right)^2}}$

3.5 Le guide circulaire

La géométrie d'un guide circulaire est décrite à la figure 3.4. Il consiste en un tube de section circulaire de rayon a , que l'on suppose infini dans la direction z .

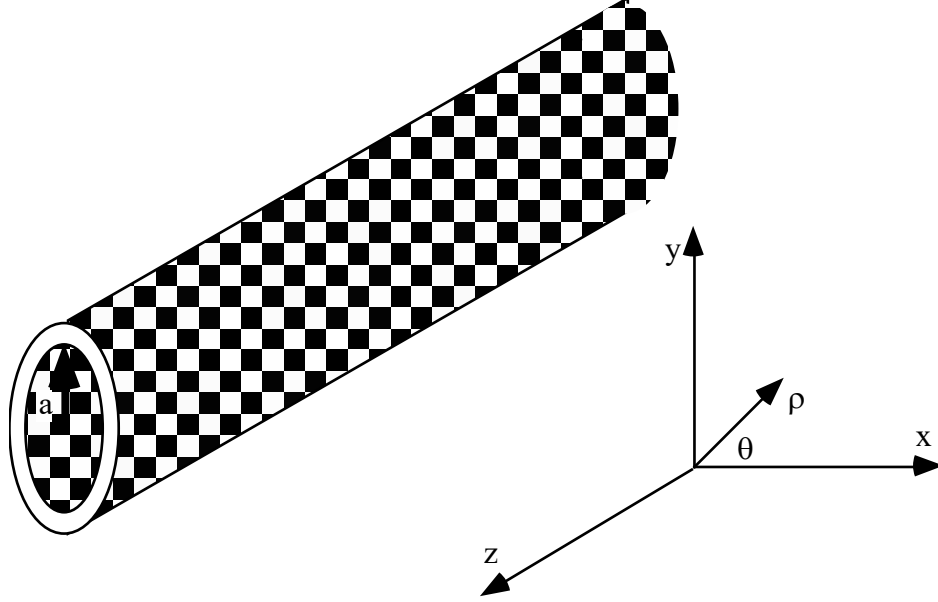


Figure 3.3 : Guide d'ondes circulaire

De même que pour le guide rectangulaire ce type de structure ne peut pas supporter de modes TEM. Avant d'étudier les modes TM et TE, il peut être judicieux d'exprimer les champs transverse (§3.1.4) en coordonnées cylindriques :

$$\begin{aligned}
 e_r &= \frac{-1}{k_c^2} \left[\gamma \frac{\partial e_z}{\partial r} + \frac{j\omega\mu}{r} \frac{\partial h_z}{\partial \phi} \right] \\
 e_\phi &= \frac{1}{k_c^2} \left[-\frac{\lambda}{r} \frac{\partial e_z}{\partial \phi} + j\omega\mu \frac{\partial h_z}{\partial r} \right] \\
 h_r &= \frac{1}{k_c^2} \left[\frac{j\omega\varepsilon}{r} \frac{\partial e_z}{\partial \phi} - \gamma \frac{\partial h_z}{\partial r} \right] \\
 h_\phi &= \frac{-1}{k_c^2} \left[j\omega\varepsilon \frac{\partial e_z}{\partial r} + \frac{\gamma}{r} \frac{\partial h_z}{\partial \phi} \right]
 \end{aligned}$$

où

$$k_c^2 = \gamma^2 + k_0^2 = k_0^2 - \beta^2$$

3.5.1 Modes TM

On résout l'équation d'ondes pour e_z en coordonnées cylindriques :

$$\left(\frac{\partial^2}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2}{\partial \varphi^2} + k_c^2 \right) e_z = 0$$

On applique à nouveau la technique de séparation des variables :

$$e_z(\rho, \varphi) = R(\rho)P(\varphi)$$

Ce qui nous donne pour l'équation d'ondes :

$$\frac{\rho^2}{R} \frac{d^2 R}{d\rho^2} + \frac{\rho}{R} \frac{dR}{d\rho} + \rho^2 k_c^2 = -\frac{1}{P} \frac{d^2 P}{d\varphi^2}$$

Le membre de gauche de cette équation ne dépendant que de ρ et le membre de droite que de φ , les deux côtés doivent être égaux à une constante k_φ^2 :

$$\begin{aligned} -\frac{1}{P} \frac{d^2 P}{d\varphi^2} &= k_\varphi^2, \quad \frac{d^2 P}{d\varphi^2} + P k_\varphi^2 = 0 \\ \frac{\rho^2}{R} \frac{d^2 R}{d\rho^2} + \frac{\rho}{R} \frac{dR}{d\rho} + \rho^2 k_c^2 &= k_\varphi^2, \quad \rho^2 \frac{d^2 R}{d\rho^2} + \rho \frac{dR}{d\rho} + R(\rho^2 k_c^2 - k_\varphi^2) = 0 \end{aligned}$$

La solution générale de l'équation en φ est de la forme

$$P(\varphi) = A \sin k_\varphi \varphi + B \cos k_\varphi \varphi$$

La solution doit être périodique en φ , ce qui veut dire que k_φ doit être entier.

$$P(\varphi) = A \sin n\varphi + B \cos n\varphi$$

L'équation en ρ prend alors la forme :

$$\rho^2 \frac{d^2 R}{d\rho^2} + \rho \frac{dR}{d\rho} + R(\rho^2 k_c^2 - n^2) = 0$$

Cette équation est une équation de Bessel, qui a des solutions de la forme suivante :

$$R(\rho) = C J_n(k_c \rho) + D Y_n(k_c \rho)$$

J_n et Y_n sont des fonctions de Bessel du premier et second type, d'ordre n . Les fonctions de Bessel du second type deviennent infinies pour un argument nul. Nous voyons donc immédiatement que D doit être égal à zéro. Nous obtenons donc finalement :

$$e_z(\rho, \varphi) = (A \sin n\varphi + B \cos n\varphi) J_n(k_c \rho)$$

où la constante C a été absorbée dans A et B . Il nous reste maintenant à déterminer le nombre d'onde de coupure k_c . Les conditions aux limites impliquent que $e_z(\rho, \varphi)$ s'annule pour $\rho=a$. Nous devons donc avoir :

$$J_n(k_c a) = 0 \quad \text{donc} \quad k_c = \frac{p_{nm}}{a}$$

où p_{nm} est le $m^{\text{ième}}$ zéro de la fonction de Bessel du premier type d'ordre n . Les zéros des fonctions de Bessel se trouvent dans toutes les bonnes tables.

La constante de propagation du mode TM_{nm} est donc donnée par :

$$\beta_{nm} = \sqrt{k_0^2 - k_c^2} = \sqrt{k_0^2 - \left(\frac{p_{nm}}{a}\right)^2}$$

et la fréquence de coupure par

$$f_{c_{nm}} = \frac{k_c}{2\pi\sqrt{\mu\epsilon}} = \frac{p_{nm}}{2\pi a\sqrt{\mu\epsilon}}$$

Le premier mode TM à se propager est donc le mode TM_{01} , obtenu pour le premier zéro de la fonction de Bessel d'ordre 1 $p_{01}=2.405$. Il n'a pas de mode TM_{10} , car $m \geq 1$. Toutes les composantes des champs électriques et magnétiques sont aisément obtenues à partir de $e_z(\rho, \varphi)$. On constate que les solutions contiennent deux constantes indépendantes A et B . Les valeurs de ces dernières vont dépendre de l'excitation du guide d'onde. Nous avons deux constantes dans le cas des guides circulaires, car à cause de la symétrie azimutale du problème les solutions peuvent avoir une dépendance sinusoïdale et/ou cosinoïdale. Il est possible de fixer A ou B nul en choisissant l'orientation du système de coordonnées cylindriques.

3.5.2 Modes TE

Le calcul est identique au cas TM. L'équation d'onde à résoudre est :

$$\left(\frac{\partial^2}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2}{\partial \varphi^2} + k_c^2 \right) h_z = 0$$

En appliquant les mêmes considérations que pour les modes TM, on obtient comme solution :

$$h_z(\rho, \varphi) = (A \sin n\varphi + B \cos n\varphi) J_n(k_c \rho)$$

Le nombre d'onde est obtenu à l'aide des conditions aux limites : $e_\varphi(\rho, \varphi)$ doit s'annuler en $\rho=a$. Nous avons d'après le début du §3.5

$$e_\varphi(\rho, \varphi) = \frac{j\omega\mu}{k_c} (A \sin n\varphi + B \cos n\varphi) J'_n(k_c \rho)$$

où J'_n est la dérivée de J_n par rapport à son argument. La condition au limite impose donc

$$J'_n(k_c a) = 0 \quad \text{donc} \quad k_{c_{nm}} = \frac{p'_{nm}}{a}$$

où p'_{nm} est le $m^{\text{ième}}$ zéro de la dérivée de la fonction de Bessel du premier type d'ordre n . Les modes TE_{nm} sont donc définis par le nombre d'onde de coupure $k_{c_{nm}}$. La constante de propagation est donnée par

$$\beta_{nm} = \sqrt{k^2 - k_{c_{nm}}^2} = \sqrt{k^2 - \left(\frac{p'_{nm}}{a}\right)^2}$$

et la fréquence de coupure correspondante par

$$f_{c_{nm}} = \frac{p'_{nm}}{2\pi a \sqrt{\mu\epsilon}}$$

De nouveau, nous avons $m \geq 1$. Il est intéressant de noter que le mode ayant la plus petite fréquence de coupure est le mode TE_{11} , correspondant à $p'_{11}=1.8141$. En effet $p'_{01}=3.832$, ce qui donne une fréquence de coupure plus élevée pour le mode TE_{01} .

Le mode TE_{11} est appelé le mode dominant du guide circulaire, car il a la fréquence de coupure la plus faible de tous les modes TE et TM.

3.5.3 Tableau récapitulatif

	Modes TM	Modes TE
k_c	$\frac{p_{nm}}{a}$	$\frac{p'_{nm}}{a}$
ω_c	$\frac{p_{nm}}{a\sqrt{\mu\varepsilon}}$	$\frac{p'_{nm}}{a\sqrt{\mu\varepsilon}}$
α	$k_c\sqrt{1-\left(\frac{\omega}{\omega_c}\right)^2}, \quad \omega < \omega_c$	$k_c\sqrt{1-\left(\frac{\omega}{\omega_c}\right)^2}, \quad \omega < \omega_c$
β	$k\sqrt{1-\left(\frac{\omega_c}{\omega}\right)^2}, \quad \omega > \omega_c$	$k\sqrt{1-\left(\frac{\omega_c}{\omega}\right)^2}, \quad \omega > \omega_c$
Ez	$(A\sin n\varphi + B\cos n\varphi)J_n(k_c\rho)e^{-\gamma z}$	0
H _z	0	$(A\sin n\varphi + B\cos n\varphi)J_n(k_c\rho)e^{-\gamma z}$
E _ρ	$\frac{-\gamma}{k_c}(A\sin n\varphi + B\cos n\varphi)J'_n(k_c\rho)e^{-\gamma z}$	$\frac{-j\omega\mu n}{k_c^2\rho}(A\cos n\varphi - B\sin n\varphi)J_n(k_c\rho)e^{-\gamma z}$
E _φ	$\frac{-\gamma n}{k_c^2\rho}(A\cos n\varphi - B\sin n\varphi)J_n(k_c\rho)e^{-\gamma z}$	$\frac{j\omega\mu}{k_c}(A\sin n\varphi + B\cos n\varphi)J'_n(k_c\rho)e^{-\gamma z}$
H _ρ	$\frac{j\omega\varepsilon n}{k_c^2\rho}(A\cos n\varphi - B\sin n\varphi)J_n(k_c\rho)e^{-\gamma z}$	$\frac{-\gamma}{k_c}(A\sin n\varphi + B\cos n\varphi)J'_n(k_c\rho)e^{-\gamma z}$
H _φ	$\frac{j\omega\varepsilon}{k_c}(A\sin n\varphi + B\cos n\varphi)J'_n(k_c\rho)e^{-\gamma z}$	$\frac{-\gamma n}{k_c^2\rho}(A\cos n\varphi - B\sin n\varphi)J_n(k_c\rho)e^{-\gamma z}$
Z	$\sqrt{\frac{\mu}{\varepsilon}}\sqrt{1-\left(\frac{\omega_c}{\omega}\right)^2}$	$\frac{\sqrt{\frac{\mu}{\varepsilon}}}{\sqrt{1-\left(\frac{\omega_c}{\omega}\right)^2}}$

3.5.4 Zéros des fonctions de Bessel du premier type

n	p_{n1}	p_{n2}	p_{n3}
0	2.405	5.520	8.654
1	3.832	7.016	10.174
2	5.135	8.417	11.620

3.5.5 Zéros des dérivées des fonctions de Bessel du premier type

n	p'_{n1}	p'_{n2}	p'_{n3}
0	3.832	7.016	10.174
1	1.841	5.331	8.536
2	3.054	6.706	9.970

4. Circuits imprimés hyperfréquences

Références:

S. Ramo, J.R. Whinnery and T. van Duzer, "Fields and Waves in Communication Electronics", Wiley, 1984, § 8.6.

F. E. Gardiol, "Hyperfréquences", volume XIII du Traité d'Electricité, Presses Polytechniques Romandes, § 2.11.

K.C. Gupta, R. Garg & I.J. Bahl, "Microstrip Lines and Slotlines", Artech House, Dedham MA, 1979

R.K Hoffmann, "Handbook of Microwave Integrated Circuits", Artech House, Dedham MA, 1987

4.1 Introduction

Depuis le début des années 50, les circuits imprimés hyperfréquences remplacent de plus en plus les câbles coaxiaux et les guides d'ondes dans les applications commerciales. Bien que présentant plus de pertes et étant sujet à une dispersion en fréquence, ces circuits sont légers, relativement bon marchés et surtout facilement productibles en série. Comme tous les circuits imprimés, ils sont en effet fabriqués à l'aide de procédés photo-lithographiques, qui assurent une grande répétabilité et une production de masse aisée.

Il existe plusieurs types de circuits imprimés hyperfréquences, dont le plus populaire est sans conteste le circuit microruban.

4.2 Ligne à ruban équilibré (stripline)

4.2.1 Définition

La structure à ruban équilibré est illustrée à la figure 4.1. Elle consiste en une piste conductrice prise en sandwich entre deux plans de masse.

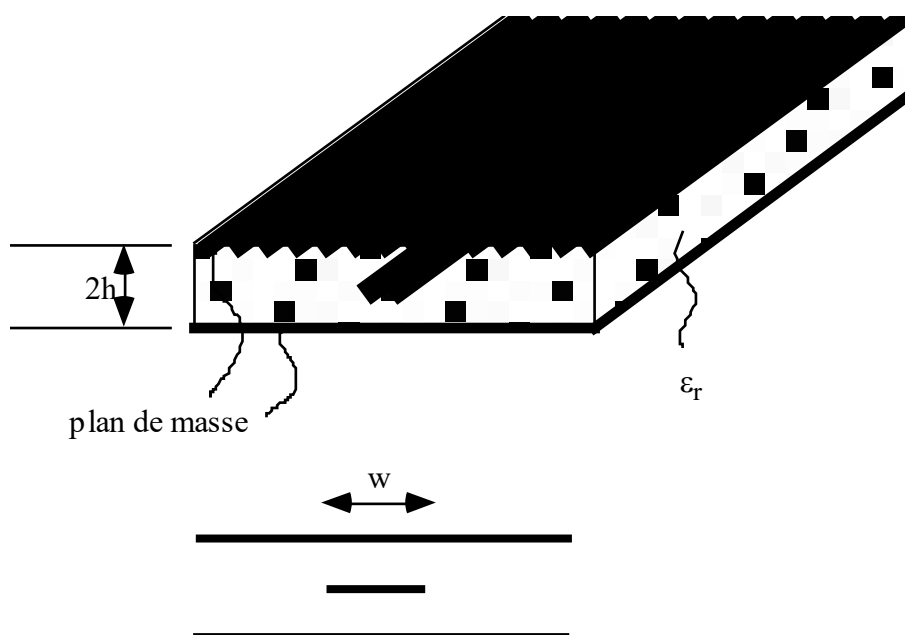


Figure 4.1 : Ligne à ruban équilibré

Cette structure est caractérisée par deux surfaces conductrices (plans de masse) séparés d'une distance $2h$. Le volume entre les plans de masse est rempli d'un diélectrique de permittivité ϵ_r . Un ruban conducteur de largeur w est situé entre les deux plans de masse.

4.2.2 Modes de propagation dans une structure à ruban équilibré

L'analyse de ce type de structure est malheureusement beaucoup plus compliquée que l'analyse des guides d'ondes vue dans les paragraphes précédent. Nous allons donc nous concentrer sur le mode dominant uniquement, car dans la très grande majorité des cas c'est le seul à être utilisé dans les lignes à ruban équilibré.

Ce mode dominant de ce type de structure est le mode TEM. En effet, le milieu entre les deux plans de masse est homogène, et le guide est formé de deux conducteurs : les deux plans de masse (qui sont au même potentiel, donc considérés comme un seul conducteur) et le conducteur central. Il faut donc résoudre l'équation de Laplace pour caractériser ce mode. Une solution exacte de cette équation peut être obtenue à l'aide de la théorie des transformations conformes (voir par exemple "Stripline Circuit Design", par H.Howe Jr., Artech House, Dedham Ma, 1974), mais cette procédure est complexe et donne des résultats sous une forme peu pratique à utiliser.

On préférera donner ici des expressions analytiques qui sont une bonne approximation de la solution rigoureuse, et qui sont surtout d'un emploi beaucoup plus commode.

4.2.3 Constante de propagation

Dans le cas d'un milieu non magnétique, la vitesse de phase d'un mode TEM est donné par :

$$v_\phi = \frac{1}{\sqrt{\epsilon_r \epsilon_0 \mu_0}} = \frac{c}{\sqrt{\epsilon_r}}$$

où c est la vitesse de la lumière dans le vide. On en déduit la constante de propagation :

$$\beta = \frac{\omega}{v_\phi} = \omega \sqrt{\epsilon_r \epsilon_0 \mu_0} = \sqrt{\epsilon_r} k_0$$

4.2.4 Impédance caractéristique

L'impédance caractéristique d'une ligne de transmission supportant un mode TEM est donnée par :

$$Z_c = \sqrt{\frac{L}{C}}$$

où L est l'inductance par unité de longueur de la ligne et C sa capacité par unité de longueur. Ces deux grandeurs sont obtenues en résolvant l'équation de Laplace numériquement et en faisant passer une courbe par les points obtenus numériquement (curve fitting). On obtient finalement pour l'impédance caractéristique (formule approchée) :

$$Z_c = \frac{30\pi}{\sqrt{\epsilon_r}} \frac{2h}{w_e + 0.441(2h)}$$

où w_e est la largeur effective du conducteur central, donnée par :

$$\frac{w_e}{2h} = \frac{w}{2h} - \begin{cases} 0 & \text{pour } \frac{w}{2h} > 0.35 \\ \left(0.35 - \frac{w}{2h}\right)^2 & \text{pour } \frac{w}{2h} < 0.35 \end{cases}$$

Ces formules sont valables pour un conducteur central d'épaisseur nulle, et sont exactes à 1% près. On note que l'impédance caractéristique de la ligne diminue lorsque la largeur du conducteur central augmente.

Dans une procédure de conception de circuit, on souhaite souvent avoir l'information inverse, c'est à dire obtenir la largeur du ruban en fonction d'une impédance caractéristique souhaitée. Ces grandeurs sont liées par la relation approchée :

$$\frac{w}{2h} = \begin{cases} x & \text{pour } \sqrt{\epsilon_r} Z_c < 120 \\ 0.85 - \sqrt{0.6 - x} & \text{pour } \sqrt{\epsilon_r} Z_c > 120 \end{cases}$$

avec

$$x = \frac{30\pi}{\sqrt{\epsilon_r} Z_c} - 0.441$$

4.2.5 Affaiblissement dans une ligne à ruban équilibré

Les pertes dans un ruban à lignes équilibré ont deux sources : les pertes diélectriques et les pertes ohmiques. Les premières sont les mêmes pour toutes les lignes TEM et sont données par

$$\alpha_d = \frac{k \tan \delta}{2} \quad [Np / m]$$

où k est le nombre d'onde dans le milieu et δ sa tangente de perte :

$$k = \omega \sqrt{\epsilon_r \epsilon_0 \mu_0} = \frac{\omega \sqrt{\epsilon_r}}{c}$$

$$\tan \delta = \frac{\epsilon''}{\epsilon'}$$

Les pertes ohmiques sont calculées à l'aide d'une méthode de perturbations. On obtient la relation approchée :

$$\alpha_c = \begin{cases} \frac{0.0027 R_s \epsilon_r Z_c}{30 \pi (2h - t)} A & \text{pour } \sqrt{\epsilon_r} Z_c < 120 \\ \frac{0.16 R_s}{Z_c 2h} B & \text{pour } \sqrt{\epsilon_r} Z_c > 120 \end{cases} \quad [Np / m]$$

$$A = 1 + \frac{2w}{2h - t} + \frac{2h + t}{\pi(2h - t)} \ln \left(\frac{4h - t}{t} \right)$$

$$B = 1 + \frac{2h}{(0.5w + 0.7t)} \left(0.5 + \frac{0.414t}{w} + \frac{1}{2\pi} \ln \frac{4\pi w}{t} \right)$$

où t est l'épaisseur du ruban et R_s est la résistance surfacique du conducteur.

$$R_s = \sqrt{\frac{\omega \mu_0}{2\sigma}}$$

où σ est la conductivité du métal.

4.3 Ligne microruban (microstrip)

4.3.1 Définition

Une ligne microruban (microstrip line en anglais) est constituée d'un conducteur métallique mince et étroit, *le ruban*, déposé sur une face d'une plaque diélectrique, *le substrat*. L'autre côté de la plaque est entièrement recouverte d'un conducteur, *le plan de masse*. Cette structure est illustrée à la figure 4.2.

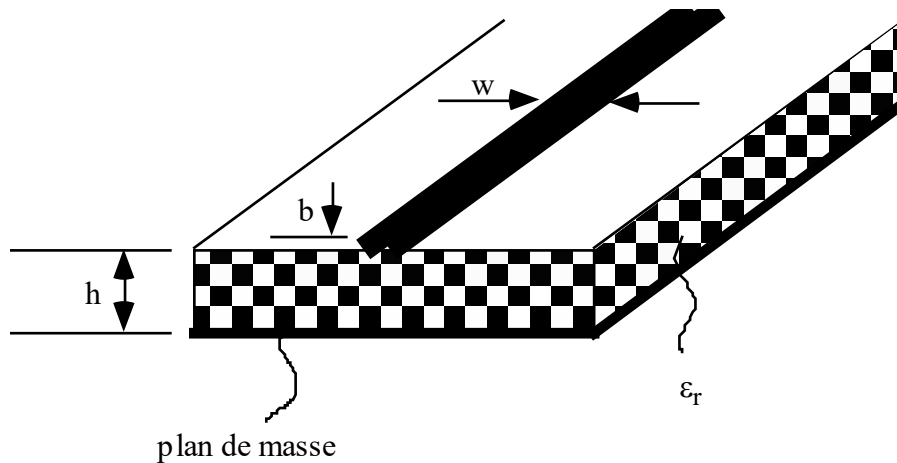


Figure 4.2: ligne microruban

Les paramètres caractéristiques de la ligne sont :

- La permittivité relative du substrat ϵ_r .
- L'épaisseur h du substrat, généralement quelques fractions de longueur d'ondes.
- La largeur w du ruban. Cette largeur est en général du même ordre de grandeur que l'épaisseur h du substrat ($0.1 h \leq w \leq 10 h$).
- L'épaisseur b du ruban, généralement petite ($b/h \ll 1$).

Ces paramètres influencent sur :

- La concentration du champ électrique dans le substrat (absence de rayonnement). Plus la constante diélectrique du substrat est élevée, plus les champs sont concentrés sous la ligne et moins elle rayonne.
- L'impédance caractéristique de la ligne, qui dépend principalement de la permittivité et du rapport w/h .

4.3.2 Modes de propagation

En première approximation, on peut considérer que la ligne microruban est la moitié d'une ligne à ruban équilibré (figure 4.3).



Figure 4.3

Si le diélectrique n'était pas présent, une ligne microruban pourrait être considérée comme une ligne bifilaire, constituée de deux conducteurs de largeur w séparés d'une distance $2h$ (le plan de masse agissant comme un miroir). Une telle structure supporterait donc un mode TEM.

La présence du diélectrique sous le ruban rend la structure inhomogène dans le plan transverse. Les modes que propagent une telle structure sont des modes hybrides. On comprend en effet bien qu'une telle structure ne peut pas propager de mode TEM : la vitesse de phase de ce mode serait en effet égale à $c / \sqrt{\epsilon_r}$ dans le diélectrique alors qu'elle devrait être égale à la vitesse de la lumière au dessus du conducteur. La désadaptation de la phase nécessite donc l'introduction des composantes longitudinales du champ électromagnétique pour satisfaire les conditions aux limites à l'interface air-diélectrique.

Dans la plupart des applications pratiques des lignes microruban, le substrat diélectrique est électriquement mince : $h < 0.05 \lambda$. Par conséquent, les composantes longitudinales du champ électromagnétique restent très faibles. On a donc un mode quasi-TEM. Cela signifie que la distribution des champs est très proche que celle obtenue dans le cas correspondant à la figure 4.4, où le ruban est dans un milieu diélectrique homogène.

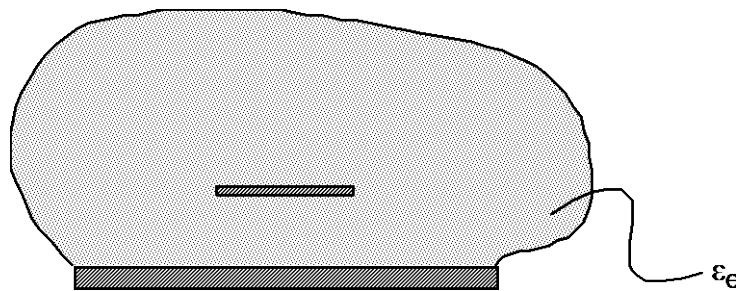


Figure 4.5 : structure homogène équivalente

Ce type de structure supporte un mode TEM, caractérisé par :

$$\begin{aligned} v_{\phi} &= \frac{c}{\sqrt{\epsilon_e}} \\ \beta &= k_0 \sqrt{\epsilon_e} \\ 1 < \epsilon_e < \epsilon_r \end{aligned}$$

Les lignes microruban ont aussi été étudiée à l'aide de techniques numériques. Les résultats obtenus ont été approchés par des formules analytiques.

4.3.3 Permittivité effective d'un circuit microruban

La permittivité effective d'un circuit microruban est la permittivité qu'aurait la structure homogène qui simule le mieux le circuit microruban. Pour une épaisseur nulle du ruban, une formule approchée de la permittivité effective est donnée par :

$$\varepsilon_e = \frac{1}{2}(\varepsilon_r + 1) + \frac{1}{2}(\varepsilon_r - 1) \left[\left(1 + 12 \frac{h}{w} \right)^{-0.5} + 0.04 \left(1 - \frac{w}{h} \right)^2 \right] \quad \text{pour } \frac{w}{h} \leq 1$$

$$\varepsilon_e = \frac{1}{2}(\varepsilon_r + 1) + \frac{1}{2}(\varepsilon_r - 1) \left(1 + 12 \frac{h}{w} \right)^{-0.5} \quad \text{pour } \frac{w}{h} \geq 1$$

L'erreur relative de ces relations approchées est inférieure à 1% lorsque

$$0.05 \leq w/h \leq 20 \quad \text{et} \quad \varepsilon_r \leq 16$$

On obtient pour la vitesse de phase et la longueur d'onde :

$$v_\phi = \frac{c}{\sqrt{\varepsilon_e}}$$

$$\lambda_g = \frac{\lambda_0}{\sqrt{\varepsilon_e}}$$

4.3.4 Impédance caractéristique

Pour une épaisseur de ruban nulle, les formules suivantes donnent une bonne approximation de l'impédance caractéristique d'une ligne microruban (erreur relative inférieure à 1% pour $0.05 \leq w/h \leq 20$:

$$Z_c = \frac{Z_0}{2\pi\sqrt{\varepsilon_e}} \ln \left(\frac{8h}{w} + \frac{w}{4h} \right) \quad \text{pour } \frac{w}{h} \leq 1$$

$$Z_c = \frac{Z_0}{\sqrt{\varepsilon_e}} \left(\frac{w}{h} + 1.393 + 0.667 \ln \left(\frac{w}{h} + 1.444 \right) \right)^{-1} \quad \text{pour } \frac{w}{h} > 1$$

où $Z_0 = 120\pi$ est l'impédance caractéristique du vide.

Dans une procédure de conception de circuit, on souhaite souvent avoir l'information inverse, c'est à dire obtenir la largeur du ruban en fonction d'une impédance caractéristique souhaitée. Ces grandeurs sont liées par la relation approchée :

$$\frac{w}{h} = 4 \left[\frac{1}{2} e^A - e^{-A} \right]^{-1} \quad \text{pour} \quad \frac{w}{h} \leq 2$$

$$\frac{w}{h} = \frac{\epsilon_r - 1}{\pi \epsilon_r} \left(\ln(B - 1) + 0.39 - \frac{0.61}{\epsilon_r} \right) + \frac{2}{\pi} (B - 1 - \ln(2B - 1)) \quad \text{pour} \quad \frac{w}{h} \geq 2$$

avec

$$A = \frac{Z_c}{Z_o} \pi \sqrt{2(\epsilon_r + 1)} + \frac{\epsilon_r - 1}{\epsilon_r + 1} \left(0.23 + \frac{0.11}{\epsilon_r} \right)$$

$$B = \frac{\pi}{2\sqrt{\epsilon_r}} \frac{Z_o}{Z_c}$$

On constate que, de même manière que pour la ligne à ruban équilibré, l'impédance caractéristique d'une ligne microruban diminue lorsque la largeur du ruban augmente.

4.3.5 Affaiblissement dans une ligne microruban

Les pertes dans un microruban ont deux sources : les pertes diélectriques et les pertes ohmiques. Les premières sont données par

$$\alpha_d = \frac{k_o \epsilon_r (\epsilon_e - 1) \tan \delta}{2\sqrt{\epsilon_e} (\epsilon_r - 1)} \quad [Np / m]$$

où k_o est le nombre d'onde dans le vide et δ la tangente de perte du diélectrique.

Les pertes ohmiques sont données approximativement par :

$$\alpha_c = \frac{R_s}{w Z_c} \quad [Np / m]$$

où R_s est la résistance surfacique du conducteur.

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

et σ est la conductivité du métal.

4.3.6 Rayonnement des lignes microruban

Le rayonnement d'une ligne microruban est lié à l'apparition de modes supérieurs non guidés. Ces derniers sont excités aux voisinage de discontinuités, comme un saut dans la largeur du

ruban, un coude, ou la fin de la ligne. Pour une impédance de ligne caractéristique de 50Ω , on peut calculer la fréquence f_m pour laquelle la proportion de puissance rayonnée reste inférieure à 1% de la puissance totale :

$$f_m [GHz] = \frac{2.14 \sqrt{\epsilon_r}}{h [mm]}$$

Pour une application à une fréquence élevée, il faut donc soit choisir un substrat diélectrique de grande permittivité, soit utiliser un substrat mince.

4.3.7 Dispersion d'un ligne microruban

L'approximation quasi-TEM utilisée dans les paragraphes précédents néglige les composantes longitudinales des champs électromagnétiques, et ne permet donc pas de prévoir la dispersion en fréquence du comportement de ces lignes. La concentration du champ électrique dans le substrat augmente en effet avec la fréquence, ce qui nous laisse penser que la permittivité effective, la constante de propagation et l'impédance caractéristique de la ligne seront aussi fonction de la fréquence. L'étude rigoureuse de ce comportement devient très complexe, et pour une utilisation pratique des microrubans on fait de nouveau usage d'une expression approchée pour la permittivité effective :

$$\epsilon_{ed}(f) = \epsilon_r - \frac{\epsilon_r - \epsilon_e}{1 + \left(\frac{f}{f_p}\right)^2} G$$

$$f_p = \frac{Z_c}{2h\mu_o}$$

$$G = 0.6 + 0.009Z_c$$

On utilise alors ϵ_{ed} plutôt que ϵ_e dans le calcul de l'impédance de la ligne, de la longueur d'onde et de la vitesse de phase. Lorsque $f \ll f_p$, cette correction n'est pas nécessaire.

4.4 Guides coplanaires

Références:

K.C. Gupta, R. Garg & I.J. Bahl, "Microstrip Lines and Slotlines", Artech House, Dedham MA, 1979.

T.Q. Deng, M.S. Leong & P.S. Kooi, "Accurate formulas for coplanar waveguide synthesis", Electronics Letters, Vol. 31, 1995, pp. 2017-2019.

4.4.1 Définition

Les guides d'ondes coplanaires connaissent un regain d'intérêt depuis quelques années, principalement comme lignes de transmission dans le domaine des ondes millimétriques (30GHz-300GHz). Elles sont en effet de fabrication beaucoup moins coûteuse que les guides d'onde traditionnels, et présentent moins de pertes haute fréquence que les lignes microruban. Un guide coplanaire est illustré à la figure 4.6. Il consiste en une ligne située du même côté que le plan de masse sur un substrat diélectrique. La ligne, de largeur s , est séparée de la masse par deux fentes de largeur w . Le substrat diélectrique a une épaisseur h .

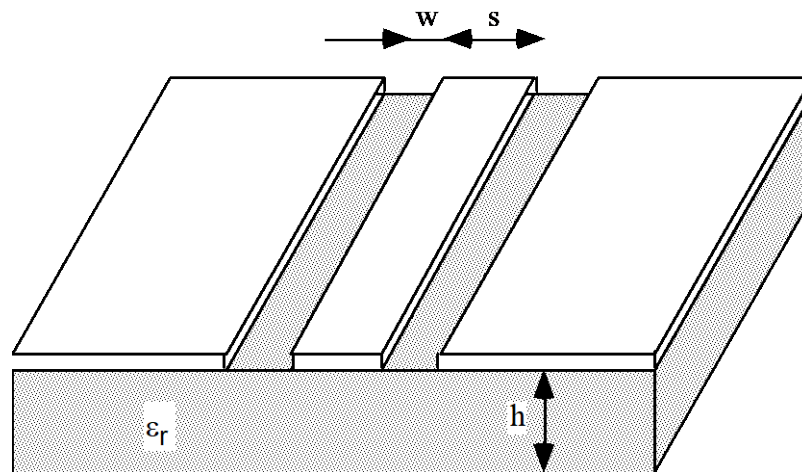


Figure 4.5 : Ligne coplanaire

4.5.2 Permittivité effective

La permittivité effective d'une telle structure est approximativement donnée par :

$$\epsilon_e = \frac{\epsilon_r + 1}{2} \left(\tanh \left(1.785 \ln \frac{h}{w} + 1.75 \right) + \frac{kw}{h} \left(0.04 - 0.7k + 0.01 \left(1 - \frac{\epsilon_r}{10} \right) (0.25 + k) \right) \right)$$
$$k = \frac{s}{s + 2w}$$

De nouveau, la permittivité effective permet d'obtenir la constante de propagation, la vitesse de phase et la longueur d'onde guidée :

$$\beta = \frac{\omega\sqrt{\epsilon_e}}{c} \quad v_\phi = \frac{c}{\sqrt{\epsilon_e}} \quad \lambda_g = \frac{\lambda_o}{\sqrt{\epsilon_e}}$$

4.5.3 Impédance caractéristique (rappel)

L'impédance caractéristique est donnée par :

$$Z_c = \frac{30\pi}{\sqrt{\epsilon_e}} \begin{cases} \frac{\pi}{\log\left(2\frac{1+\sqrt{k}}{1-\sqrt{k}}\right)} & 0 \leq k \leq 0.707 \\ \frac{\log\left(2\frac{1+\sqrt{k}}{1-\sqrt{k}}\right)}{\pi} & 0.707 < k \leq 1 \end{cases}$$

4.6 Tableau comparatif

Caractéristique	câble coaxial	Guides d'ondes	Ruban équilibré	Microruban
Mode dominant	TEM	TE ₁₀	TEM	Quasi TEM
Autres modes	TM, TE	TM, TE	TM, TE	Hybrides
Dispersion	Aucune	Moyenne	Aucune	Faible
Bande passante	Elevée	Faible	Elevée	Elevée
Pertes	Moyennes	Faibles	Elevées	Elevées
Puissance max.	Moyenne	Elevée	Faible	Faible
Dimensions	Grandes	Grandes	Moyennes	Petites
Facilité de fabr.	Moyenne	Moyenne	Facile	Facile
Intégration de composants	Difficile	Difficile	Moyenne	Facile

5. Caractérisation des circuits hyperfréquences

Références:

F. E. Gardiol, "Hyperfréquences", volume XIII du Traité d'Electricité, Presses Polytechniques Romandes, Chap 6.
R.E. Collin, "Foundations for Microwave Engineering", Mc Graw Hill, 1992, Chap. 4.

5.1 Introduction

5.2 Courant, tension et impédance

Dans la bande des hyperfréquences, les courants, tensions et impédances sont des grandeurs difficiles à définir, et à l'exception du cas des lignes supportant exclusivement un mode TEM, elles ne sont pas définies de manière univoque. Néanmoins, le modèle de Kirchhoff permet une description très aisée d'un système, à laquelle on ne souhaite pas renoncer. On définit donc des courants et des tensions équivalents sur les guides et lignes de transmission, on se rappelant que hormis le cas TEM ces grandeurs n'ont qu'une signification formelle.

Chaque mode propagé par un guide sera représenté par une onde de courant et de tension équivalents.

5.2.1 Modes TEM

La mesure des courants et des tensions est difficile, voire impossible, dans la bande des hyperfréquences, à moins de pouvoir définir clairement une paire de portes. Ceci n'est le cas que pour des lignes supportant un mode TEM ou quasi TEM.

La figure 5.1 illustre les champs électriques et magnétiques d'une ligne de transmission TEM arbitraire.

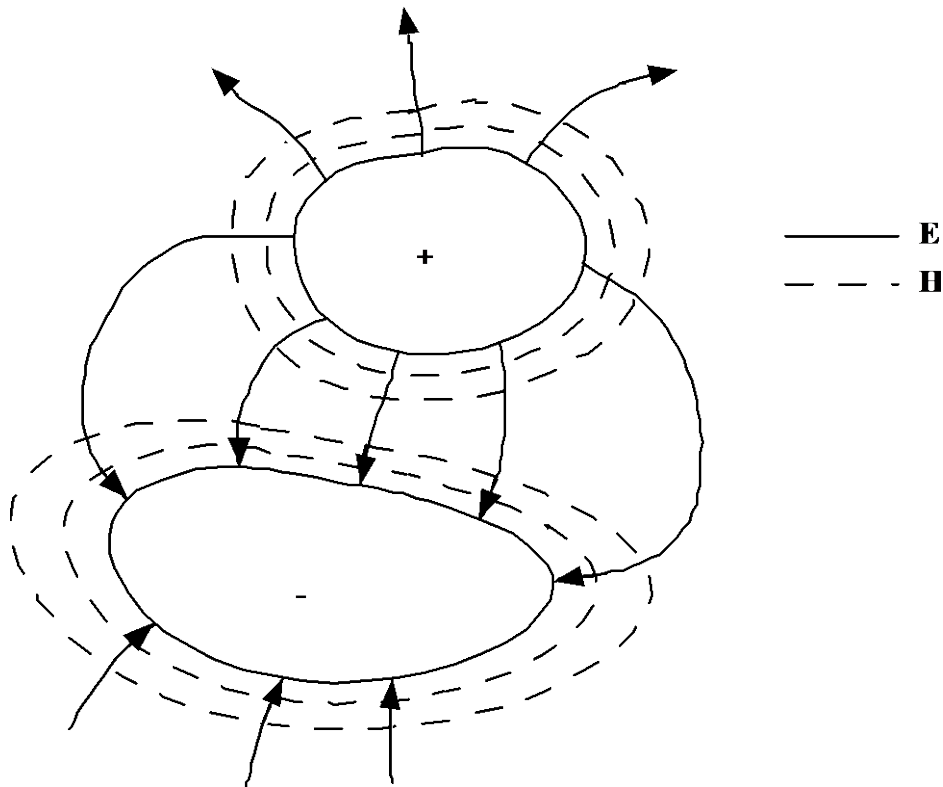


Figure 5.1 : Ligne TEM arbitraire

La tension du conducteur + par rapport au conducteur - est donnée par :

$$V = \int_{+}^{-} \mathbf{E} \cdot d\mathbf{l}$$

Dans le cas d'un mode TEM, le champ a un comportement statique, et la tension définie ci-dessus ne dépend pas du chemin d'intégration, pourvu que ce chemin aille du conducteur + vers le conducteur -. Dans ce cas, la définition est univoque, et il n'y a pas d'ambiguïté sur la tension. Le courant total du conducteur + peut être obtenu par la loi d'ampère, qui pour un mode TEM s'écrit :

$$I = \oint_{C+} \mathbf{H} \cdot d\mathbf{l}$$

où $C+$ est un chemin fermé entourant le conducteur + mais pas le conducteur -. L'impédance caractéristique s'écrit alors :

$$Z_c = \frac{V}{I} = \sqrt{\frac{L}{C}}$$

où L est l'inductance par unité de longueur de la ligne TEM et C sa capacité par unité de longueur.

5.2.2 Modes non-TEM

La situation est beaucoup moins claire pour les modes non-TEM, comme le montre cet exemple très simple.

Les champs transverses du mode TE₁₀ dans un guide d'ondes rectangulaires sont donnés par

$$E_y(x, y, z) = E_0 \frac{j\omega\mu a}{\pi} \sin \frac{\pi x}{a} e^{-j\beta z} = E_0 e_y(x, y) e^{-j\beta z}$$
$$H_x(x, y, z) = E_0 \frac{j\beta a}{\pi} \sin \frac{\pi x}{a} e^{-j\beta z} = E_0 h_y(x, y) e^{-j\beta z}$$

La tension est donc donnée par

$$V = E_0 \frac{-j\omega\mu a}{\pi} \sin \frac{\pi x}{a} e^{-j\beta z} \int_y dy$$

Cette tension dépend donc de la position x où l'on se place dans le guide, ainsi que du contour d'intégration y. Le résultat obtenu est donc différent si l'on intègre y de 0 à b en x=a/2 ou en x=0. Quel est donc la tension ?

La réponse est que dans ce cas, il n'y a pas de tension "correcte", mesurable, unique et pertinente pour toutes les applications.

On peut par conséquent définir une tension, un courant et une impédance équivalente de beaucoup de manières différentes pour un mode non-TEM. Pour obtenir des résultats utiles, il est néanmoins bon de suivre les règles suivantes :

- La tension et le courant ne sont définis que pour un mode. On choisit une tension proportionnelle au champ électrique transverse et un courant proportionnel au champ magnétique transverse.
- Pour permettre l'utilisation du modèle de Kirchhoff, le produit du courant et de la tension équivalente doit donner le flux de puissance du mode considéré.
- Le quotient de la tension par le courant correspondant à une onde doit être égal à l'impédance caractéristique de la ligne. Cette dernière peut théoriquement être choisie arbitrairement, mais on a tout avantage à la poser égale à l'impédance d'onde du mode considéré.

Dans un guide arbitraire, les champs transverses s'expriment en fonction d'une onde progressive et d'une onde rétrograde. Le potentiel et le courant définis selon les critères ci-dessus s'expriment donc de la même manière :

$$\begin{aligned}
\mathbf{E}_t(x, y, z) &= \mathbf{e}_t(x, y) \left(E_o^+ e^{-j\beta z} + E_o^- e^{j\beta z} \right) \\
&= \frac{\mathbf{e}_t(x, y)}{C_1} \left(V^+ e^{-j\beta z} + V^- e^{j\beta z} \right) \\
\mathbf{H}_t(x, y, z) &= \mathbf{h}_t(x, y) \left(E_o^+ e^{-j\beta z} - E_o^- e^{j\beta z} \right) \\
&= \frac{\mathbf{h}_t(x, y)}{C_2} \left(I^+ e^{-j\beta z} - I^- e^{j\beta z} \right)
\end{aligned}$$

On écrit donc, pour la tension et le courant :

$$\begin{aligned}
V(z) &= V^+ e^{-j\beta z} + V^- e^{j\beta z} \\
I(z) &= I^+ e^{-j\beta z} - I^- e^{j\beta z}
\end{aligned}$$

L'impédance caractéristique de cette onde est définie (par analogie au cas TEM) comme

$$Z_c = \frac{V^+}{I^+} = \frac{V^-}{I^-} = \frac{C_1 E_o^+}{C_2 E_o^+} = \frac{C_1}{C_2}$$

Si de plus on souhaite que l'impédance caractéristique soit égale à l'impédance d'onde du mode, on obtient :

$$\frac{C_1}{C_2} = Z_{\text{mod}}$$

5.2.3 Concept d'impédance

Il convient de distinguer :

- L'impédance intrinsèque du milieu. Elle ne dépend que du matériau constituant le milieu et de ses paramètres

$$Z_o = \sqrt{\frac{\mu}{\varepsilon}}$$

- L'impédance d'onde d'un mode. Il dépend du type de mode (TEM, TE et TM) ainsi que du type de guide, des matériaux utilisés, de la géométrie et de la fréquence

$$Z_{\text{mod}} = \frac{|\mathbf{E}_t|}{|\mathbf{H}_t|}$$

- L'impédance caractéristique, définie comme le rapport de la tension sur le courant. Elle n'est par définition univoque que pour un mode TEM.

$$Z_c = \frac{V^+}{I^+} = \frac{V^-}{I^-} = \sqrt{\frac{L}{C}}$$

5.3 La matrice d'impédance

Les concepts de tension, courant et impédance définis pour les lignes de transmissions aux sections précédentes peuvent être utilisés pour caractériser des composants hyperfréquences. Ces derniers sont alors caractérisés par une *matrice d'impédance*, obtenus à partir des courants et des tensions définies sur les lignes de transmission reliés aux accès du composants.

Références:

R.E. Collin, "Foundations for Microwave Engineering", Mc Graw Hill, 1992.

5.3.1 Impédance d'un monoporte

Le composant hyperfréquence le plus simple ne comporte qu'un seul accès. Sa matrice d'impédance se réduit alors à un scalaire, défini par le quotient de la tension par le courant "mesurés" à l'accès du monoporte, le *plan de référence*.

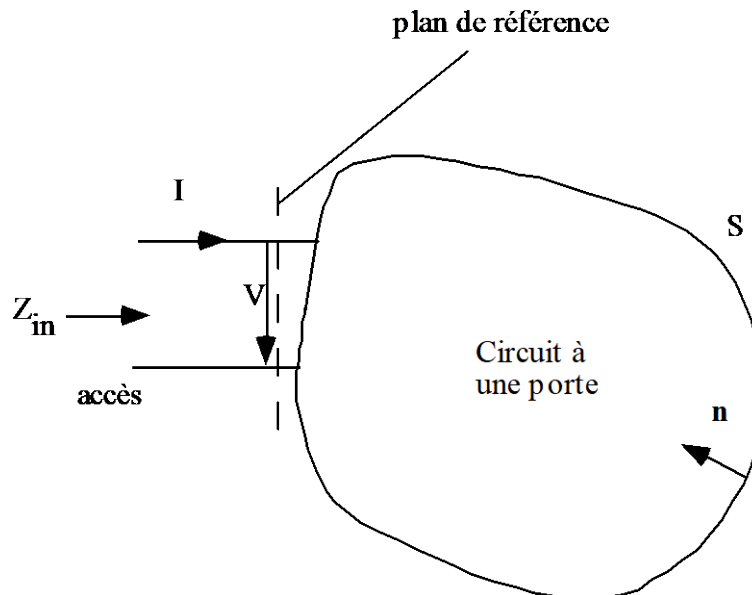


Figure 5.2 : monoporte

$$Z_{in} = \frac{V}{I}$$

5.3.2 Caractéristiques de l'impédance d'un monoporte

La puissance complexe fournie au monoporte est donnée par le vecteur de Poynting :

$$P = \frac{1}{2} \oint_S \mathbf{E} \times \mathbf{H}^* \cdot d\mathbf{s} = P_r + 2j\omega(W_m - W_e)$$

Les champs E et H sont liés, sur la ligne de transmission, au courant et la tension :

$$\mathbf{E}_t(x, y, z) = V(z) \frac{\mathbf{e}_t(x, y)}{C_1} e^{-j\beta z}$$

$$\mathbf{H}_t(x, y, z) = I(z) \frac{\mathbf{h}_t(x, y)}{C_2} e^{-j\beta z}$$

Avec, par la définition choisie du courant et de la tension

$$\frac{1}{C_1 C_2} \int_s \mathbf{e}_t \times \mathbf{h}_t \cdot \mathbf{ds} = 1$$

Donc

$$P = \frac{1}{2 C_1 C_2} \int_s V I^* \mathbf{e}_t \times \mathbf{h}_t \cdot \mathbf{ds} = \frac{1}{2} V I^*$$

Or l'impédance d'entrée peut s'écrire en fonction de la puissance moyenne:

$$Z_{in} = R + jX = \frac{V}{I} = \frac{V I^*}{|I|^2} = \frac{2P}{|I|^2}$$

$$= \frac{2(P_r + 2j\omega(W_m - W_e))}{|I|^2}$$

Où P_r est la puissance réelle moyenne, W_m l'énergie magnétique emmagasinée et W_e l'énergie électrique emmagasinée. On peut déduire de la relation ci-dessus que

- R est proportionnel à la puissance réelle dissipée par le système (pertes)
- X est proportionnel à l'énergie réactive moyenne emmagasinée dans le système

5.3.3 Matrices d'impédance et d'admittance

Considérons le composant quelconque représenté à la figure 5.3. Il est caractérisé par un certain nombre de portes d'accès, définis par des plans sur les lignes de transmission reliant le composant au monde extérieur. Ces plans, notés t_n sur la figure sont les *plans de référence* entre lesquels un composant est défini.

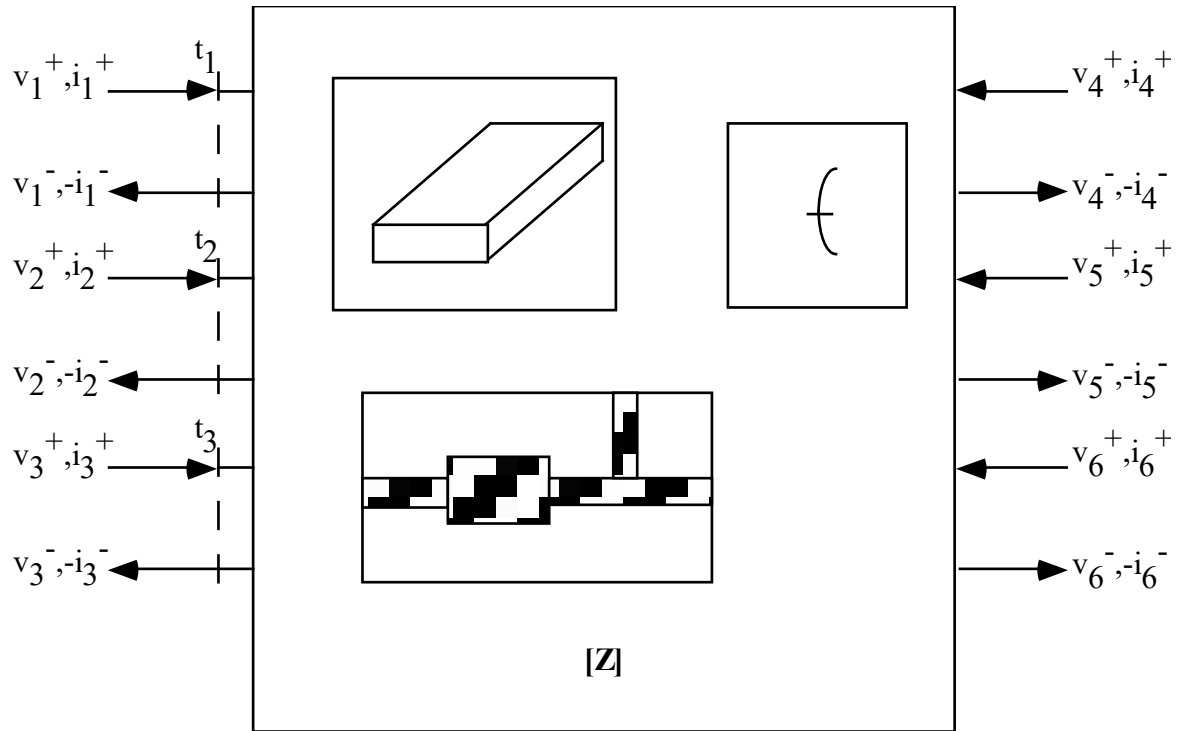


Figure 5.3 : composant hyperfréquence multiporte avec ses portes d'accès

Un axe de coordonnée z_i est lié à chaque ligne de transmission i . Par définition, l'origine de cet axe est placé au plan de référence. Nous avons donc aux portes $t_1, t_2, \dots, t_n$

$$V_n = V_n^+ + V_n^-$$

$$I_n = I_n^+ - I_n^-$$

Les matrices d'impédance et d'admittance caractérisant le composant sont définis par les relations

$$[V] = [Z][I]$$

$$[I] = [Y][V]$$

avec

$$Z_{ij} = \left. \frac{V_i}{I_j} \right|_{I_k=0 \text{ pour } k \neq j}$$

$$Y_{ij} = \left. \frac{I_i}{V_j} \right|_{V_k=0 \text{ pour } k \neq j}$$

Par conséquent,

- La matrice d'impédance est obtenue en circuit-ouvert
- La matrice d'admittance est obtenue en court circuit

La matrice d'impédance est l'inverse de la matrice d'admittance

$$[Y] = [Z]^{-1}$$

5.3.4 Propriétés des matrices d'impédance et d'admittance

5.3.4.1 Réciprocité

Considérons le cas où les hypothèses de base du théorème de réciprocité de Lorentz sont respectées, soit un cas où le composant ne comporte pas d'élément anisotrope ou non-linéaire. Dans le cas d'un composant hyperfréquence, cela signifie que le composant ne comporte ni élément actif, ni ferrite (élément magnétique anisotrope), ni plasma. On étudie alors le cas du composant illustré à la figure 5.4, où tous les accès sauf deux sont court-circuité. Soient maintenant les champs E_a , H_a , E_b , et H_b dus à deux sources indépendantes situées quelque part dans le circuit. Le théorème de réciprocité de Lorentz nous dit alors

$$\oint_S \mathbf{E}_a \times \mathbf{H}_b \cdot d\mathbf{s} = \oint_S \mathbf{E}_b \times \mathbf{H}_a \cdot d\mathbf{s}$$

où S est une surface d'intégration fermée comprenant le composant.

On choisit la surface fermée S comme la limite extérieure du circuit passant par les plans de référence des lignes d'alimentation, telle que $E_{\tan}=0$, sauf pour les plans de références 1 et 2

(Si le circuit et les lignes de transmissions sont en métal, ceci est toujours vrai. Si ce n'est pas le cas, on peut toujours choisir une surface suffisamment éloignée pour que $E_{\tan}$ soit négligeable).

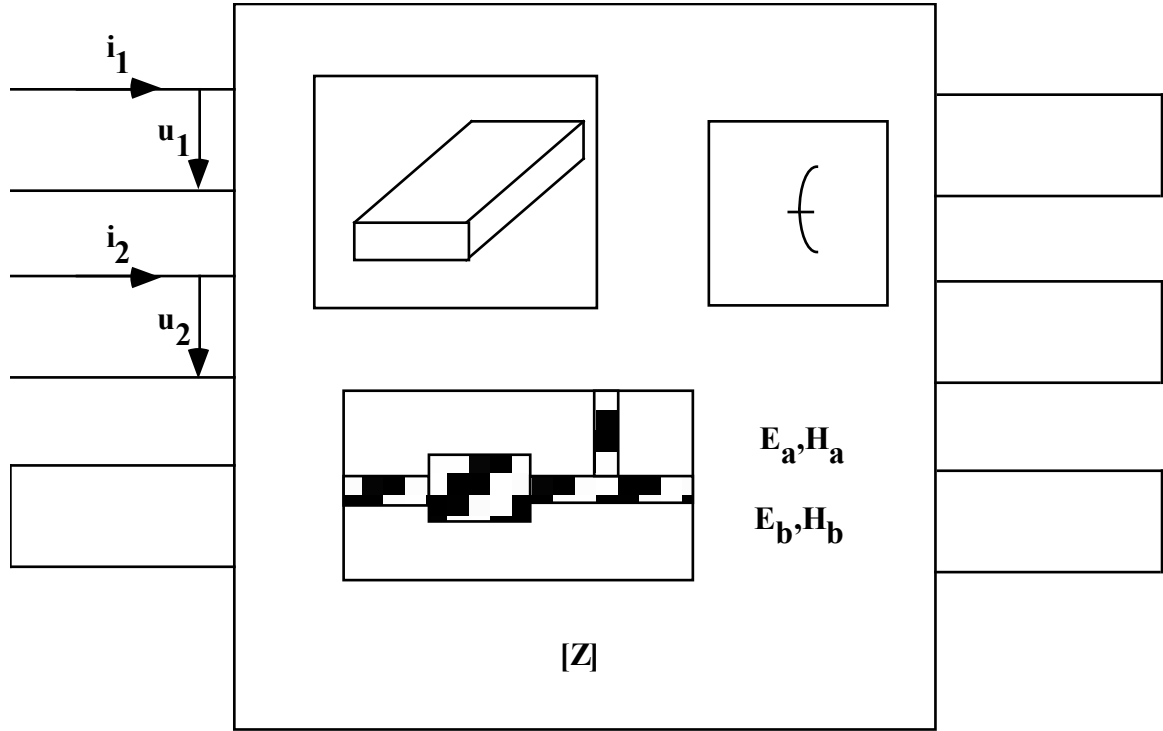


Figure 5.4 : illustration du principe de réciprocité

Les seules contributions aux intégrales proviennent alors des plan de référence 1 et 2, les seuls à ne pas être court-circuités.

On écrit dans ces plans :

$$\begin{aligned} \mathbf{E}_{1a} &= V_{1a} \frac{\mathbf{e}_1}{C_1} & \mathbf{H}_{1a} &= I_{1a} \frac{\mathbf{h}_1}{K_1} \\ \mathbf{E}_{1b} &= V_{1b} \frac{\mathbf{e}_1}{C_1} & \mathbf{H}_{1b} &= I_{1b} \frac{\mathbf{h}_1}{K_1} \\ \mathbf{E}_{2a} &= V_{2a} \frac{\mathbf{e}_2}{C_2} & \mathbf{H}_{2a} &= I_{2a} \frac{\mathbf{h}_2}{K_2} \\ \mathbf{E}_{2b} &= V_{2b} \frac{\mathbf{e}_2}{C_2} & \mathbf{H}_{2b} &= I_{2b} \frac{\mathbf{h}_2}{K_2} \end{aligned}$$

Et le théorème de réciprocité devient :

$$\begin{aligned} (V_{1a}I_{1b} - V_{1b}I_{1a}) \int_{s_1} \frac{1}{C_1 K_1} \mathbf{e}_1 \times \mathbf{h}_1 \cdot d\mathbf{s} + \\ (V_{2a}I_{2b} - V_{2b}I_{2a}) \int_{s_2} \frac{1}{C_2 K_2} \mathbf{e}_2 \times \mathbf{h}_2 \cdot d\mathbf{s} = 0 \end{aligned}$$

Or, par définition

$$\int_{s_1} \frac{1}{C_1 K_1} \mathbf{e}_1 \times \mathbf{h}_1 \cdot d\mathbf{s} = \int_{s_1} \frac{1}{C_2 K_2} \mathbf{e}_2 \times \mathbf{h}_2 \cdot d\mathbf{s} = 1$$

Donc

$$V_{1a}I_{1b} - V_{1b}I_{1a} + V_{2a}I_{2b} - V_{2b}I_{2a} = 0$$

Or

$$I_1 = Y_{11}V_1 + Y_{12}V_2$$

$$I_2 = Y_{21}V_1 + Y_{22}V_2$$

Donc

$$(V_{1a}V_{2b} - V_{1b}V_{2a})(Y_{12} - Y_{21}) = 0$$

Cette relation doit être valable de manière indépendante des sources, donc des tensions. On en déduit :

$$Y_{12} = Y_{21}$$

Cette relation peut être aisément généralisée à toutes les portes du composant. On écrit donc de manière générale, pour un circuit ne comportant ni éléments actifs, ni ferrites, ni plasmas

$$Y_{ij} = Y_{ji}$$

$$Z_{ij} = Z_{ji}$$

On en déduit donc que la matrice d'admittance (par conséquent aussi la matrice d'impédance) est symétrique lorsque le composant est réciproque.

5.3.4.2 Circuits sans pertes

Considérons un composant sans pertes à N accès. On peut écrire que pour ce composant, la puissance réelle moyenne fournie au circuit est nulle

$$\text{Re}\{P_{av}\} = 0$$

Par définition des courants et des tensions aux accès, la puissance moyenne délivrée au composant est donnée par :

$$P_{av} = \frac{1}{2} [V]^t [I]^*$$

Ce qui peut être écrit en terme de la matrice d'impédance du composant comme

$$\begin{aligned} P_{av} &= \frac{1}{2} ([Z][I])^t [I]^* \\ &= \frac{1}{2} [I]^t [Z][I]^* \\ &= \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N I_m Z_{mn} I_n^* \end{aligned}$$

Les courants I_n sont indépendants entre eux, donc la partie réelle de chaque terme de type $m=n$ doit être nulle :

$$\operatorname{Re} \{ I_n Z_{nn} I_n^* \} = |I_n|^2 \operatorname{Re} \{ Z_{nn} \} = 0$$

On en déduit que la partie réelle des termes diagonaux de la matrice d'impédance doivent être nuls.

$$\operatorname{Re} \{ Z_{nn} \} = 0$$

On suppose maintenant que tous les courants entrant dans le circuit sont nuls sauf I_n et I_m . On écrit alors

$$\operatorname{Re} \{ (I_n I_m^* + I_m I_n^*) Z_{mn} \} = (I_n I_m^* + I_m I_n^*) \operatorname{Re} \{ Z_{mn} \} = 0$$

On en déduit que

$$\operatorname{Re} \{ Z_{mn} \} = 0$$

Nous avons donc démontré qu'un circuit sans pertes aura une matrice d'impédance (d'admittance) purement imaginaire

5.3.5 Exemples de matrices d'impédances

1) ligne de transmission

Soit le tronçon de ligne de transmission illustré à la figure 5.5

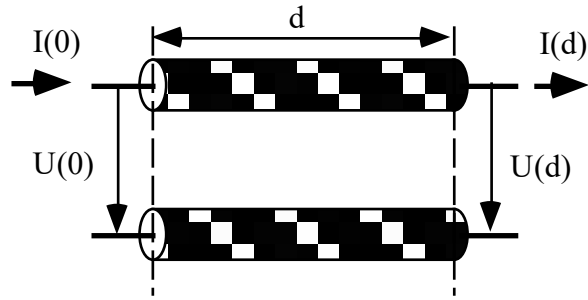


Figure 5.5 : tronçon de ligne de transmission de longueur d

Il peut être représenté par le biporte équivalent suivant

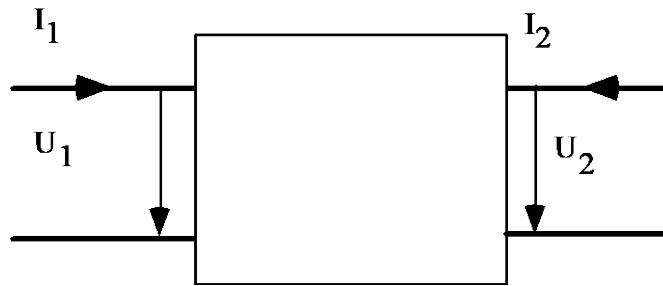


Figure 5.6 : biporte équivalent

où, par définition,

$$\begin{aligned} U_1 &= U(0) , \quad I_1 = I(0) , \\ U_2 &= U(d) , \quad I_2 = -I(d) \end{aligned}$$

En sachant que

$$\begin{aligned} U(z) &= U_+ e^{-\gamma z} + U_- e^{+\gamma z} \\ I(z) &= I_+ e^{-\gamma z} - I_- e^{+\gamma z} \end{aligned}$$

On peut écrire

$$\begin{aligned} U(0) &= U_+ + U_- & U(d) &= U_+ e^{-\gamma d} + U_- e^{+\gamma d} \\ I(0) &= I_+ - I_- & I(d) &= I_+ e^{-\gamma d} - I_- e^{+\gamma d} \end{aligned}$$

Les matrices d'impédance et d'admittance du biporte s'écrivent alors

$$\begin{aligned}
\begin{bmatrix} U_1 \\ U_2 \end{bmatrix} &= \begin{bmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{bmatrix} \begin{bmatrix} I_1 \\ I_2 \end{bmatrix} \\
&= Z_c \begin{bmatrix} \coth(\gamma d) & \frac{1}{\sinh(\gamma d)} \\ \frac{1}{\sinh(\gamma d)} & \coth(\gamma d) \end{bmatrix} \begin{bmatrix} I_1 \\ I_2 \end{bmatrix} \\
\begin{bmatrix} I_1 \\ I_2 \end{bmatrix} &= \begin{bmatrix} Y_{11} & Y_{12} \\ Y_{21} & Y_{22} \end{bmatrix} \begin{bmatrix} U_1 \\ U_2 \end{bmatrix} \\
&= Y_c \begin{bmatrix} \coth(\gamma d) & \frac{-1}{\sinh(\gamma d)} \\ \frac{-1}{\sinh(\gamma d)} & \coth(\gamma d) \end{bmatrix} \begin{bmatrix} U_1 \\ U_2 \end{bmatrix}
\end{aligned}$$

2) Circuit équivalent en T d'un biporte réciproque

Un biporte réciproque est caractérisé par la matrice d'impédance suivante :

$$\begin{bmatrix} Z_{11} & Z_{12} \\ Z_{12} & Z_{22} \end{bmatrix}$$

Un tel biporte peut être représenté par un circuit équivalent en T

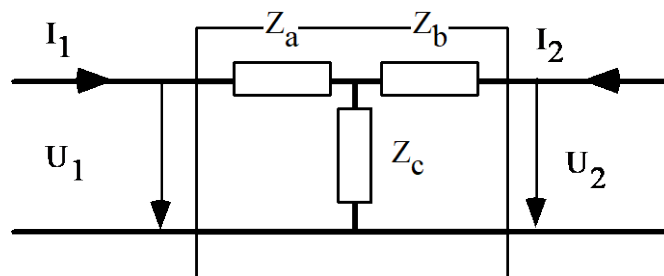


Figure 5.7 : équivalent en T d'un biporte réciproque

où

$$\begin{aligned}
Z_a &= Z_{11} - Z_{12} \\
Z_b &= Z_{22} - Z_{12} \\
Z_c &= Z_{12}
\end{aligned}$$

exemple : circuit équivalent en T d'un tronçon de ligne :

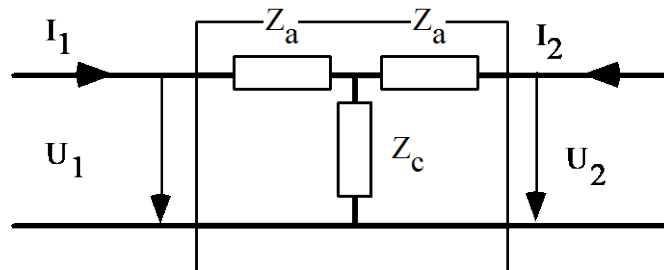


Figure 5.8 : équivalent en T d'un tronçon de ligne

avec

$$\begin{aligned}
 Z_a &= Z_{caract} \left(\coth(\gamma d) - \frac{1}{\sinh(\gamma d)} \right) \\
 &= Z_{caract} \left(\frac{ch(\gamma d) - 1}{\sinh(\gamma d)} \right) = Z_{caract} \tanh\left(\frac{\gamma d}{2}\right) \\
 Z_c &= \frac{Z_{caract}}{\sinh(\gamma d)}
 \end{aligned}$$

3) Circuit équivalent en Π d'un biporte réciproque

Un biporte réciproque est caractérisé par la matrice d'admittance suivante :

$$\begin{bmatrix} Y_{11} & Y_{12} \\ Y_{12} & Y_{22} \end{bmatrix}$$

Un tel biporte peut être représenté par un circuit équivalent en Π

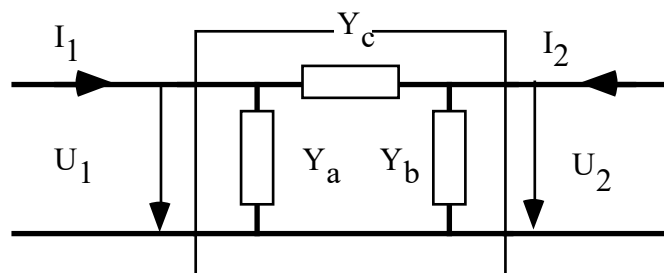


Figure 5.9 : équivalent en Π d'un biporte réciproque

avec

$$\begin{aligned} Y_a &= Y_{11} + Y_{12} \\ Y_b &= Y_{22} + Y_{12} \\ Y_c &= -Y_{12} \end{aligned}$$

Application : circuit équivalent en Π d'un tronçon de ligne :

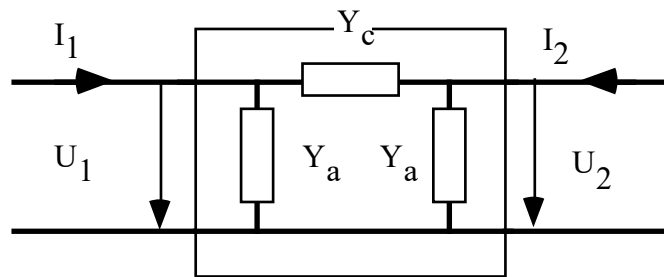


Figure 5.10 : équivalent en Π d'un tronçon de ligne

avec

$$\begin{aligned} Y_a &= Y_{caract} \left(\coth(\gamma d) - \frac{1}{\sinh(\gamma d)} \right) \\ &= Y_{caract} \left(\frac{\cosh(\gamma d) - 1}{\sinh(\gamma d)} \right) = Y_{caract} \tanh\left(\frac{\gamma d}{2}\right) \\ Y_c &= \frac{Y_{caract}}{\sinh(\gamma d)} \end{aligned}$$

5.4 La matrice de répartition

Références:

R.E. Collin, "Foundations for Microwave Engineering", Mc Graw Hill, 1992.

F. E. Gardiol, "Hyperfréquences", volume XIII du Traité d'Electricité, Presses Polytechniques Romandes, chap. 6

Nous avons vu dans les paragraphes précédents que les courants et les tensions étaient des grandeurs équivoques en hyperfréquences. Une des conséquences directe de cette non univocité est le fait que la tension et le courant ne sont souvent pas mesurables en hyperfréquences. Ils permettent donc la caractérisation théorique d'un composant par la matrice d'impédance, mais ces impédances ne peuvent pas être mesurées comme en basse fréquence. Un composant hyperfréquences ne peut souvent pas être caractérisé par la mesure des courants et des tensions à ses bornes. Pour pallier à cet inconvénient majeur, on introduit les notions de matrice de répartition et d'amplitudes normalisées pour caractériser les circuits hyperfréquences.

5.4.1 Amplitudes normalisées

On définit les amplitudes normalisées **a** et **b** comme

$$a_i = \frac{v_i + Z_{ci} i_i}{2\sqrt{Z_{ci}}} \quad , \quad b_i = \frac{v_i - Z_{ci} i_i}{2\sqrt{Z_{ci}}}$$

Note : ces amplitudes normalisées ont la dimension de racines carrées de puissances, et la puissance sur une ligne de transmission est facilement mesurable en hyperfréquences.

La relation inverse est donnée par

$$v_i = \sqrt{Z_{ci}} (a_i + b_i) \quad , \quad i_i = \frac{(a_i - b_i)}{\sqrt{Z_{ci}}}$$

Ces amplitudes normalisées sont définies sur les lignes de transmission reliant les accès du composant. Or, sur ces lignes de transmission :

$$\begin{aligned} v_i &= v_i^+ e^{-j\beta z} + v_i^- e^{+j\beta z} \\ i_i &= i_i^+ e^{-j\beta z} + i_i^- e^{+j\beta z} \end{aligned}$$

on peut en déduire

$$a_i = \frac{v_i^+}{\sqrt{Z_{ci}}} e^{-j\beta z}$$

$$b_i = \frac{v_i^-}{\sqrt{Z_{ci}}} e^{+j\beta z}$$

donc :

- a_i : onde purement progressive, donnant le signal (racine carrée de la puissance) entrant dans l'accès i
- b_i : onde purement rétrograde donnant le signal (racine carrée de la puissance) sortant de l'accès i

5.4.2 Plans de référence

Un composant hyperfréquence est défini entre ses accès, qui sont des plans transverses aux lignes de transmissions le reliant au monde extérieur. Ce plan sert d'origine à l'axe des coordonnées longitudinales z_i liée à la ligne de transmission i (figure 5.11)

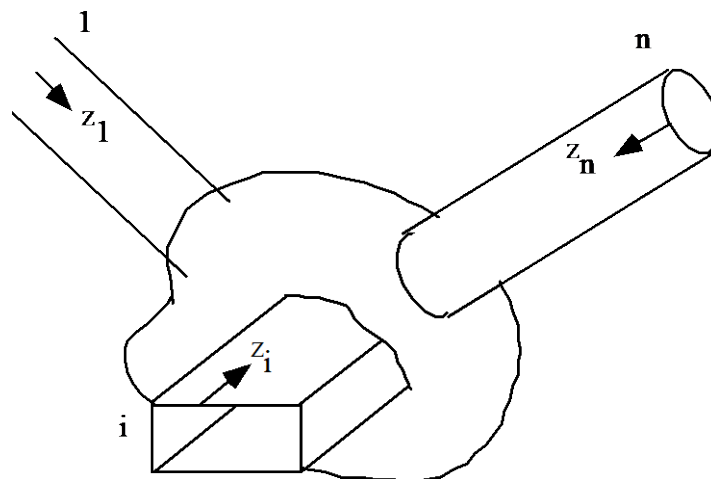


Figure 5.11 : composant hyperfréquence avec ses plans de référence

Par définition, les plans de références répondent aux critères suivants

- Les plans de références sont suffisamment éloignés du circuit pour que les modes évanescents soit atténués
- Les lignes de transmission ne supportent que le mode dominant
- Les lignes de transmission sont sans pertes

On peut remarquer que la puissance active à l'accès i d'un composant est donnée par

$$P_i = \operatorname{Re} [v_i i_i^*] = \operatorname{Re} [(a_i + b_i)(a_i^* - b_i^*)] = |a_i|^2 - |b_i|^2$$

$|a_i|^2$ est donc la puissance active entrant dans le composant à l'accès i et $|b_i|^2$ est la puissance active sortant du composant à l'accès i

5.4.3 Matrice de répartition d'un composant

Un composant hyperfréquence peut être caractérisé en fonction des amplitudes généralisées sur les lignes des transmissions à ses plans de références (figure 5.12). Il est alors caractérisé par sa matrice de répartition

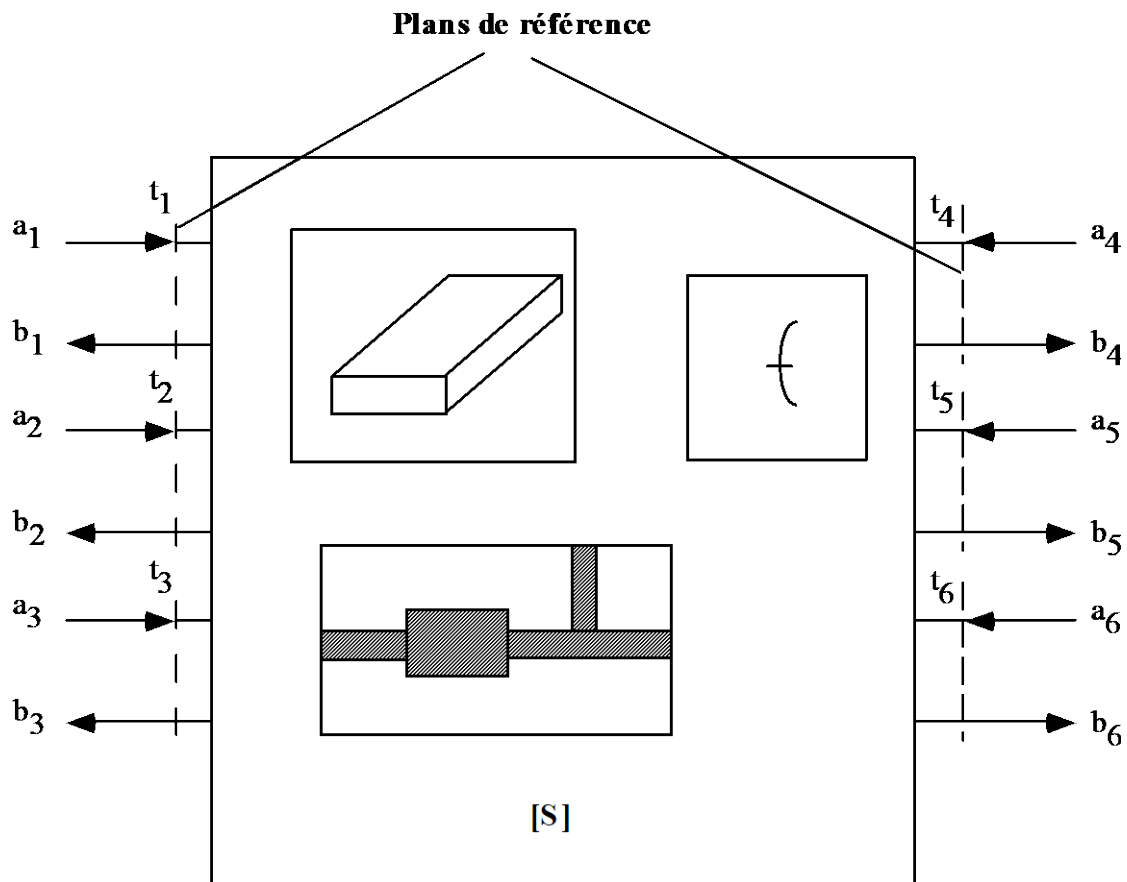


Figure 5.12 : composant hyperfréquences avec ses plans de référence

La matrice de répartition du composant est alors définie par

$$[b] = [S][a]$$

avec

$$s_{ij} = \left. \frac{b_i}{a_j} \right|_{a_k=0, k \neq j}$$

5.4.4 Propriétés de la matrice de répartition

- Si la matrice d'impédance caractérise un composant en circuit-ouvert ($Z_{ij} = v_i/i_j$, $i_k=0$ pour $k \neq j$) et la matrice d'admittance caractérise le même composant en court-circuit ($Y_{ij} = i_i/v_j$, $i_k=0$ pour $k \neq j$), alors la matrice de répartition caractérise ce composant terminé par des charges adaptées ($S_{ij} = b_i/a_j$, $a_k=0$ pour $k \neq j$).
- Le terme s_{ij} est la fonction de transfert du signal entre la porte j et la porte i .
- La matrice de répartition dépend du composant et de l'environnement par les lignes de transmissions.
- Changer l'impédance caractéristiques des lignes de transmission revient à changer la matrice de répartition

5.4.4.1 Réciprocité

dans le cas d'un circuit ne comportant pas d'éléments actifs, anisotropes ou non-linéaires (hypothèses de réciprocité), nous avons vu que :

$$z_{ij} = z_{ji}$$

On démontre facilement par transformation matricielle que pour un circuit réciproque :

$$s_{ij} = s_{ji}$$

Un circuit réciproque a une matrice de répartition symétrique

5.4.4.2 Circuit sans pertes

Un circuit sans pertes est un circuit qui ne dissipe pas de puissance active. Pour un tel circuit, nous avons donc que la somme des puissances actives sortant de tous les accès doit être égale à la somme des puissances actives rentrée par tous les accès :

$$\sum |a_i|^2 = \sum |b_i|^2$$

Cette relation peut s'écrire matriciellement

$$[\tilde{a}][a] - [\tilde{b}]b = 0$$

où le tilde signifie la matrice transposée et complexe conjuguée :

$$[\tilde{a}] = [a^*]$$

De plus, par définition,

$$\begin{aligned} [b] &= [S][a] \\ [\tilde{b}] &= [\tilde{a}][\tilde{S}] \end{aligned}$$

Donc

$$\begin{aligned} [\tilde{a}][a] - [\tilde{a}][\tilde{S}][S][a] &= 0 \\ [\tilde{a}]\{1 - [\tilde{S}][S]\}a &= 0 \\ [\tilde{S}][S] &= [1] \end{aligned}$$

Ce qui peut aussi s'écrire

$$\sum_{i=1}^N s_{ij}^* s_{ik} = \delta_{jk} \quad \delta_{jk} = \begin{cases} 1 & \text{si } j = k \\ 0 & \text{si } j \neq k \end{cases}$$

5.4.4.3 Déplacement d'un plan de référence

L'origine des axes z_i , donc la position des plans de référence, sont définies arbitrairement, pour autant que l'on soit en propagation monomodale. Il peut donc être intéressant d'étudier l'effet sur la matrice de répartition d'une translation du plan de référence le long de la ligne de transmission (figure 5.13).

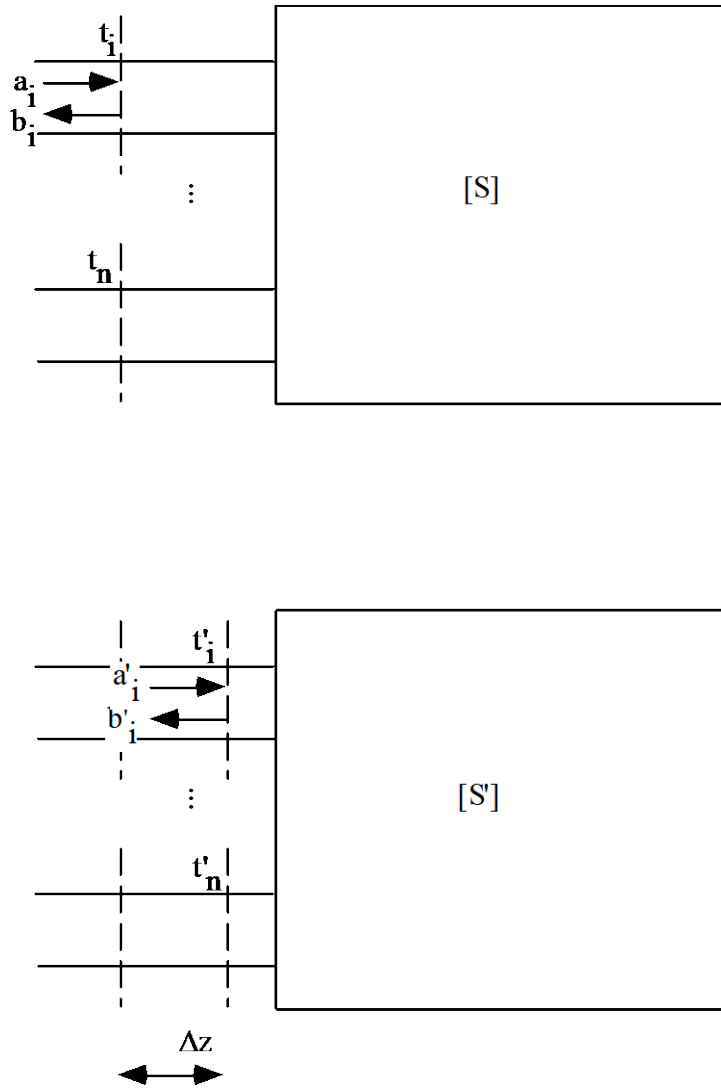


Figure 5.13 : effet du déplacement du plan de référence

Les amplitudes normalisées a'_i et b'_i liées au système de coordonnées déplacé peuvent s'exprimer en fonction des amplitudes normalisées a_i et b_i liées au système de coordonnées original par

$$a'_i = a_i e^{-j\varphi_i}$$

$$b'_i = b_i e^{j\varphi_i}$$

$$\varphi_i = -\beta_i \Delta z_i$$

Or, puisque nous avons

$$[b'] = [S'] [a'] \quad \text{et} \quad [b] = [S] [a]$$

nous pouvons écrire

$$s'_{ii} = s_{ii} e^{j2\varphi_i}$$

De manière plus générale, on a les relations

$$\begin{aligned}[a] &= \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} d \\[d] &= \begin{bmatrix} diag & e^{-j\varphi} \end{bmatrix} a \\[b] &= \begin{bmatrix} diag & e^{-j\varphi} \end{bmatrix} b' \\[b'] &= \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} b\end{aligned}$$

avec

$$\begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} = \begin{bmatrix} e^{j\varphi_1} & 0 & \dots & 0 \\ 0 & e^{j\varphi_2} & & : \\ : & & & \\ 0 & \dots & & e^{j\varphi_n} \end{bmatrix}$$

donc

$$\begin{aligned}[b'] &= \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} S \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} d \\[S'] &= \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} S \begin{bmatrix} diag & e^{j\varphi} \end{bmatrix} \\s'_{ij} &= s_{ij} e^{j(\varphi_i + \varphi_j)}\end{aligned}$$

5.4.4.4 Passage de la matrice d'impédance à la matrice de répartition

A l'aide des deux matrices diagonales

$$\begin{aligned}[G] &= \begin{bmatrix} diag & Z_{ci} \end{bmatrix} = \begin{bmatrix} Z_{c1} & 0 & \dots & 0 \\ 0 & Z_{c2} & & : \\ : & & & \\ 0 & \dots & & Z_{cn} \end{bmatrix} \\[F] &= \begin{bmatrix} diag & \frac{1}{2\sqrt{Z_{ci}}} \end{bmatrix} = \begin{bmatrix} \frac{1}{2Z_{c1}} & 0 & \dots & 0 \\ 0 & \frac{1}{2Z_{c2}} & & : \\ : & & & \\ 0 & \dots & & \frac{1}{2Z_{cn}} \end{bmatrix}\end{aligned}$$

et de la définition des amplitudes normalisées, on montre aisément que

$$\begin{aligned}[S] &= [F][[Z] - [G]][[F][[Z] + [G]]]^{-1} \\ &= [F][[Z] - [G]][[Z] + [G]]^{-1}[F]^{-1}\end{aligned}$$

et

$$[Z] = [F]^{-1}[[1] + [s]][[1] - [s]]^{-1}[F][G]$$

où $[1]$ est la matrice unité.

5.4.5 Graphes de fluence

Les termes de la matrices de répartition sont des fonctions de transfert, liant une entrée à une sortie. On peut les représenter graphiquement par les *graphes de fluence*

exemple 1 : biporte

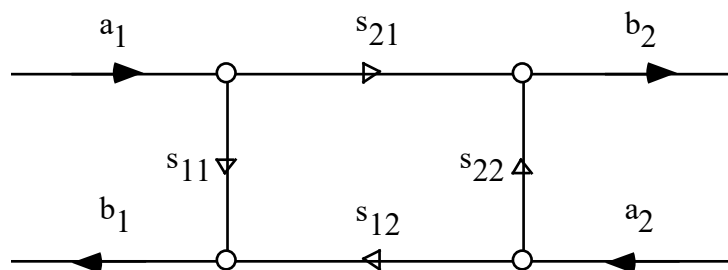


Figure 5.14 : graphe de fluence d'un biporte

exemple 2 : triporte

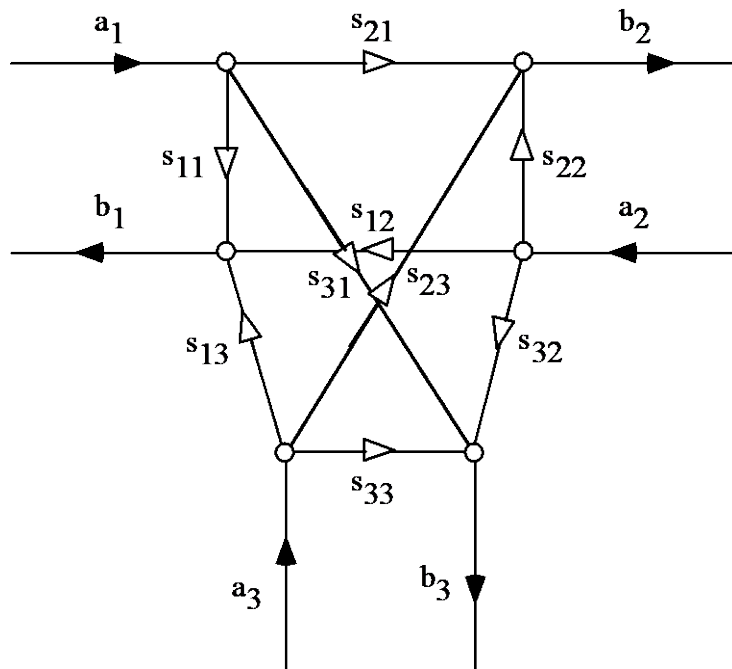


Figure 5.15 : graphe de fluence d'un triporte

exemple 3 : mise en cascade de biportes

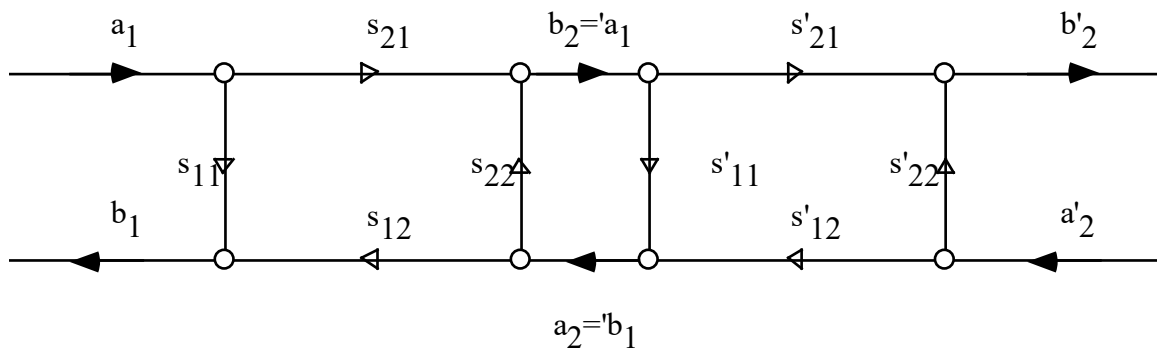


Figure 5.16 : mise en cascade de deux biportes

5.4.5.1 Règles de réduction des graphes de fluence

1) multiplication

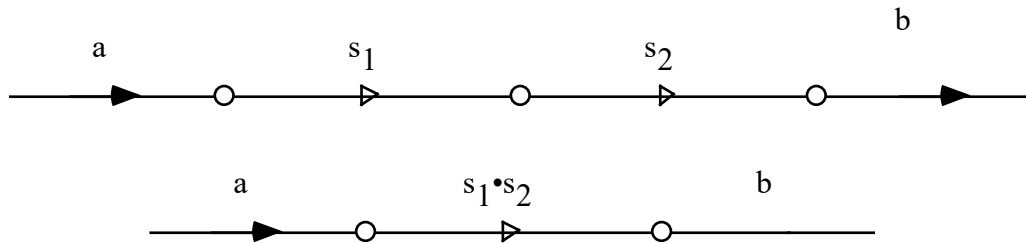


Figure 5.17 : mise en série de deux graphes de fluence

2) addition

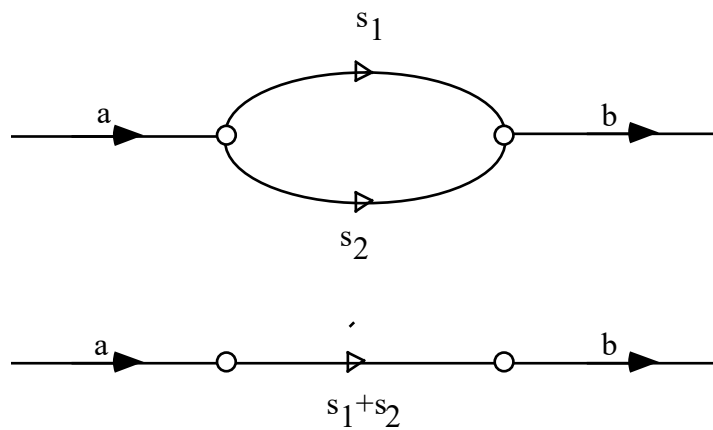


Figure 5.18: mise en parallèle de deux graphes de fluence

3) rétroaction

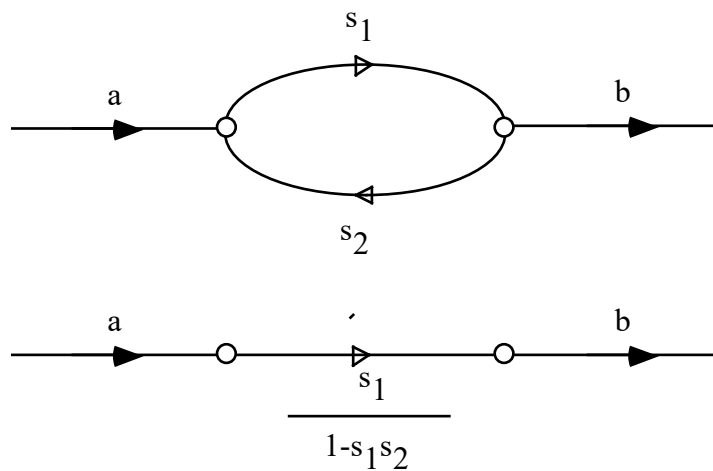


Figure 5.19 : rétroaction de deux graphes de fluence

5.4.5.2 Exemple

Trouver le coefficient de réflexion à l'entrée d'un biporte réciproque terminé par un court-circuit

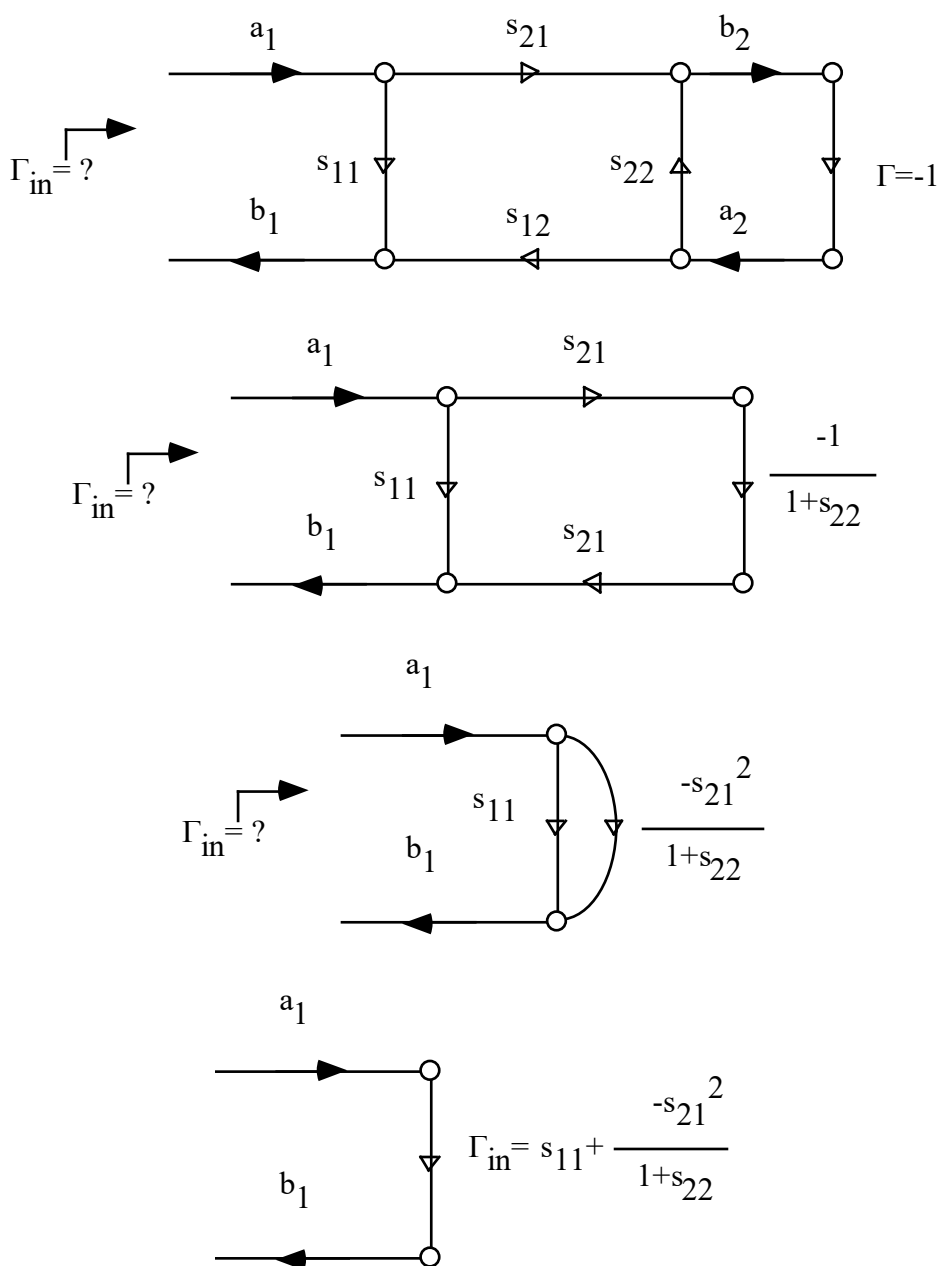
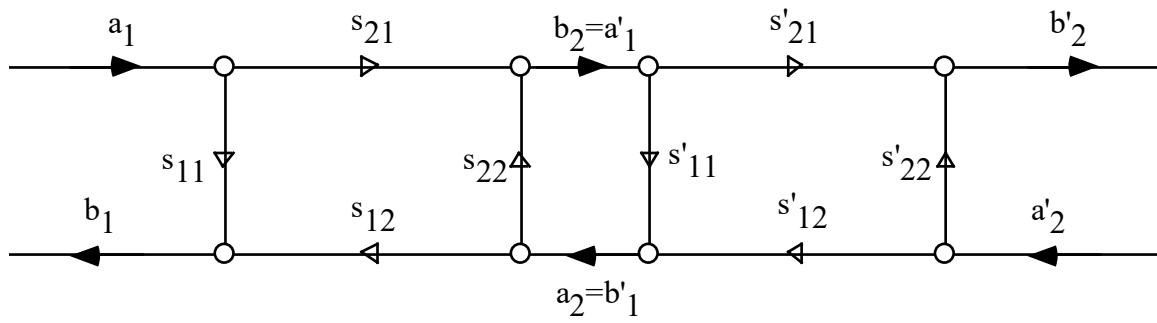
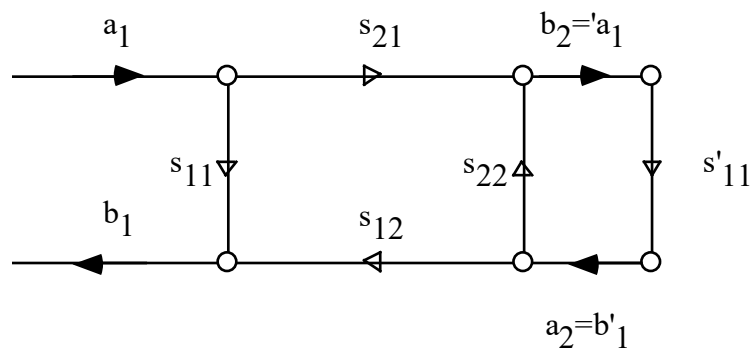


Figure 5.20 : réduction d'un graphe de fluence

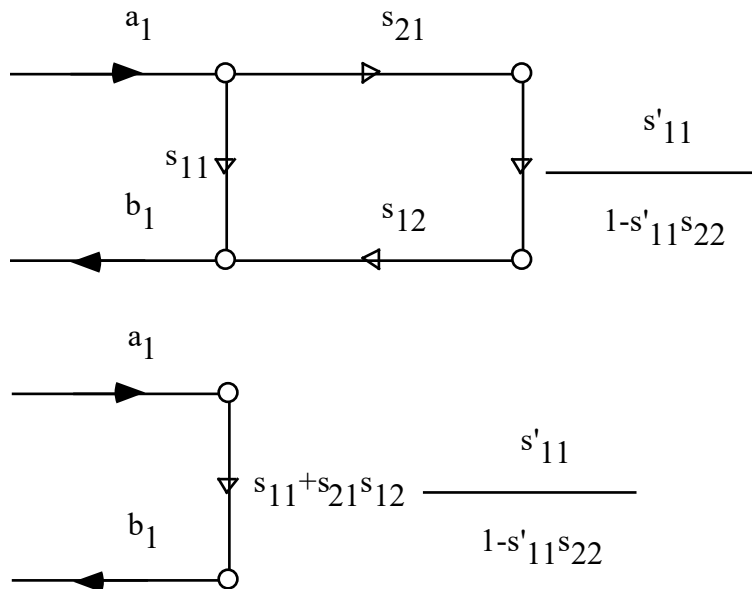
5.4.5.3 Exemple : mise en cascade de deux biportes



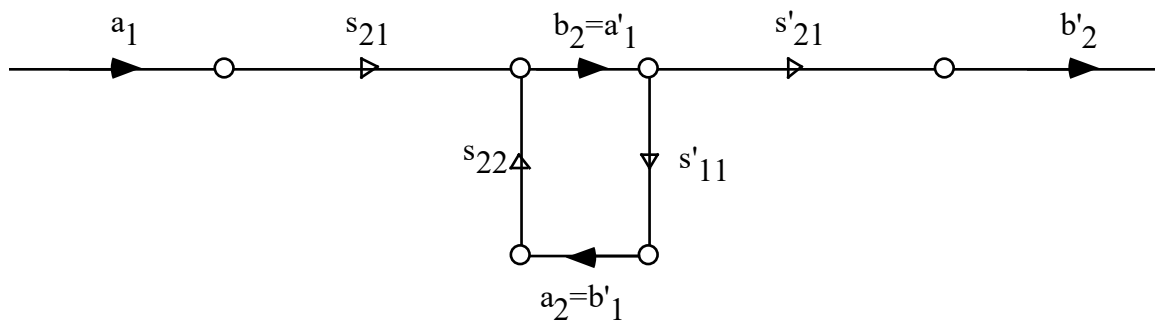
Etape 1 : on recherche les chemins possibles menant de a_1 à b_1 :



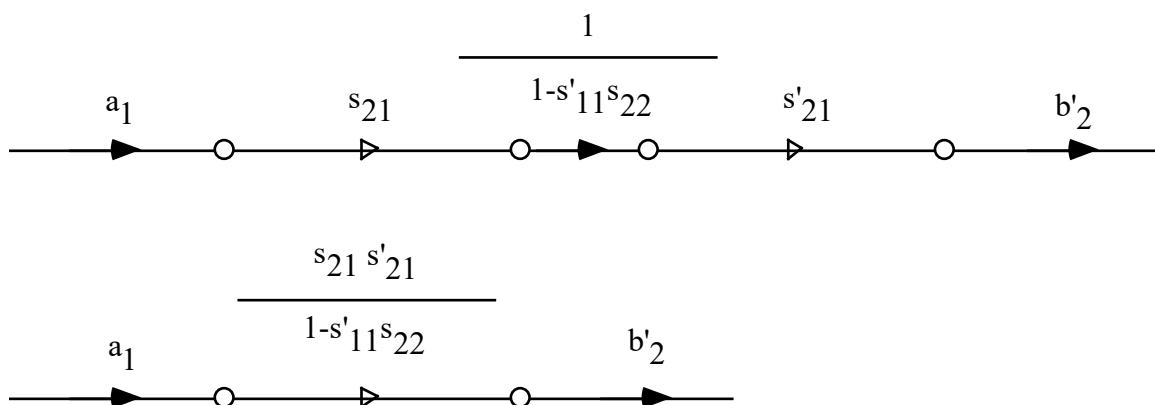
Etape 2 : on effectue la réduction



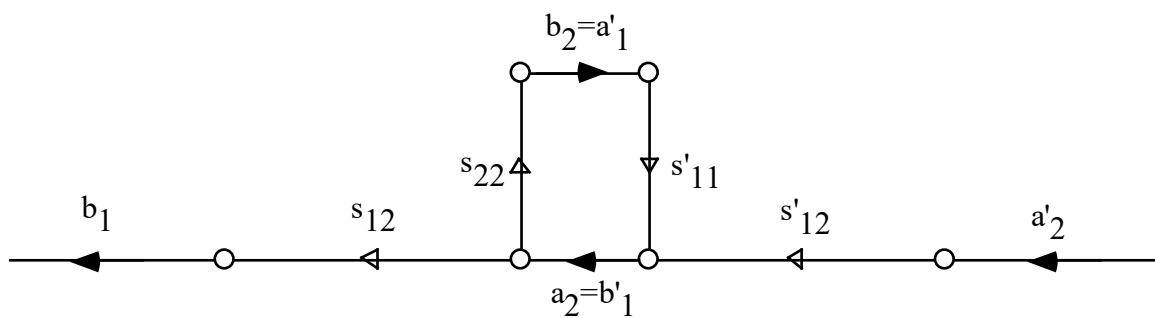
Etape 3 : on recherche les chemins possibles menant de a_1 à b'_2 :



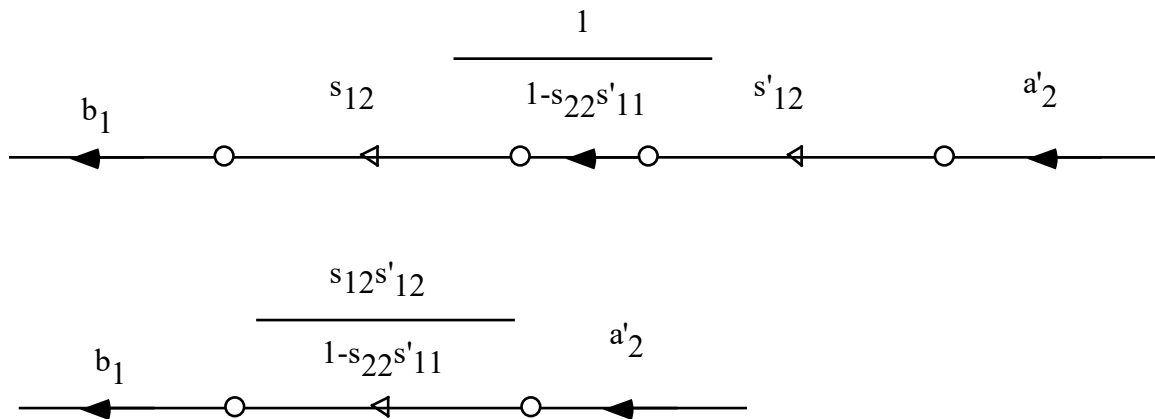
Etape 4 : on effectue la réduction



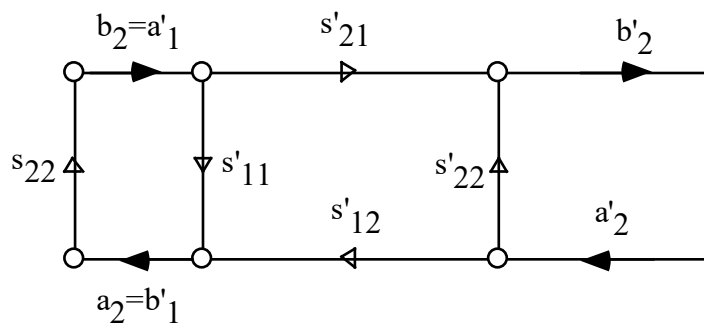
Etape 5 : on recherche les chemins possibles menant de a'2 à b1 :



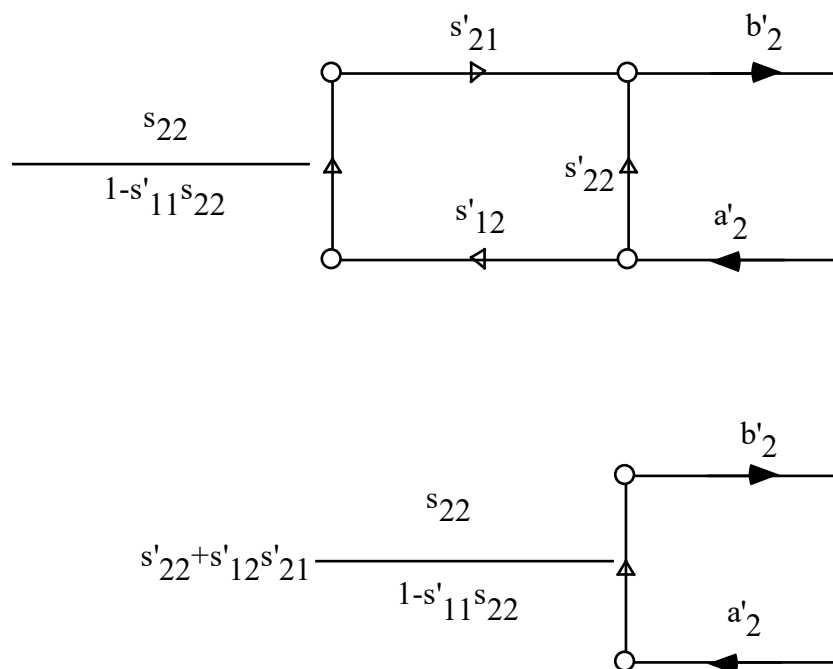
Etape 6: on réduit



Etape 7 : on recherche les chemins possibles menant de a'_2 à b'_2 :



Etape 8 : on réduit :



On a finalement obtenu la matrice de répartition de la mise en cascade de deux biportes :

$$\begin{bmatrix} b_1 \\ b_2 \end{bmatrix} = \begin{bmatrix} s_{11} + \frac{s_{21}s_{12}s'_{11}}{1 - s'_{11}s_{22}} & \frac{s_{12}s'_{12}}{1 - s'_{11}s_{22}} \\ \frac{s_{21}s'_{21}}{1 - s'_{11}s_{22}} & s'_{22} + \frac{s'_{12}s_{21}s_{22}}{1 - s'_{11}s_{22}} \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \end{bmatrix}$$

5.4.6 Résumé des caractéristiques générales de la matrice de répartition

- s_{ij} : fonction de transfert entre la porte j et i
- s_{ii} : coefficient de réflexion à l'accès i
- $|s_{ij}|^2 = \frac{P_i}{P_j}$ transfert de puissance entre l'accès j et i
- La matrice de répartition d'un circuit réciproque est symétrique
- La matrice de répartition d'un circuit sans pertes est de type $[S][\tilde{S}] = [I]$
- La matrice de répartition d'un circuit adapté a la diagonale principale nulle

5.4.7 Monoportes

La matrice d'un répartition d'un monoporte (ou bipôle) se réduit à un seul terme, le coefficient de réflexion

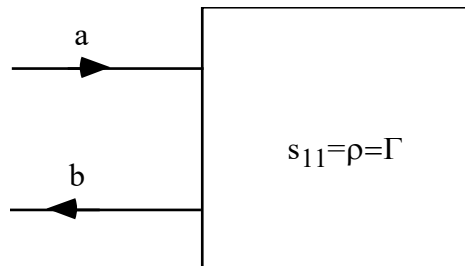


Figure 5.21 : composant à un accès, le monoporte

Son graphe de fluence est élémentaire (fig. 5.22)

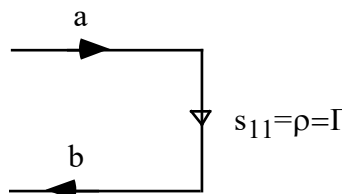


Figure 5.22 : graphe de fluence d'un monoporte

5.4.7.1 Monoportes sans pertes

Pour qu'un monoporte soit sans pertes, nous devons avoir :

$$|a|^2 = |b|^2$$

Ce qui revient à dire

$$\frac{|b|}{|a|} = |s| = 1$$

Nous avons donc qu'un monoporte sans pertes est un élément à réflexion totale, avec

$$s = e^{j\varphi}$$

Il existe deux cas particuliers de réflexion totale, soit le court-circuit où

$$b = -a \quad \text{donc} \quad s = -1$$

et le circuit-ouvert, avec

$$b = a \quad \text{donc} \quad s = 1$$

5.4.7.2 Monoporte adapté

Un monoporte adapté absorbe toute la puissance incidente et est caractérisé par un coefficient de réflexion nul :

$$s = 0$$

Remarque : il n'est manifestement pas possible d'adapter un monoporte sans pertes, les deux exigences étant incompatibles.

5.4.8 Biportes : composants à deux accès

La matrice de répartition d'un biporte comporte quatre termes :

$$[S] = \begin{bmatrix} s_{11} & s_{12} \\ s_{21} & s_{22} \end{bmatrix}$$

où les deux termes diagonaux sont les coefficients de réflexion aux deux portes et les deux termes hors diagonale les fonctions de transfert entre les deux portes. Le graphe de fluence d'un biporte est illustré à la figure 5.23

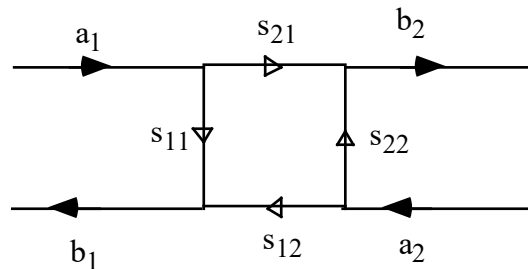


fig. 5.23 : graphe de fluence d'un biporte

5.4.8.1 Propriétés des biportes

- Pour un biporte réciproque, $s_{21}=s_{12}$

- Pour un biporte sans pertes

$$|s_{11}|^2 + |s_{21}|^2 = 1$$

$$|s_{12}|^2 + |s_{22}|^2 = 1$$

$$s_{11}^* s_{12} + s_{21}^* s_{22} = 0$$

- Pour un biporte à la fois réciproque et sans pertes, nous avons en plus

$$|s_{11}| = |s_{22}|$$

- Lorsque

$$s_{11} = s_{22}$$

on dit que le biporte est symétrique

- Pour un biporte adapté

$$s_{11} = s_{22} = 0$$

5.4.8.2 Exemples de biportes

- a) Affaiblisseur adapté

$$[S] = \begin{bmatrix} 0 & s_{12} \\ s_{12} & 0 \end{bmatrix}$$

Il s'agit d'un élément réciproque qui atténue la puissance entre l'entrée et la sortie. On définit le niveau d'affaiblissement comme

$$LA = 10 \log \frac{P_1}{P_2} = 10 \log \frac{|a_1|^2}{|b_2|^2} = 10 \log \frac{1}{|s_{12}|^2} = -20 \log |s_{12}|$$

Un affaiblisseur adapté absorbe de la puissance, il s'agit donc d'un élément avec pertes.

b) Déphaseur

La matrice de répartition d'un déphaseur réciproque adapté et sans pertes est donnée par :

$$[S] = \begin{bmatrix} 0 & e^{j\varphi} \\ e^{j\varphi} & 0 \end{bmatrix}$$

c) Isolateur

Un isolateur est un élément non réciproque qui laisse passer le signal dans un sens et le bloque dans l'autre. Cet élément est très utile pour protéger un élément (un générateur par exemple) contre une réflexion parasite. La matrice de répartition d'un isolateur idéal est donnée par :

$$[S] = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}$$

d) Biportes à caractéristiques dépendant de la fréquence (filtres, etc.)

La matrice de répartition d'un filtre est donnée par

$$[S] = \begin{bmatrix} s_{11}(\omega) & s_{12}(\omega) \\ s_{12}(\omega) & s_{22}(\omega) \end{bmatrix}$$

5.4.9 Triportes

La matrice de répartition d'un triporte est donnée par

$$[S] = \begin{bmatrix} s_{11} & s_{12} & s_{13} \\ s_{21} & s_{22} & s_{23} \\ s_{31} & s_{32} & s_{33} \end{bmatrix}$$

et son graphe de fluence par

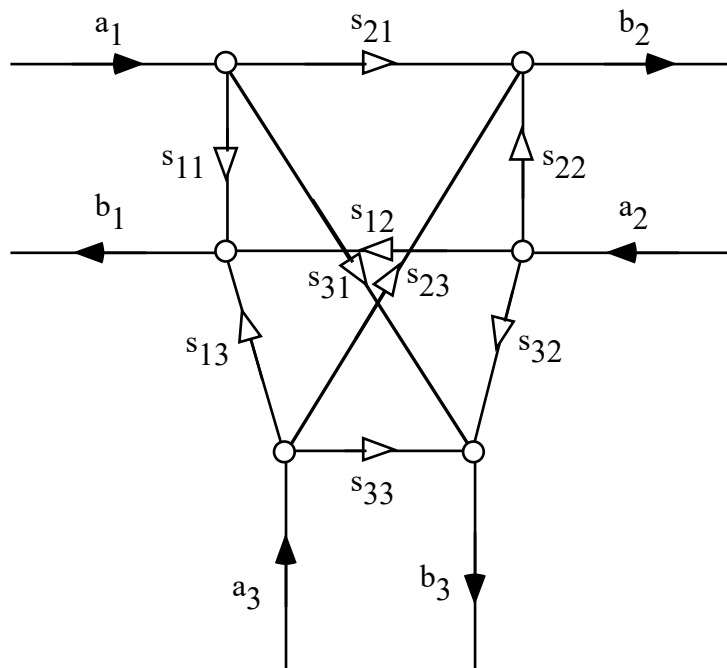


Figure 5.24 : graphe de fluence d'un triporte

5.4.9.1 Caractéristiques d'un triporte

- Pour un biporte réciproque, $s_{21}=s_{12}$, $s_{23}=s_{32}$, $s_{13}=s_{31}$

- Pour un biporte sans pertes

$$\sum_{i=1}^3 s_{ij}^* s_{ik} = \delta_{jk} \quad j, k = 1, 2, 3$$

- Pour un biporte adapté

$$s_{11} = s_{22} = s_{33} = 0$$

- Un triporte réciproque et sans pertes ne peut pas être adapté à ses trois accès. En effet, si cet élément existait, les relations impliquant qu'il est sans pertes s'écriraient :

$$|s_{12}|^2 + |s_{13}|^2 = 1$$

$$|s_{12}|^2 + |s_{23}|^2 = 1$$

$$|s_{23}|^2 + |s_{13}|^2 = 1$$

$$s_{13}^* s_{23} = 0$$

$$s_{12}^* s_{23} = 0$$

$$s_{12}^* s_{13} = 0$$

Supposons que s_{13} soit non nul. il en résulte que $s_{23}=0$ et $s_{12}=0$, ce qui est incompatible avec la deuxième relation. On peut refaire le même raisonnement en partant avec s_{23} ou s_{12} non nul.

5.4.9.2 exemples de triportes

a) Le circulateur

Un triporte non réciproque sans pertes peut être adapté à ses trois accès. Dans ce cas, les relations de conservation d'énergie s'écrivent :

$$|s_{21}|^2 + |s_{31}|^2 = 1$$

$$|s_{12}|^2 + |s_{32}|^2 = 1$$

$$|s_{23}|^2 + |s_{13}|^2 = 1$$

$$s_{12}^* s_{13} = 0$$

$$s_{21}^* s_{23} = 0$$

$$s_{31}^* s_{32} = 0$$

On suppose à nouveau que s_{13} est non nul. On obtient alors

$$s_{13} \neq 0 \Rightarrow s_{12} = 0 \Rightarrow |s_{32}| = 1 \Rightarrow s_{31} = 0 \Rightarrow |s_{21}| = 1 \Rightarrow s_{23} = 0 \Rightarrow |s_{13}| = 1$$

On peut choisir les plans de référence de manière à ce que les termes non nuls de la matrice de répartition soient réels. On obtient alors la matrice de répartition suivante :

$$[S] = \begin{bmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix}$$

L'élément obtenu est un circulateur, et son graphe de fluence est illustré à la figure 5.25

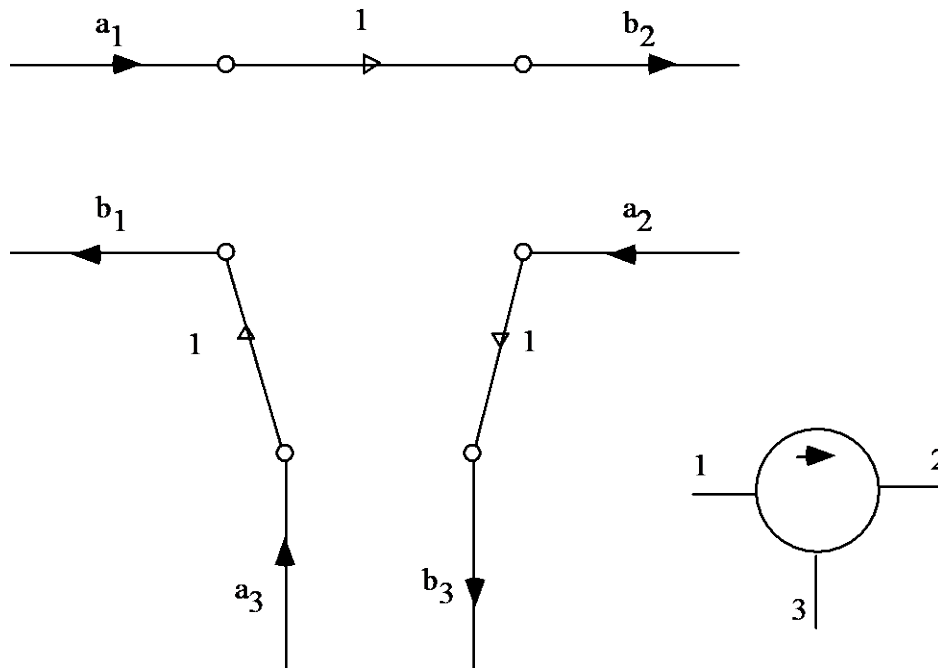


Figure 5.25 : graphe de fluence d'un circulateur idéal

b) triporte réciproque et adapté à deux accès

On démontre facilement que le seul triporte réciproque sans pertes adapté à deux accès ,n'est pas très intéressant, car il est complètement découplé : les termes de sa matrice de répartition sont donnés par :

$$|s_{ij}| = |s_{ji}| = |s_{kk}| \quad \text{avec} \quad i, j, k = 1, 2, 3$$

et un exemple de graphe de fluence est donné à la figure 5.26

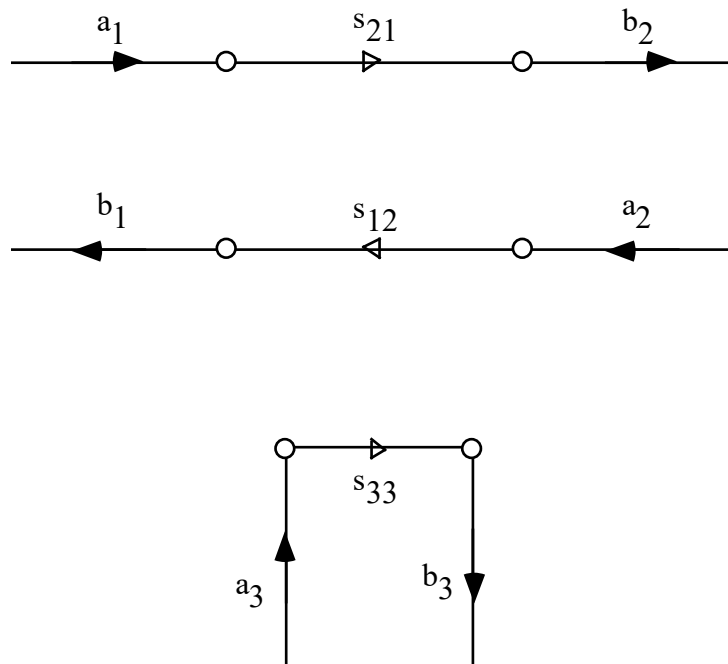


Figure 5.26 : triporte réciproque sans pertes adapté à deux de ses accès

5.4.10 Le quadriporte

La matrice de répartition et le graphe de fluence d'un quadriporte sont donnés par :

$$[S] = \begin{bmatrix} s_{11} & s_{12} & s_{13} & s_{14} \\ s_{21} & s_{22} & s_{23} & s_{24} \\ s_{31} & s_{32} & s_{33} & s_{34} \\ s_{41} & s_{42} & s_{43} & s_{44} \end{bmatrix}$$

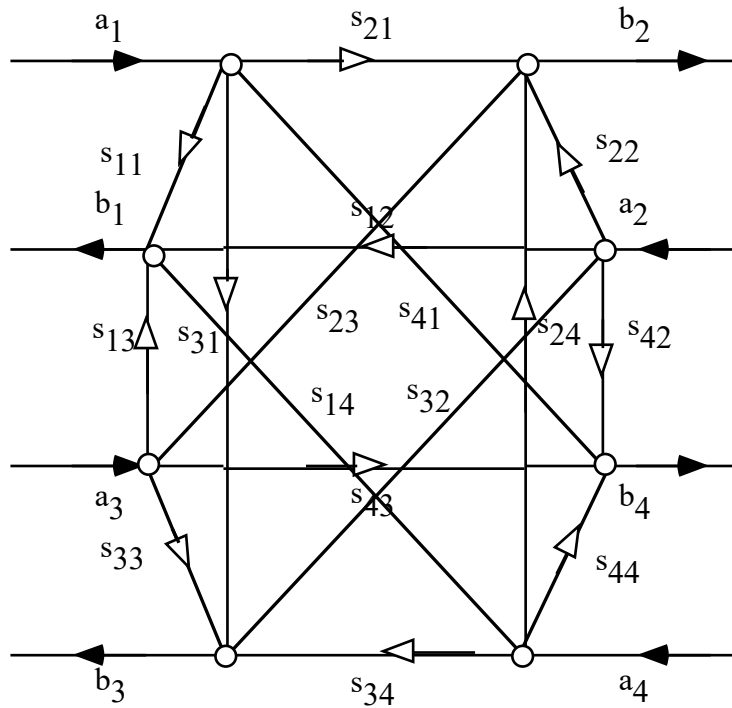


Figure 5.27 : graphe de fluence d'un quadriporte

5.4.10.1 Le coupleur directif

On démontre en considérant les équations de conservation de puissance que le seul quadriporte réciproque, adapté est sans perte à une matrice de répartition de type

$$[S] = \begin{bmatrix} 0 & s_{12} & 0 & s_{14} \\ s_{12} & 0 & s_{23} & 0 \\ 0 & s_{23} & 0 & s_{34} \\ s_{14} & 0 & s_{34} & 0 \end{bmatrix}$$

où de plus, nous devons avoir

$$|s_{12}|^2 + |s_{14}|^2 = 1$$

$$|s_{12}|^2 + |s_{23}|^2 = 1$$

$$|s_{23}|^2 + |s_{34}|^2 = 1$$

$$|s_{14}|^2 + |s_{34}|^2 = 1$$

$$s_{12}^* s_{23} + s_{14}^* s_{34} = 0$$

$$s_{12}^* s_{23} + s_{14}^* s_{34} = 0$$

Ce type d'élément qui lie à chaque fois une entrée à deux sorties, la troisième sortie étant isolée, est appelé un coupleur directif. En choisissant les plans de référence de manière judicieuse, on obtient finalement la matrice de répartition suivante:

$$[S] = \begin{bmatrix} 0 & \alpha & 0 & \beta e^{j\psi} \\ \alpha & 0 & \beta e^{j\theta} & 0 \\ 0 & \beta e^{j\theta} & 0 & \alpha \\ \beta e^{j\psi} & 0 & \alpha & 0 \end{bmatrix}$$

avec

$$\alpha^2 + \beta^2 = 1$$

$$\psi + \theta = \pi + 2n\pi$$

Démonstration :

La conservation de l'énergie donne

$$\begin{bmatrix} 0 & s_{12} & s_{13} & s_{14} \\ s_{12} & 0 & s_{23} & s_{24} \\ s_{13} & s_{23} & 0 & s_{34} \\ s_{14} & s_{23} & s_{34} & 0 \end{bmatrix}^* \begin{bmatrix} 0 & s_{12} & s_{13} & s_{14} \\ s_{12} & 0 & s_{23} & s_{24} \\ s_{13} & s_{23} & 0 & s_{34} \\ s_{14} & s_{23} & s_{34} & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

La troisième ligne * la première colonne et la première ligne * la deuxième colonne s'écrivent :

$$s_{13}^* s_{14} + s_{23}^* s_{24} = 0$$

$$s_{13}^* s_{23} + s_{14}^* s_{24} = 0$$

On multiplie la expression ligne par s_{14}^* et la deuxième par s_{23}^* puis on la soustrait à la première. On obtient

$$s_{13}^* (s_{14}^2 - s_{23}^2) = 0$$

Cette équation a deux solutions possibles :

a) $s_{13} = 0$

Les équations de conservation de l'énergie et les deux relations ci-dessus nous donnent alors

$$s_{14} \neq 0, s_{23} \neq 0, s_{24} = 0$$

On obtient la matrice de répartition

$$[S] = \begin{bmatrix} 0 & s_{12} & 0 & s_{14} \\ s_{12} & 0 & s_{23} & 0 \\ 0 & s_{23} & 0 & s_{34} \\ s_{14} & 0 & s_{34} & 0 \end{bmatrix}$$

b) $|s_{14}| = |s_{23}|$

On choisit alors les plans de référence de manière à ce qu'il y ait des termes imaginaires :

$$s_{14} = s_{23} = j\delta$$

On utilise en suite deux autres relations de la conservation de l'énergie : ligne 1 * colonne 1 et ligne 3 * colonne 3

$$\begin{aligned} |s_{12}|^2 + |s_{13}|^2 + |s_{14}|^2 &= 1 \\ |s_{13}|^2 + |s_{23}|^2 + |s_{34}|^2 &= 1 \end{aligned}$$

On obtient

$$|s_{12}| = |s_{34}|$$

On définit ces plans de référence de manière à ce que ces termes soient réels :

$$s_{12} = s_{34} = \gamma$$

On reprend deux relations de conservation de l'énergie : ligne 1 * colonne 4 et ligne 3 * colonne 4

$$\begin{aligned} s_{12}^* s_{24} + s_{13}^* s_{34} &= 0 = \gamma (s_{24} + s_{31}^*) \\ s_{13}^* s_{14} + s_{23} s_{24} &= 0 = \delta (s_{31}^* - s_{24}) \end{aligned}$$

Ce système d'équations admet deux solutions :

a) $s_{13} = s_{24} = 0$

Nous obtenons alors une matrice de répartition du quadriporte du même type que précédemment

$$[S] = \begin{bmatrix} 0 & s_{12} & 0 & s_{14} \\ s_{12} & 0 & s_{23} & 0 \\ 0 & s_{23} & 0 & s_{34} \\ s_{14} & 0 & s_{34} & 0 \end{bmatrix} = \begin{bmatrix} 0 & \gamma & 0 & j\delta \\ \gamma & 0 & j\delta & 0 \\ 0 & j\delta & 0 & \gamma \\ j\delta & 0 & \gamma & 0 \end{bmatrix}$$

b) $\gamma = \delta = 0$

Cette solution correspond plus à un quadriporte, mais à deux biportes découplés :

$$[S] = \begin{bmatrix} 0 & 0 & s_{13} & 0 \\ 0 & 0 & 0 & s_{24} \\ s_{13} & 0 & 0 & 0 \\ 0 & s_{24} & 0 & 0 \end{bmatrix}$$

La seule solution possible pour un quadriporte réciproque, adapté et sans pertes a donc une matrice de répartition de la forme :

$$[S] = \begin{bmatrix} 0 & s_{12} & 0 & s_{14} \\ s_{12} & 0 & s_{23} & 0 \\ 0 & s_{23} & 0 & s_{34} \\ s_{14} & 0 & s_{34} & 0 \end{bmatrix}$$

La démonstration continue par la réécriture des relations de conservation d'énergie pour cette matrice de répartition :

$$\begin{aligned} |s_{12}|^2 + |s_{14}|^2 &= 1 \\ |s_{12}|^2 + |s_{23}|^2 &= 1 \\ |s_{23}|^2 + |s_{34}|^2 &= 1 \\ |s_{14}|^2 + |s_{34}|^2 &= 1 \\ s_{12}s_{23}^* + s_{14}s_{34}^* &= 0 \\ s_{12}s_{14}^* + s_{23}s_{34}^* &= 0 \end{aligned}$$

On en déduit

$$\begin{aligned} |s_{12}| &= |s_{34}| = \alpha \\ |s_{14}| &= |s_{23}| = \beta \\ \alpha^2 + \beta^2 &= 1 \end{aligned}$$

On écrit ces termes sous forme polaire :

$$s_{12} = \alpha e^{j\varphi}$$

$$s_{34} = \alpha e^{j\eta}$$

$$s_{14} = \beta e^{j\psi}$$

$$s_{23} = \beta e^{j\theta}$$

On déduit des deux dernières relations de conservation de l'énergie que

$$(\varphi + \eta) = (\psi + \theta) + \pi + 2n\pi$$

On choisit alors les plans de référence pour que

$$\varphi = \eta = 0$$

Et on obtient le résultat final pour la matrice de répartition d'un quadriporte réciproque, adapté et sans pertes :

$$[S] = \begin{bmatrix} 0 & \alpha & 0 & \beta e^{j\psi} \\ \alpha & 0 & \beta e^{j\theta} & 0 \\ 0 & \beta e^{j\theta} & 0 & \alpha \\ \beta e^{j\psi} & 0 & \alpha & 0 \end{bmatrix}$$

Avec

$$\alpha^2 + \beta^2 = 1$$

$$\psi + \theta = \pi + 2n\pi$$

Le graphe de fluence de cet élément est donné à la figure 5.28

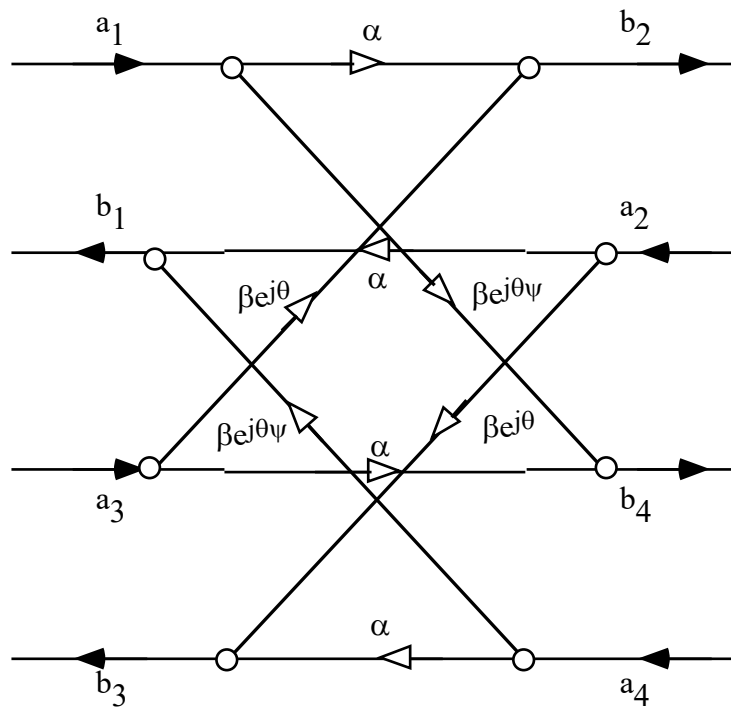


Figure 5.28 : graphe de fluence d'un quadriporte adapté, réciproque et sans pertes

Cas particulier : le coupleur symétrique

On choisit

$$\psi = \theta = \frac{\pi}{2}$$

et on obtient

$$[S] = \begin{bmatrix} 0 & \alpha & 0 & j\beta \\ \alpha & 0 & j\beta & 0 \\ 0 & j\beta & 0 & \alpha \\ j\beta & 0 & \alpha & 0 \end{bmatrix}$$

Son graphe de fluence est illustré à la figure 5.29

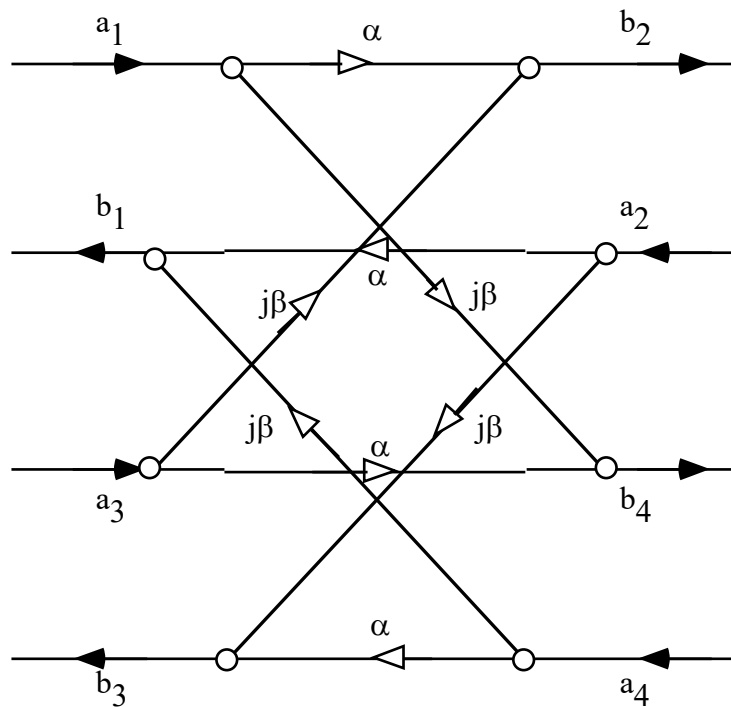


Figure 5.29 : graphe de fluence d'un coupleur symétrique

Exemple : le coupleur hybride $\alpha = \beta = \frac{1}{\sqrt{2}}$

$$[S] = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 & 1 & 0 & j \\ 1 & 0 & j & 0 \\ 0 & j & 0 & 1 \\ j & 0 & 1 & 0 \end{bmatrix}$$

Cas particulier : le coupleur asymétrique

On choisit

$$\psi = 0, \theta = \pi$$

et on obtient

$$[S] = \begin{bmatrix} 0 & \alpha & 0 & \beta \\ \alpha & 0 & -\beta & 0 \\ 0 & -\beta & 0 & \alpha \\ \beta & 0 & \alpha & 0 \end{bmatrix}$$

Son graphe de fluence est illustré à la figure 5.30

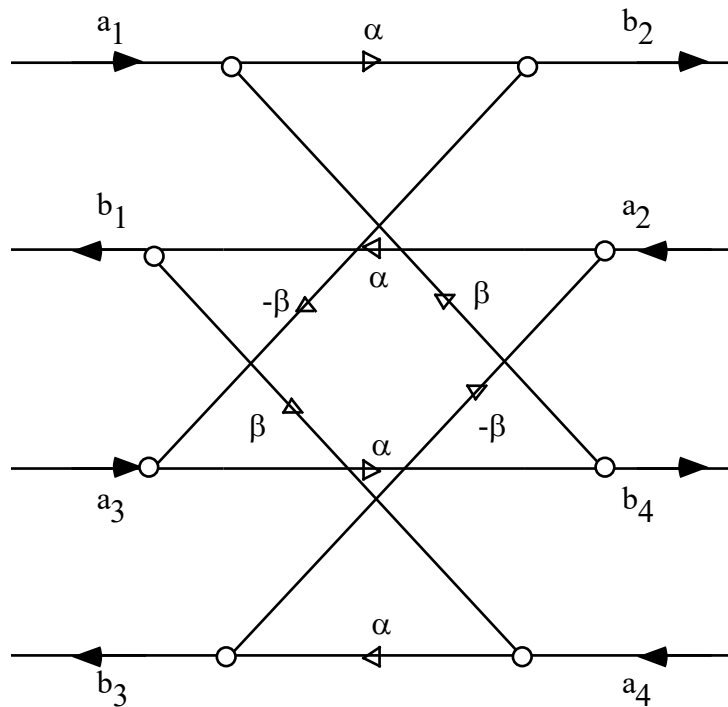


Figure 5.30 : graphe de fluence d'un coupleur asymétrique

Exemple : le coupleur en T hybride, ou cercle hybride (magic T, rat-race) $\alpha = \beta = \frac{1}{\sqrt{2}}$

$$[S] = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 & 1 & 0 & 1 \\ 1 & 0 & -1 & 0 \\ 0 & -1 & 0 & 1 \\ 1 & 0 & 1 & 0 \end{bmatrix}$$

5.10.4.2 Coupleur directif réel

On numérote les accès de manière à avoir

$$\alpha \geq \beta$$

On définit les grandeurs suivantes :

- Niveau d'affaiblissement : $LA = -20 \log \alpha$
- Niveau de couplage : $LC = -20 \log \beta$

Dans le cas d'un coupleur réel, nous avons de plus que

$$s_{ii}, s_{13}, s_{24} \text{ petits mais } \neq 0$$

Le coupleur est alors en plus caractérisé par les coefficients de réflexion à ses accès et par ses niveaux d'isolation :

- $LI_{13} = -20 \log |s_{13}|$
- $LI_{24} = -20 \log |s_{24}|$

La qualité d'un coupleur est donnée par sa directivité :

- $LD_{13} = LI_{13} - LC = -20 \log \frac{|s_{13}|}{\beta}$
- $LD_{24} = LI_{24} - LC = -20 \log \frac{|s_{24}|}{\beta}$

Plus la directivité est élevée, meilleur est le coupleur.

5.10.4.3 Symbole électrique du coupleur directif

Le symbole électrique du coupleur directif est illustré à la figure 5.31

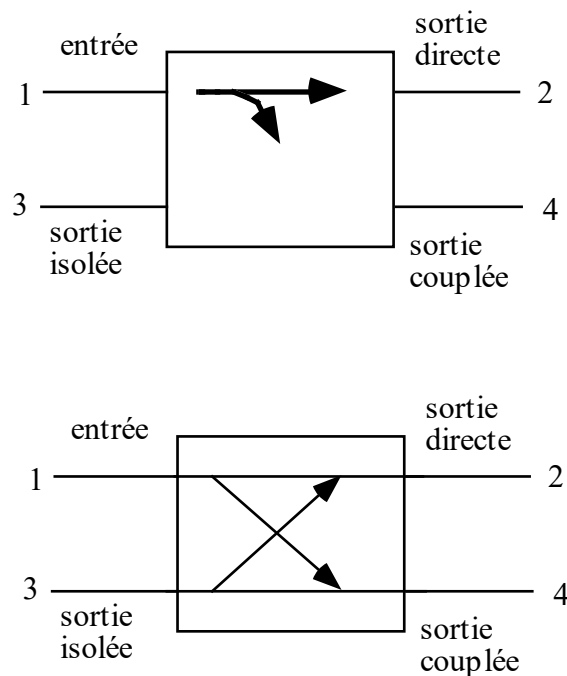


Figure 5.31 : Symbole du coupleur directif

6. La propagation d'ondes électromagnétiques dans l'espace libre

6.1 Propagation dans l'espace libre : Approche intuitive

Considérons une source de rayonnement électromagnétique délivrant une puissance P_e (figure 6.1).

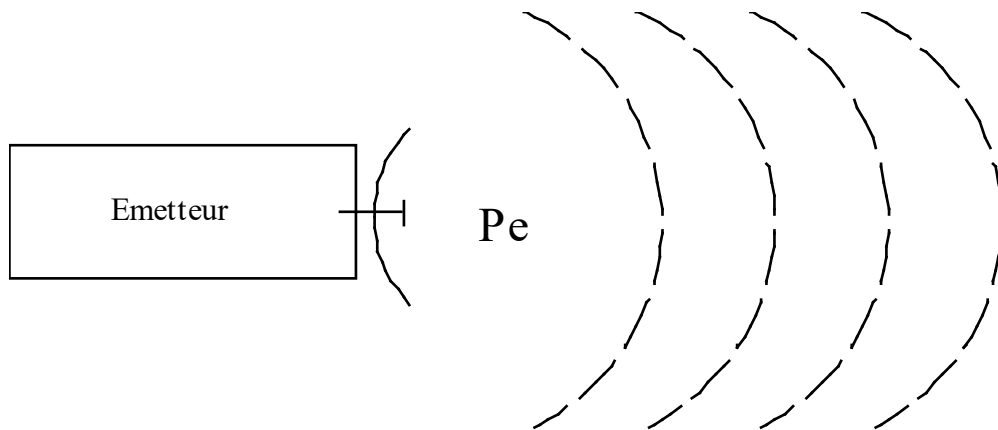


Figure 6.1

Supposons de plus que cette source soit isotrope, c'est à dire qu'elle émet uniformément dans toutes les directions de l'espace.

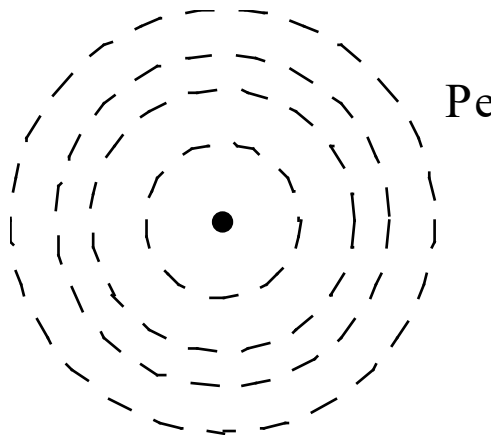


Figure 6.2

Alors, la densité de puissance à une distance R de la source sera de (figure 6.2) :

$$S = \frac{P_e}{4\pi R^2}$$

Plaçons maintenant un capteur ponctuel de surface ds à une distance R de la source. La puissance reçue par ce capteur sera de

$$P_r = \frac{P_e ds}{4\pi R^2}$$

Donc

$$\frac{P_r}{P_e} = \frac{ds}{4\pi R^2}$$

On en déduit que la puissance transmise entre cet émetteur isotrope et ce capteur ponctuel est inversement proportionnelle au carré de la distance qui les sépare.

6.2 Formule de Friis

Si on considère un émetteur et un récepteur réels, possédant des antennes physiquement réalisable en lieu et place du radiateur omnidirectionnel et du capteur infinitésimal, on écrit

$$\frac{P_r}{P_e} = \left(\frac{\lambda}{4\pi R} \right)^2 g_e g_r$$

avec

λ la longueur d'onde du signal transmis
 g_e le gain de l'antenne d'émission
 g_r le gain de l'antenne de réception

Cette relation est **la formule de Friis**, et son développement rigoureux, ainsi que les notions de gain d'antenne seront vus au chapitre 3.

Le point qui nous intéresse pour le moment est le terme entre parenthèses, soit $\left(\frac{\lambda}{4\pi R} \right)^2$, qui est le **facteur de pertes dans l'espace libre** et tient compte des pertes dues au caractère sphérique de la propagation de l'énergie fournie par la source.

6.3 Comparaison avec la propagation guidée

La figure 6.3 donne l'atténuation en fonction de la distance d'une onde se propageant dans un câble coaxial, un guide circulaire, une fibre optique ou dans l'espace libre (à 300 MHz dans ce cas).

On suppose une atténuation de 60 dB/km pour le câble coaxial, de 2 dB/km pour le guide circulaire et de 0.5 dB/km pour la fibre optique.

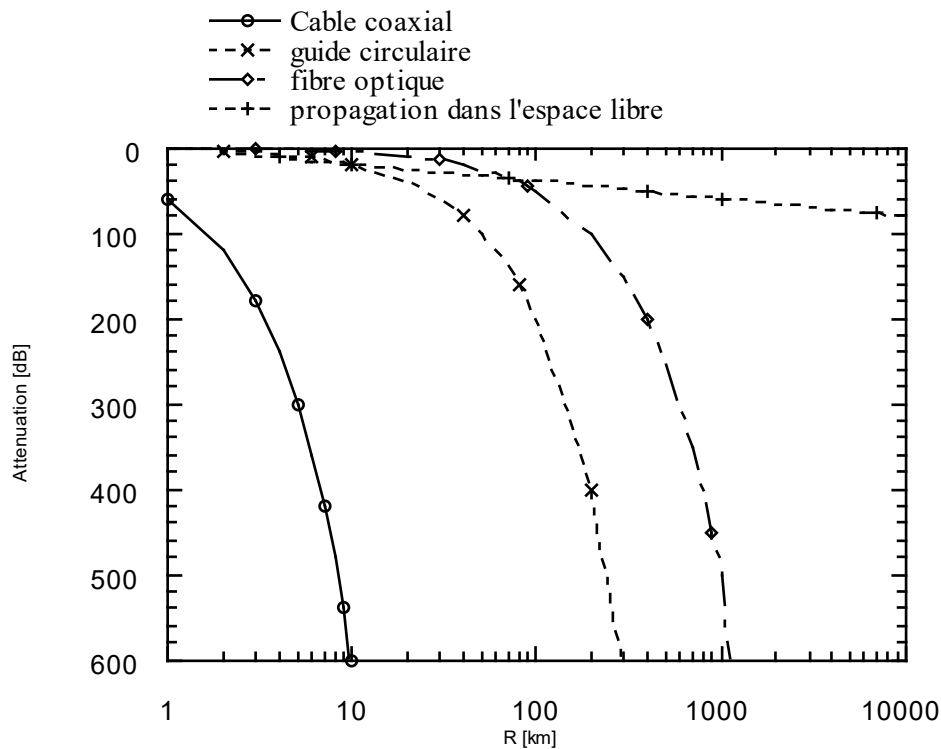


Figure 6.3 : Propagation dans différents milieux

Cette figure illustre clairement que pour une propagation à grande distance, la propagation d'ondes dans l'espace libre est beaucoup plus avantageuse que la propagation guidée : Cette dernière présente en effet un affaiblissement exponentiel avec la fréquence, alors que l'atténuation n'augmente qu'avec le carré de la distance dans le cas de la propagation dans l'espace libre.

6.4 Propagation dans l'atmosphère

La figure 6.3 représente l'atténuation d'une onde dans l'espace libre. L'atmosphère terrestre est fréquemment assimilée à l'espace libre, mais elle présente néanmoins quelques particularités qui vont avoir une certaine influence sur la propagation des ondes électromagnétiques : L'atmosphère présente des pertes et elle n'est pas homogène. De plus, une de ses couches est ionisée par la radiation du soleil, et influence fortement les ondes en dessous d'une certaine fréquence.

6.4.1 Pertes dans l'atmosphère

Les pertes dans l'atmosphère sont causées principalement par l'absorption de l'énergie microonde par la vapeur d'eau et les molécules d'oxygène. Un maximum a lieu dans l'absorption lorsque la fréquence du signal coïncide avec une résonance moléculaire de l'eau ou de l'oxygène, et la courbe d'atténuation due à l'atmosphère présente donc des pics distincts à ces fréquences (figure 2.4). Aux fréquences inférieures à 10 GHz, l'atmosphère n'a que très peu d'effet sur le signal. Des pics de résonance dus à la vapeur d'eau se trouvent à 22.2 GHz et 183.3 GHz, d'autre dus à l'oxygène à 60 et 120 GHz. Il existe donc des fenêtres dans les ondes millimétriques, aux environs de 35, 94 et 135 GHz où des systèmes de transmission peuvent exister avec des pertes atmosphériques raisonnables.

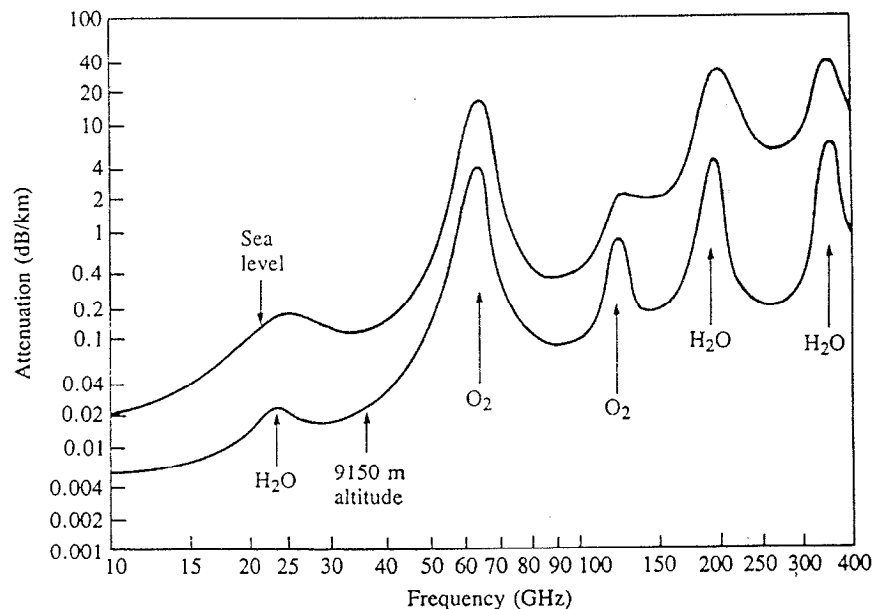


Figure 6.4 : Pertes dues à l'atmosphère terrestre

Dans certaines applications, les fréquences où les pertes atmosphériques sont élevées sont recherchées :

- Pour la surveillance de l'atmosphère terrestre par télémétrie pour des applications météorologiques par exemple, on utilise des radiomètres à 20 ou 55 GHz, de manière à maximiser les effets de l'atmosphère.
- Pour les communications entre satellites, qui se fait souvent à 60 GHz. Ceci présente l'avantage d'offrir une grande largeur de bande et de petites antennes, tout en minimisant le risque d'interférence, d'écoute ou de brouillage depuis la terre.
- Les militaires cherchent souvent une portée limitée pour garantir la confidentialité.

6.4.2 Inhomogénéité de l'atmosphère

Les propriétés électromagnétiques de l'air ($\epsilon_r \approx 1$, $\mu_r \approx 1$) sont très proches de celles du vide, mais elles ne sont pas rigoureusement identiques. Lorsqu'on étudie un trajet couvrant une grande distance dans l'atmosphère, il faut tenir compte des variations de la permittivité dues à la pression p , la température T et à l'humidité v . La permittivité relative en hyperfréquences est donnée par une relation empirique

$$\epsilon_r = n^2 \cong \left[1 + \left(\frac{79p}{T} - \frac{11v}{T} + \frac{3.8 \cdot 10^5 v}{T^2} \right) 10^{-6} \right]^2$$

où p est la pression barométrique en millibar, T la température en Kelvin, v la pression de vapeur d'eau en millibar et n l'indice de réfraction. Des mesures ont permis d'établir un profil moyen pour ϵ_r , en fonction de l'altitude h , que l'on appelle l'atmosphère standard (figure 6.5).

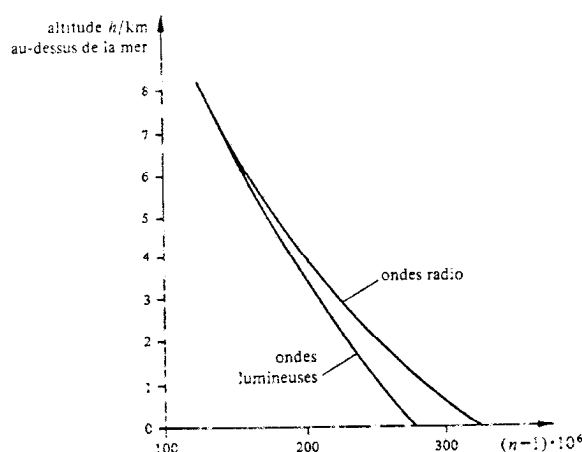


Figure 6.5 : Profil de l'atmosphère standard [fig. 8.14, Traité d'Electricité, vol XIII, Hyperfréquences, F. Gardiol, PPR, 1981]

L'indice de réfraction n est une fonction de l'altitude h et donc de la distance r jusqu'au centre de la terre. Considérons la figure 2.6, où le problème de la propagation dans l'atmosphère est simplifié à un problème à trois couches homogènes d'indice n_0 , n_1 et n_2 .

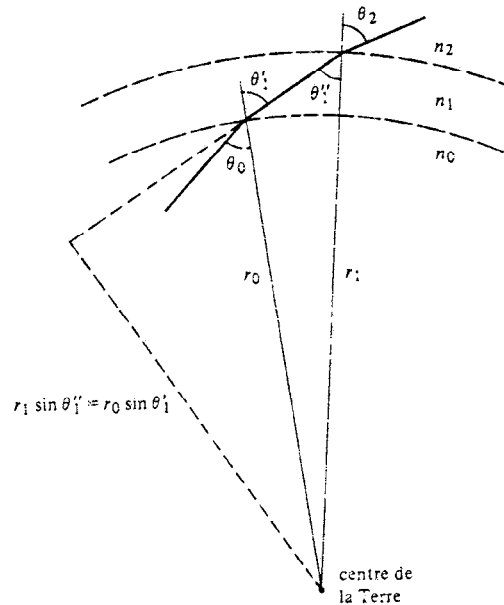


Figure 6.6 : Propagation dans des milieux sphériquement stratifiés [fig 8.15, Traité d'Electricité, vol XIII, Hyperfréquences, F. Gardiol, PPR, 1981]

La loi de Snell nous dit que

$$n_0 \sin \theta_0 = n_1 \sin \theta_1'$$

$$n_1 \sin \theta_1'' = n_2 \sin \theta_2$$

De plus, des considérations géométriques montrent que

$$r_1 \sin \theta_1'' = r_0 \sin \theta_1'$$

En combinant ces deux relations, on obtient

$$n_0 r_0 \sin \theta_0 = n_2 r_1 \sin \theta_2$$

Soit en généralisant à une variation continue de l'indice

$$nr \sin \theta = n_0 r_0 \sin \theta_0$$

La propagation d'une onde électromagnétique n'est donc pas rectiligne dans l'atmosphère, mais suit une courbe (figure 2.7). Pour la planification d'une liaison hertzienne (soit d'une liaison terrestre), il est néanmoins plus pratique d'admettre malgré tout une propagation rectiligne et de corriger en conséquence la courbure terrestre en augmentant artificiellement le rayon terrestre R . Pour une zone tempérée, ce rayon fictif a été déterminé par l'observation et vaut

$$R' = \frac{4R}{3} = 8500 \text{ km}$$



Figure 6.7 : rayon terrestre fictif

Le profil moyen de l'atmosphère donné à la figure 6.6 n'est qu'une valeur moyenne de valeurs effectivement mesurées. Lors de fluctuations atmosphériques (inversions thermiques par exemple) on peut avoir localement une région où l'indice de réfraction croît avec l'altitude. Des ondes guidées apparaissent alors; Au voisinage de ces canaux se situent des zones d'ombre où il est apparemment impossible de transmettre ou de recevoir. Ces anomalies se produisent plus fréquemment au dessus des côtes (par exemple dans le triangle des Bermudes).

6.4 3 Effet de la ionosphère

La ionosphère est une des couches extérieures de l'atmosphère, et contient des particules ionisées, donc chargées par le rayonnement solaire. Cette région forme un plasma, donc un milieu non linéaire, caractérisé par N , le nombre de ions par unité de volume qu'il contient.

Une onde électromagnétique arrivant dans un plasma sera réfléchiée, transmise ou absorbée, en fonction de la densité du plasma et de la fréquence de l'onde.

On associe une permittivité effective à une région contenant un plasma uniforme

$$\varepsilon_e = \varepsilon_0 \left(1 - \frac{\omega_p^2}{\omega^2} \right) \quad \omega_p = \sqrt{\frac{Nq^2}{m\varepsilon_0}}$$

avec

N : nombre de ions/volume

q : charge de l'électron

m : masse de l'électron

$\varepsilon_0 = 8.854 \cdot 10^{-12} \text{ As/Vm}$: permittivité du vide

Une onde électromagnétique incidente sur un plasma peut, en première approximation, présenter trois types de comportement :

1) $\omega \ll \omega_p$

La permittivité effective du milieu tend vers $-\infty$, et l'impédance caractéristique du milieu

$$Z_{milieu} = \sqrt{\frac{\mu}{\epsilon}}$$

est imaginaire et proche de 0. Il y aura donc réflexion totale de l'onde incidente (figure 6.8).

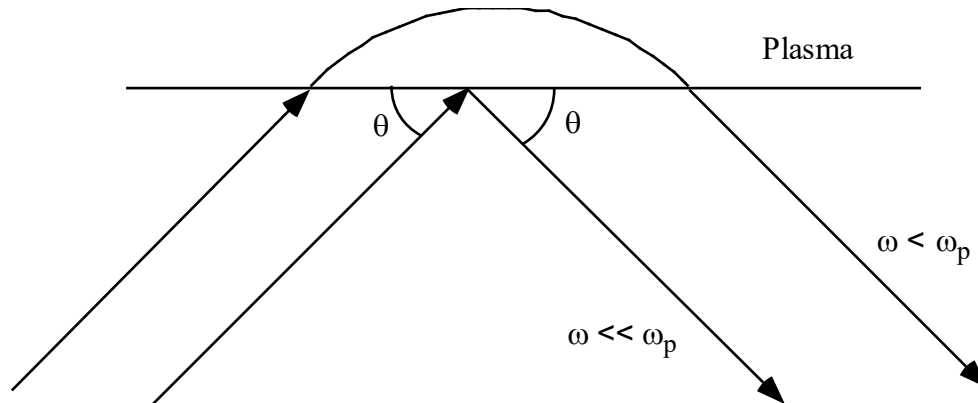


Figure 6.8 : Onde réfléchi par le plasma

II) $\omega = \omega_p$

La permittivité effective du plasma est nulle, ainsi que le facteur de propagation k de l'onde

$$k = \omega \sqrt{\epsilon \mu}$$

L'onde est absorbée par le plasma et entretient la pulsation de ce dernier (figure 6.9).

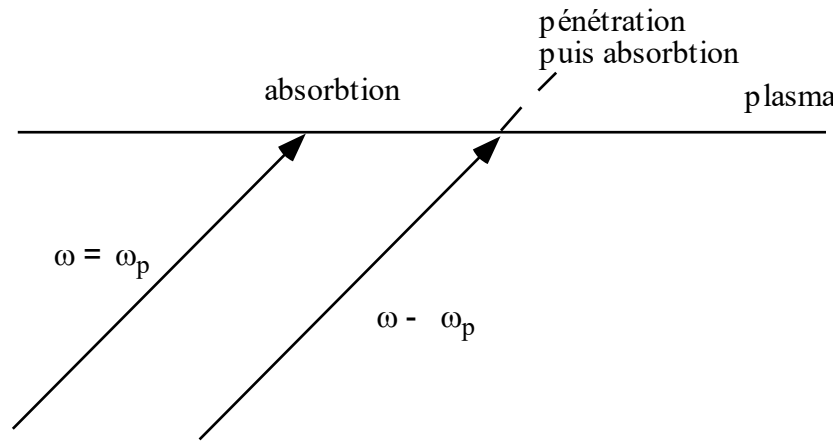


Figure 6.9 : onde absorbée par le plasma

III) $\omega \gg \omega_p$

La permittivité effective du plasma tend vers l'unité, et l'onde n'est pas perturbée par le plasma (figure 2.10)

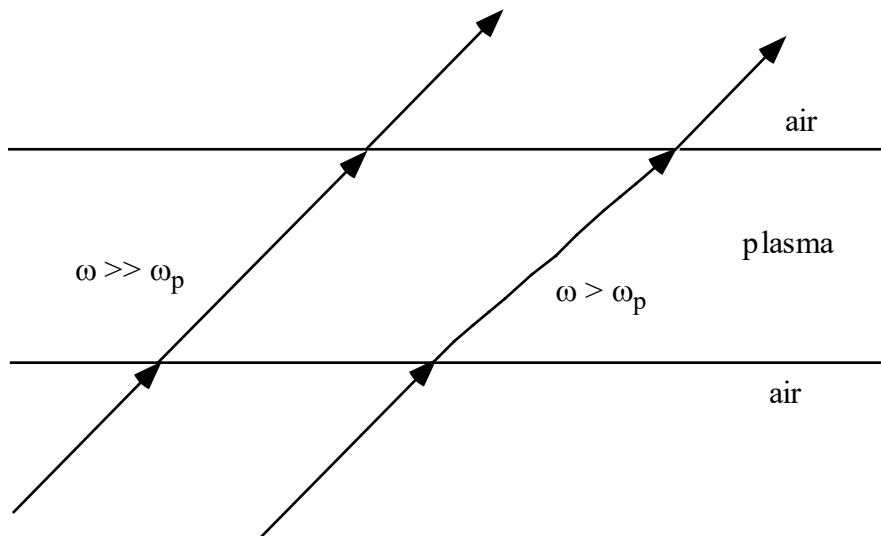


Figure 6.10 : onde non perturbée par le plasma.

Pour le plasma formant la ionosphère, la fréquence de pulsation propre est aux environs de

$$f_p = 0,2 \pi \approx 8 \text{ MHz}$$

Il existe d'autres types de plasmas, créés par exemple par un corps entrant dans l'atmosphère (satellite, météorite), la foudre, une explosion nucléaire, ...

La navette spatiale, par exemple, génère autour d'elle un plasma très dense lorsqu'elle rentre dans l'atmosphère. Ce plasma est dû à la très haute température causée localement par le frottement. Ce phénomène est à l'origine du black-out entre la navette et la terre, qui dure quelques minutes.

6.5 Réflexion sur un obstacle

Une onde électromagnétique qui atteint un obstacle est en partie réfléchi. Le récepteur peut ainsi recevoir une onde directe et une onde indirecte, réfléchi par l'obstacle. La combinaison de ces deux ondes peut être constructive ou destructive, selon leur déphasage relatif.

Deux signaux de même fréquence sont en opposition de phase lorsque la différence de leurs trajets est un multiple impair d'une demi-longueur d'onde (figure 6.11),

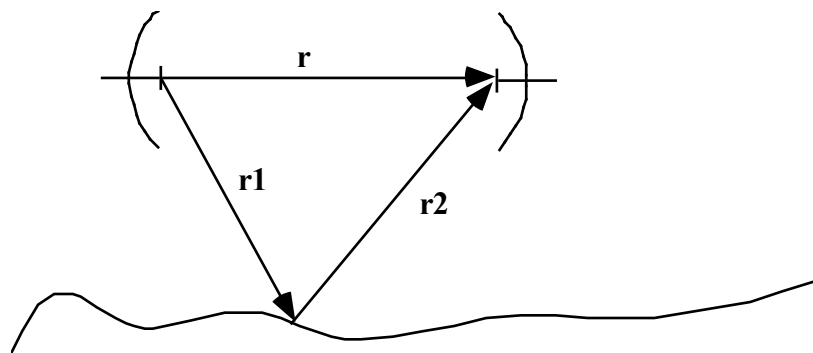


Figure 6.11 : onde directe et indirecte

donc lorsque

$$|r_1| + |r_2| - |r| = n \frac{\lambda}{2}, \quad n \text{ impair}$$

Le lieu des points qui satisfont à cette relation se trouve sur une famille d'ellipsoïdes, dont les foyers coïncident avec les deux antennes (figure 6.12), les ellipsoïdes de Fresnel. En un point donné entre les deux antennes, le rayon h_0 de la section droite du premier ellipsoïde vaut

$$h_0 = \sqrt{\frac{\lambda L_1 L_2}{L}}$$

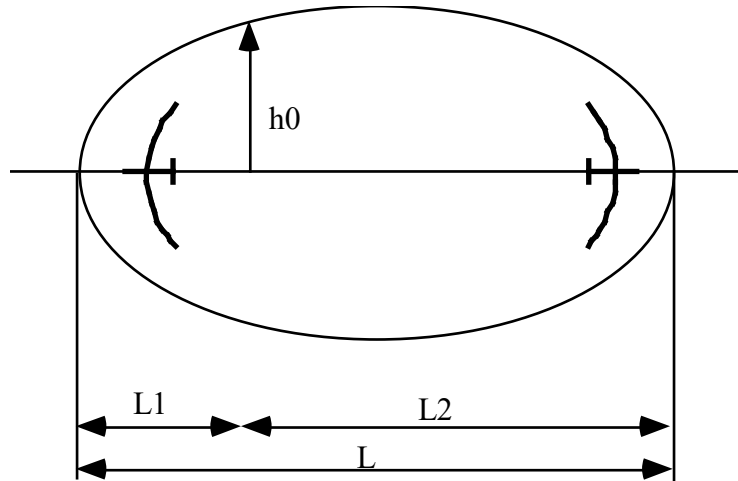


Figure 6.12 : Première ellipsoïde de Fresnel

On constate qu'en pratique seuls les signaux réfléchis dans la première zone de Fresnel ont une amplitude suffisante pour provoquer des interférences. Dans la mesure du possible, on veille donc à ce que cette zone soit libre d'obstacles.

Ceci n'est malheureusement pas toujours possible, notamment dans le cas des communications intérieures et des communications mobiles. Le problème posé par les interférences est alors résolu en utilisant une diversité spatiale ou spectrale à la réception (cf.. chapitre antennes).

6.5.1 Application : Dégagement du trajet d'une liaison hertzienne

Il faut prendre en compte deux points lors du calcul du dégagement d'une liaison hertzienne :

- 1) De la courbure terrestre
- 2) Du premier ellipsoïde de Fresnel

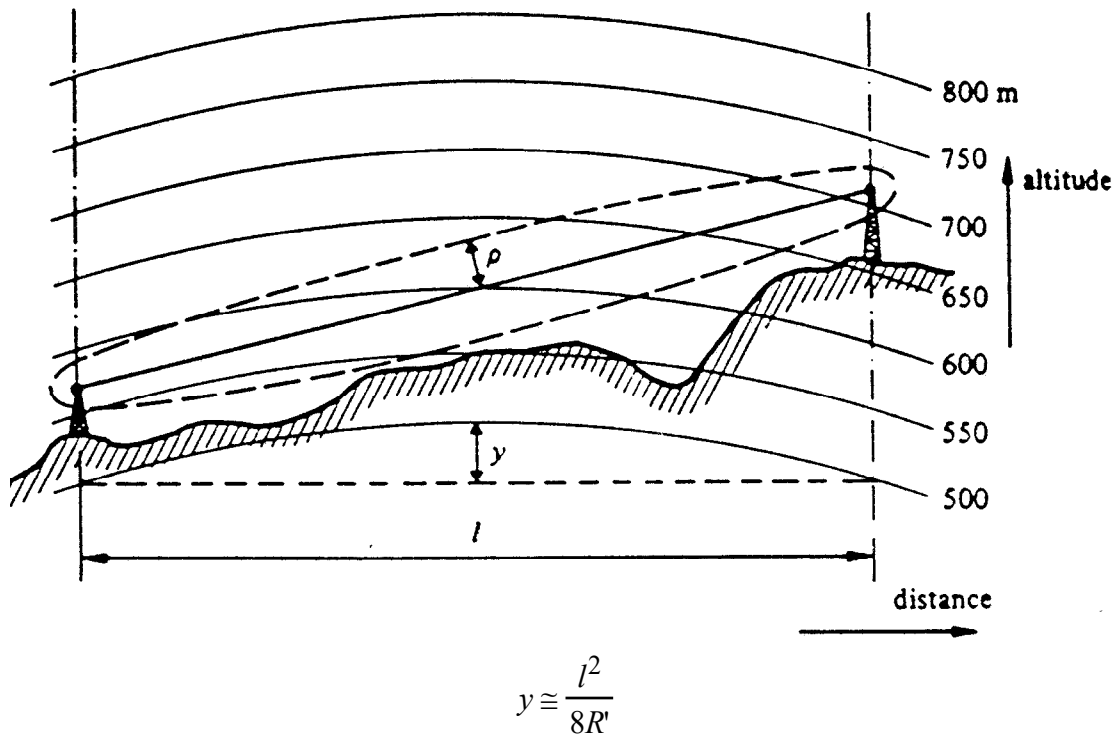


Figure 6.13 : Dégagement d'une liaison hertzienne [tirée du vol XVIII du Traité d'Electricité, "Systèmes de Télécommunications", par P.G. Fontollet, PPR 1983, fig. 12.4]

6.6 Effet Doppler

Un récepteur perçoit le signal d'un émetteur mobile à une fréquence différente de celle d'émission. Si le signal d'émission est monochromatique de fréquence f_0 , et que le véhicule se déplace à la vitesse v , la fréquence f_r du signal reçu sera égale à

$$f_r = f_0 \left(1 + \frac{v}{c} \cos \alpha \right)$$

avec

v : vitesse du mobile

c : vitesse de la lumière

α : angle formé entre le vecteur vitesse du mobile $\mathbf{v}$, et le vecteur de propagation d'onde (donnant la direction de propagation de l'onde) $\mathbf{k}$ (figure 2.14)

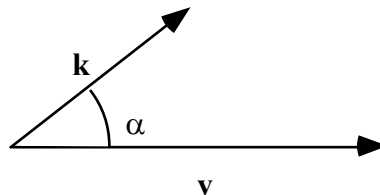


Figure 6.14 : Effet Doppler

Pour un mobile qui s'approche du récepteur, $f_r > f_0$, alors que pour un mobile qui s'éloigne $f_r < f_0$.

Exemple 1 : Radar Doppler

L'effet Doppler est utilisé par les radars pour plusieurs applications : La mesure de la vitesse et la détection de mouvement par exemple.

La vitesse radiale (figure 2.15) d'un objet est obtenue par la relation

$$v_r = \frac{c \Delta f}{2 f_0} \quad \Delta f = f_0 - f_r$$

Figure 6.15 : Principe d'un radar Doppler

Exemple 2 : Communications avec des mobiles

L'effet Doppler affecte les transmissions avec un mobile, qu'il soit aéronautique, maritime ou terrestre, et se traduit par un décalage en fréquence. Deux situations peuvent se présenter :

1) Le mobile est bien dégagé, et ne reçoit qu'une onde directe. Dans ce cas, un signal émis à la fréquence f_0 sera perçu à la fréquence f_r (cf. ci-dessus), mais un signal monochromatique restera monochromatique.

2) Le mobile n'est pas dans un espace dégagé, et l'onde est réfléchiée et diffractée avant d'atteindre le mobile. Ce dernier reçoit alors la superposition d'un grand nombre d'ondes provenant de directions $\mathbf{a}_i$ pratiquement quelconques, et subissant un décalage différent. Une raie spectrale f est ainsi transformée en une répartition d'énergie sur l'intervalle de fréquence :

$$\left[f_0 - \Delta_f; f_0 + \Delta_f \right] \quad \Delta_f = \frac{f_0 v}{c}$$

Pour les mobiles terrestres, l'élargissement Doppler est très faible : Dans le cas d'un système de téléphone cellulaire GSM par exemple ($f_{\max} = 960$ MHz), en considérant une voiture circulant à 140 km/h, nous obtenons

$$\Delta_f = 124 \text{ Hz}$$

Donc un étalement du spectre maximum de 248 Hz.

L'effet Doppler impose l'utilisation d'un système d'asservissement en fréquence lorsque certains types de modulation sont utilisés.

Exemple 3 : Communication par satellites et orbite géostationnaire

L'effet Doppler affecte évidemment aussi les liaisons avec les vaisseaux spatiaux et les satellites.

L'orbite d'un satellite tournant autour de la terre est une ellipse dont l'un des foyers coïncide avec le centre de la terre. La vitesse de rotation du satellite autour de la terre est régie par la troisième loi de Kepler :

$$T = 2\pi \sqrt{\frac{a^3}{GM}} = 2\pi \sqrt{\frac{(h+R)^3}{GM}}$$

avec :

T : La période de rotation du satellite

a : le demi-grand axe de l'ellipse

$G = 6.67 \cdot 10^{-11} \text{ Nm}^2/\text{kg}$: constante de gravitation universelle

$M = 6 \cdot 10^{24} \text{ kg}$: la masse de la terre (supposée homogène)

$R = 6370 \text{ km}$: le rayon terrestre

Il existe une orbite où le satellite paraît immobile vu de la terre : elle répond aux conditions suivantes :

- Synchronisme avec la rotation terrestre. La période de rotation est égale à un jour sidéral, soit 23h 56' 4".
- Orbite circulaire.
- Orbite équatoriale.

Cette orbite est dite **géostationnaire**. Son altitude est déterminée à l'aide de la loi de Kepler, et vaut

$$h = 35'786 \text{ km}$$

Avantages d'un satellite géostationnaire par rapport à un satellite à défilement :

- L'antenne au sol est pointée vers un point fixe dans le ciel (pas de tracking)
- Pas d'effet Doppler : le satellite paraît immobile depuis le sol
- le satellite n'est pas ralenti par les couches supérieures de l'atmosphère

Inconvénients d'un satellite géostationnaire par rapport à un satellite à défilement :

- Le temps nécessaire au trajet terre-satellite est élevé. Une transmission téléphonique aller et retour prend presque une demi seconde, ce qui est gênant pour la conversation.
- La distance à parcourir est grande, l'équipement est donc cher.

6.7 Ondes Electromagnétiques dans l'espace libre

Une onde électromagnétique est un flux d'énergie se propageant dans un milieu, sous la forme d'un champ électrique et d'un champ magnétique. Ces derniers sont liés par les équations de Maxwell, et répondent à une équation d'onde. Nous ne considérerons ici que le cas harmonique, soit un champ présentant une dépendance sinusoïdale avec le temps. De plus, on utilisera une notation de phaseurs :

$$\mathcal{E}(\mathbf{r}, t) = \sqrt{2} \operatorname{Re} \left[\mathbf{E}(\mathbf{r}) e^{j\omega t} \right]$$

Le facteur $\sqrt{2}$ est introduit dans la définition pour que le module du phaseur corresponde à la valeur efficace du signal, conformément à la convention adoptée dans le Traité d'Electricité.

(Note : Dans plusieurs ouvrages, ce facteur ne fait partie de la définition du phaseur, dont la norme est alors la valeur de crête du signal. Dans ce dernier cas, les relations associées aux énergies et aux puissances contiennent un facteur 1/2, qui n'existe pas avec notre définition.)

6.7.1 Equation d'onde

En prenant comme point de départ les équations de Maxwell sous leur forme harmonique dans un milieu sans sources :

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H}$$

$$\nabla \times \mathbf{H} = j\omega\varepsilon\mathbf{E}$$

$$\nabla \cdot \mathbf{E} = 0$$

$$\nabla \cdot \mathbf{H} = 0$$

et en y appliquant les principes de l'analyse vectorielle, on obtient les équations d'onde pour le champ électrique et le champ magnétique

$$\nabla^2 \mathbf{E} + k^2 \mathbf{E} = 0$$

$$\nabla^2 \mathbf{H} + k^2 \mathbf{H} = 0$$

Où k dépend du milieu et est appelé le nombre d'onde.

$$k = \omega\sqrt{\varepsilon\mu} = \omega\sqrt{\varepsilon_0\mu_0}\sqrt{\varepsilon_r\mu_r} = \frac{\omega}{c}\sqrt{\varepsilon_r\mu_r}$$

Parmi l'infinité de solutions qui satisfont à cette équation d'ondes, nous allons nous intéresser aux ondes planes et aux ondes sphériques : L'onde plane, car il s'agit de l'onde la plus simple, et l'onde sphérique, car c'est elle que l'on rencontre le plus fréquemment dans la nature. En effet, l'onde électromagnétique émise par une antenne devient une onde sphérique dès que l'on s'éloigne suffisamment de l'antenne. Du point de vue de la réception, l'onde captée par l'antenne du récepteur peut-être localement assimilée à une onde plane (figure 6.16).

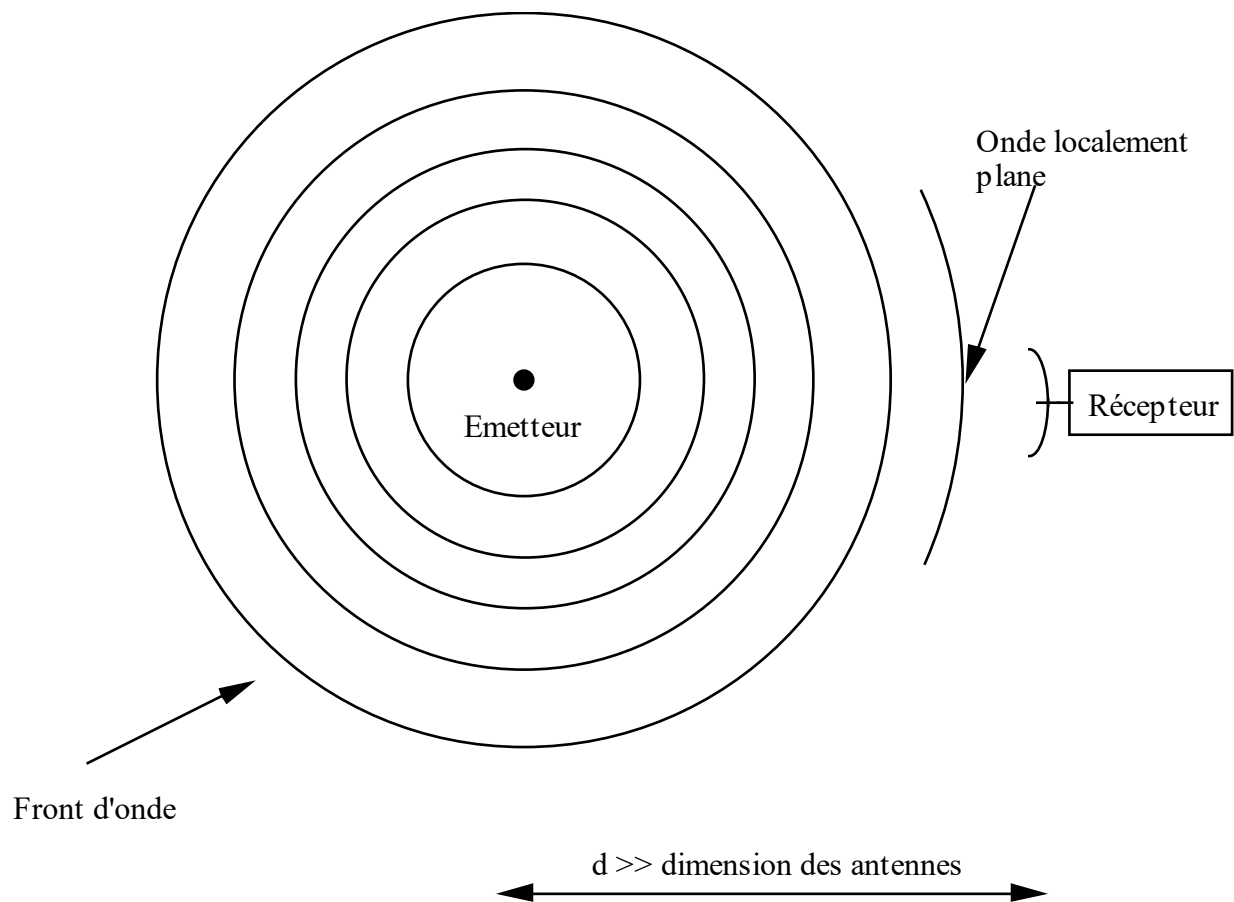


Figure 2.16 : Onde sphérique

6.7.2 Ondes planes

Les ondes planes sont caractérisées par le fait que les vecteurs $\mathbf{E}$, $\mathbf{H}$ et $\mathbf{k}$ (vecteur de propagation) sont mutuellement perpendiculaires et forment un repère direct : $\mathbf{E} \times \mathbf{H}$ est dans le même sens que $\mathbf{k}$. De plus, la phase de $\mathbf{E}$ et de $\mathbf{H}$ est constante sur des plans, appelés plans équiphases, qui sont perpendiculaires à la direction de propagation de l'onde $\mathbf{k}$. On pourrait, de manière équivalente dire qu'une onde plane est une onde où les champs ne varient que dans une seule direction d'un repère cartésien, par exemple $\hat{\mathbf{z}}$

Exemple :

Si l'on fait coïncider un repère cartésien avec $\mathbf{E}$, $\mathbf{H}$ et $\mathbf{k}$ respectivement, de manière à ce que les seules composantes non nulles de ces vecteurs sont E_x , H_y et k_z . L'équation d'onde devient

$$\frac{d^2 E_x}{dz^2} + k^2 E_x = 0$$

Les solutions à cette équation sont des combinaisons linéaires de e^{jkz} et e^{-jkz} . En particulier,

considérons la solution

$$E_x = E_0 e^{-jkz}$$

qui est une onde se propageant dans la direction positive de l'axe z. Le champ magnétique associé est obtenu par

$$j\omega\mu\mathbf{H} = -\nabla \times \mathbf{E} = \hat{\mathbf{y}} \ jk \ E_x$$

$$E_x = \sqrt{\frac{\mu}{\varepsilon}} H_y$$

Le facteur de proportionnalité entre les deux champs a la dimension d'une impédance : c'est l'impédance caractéristique du milieu :

$$Z_c = \sqrt{\frac{\mu}{\varepsilon}} \quad \text{Dans le vide : } Z_c = 120 \pi \cong 377 \Omega$$

Cette grandeur joue pour les ondes le même rôle que l'impédance caractéristique pour les lignes de transmission.

Les champs instantanés sont donnés par

$$E_x = E_0 \cos(\omega t - kz)$$

$$H_y = \frac{E_0}{Z_c} \cos(\omega t - kz)$$

Dans un milieu sans pertes (propagation dans l'espace libre), le vecteur de Poynting est réel et dirigé selon $\mathbf{k}$, et la densité de puissance transmise est égale à :

$$P = \frac{E_0^2}{Z_c} = Z_c H_0^2$$

6.7.3 Onde sphérique

Dans une onde sphérique, les plans équiphase sont des sphères. Ceci veut dire que, si on place l'origine d'un système de coordonnées sphériques à la source de l'onde (centre des sphères équiphase), la propagation de l'onde est radiale. Les champs électrique et magnétique et le vecteur de propagation étant perpendiculaires entre eux selon un système direct ($\mathbf{E} \times \mathbf{H}$ est dans le même sens que $\mathbf{k}$), on écrit donc par exemple :

$$\mathbf{k} = k\hat{\mathbf{r}}$$

$$\mathbf{E} = E_{\theta}\hat{\boldsymbol{\theta}}$$

$$\mathbf{H} = H_{\varphi}\hat{\boldsymbol{\varphi}}$$

$$E_{\theta} = Z_c H_{\varphi} = \sqrt{\frac{\mu}{\varepsilon}} H_{\varphi}$$

Une onde sphérique est de la forme

$$E_{\theta} = E_0 \frac{e^{-jkr}}{r}$$

Pour une onde s'éloignant de l'origine.

6.7.3 Polarisation de l'onde

L'orientation du champ électrique est appelée polarisation de l'onde électromagnétique. Elle peut être de trois types : linéaire, circulaire ou elliptique.

Polarisation linéaire

La direction du champ électrique reste constante en fonction du temps à un point fixé de l'espace. Pour une onde se propageant près de la surface terrestre, on utilise fréquemment les termes de polarisation verticale ou horizontale pour un champ électrique vertical ou horizontal.

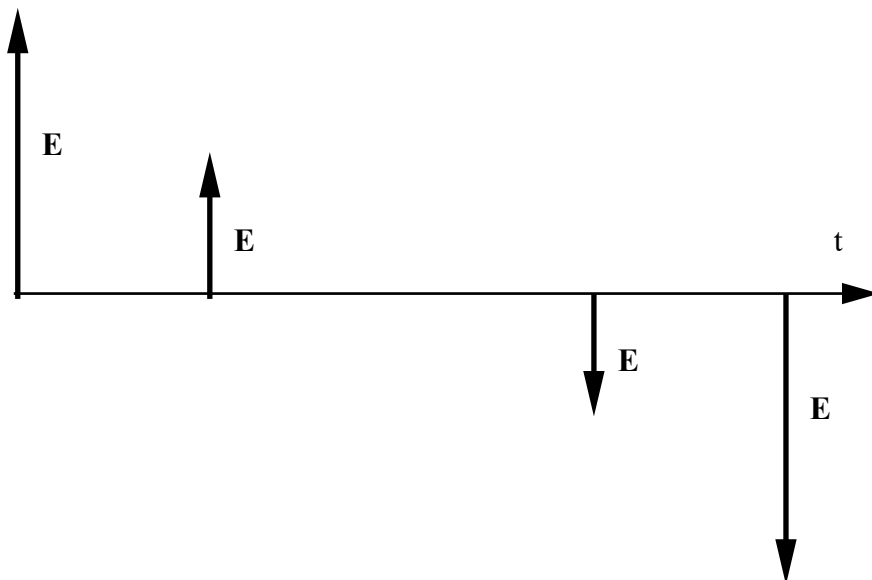


Figure 2.17 : polarisation linéaire

Polarisation circulaire

On dit que la polarisation d'une onde est circulaire lorsqu'à un point fixe de l'espace, l'extrémité du phaseur représentant le champ électrique décrit un cercle. On parle de polarisation circulaire droite lorsque le phaseur tourne dans le sens des aiguilles d'une montre pour une onde s'éloignant de l'observateur, d'une polarisation circulaire gauche dans le cas inverse.

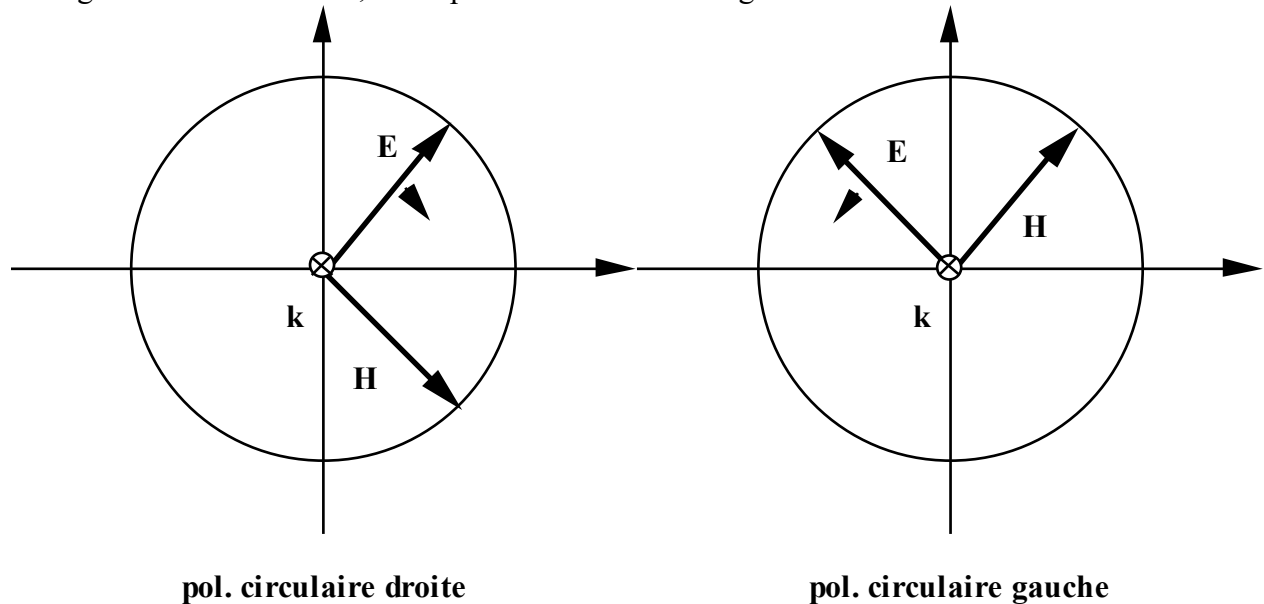


Figure 6.18 : Polarisation circulaire

Polarisation elliptique

La polarisation est dite elliptique lorsqu'à un point fixe de l'espace, l'extrémité du phaseur représentant le champ électrique décrit une ellipse. C'est le cas le plus général.

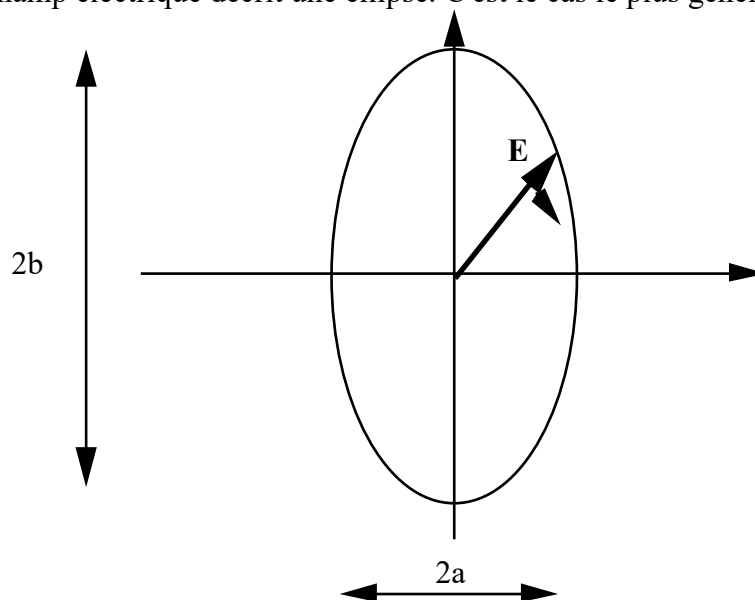


Figure 6.19 : Polarisation elliptique

6.7.4 Caractéristique de polarisation d'un champ

Un champ électrique à dépendance sinusoïdale du temps est défini par

$$\mathbf{E}(t) = \sqrt{2} \left[\hat{\mathbf{x}} E_{0x} \cos(\omega t + \varphi_x) + \hat{\mathbf{y}} E_{0y} \cos(\omega t + \varphi_y) + \hat{\mathbf{z}} E_{0z} \cos(\omega t + \varphi_z) \right]$$

ce qui peut s'écrire

$$\mathbf{E}(t) = \mathbf{E}(0) \cos(\omega t) + \mathbf{E}(T/4) \sin(\omega t)$$

où

$$\begin{aligned} \mathbf{E}(0) &= \sqrt{2} \left[\hat{\mathbf{x}} E_{0x} \cos(\varphi_x) + \hat{\mathbf{y}} E_{0y} \cos(\varphi_y) + \hat{\mathbf{z}} E_{0z} \cos(\varphi_z) \right] \\ \mathbf{E}(T/4) &= -\sqrt{2} \left[\hat{\mathbf{x}} E_{0x} \sin(\varphi_x) + \hat{\mathbf{y}} E_{0y} \sin(\varphi_y) + \hat{\mathbf{z}} E_{0z} \sin(\varphi_z) \right] \end{aligned}$$

ou encore, en terme de vecteur phaseur

$$\begin{aligned} \mathbf{E}(0) &= \text{Re}[\sqrt{2} \mathbf{E}] \\ \mathbf{E}(T/4) &= -\text{Im}[\sqrt{2} \mathbf{E}] \end{aligned}$$

Les vecteurs $\mathbf{E}(0)$ et $\mathbf{E}(T/4)$ sont deux vecteurs conjugués de l'ellipse de polarisation. Dans le cas d'une polarisation linéaire, ils sont équipollents, ce qui revient à écrire

$$\begin{aligned} \mathbf{E}(0) \times \mathbf{E}(T/4) &= 0 \\ \langle E^2 \rangle &\neq 0 \end{aligned}$$

Dans le cas d'une polarisation circulaire, les demi-axes de l'ellipse sont orthogonaux et ont la même amplitude. On peut donc écrire

$$\begin{aligned} \mathbf{E}(0) \cdot \mathbf{E}(T/4) &= 0 \\ |\mathbf{E}(0)| &= |\mathbf{E}(T/4)| \neq 0 \end{aligned}$$

Ce qui en termes de vecteur phaseur s'écrit

$$\mathbf{E} \cdot \mathbf{E} = 0$$

7. Antennes

7.1 Introduction

"Une antenne est un transducteur servant à transformer une énergie électromagnétique guidée en énergie électromagnétique rayonnée, et réciproquement".

Autrement dit, une antenne accepte une puissance électrique fournie par un générateur sous forme courant/tension et l'émet dans l'espace environnant sous forme d'onde électromagnétique (émission). Mais elle peut aussi recevoir une onde électromagnétique de l'espace environnant et la transformer en puissance fournie au récepteur (réception). Cet aspect dual du fonctionnement des antennes est de la plus grande importance : On peut en effet démontrer qu'une antenne est **réciproque**, c'est à dire que ses caractéristiques fondamentales sont **identiques pour l'émission et la réception**. Les aspects transducteurs d'énergie des antennes sont caractérisées par certaines grandeurs, comme l'impédance d'entrée, la résistance de rayonnement ou le rendement.

Une antenne peut aussi être considérée comme un filtre spatial : à l'émission, elle distribue une puissance dans l'espace en privilégiant certaines directions par rapport à d'autres. A la réception, elle est beaucoup plus sensible aux ondes provenant de ces mêmes directions privilégiées. Le côté filtre des antennes est caractérisé par leur diagramme de rayonnement.

7.2 Exemples d'antennes

Différents types d'antennes sont illustrés aux figures 7.1 - 7.4. On y reconnaît des antennes très courantes, de type dipôle ou cornet, mais aussi des radiateurs plus exotiques, comme le patch microruban.

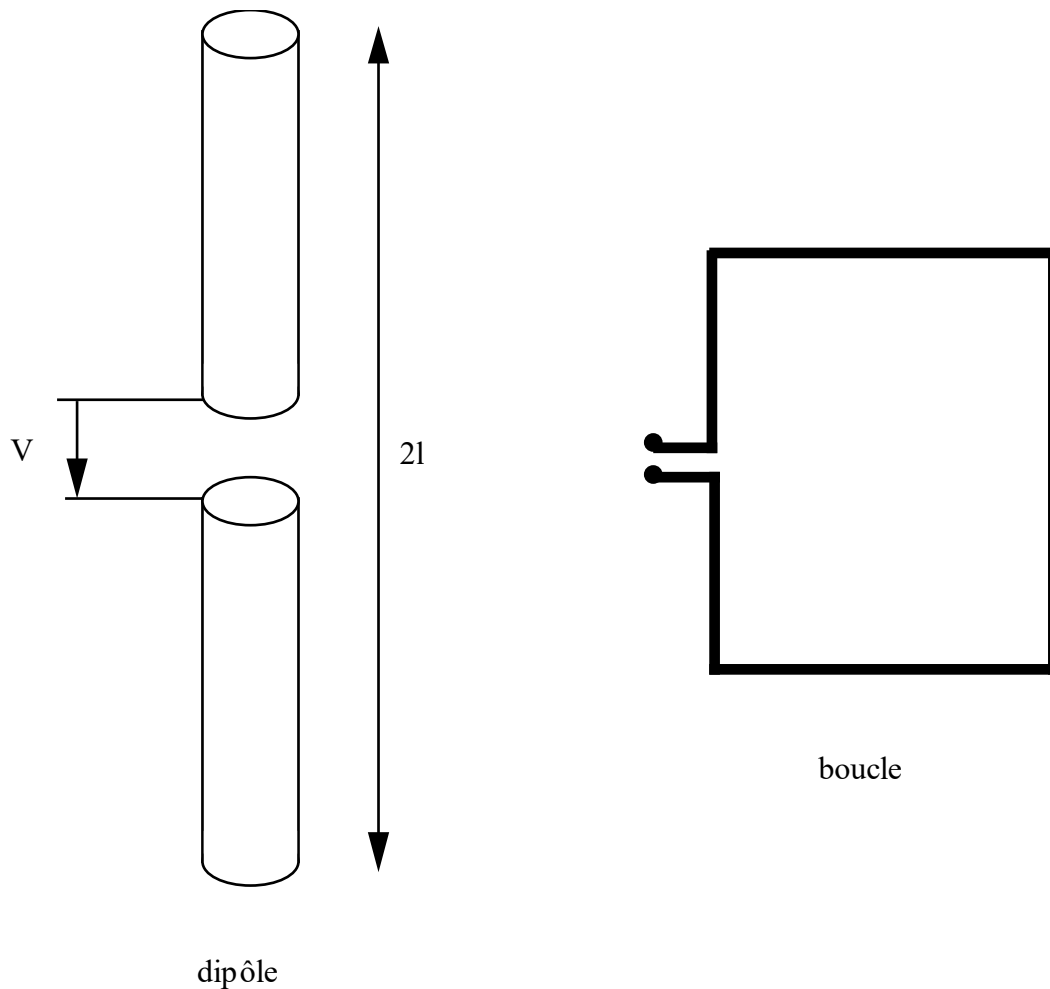
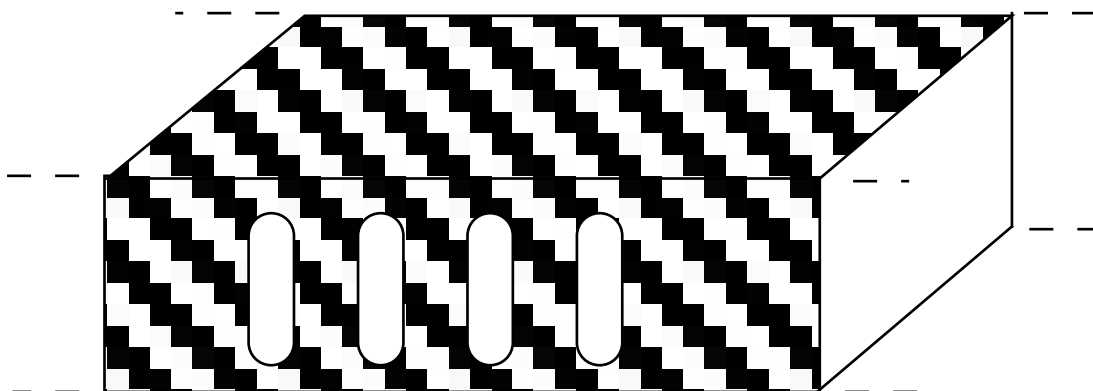
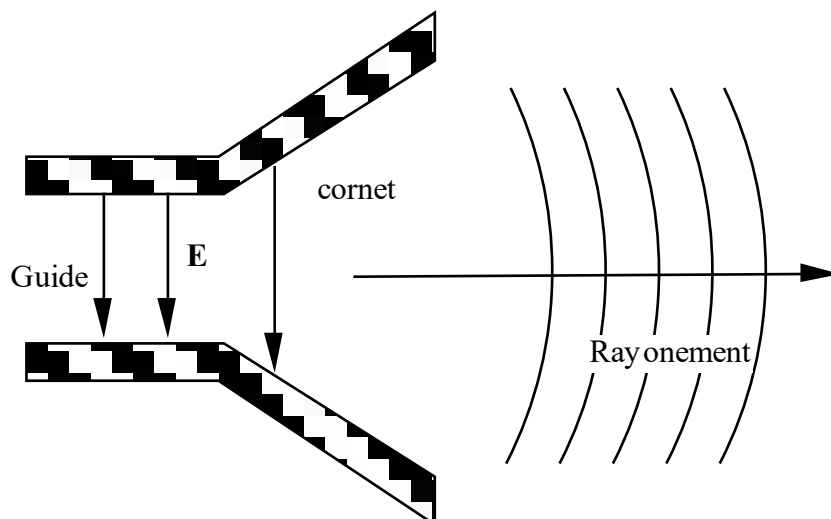


Fig. 7.1 : exemples d'antenne filaires



Fentes dans un guide d'onde

Fig. 7.2 : exemples d'antennes à ouverture

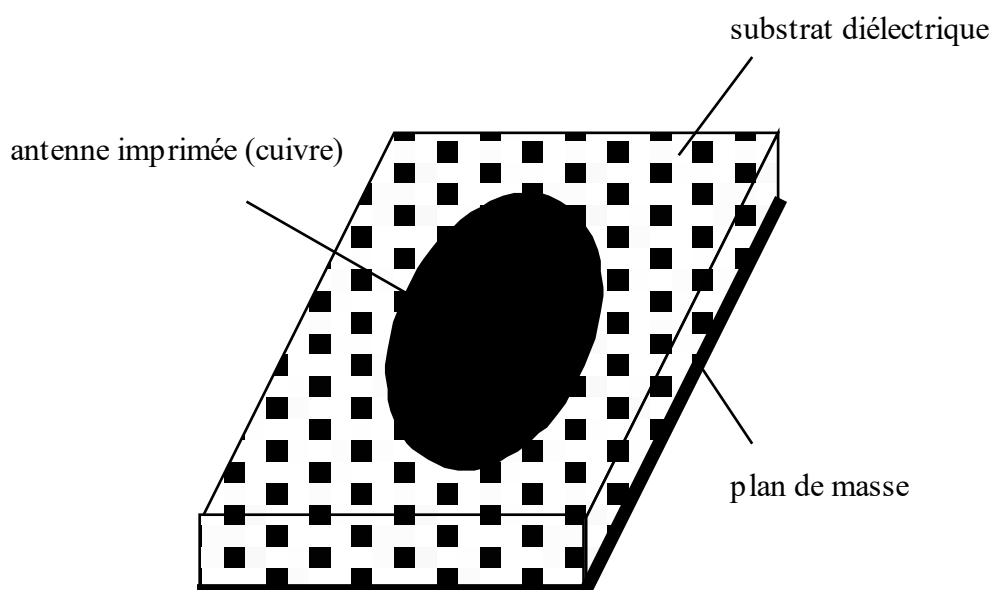


Fig. 7.3 : exemple d'antennes imprimées : l'antenne microruban

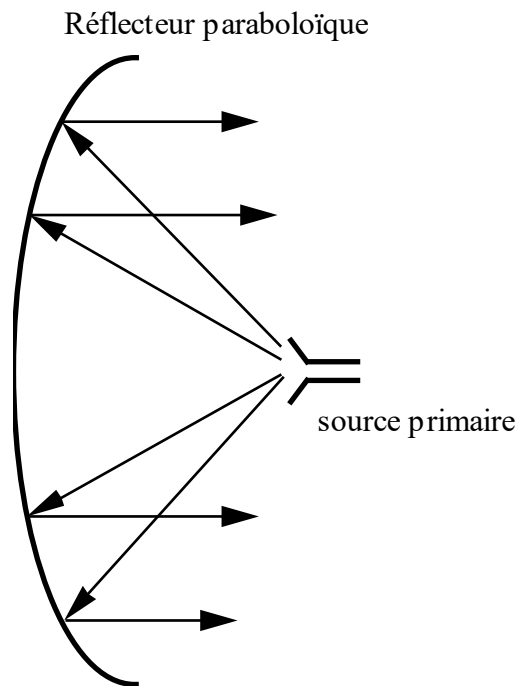


Fig. 7.4 : Antenne à réflecteur

Toutes ces configurations vont avoir des comportements spécifiques, dépendant de leur géométrie et des matériaux utilisés. La détermination des caractéristiques d'une antenne particulière peut, selon les cas, être relativement complexe (recours aux équations de Maxwell) et sort du contexte de ce cours. Cette caractérisation est en général faite par le fabricant de l'antenne, qui fournit à l'ingénieur utilisateur les "grandeurs système" lui permettant d'intégrer l'élément rayonnant dans son design. Ces grandeurs sont l'impédance d'entrée et la résistance de rayonnement de l'antenne, son diagramme de rayonnement, son gain et son rendement. Elles seront définies dans les paragraphes qui suivent, avec des exemples pour certaines classes d'antennes.

7.3 L'antenne élémentaire : le dipôle de Hertz

Considérons l'élément de courant infinitésimal d'intensité I et de longueur dl orienté selon $\hat{z}$ représenté à la figure 3.5. Il s'agit de l'antenne élémentaire, à partir du comportement de laquelle on peut déduire les caractéristiques de toutes les antennes en faisant usage du théorème de superposition. Cet élément est appelé le dipôle de Hertz. Le terme dipôle vient du fait que l'équation de continuité entre le courant et la charge impose deux charge égales mais de signe opposé à chaque extrémité de l'élément.

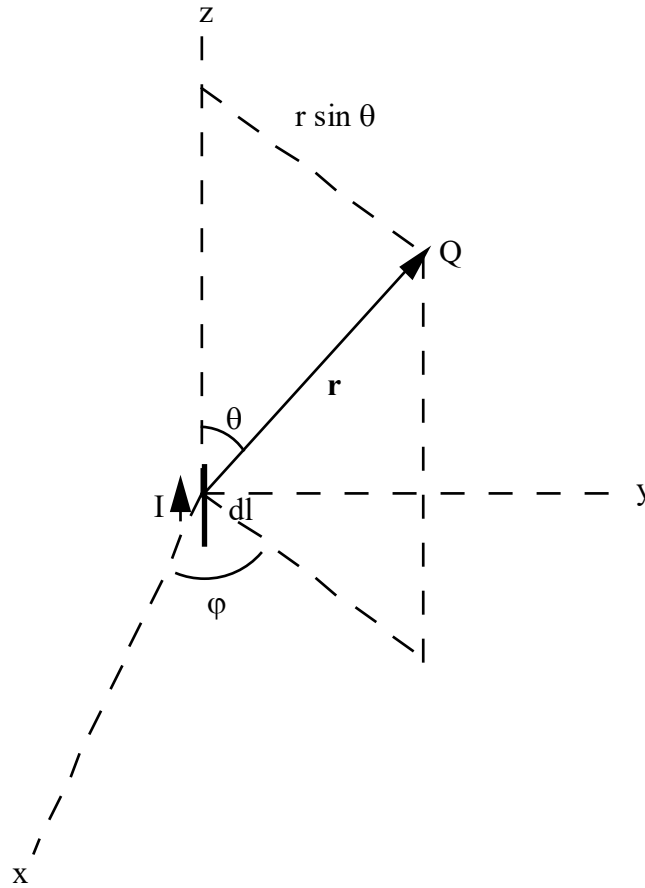


Fig. 7.5 : le dipôle de Hertz

Le champ électromagnétique rayonné par cet élément est donné par

$$H_{\varphi} = \frac{Idl}{4\pi} e^{-jkr} \left(\frac{jk}{r} + \frac{1}{r^2} \right) \sin \theta$$

$$E_r = \frac{Idl}{4\pi} e^{-jkr} \left(\frac{2Z_c}{r^2} + \frac{2}{j\omega\epsilon r^3} \right) \cos \theta$$

$$E_{\theta} = \frac{Idl}{4\pi} e^{-jkr} \left(\frac{j\omega\mu}{r} + \frac{Z_c}{r^2} + \frac{1}{j\omega\epsilon r^3} \right) \sin \theta$$

Dans ces relations, les termes en $1/r^2$ et $1/r^3$ vont décroître beaucoup plus rapidement que les termes en $1/r$ à mesure que l'on s'éloigne. Suffisamment loin de l'antenne, le champ électromagnétique se propagera radialement et aura les composantes

$$H_{\varphi} \approx jk \frac{Idl}{4\pi r} e^{-jkr} \sin \theta$$

$$E_r \approx 0$$

$$E_{\theta} \approx \frac{j\omega\mu Idl}{4\pi r} e^{-jkr} \sin \theta$$

Il s'agit d'une onde sphérique.

7.3 Champ lointain d'une antenne

Une antenne est constituée de différentes parties métalliques supportant un courant de surface $\mathbf{J}_s$, source du rayonnement électromagnétique. Ce courant de surface peut être décomposé en une infinité de sources de courant élémentaires, produisant chacune un rayonnement élémentaire. En appliquant le théorème de superposition pour trouver le rayonnement total de l'antenne et en utilisant le résultat du paragraphe précédent, on déduit que le champ électromagnétique émis par une antenne est, dès que l'on se place "suffisamment loin" de cette dernière, une onde sphérique de forme

$$\mathbf{E}(r, \theta, \varphi) = E_\theta(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\theta} + E_\varphi(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\varphi}$$

$$\mathbf{H}(r, \theta, \varphi) = H_\theta(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\theta} + H_\varphi(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\varphi}$$

avec

$$\mathbf{E} = -Z_c \hat{\mathbf{r}} \times \mathbf{H}$$

$$\mathbf{H} = \frac{1}{Z_c} \hat{\mathbf{r}} \times \mathbf{E}$$

Le domaine où cette approximation est valable est appelé le **champ lointain** de l'antenne.

Cette condition est en général remplie à une distance L

$$L \geq \frac{2d^2}{\lambda}$$

où d est la plus grande dimension de l'antenne. Cette relation est empirique, et n'est valable que pour une antenne dont les dimensions ne sont pas trop petites par rapport à la longueur d'onde.

Dans les paragraphes suivants, nous supposons toujours être dans le champ lointain de l'antenne.

7.4 Diagramme de rayonnement

Le champ lointain d'une antenne est souvent décrit de la manière suivante, où les dépendances angulaires des champs sont mises en évidence à l'aide d'une fonction $\mathbf{f}(\theta, \varphi)$, qui est elle aussi un phaseur :

$$\begin{aligned}\mathbf{E}(r, \theta, \varphi) &= -\frac{jZ_c}{2\lambda} f_\theta(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\theta} - \frac{jZ_c}{2\lambda} f_\varphi(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\varphi} \\ \mathbf{H}(r, \theta, \varphi) &= \frac{j}{2\lambda} f_\varphi(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\theta} - \frac{j}{2\lambda} f_\theta(\theta, \varphi) \frac{e^{-jkr}}{r} \hat{\varphi}\end{aligned}$$

Le vecteur de Poynting est purement radial et réel :

$$\mathbf{S} = \mathbf{E} \times \mathbf{H}^* = \frac{1}{Z_c} |\mathbf{E}|^2 \hat{\mathbf{r}}$$

La densité de puissance rayonnée vers l'extérieur est donc égale à :

$$\begin{aligned}p(r, \theta, \varphi) &= \mathbf{S} \cdot \hat{\mathbf{r}} = \frac{1}{Z_c} |\mathbf{E}|^2 = \frac{1}{Z_c} \left(|E_\theta(\theta, \varphi)|^2 + |E_\varphi(\theta, \varphi)|^2 \right) \\ &= \frac{Z_c}{4\lambda^2 r^2} \left(|f_\theta(\theta, \varphi)|^2 + |f_\varphi(\theta, \varphi)|^2 \right) \quad [W / m^2]\end{aligned}$$

Nous retrouvons ainsi le résultat de la formule de Friis : La densité de puissance rayonnée par une antenne décroît avec le carré de la distance. On définit alors l'intensité de rayonnement U comme :

$$U(\theta, \varphi) = r^2 p(r, \theta, \varphi) = \frac{Z_c}{4\lambda^2} \left(|f_\theta(\theta, \varphi)|^2 + |f_\varphi(\theta, \varphi)|^2 \right) \quad [W / \text{stéradian}]$$

7.4.1 Diagrammes de champ

Un diagramme de rayonnement de champ est une représentation graphique des fonctions $f_\theta(\theta, \varphi)$ (associée à E_θ et H_φ) et $f_\varphi(\theta, \varphi)$ (associée à E_φ et H_θ). Il ne faut pas oublier que l'on travaille avec des phaseurs. Les composantes du champ sont donc des valeurs complexes, de même que $f_\theta(\theta, \varphi)$ et $f_\varphi(\theta, \varphi)$. Dans la plupart des problèmes d'antennes, on s'intéresse surtout à l'amplitude du champ rayonné. La seule information à retenir est alors contenue dans les normes $|f_\theta(\theta, \varphi)|$ et $|f_\varphi(\theta, \varphi)|$.

Le diagramme de champ normalisé est donné par

$$D_{E\theta} = \frac{|E_\theta(\theta, \varphi)|}{|E_\theta(\theta_{\max}, \varphi_{\max})|} = \frac{|f_\theta(\theta, \varphi)|}{|f_\theta(\theta_{\max}, \varphi_{\max})|}$$

pour la composante θ et par

$$D_{E\varphi} = \frac{|E_\varphi(\theta, \varphi)|}{|E_\varphi(\theta_{\max}, \varphi_{\max})|} = \frac{|f_\varphi(\theta, \varphi)|}{|f_\varphi(\theta_{\max}, \varphi_{\max})|}$$

pour la composante φ .

La direction $(\theta_{\max}, \varphi_{\max})$ correspond à la valeur maximum de $|E_\theta(\theta, \varphi)|$ pour $D_{E\theta}$ respectivement de $|E_\varphi(\theta, \varphi)|$ pour $D_{E\varphi}$. Attention, ces directions ne sont pas les mêmes.

Les diagrammes de champ représentent aussi bien les champs électriques que les champs magnétiques, les composantes H_θ et H_φ ayant respectivement les mêmes dépendances angulaires que les composantes E_φ et E_θ .

Les diagrammes ainsi définis ont toujours des valeurs comprises entre 0 et 1. Des valeurs en décibels sont souvent utilisées :

$$D_{E\theta, \varphi} \text{ (dB)} = 20 \log_{10} |D_{E\theta, \varphi}|$$

7.4.2 Diagrammes de puissance

Le diagramme de rayonnement de puissance (normalized power pattern) est une représentation graphique de la densité de puissance $p(r, \theta, \varphi)$ ou de l'intensité $U(\theta, \varphi)$ en fonction des angles θ, φ .

$$D_p(\theta, \varphi) = \frac{p(r, \theta, \varphi)}{p(r, \theta_{\max}, \varphi_{\max})} = \frac{U(\theta, \varphi)}{U(\theta_{\max}, \varphi_{\max})}$$

Ce diagramme est beaucoup plus utilisé que les diagrammes de champs, car il est d'une part plus facilement mesurable (la puissance est une grandeur plus aisée et moins coûteuse à mesurer que le champ), et d'autre part la densité de puissance est un scalaire réel. Il n'y a donc pas de composantes à distinguer et de valeur absolue à prendre. La normalisation est unique et se fait

selon la direction où la densité de puissance est maximale, donnée par les angles $\theta_{\max}, \varphi_{\max}$. Pour les diagrammes de rayonnement de puissance on utilise souvent des valeurs en décibels, données par :

$$D_p(\theta, \varphi) \text{ (dB)} = 10 \log_{10}(D_p(\theta, \varphi))$$

7.4.3 Paramètres caractéristiques d'un diagramme de rayonnement de puissance

Le diagramme de rayonnement dépend des deux angles θ et φ . Il s'agit donc d'un diagramme à trois dimensions. On représente en général des coupes à deux dimensions de ce diagramme, en posant soit $\theta = \text{cte}$, ou , plus fréquemment $\varphi = \text{cte}$.

Un diagramme de rayonnement comporte en général plusieurs lobes, séparés par des zéros. La figure 7.6 représente une coupe $\varphi = \text{cte}$ typique d'un diagramme de rayonnement en représentation linéaire. La figure 7.7 représente la même coupe en échelle logarithmique.

On appelle lobe principal (main lobe) le lobe contenant la direction principale de rayonnement, les autres étant des lobes secondaires (minor lobe). Un lobe latéral (side lobe) est un lobe dans une direction autre que celle souhaitée pour le rayonnement de l'antenne. Dans la plupart des cas, l'antenne est conçue pour exploiter son rayonnement dans la direction du lobe principal, et tous les lobes secondaires sont latéraux et vice versa.

La largeur du lobe principal ou largeur du faisceau (Beam width, θ_{BW}) est l'angle entre les deux zéros entourant ce lobe. Il n'est pas toujours facile à déterminer, et on utilise aussi la largeur de faisceau à mi-puissance (Half Power Beam Width, θ_{HPBW}), qui est l'angle formé entre les deux directions où la densité de puissance est la moitié de la valeur maximale.

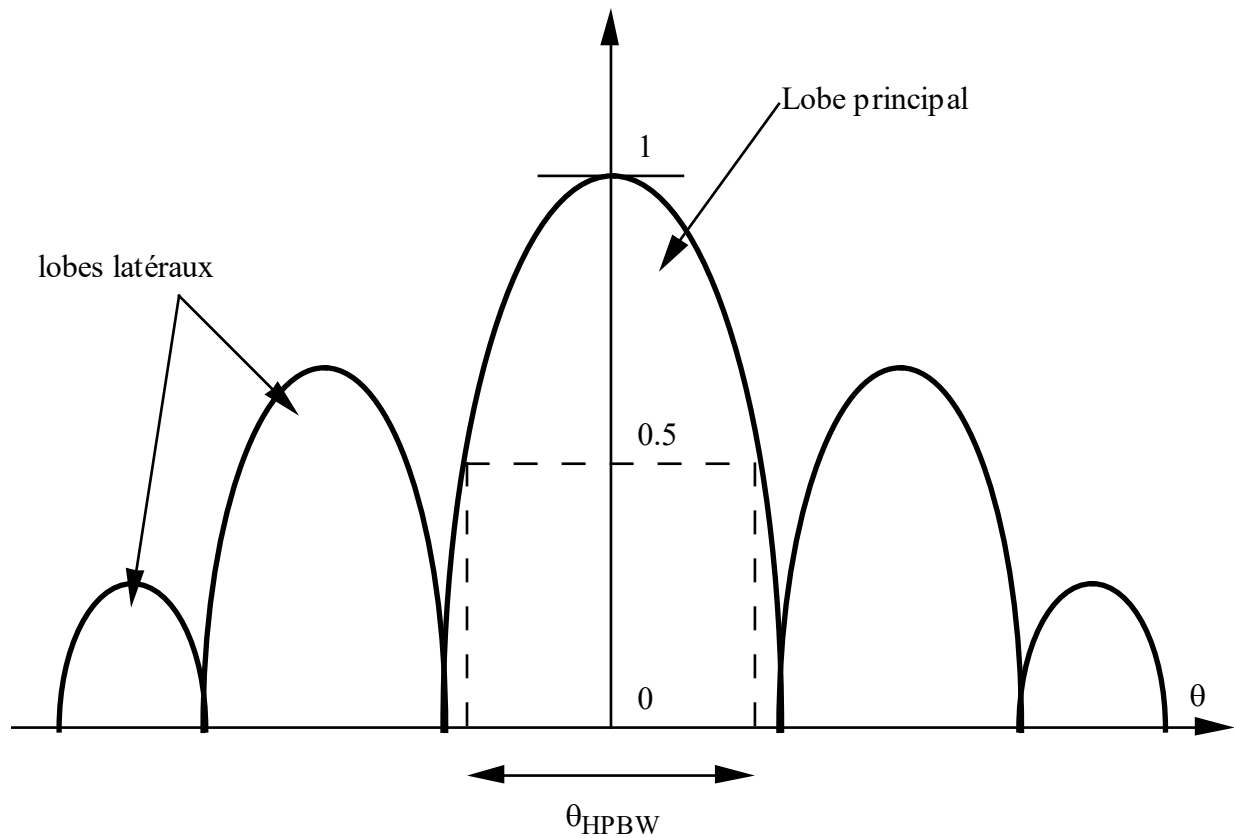


Figure 7.6 : Diagramme de rayonnement linéaire

Le niveau des lobes latéraux (side lobe level SLL) est défini comme (figure 3.7) :

$$SLL = 10 \log_{10} \left(\frac{P_{\max}(\text{lobe principal})}{P_{\max}(\text{lobes latéraux})} \right)$$

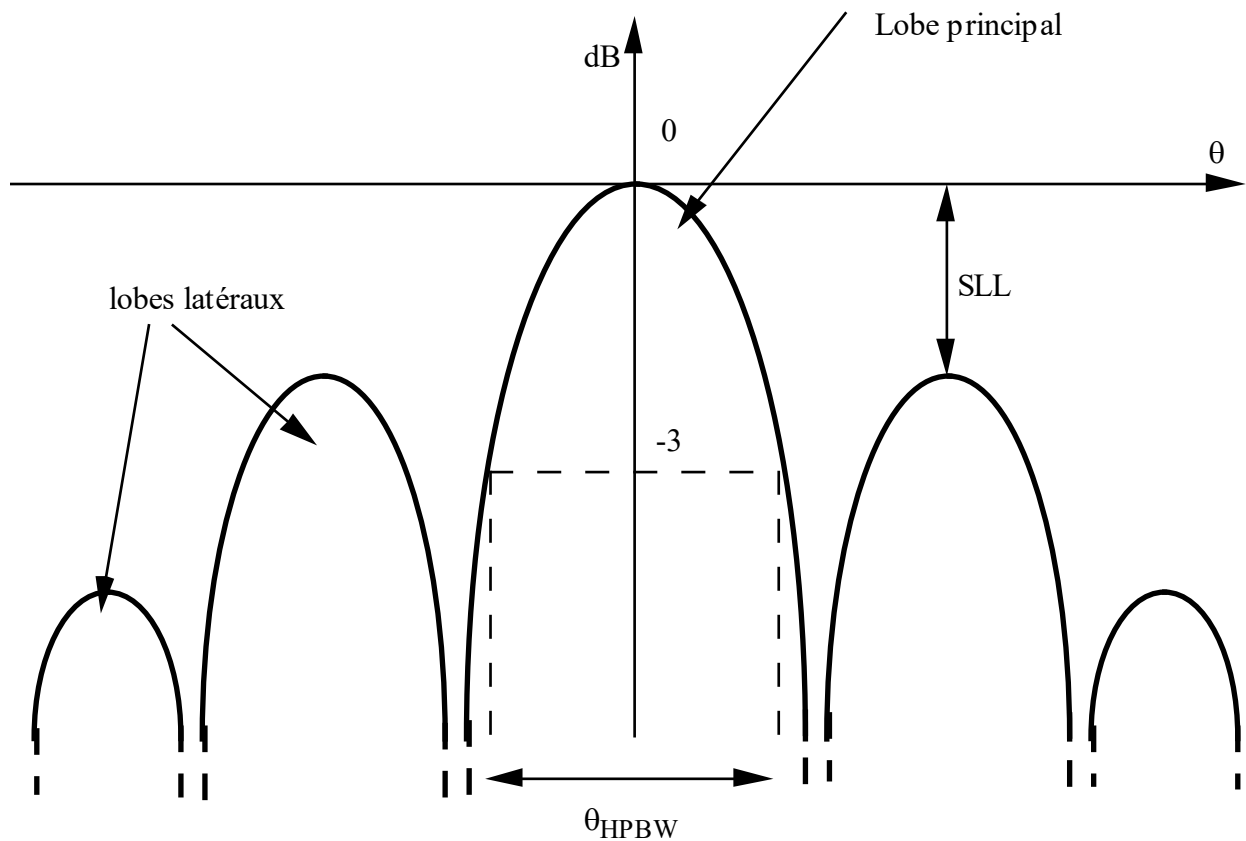


Figure 7.7 : diagramme de rayonnement en dB

7.5 Directivité, gain et rendement d'une antenne

Directivité

La directivité d'une antenne est définie comme l'intensité de rayonnement $U(\theta, \varphi)$ normalisée par l'intensité moyenne rayonnée par l'antenne (figure 3.8) :

$$d(\theta, \varphi) = \frac{U(\theta, \varphi)}{P_{\text{rayonnée}} / 4\pi} = 4\pi \frac{U(\theta, \varphi)}{P_{\text{rayonnée}}}$$

La directivité est en général exprimée en dB :

$$D(\theta, \varphi) = 10 \log_{10}(d(\theta, \varphi))$$

On donne souvent le nom de directivité à la valeur maximale de $D(\theta, \varphi)$, qui est obtenue pour la direction principale de rayonnement $(\theta_{\max}, \varphi_{\max})$.

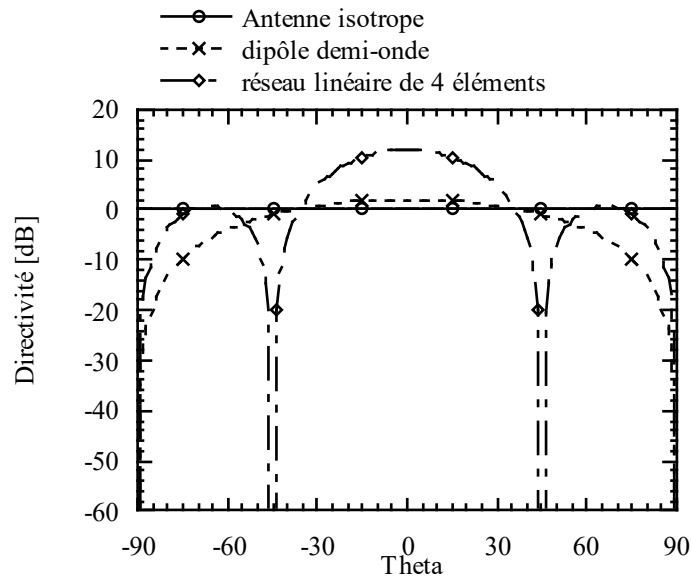


Figure 7.8 : Exemples de directivité d'antennes en dB

Le diagramme de rayonnement de puissance n'est rien d'autre que la directivité normalisée par rapport à la directivité maximale. La figure 3.8 montre des exemples de diagrammes de directivité. On constate que plus le lobe principal est étroit, plus la directivité est élevée, car plus la puissance rayonnée est concentrée dans une portion restreinte de l'espace.

Gain

Le gain d'une antenne est défini comme l'intensité de rayonnement $U(\theta, \varphi)$ normalisée par l'intensité de rayonnement d'une antenne isotrope sans pertes :

$$g(\theta, \varphi) = \frac{U(\theta, \varphi)}{U_{\text{isotrope}}} = \frac{U(\theta, \varphi)}{P_{\text{fournie}} / 4\pi} = 4\pi \frac{U(\theta, \varphi)}{P_{\text{fournie}}}$$

Le gain est en général exprimé en dB :

$$G(\theta, \varphi) = 10 \log_{10}(g(\theta, \varphi))$$

On donne souvent le nom de gain à la valeur maximale de $G(\theta, \varphi)$, qui est obtenue pour la direction principale de rayonnement $(\theta_{\text{max}}, \varphi_{\text{max}})$.

On précise parfois qu'il s'agit de décibels isotropes, dBi, pour indiquer le fait que le rayonnement de l'antenne est comparé au rayonnement d'une antenne isotrope. Il faut noter que cette dernière est une vue de l'esprit, irréalisable physiquement, mais bien utile en temps que concept mathématique.

Le gain d'une antenne est, contrairement à la directivité, une grandeur mesurable. La technique habituelle est de comparer le gain de l'antenne à mesurer au gain d'une antenne de référence (mesure substitutive). Il suffit de mesurer le gain maximum de l'antenne obtenu pour la direction de rayonnement principal. Cette information, couplée au diagramme de rayonnement normalisé, caractérise entièrement l'antenne du point de vue du rayonnement.

L'antenne de référence utilisée pour la mesure du gain est souvent un dipôle demi-onde (longueur totale $\lambda/2$), surtout dans la bande UHF. C'est pourquoi on trouve quelquefois le gain d'une antenne spécifié en dB_D , c'est à dire comparé à la puissance rayonnée dans la direction $(\theta_{\max}, \varphi_{\max})$ par un dipôle.

Puisque le gain d'un dipôle demi-onde est de :

$$g_D = 1.64$$

$$G_D = 2.15 \text{ dB}_i$$

nous avons la correspondance

$$x \text{ dB}_i = x + 2.15 \text{ dB}_D$$

Rendement

Les formules donnant le gain et la directivité sont très similaires : elle ne diffèrent que par l'utilisation de la puissance totale fournie à l'antenne (gain) au lieu de la puissance totale rayonnée par l'antenne (directivité). Le gain prend donc en compte les pertes de l'antenne, la directivité ne le fait pas.

Soit η le rendement de l'antenne, défini par

$$\eta = \frac{P_{\text{rayonnée}}}{P_{\text{fournie}}} \quad \eta \in [0;1]$$

Nous obtenons alors

$$g(\theta, \varphi) = \eta \text{ } d(\theta, \varphi)$$

Le gain et la directivité sont identiques pour une antenne sans pertes.

7.6 Surface effective

Considérons une antenne idéale et sans pertes à la réception, illuminée par une onde plane de densité de puissance p (figure 3.9). Si cette antenne était capable de capter toute la puissance "passant à sa portée", la puissance qu'elle transmettrait à la charge adaptée la terminant serait égale à :

$$P_r = p A = \frac{E_0^2}{Z_c} A$$

où A est la surface de l'antenne.

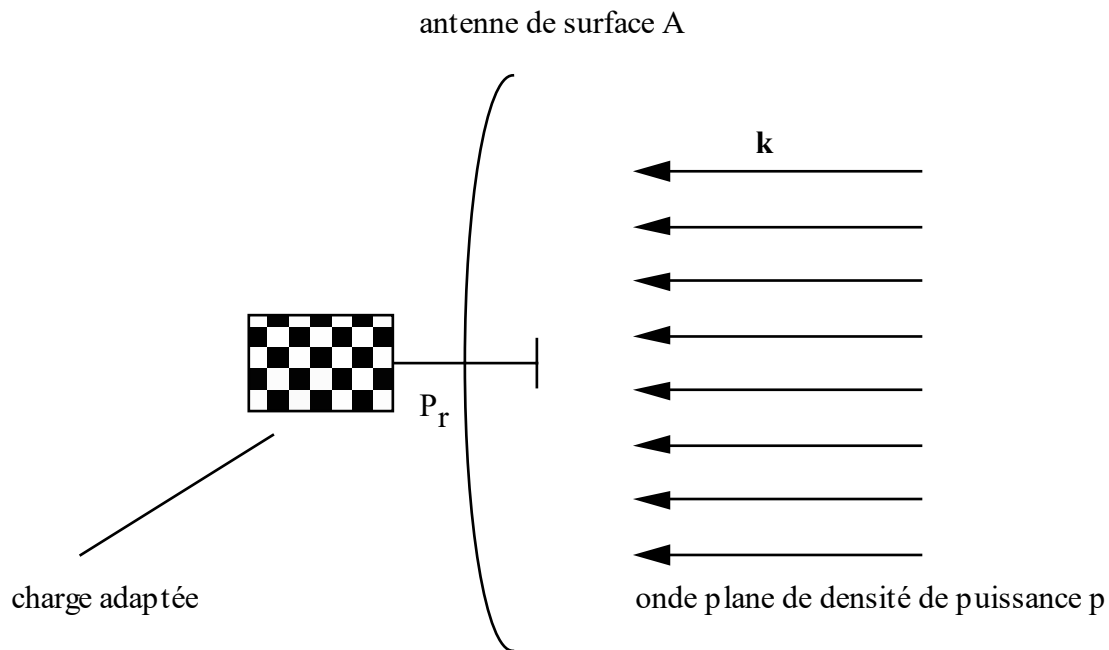


Figure 7.9 : puissance maximale captée par une antenne.

Cette puissance est un maximum théorique de la puissance qu'une antenne peut capter. Elle est diminuée par deux facteurs :

- un rendement de surface de l'antenne, qui fait que toute la surface du radiateur ne participe de manière identique à la réception du signal ;
- le fait que l'antenne est sensible à la direction d'où provient le signal, comme le montre son diagramme de rayonnement.

On définit donc comme **la surface effective** (ou surface de captation) $A_e(\theta, \varphi)$ d'une antenne la grandeur, dépendant de la direction (θ, φ) de l'onde incidente, qui multipliée par la densité de puissance de l'onde donne la puissance transmise à une charge adaptée par l'antenne. Cette grandeur a effectivement la dimension d'une surface.

$$P_r(\theta, \varphi) = A_e(\theta, \varphi)p$$

Cette surface effective n'est pas indépendante du gain de l'antenne : Si on considère un système de transmission formé de deux antennes A et B (figure 7.10), séparées par une distance L, le rapport entre puissances transmises et reçues est donné par

$$\left. \frac{P_r}{P_{fourni}} \right|_{A \rightarrow B} = \frac{1}{4\pi L^2} g_A(\theta_A, \varphi_A) A_{eB}(\theta_B, \varphi_B)$$

$$\left. \frac{P_r}{P_{fourni}} \right|_{B \rightarrow A} = \frac{1}{4\pi L^2} g_B(\theta_B, \varphi_B) A_{eA}(\theta_A, \varphi_A)$$

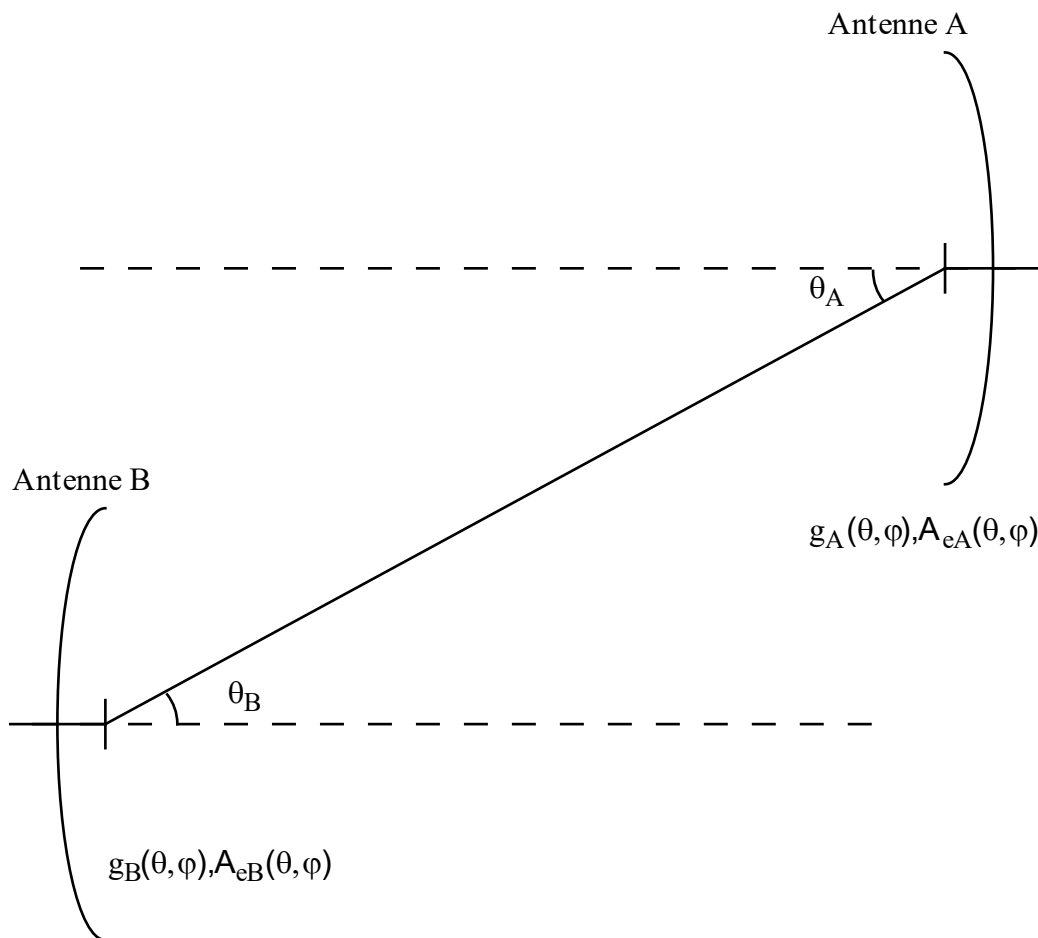


Figure 7.10 : transmission entre deux antennes

Or, à cause du théorème de réciprocité, ces deux rapports de puissances doivent être égaux. Nous obtenons donc

$$\frac{g_A(\theta_A, \varphi_A)}{A_{eA}(\theta_A, \varphi_A)} = \frac{g_B(\theta_B, \varphi_B)}{A_{eB}(\theta_B, \varphi_B)}$$

La valeur de ce rapport peut être déterminée par l'étude d'un système formé par deux antennes simples et bien connues, par exemple un dipôle et un radiateur isotrope. On obtient finalement :

$$g(\theta, \varphi) = \frac{4\pi}{\lambda^2} A_e(\theta, \varphi)$$

Commentaire

Le gain d'une antenne est directement proportionnel à sa surface effective. Une antenne de grandes dimensions aura donc un gain plus élevé qu'une antenne de petites dimensions, et son rayonnement sera plus directif.

7.7 Résistance de rayonnement

Du point de vue de la théorie des circuits, une antenne à l'émission correspond à un élément passif linéaire dissipatif, possédant une impédance d'entrée complexe Z_{in} , qui est une fonction de la fréquence (figure 3.11). Si l'antenne est excitée par un courant de valeur efficace I , la puissance fournie à l'antenne est

$$P_f = I^2 \operatorname{Re}[Z_{in}]$$

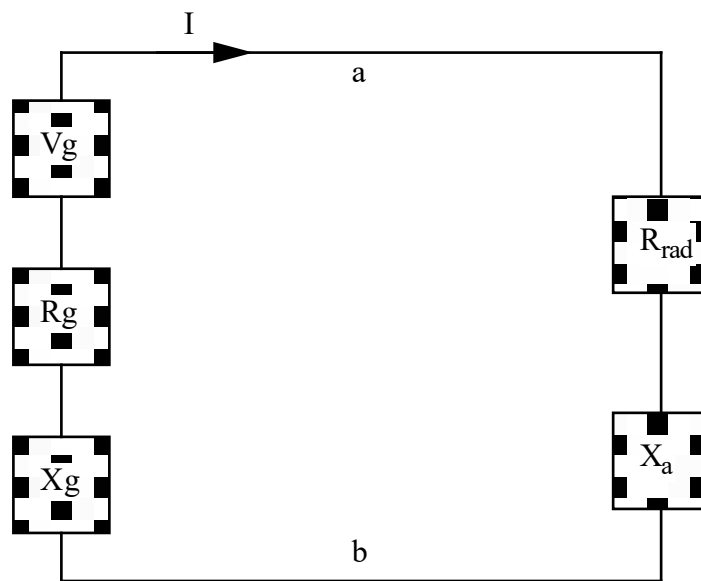
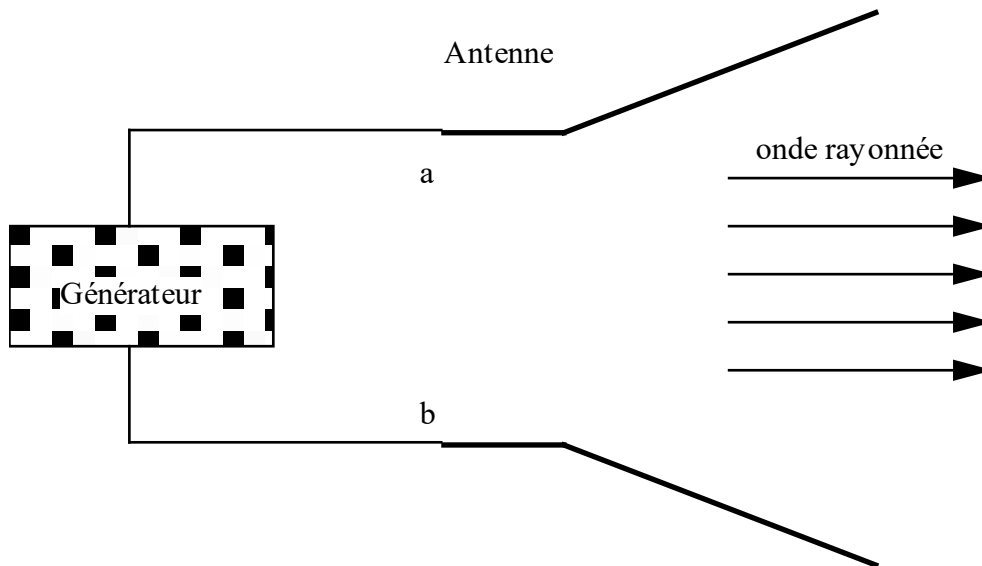


Figure 7.11 : circuit équivalent de Thévenin d'une antenne

Si on admet que l'antenne est faite avec des matériaux sans pertes, la conservation de l'énergie implique que cette puissance doit être égale à la puissance électromagnétique rayonnée P_{rad} .

On peut donc obtenir la partie réelle de l'impédance d'entrée d'une antenne idéale, dite résistance de rayonnement R_{rad} , comme

$$R_{rad} = \text{Re}[Z_{in}] = \frac{P_{rad}}{I^2}$$

Le calcul de la partie imaginaire de l'impédance d'entrée est beaucoup plus délicat. La réactance n'est pas liée au rayonnement, mais plutôt au champ proche, et est très sensible à la géométrie. On se contentera souvent d'une estimation empirique approchée.

7.8 Désadaptation

La puissance fournie à l'antenne P_f n'est souvent pas la puissance maximum disponible dans le générateur, P_g . Considérons le générateur de tension V_g de la figure 3.11, et supposons que son impédance soit réelle R_g . Ce générateur fournit sa puissance maximale pour l'adaptation conjuguée, c'est à dire dans ce cas lorsque X_a est nul et $R_{rad} = R_g$. Cette puissance vaut alors :

$$P_g = \frac{1}{2} V_g I_g^* = \frac{|V_g|^2}{4R_g}$$

Si l'impédance de l'antenne est différente (pas d'adaptation conjuguée), la puissance fournie à l'antenne P_f est inférieure à P_g , et vaut

$$P_f = \left(1 - |\Gamma_g|^2\right) P_g \quad \text{où} \quad \Gamma_g = \frac{Z_{in} - R_g}{Z_{in} + R_g}$$

Γ_g étant le coefficient de réflexion entre générateur et antenne.

7.8.1 Antenne réelle à l'émission

En combinant les effets de la désadaptation et des pertes ohmiques, on trouve que la puissance rayonnée P_{rad} par une antenne émettrice vaut, pour une antenne non idéale :

$$P_{rad} = \eta_e P_f = \eta_e \left(1 - |\Gamma_g|^2\right) P_g$$

Dans le cas idéal, $P_{rad} = P_g$.

7.8.2 Antenne réelle à la réception

Les mêmes considérations sont valables pour une antenne à la réception : La puissance P_{load} délivrée à la charge R_L n'est pas égale à la puissance disponible à la réception P_r : En général on a, en tenant compte des pertes et d'une désadaptation éventuelle

$$P_{load} = \eta_r (1 - |\Gamma_L|^2) P_r \quad \text{avec} \quad \Gamma_L = \frac{Z_{in} - R_L}{Z_{in} + R_L}$$

7.9 Polarisation

Emission

La polarisation d'une antenne est la polarisation du champ électromagnétique émis par cette dernière, lorsqu'elle est excitée. Une antenne peut donc avoir une polarisation linéaire, circulaire ou elliptique suivant que le champ qu'elle émet est linéaire, circulaire ou elliptique.

Réception

La polarisation d'une antenne est la polarisation de l'onde électromagnétique incidente qui délivre le maximum de puissance à la charge.

Par le théorème de réciprocité, on démontre que la polarisation d'une antenne est la même à l'émission ou à la réception.

7.9.1 Etat de polarisation d'une antenne et facteur de dépolarisation

La polarisation d'une antenne peut être définie par un vecteur phaseur $\mathbf{e}$ dit état de polarisation et définit comme le champ électrique normalisé émis ou reçu par l'antenne

$$\mathbf{e} = \frac{\mathbf{E}}{|\mathbf{E}|} = \frac{\mathbf{E}}{\sqrt{\mathbf{E} \cdot \mathbf{E}^*}}$$

Considérons maintenant deux antennes ayant respectivement les états de polarisation $\mathbf{e}_1$ et $\mathbf{e}_2$. On construit un système de transmission à l'aide de ces deux antennes en définissant l'axe de transmission z comme allant de l'antenne 1 vers l'antenne 2. Cette dernière pointe donc vers les z négatifs, et son état de polarisation devint donc $\mathbf{e}_2^*$ dans ce repère de coordonnées.

Du champ $\mathbf{E}_1$ émis par l'antenne 1, l'antenne 2 ne peut recevoir que la projection sur $\mathbf{e}_2^*$, soit

$$\mathbf{E}_1 \cdot \mathbf{e}_2^*$$

Les pertes de puissance par rapport à la situation optimale sont alors données par le Facteur De Dépolarisation (FDP) :

$$FDP = \frac{|\mathbf{E}_1 \cdot \mathbf{e}_2^*|^2}{|\mathbf{E}_1|^2} = |\mathbf{e}_1 \cdot \mathbf{e}_2^*|^2$$

En plus des facteurs ohmiques et de désadaptation, la **dépolarisation** peut réduire la puissance reçue par une antenne.

7.9.2 Formule de Friis pour les antennes réelles

La formule de Friis pour les antennes réelles devient donc

$$\begin{aligned} \frac{P_{Load}}{P_{fournie}} &= FDP \left(1 - |\Gamma_L|^2\right) \left(1 - |\Gamma_g|^2\right) \eta_r \eta_e d_e d_r \left(\frac{\lambda^2}{4\pi R}\right) \\ &= FDP \left(1 - |\Gamma_L|^2\right) \left(1 - |\Gamma_g|^2\right) g_e g_r \left(\frac{\lambda^2}{4\pi R}\right) \end{aligned}$$

7.10 Température de bruit d'une antenne

Dans une chaîne de réception (figure 3.12), le bruit introduit par les éléments avant le premier étage d'amplification sera prépondérant pour le rapport signal sur bruit du récepteur, car ce bruit est amplifié par tous les amplificateurs de la chaîne. Pour un récepteur hyperfréquences, ce bruit comprend trois éléments : Le bruit propre de l'antenne, le bruit introduit par la ligne de transmission et le bruit propre de l'étage d'entrée de l'amplificateur. Pour ne pas trop dégrader le signal reçu, il faut, dans la mesure du possible, minimiser ces trois niveaux de bruit. Plusieurs points doivent être pris en considération :

- Il est indispensable de minimiser les pertes dans la ligne de transmission reliant l'antenne au récepteur. Deux mesures s'imposent : choisir des lignes de transmissions à faibles pertes, et placer le récepteur aussi près que possible de l'antenne.
- Le premier amplificateur du récepteur doit être un *amplificateur à faible bruit* (low noise amplifier).
- Le "bruit ambiant" capté par l'antenne doit être aussi faible que possible. Ce bruit est constitué de rayonnements parasites appartenant à d'autres transmissions (important

par exemple pour des transmission mobiles en environnement urbain), du rayonnement solaire, du rayonnement galactique, etc. Pour cela, il est nécessaire de connaître les caractéristiques de ces différents bruits.

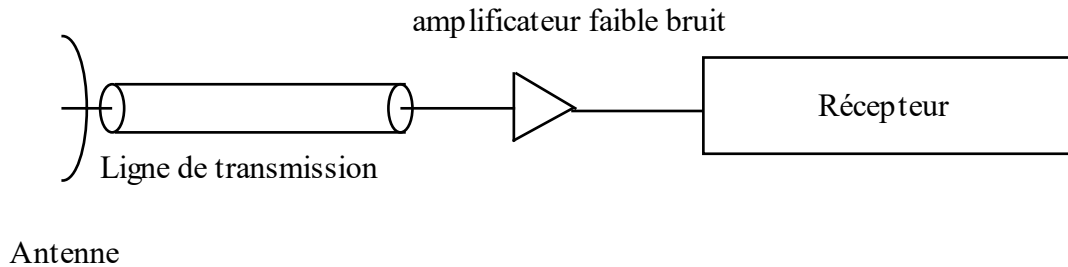


Figure 7.12 : Puissance de bruit dans un récepteur

Température de bruit équivalente du ciel

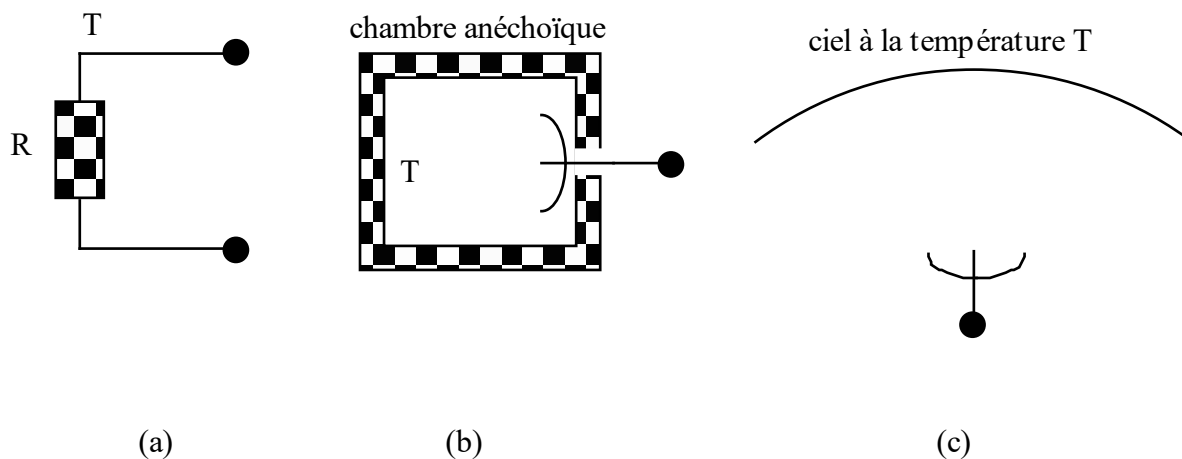


Figure 7.13 : Puissance de bruit

La puissance due au bruit aux bornes d'une résistance R (figure 7.13 a) à la température T est donnée par :

$$P_n = kT\Delta f \quad k = 1.38 \cdot 10^{-23} \text{ [J K}^{-1}\text{]}$$

où k est la constante de Boltzmann. Une antenne de résistance de rayonnement R , placée dans une chambre anéchoïque à la température T (figure 3.13 b), présente la même puissance de bruit à ses bornes. Si maintenant la même antenne est sortie de la chambre anéchoïque et pointée vers un ciel de température T (figure 3.13 c), la puissance due au bruit restera inchangée à ses bornes.

La température équivalente de bruit d'une antenne dépend de plusieurs facteurs : la fréquence, la direction de pointage, l'altitude. Une courbe de cette température en fonction de la fréquence est donnée à la figure 7.14.

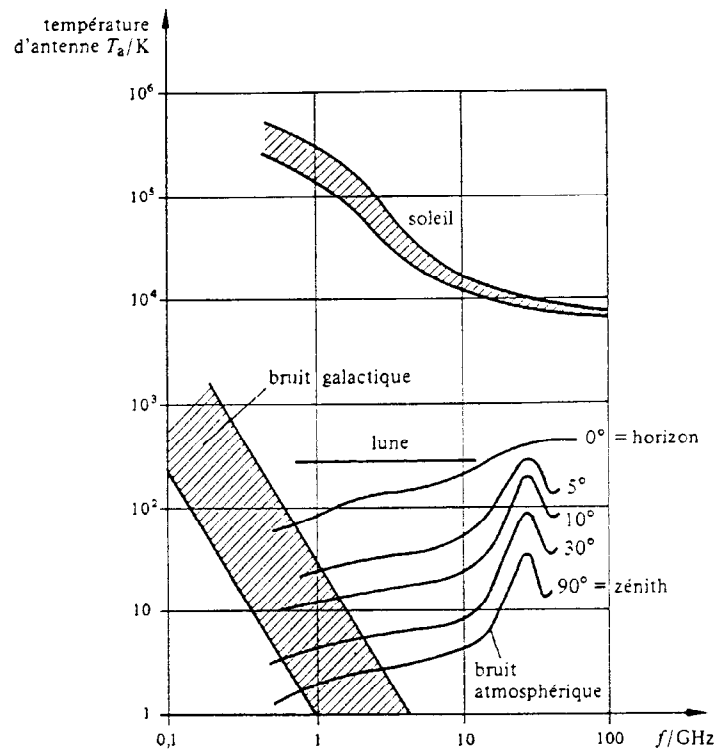


Figure 7.14 : Température de bruit d'une antenne en fonction du pointage (figure 7.40 du volume XIII du Traité d'Electricité, Hyperfréquences, par F.E. Gardiol, PPR, 1981)

On constate qu'il existe une fenêtre entre 1 et 10 GHz où la température de bruit d'une antenne pointée vers le ciel est faible. Ceci explique pourquoi la plupart des liaisons avec des satellites est faite dans cette bande de fréquences. Cette dernière commence malheureusement à être saturée, et des bandes à des fréquences plus élevées ont été ouvertes pour des communications avec des satellites. Ces nouvelles liaisons auront un rapport signal sur bruit moins favorable, mais à cause des fréquences plus élevées la bande passante est plus large et permet un débit d'information plus élevé.

7.11 Types d'antennes et leur caractéristiques

7.11.1 Bande passante

Nous avons jusqu'à présent toujours considéré les antennes en régime sinusoïdal, soit en ne fonctionnant qu'à une seule fréquence. Il est clair que pour transmettre de l'information, on travaillera sur une bande de fréquence, la bande passante du système, pour lesquelles on souhaite que l'antenne possède des caractéristiques satisfaisantes du point de vue du diagramme de rayonnement, de la polarisation et de l'adaptation.

On appelle "bande passante" de l'antenne les fréquences pour lesquelles ses caractéristiques sont jugées satisfaisantes. Cette bande est généralement donnée en % par rapport à la fréquence centrale de la bande :

$$B = \frac{f_{\text{sup}} - f_{\text{inf}}}{f_{\text{central}}} * 100\%$$

Cette définition est relative et subjective, et il est souvent nécessaire de la chiffrer plus précisément.

Bande passante (adaptation)

La bande passante en adaptation se définit en général à l'aide du Rapport d'Ondes Stationnaires, qui est une mesure de l'adaptation de l'antenne :

$$ROS = \frac{1 + |\rho|}{1 - |\rho|}$$

où ρ est le coefficient de réflexion aux bornes de l'antenne.

On dira par exemple qu'une antenne a une bande passante de 10% pour un rapport d'ondes stationnaires de 2, ce qui veut dire que la bande de fréquence pour lesquelles le $ROS < 2$ est de 10%.

Parfois, on définit la bande passante directement à partir du coefficient de réflexion.

7.11.2 Petites antennes

On dit qu'une antenne qu'elle est petite lorsque sa dimension la plus grande est beaucoup plus petite que la longueur d'onde :

$$d \ll \lambda$$

Ces antennes se caractérisent par une faible directivité et une petite bande passante

7.11.3 Antennes résonantes

Les antennes résonnantes ont une dimension de l'ordre d'une demi longueur d'onde :

$$d \approx \frac{\lambda}{2}$$

Ces antennes ont un gain qui peut aller jusqu'à une dizaine de dB et une bande passante (adaptation) pouvant atteindre 20 -25 %

7.11.4 Grandes antennes

Une antenne est dite grande lorsque sa dimension la plus grande est beaucoup plus grande que la longueur d'onde :

$$d \gg \lambda$$

Ce sont des antennes que l'on peut optimiser pour obtenir de très grandes bandes passantes ou des gains élevés. Une antenne log-périodique par exemple présente une bande passante de plusieurs octaves, alors qu'une parabole peut avoir un gain de 80 dB.

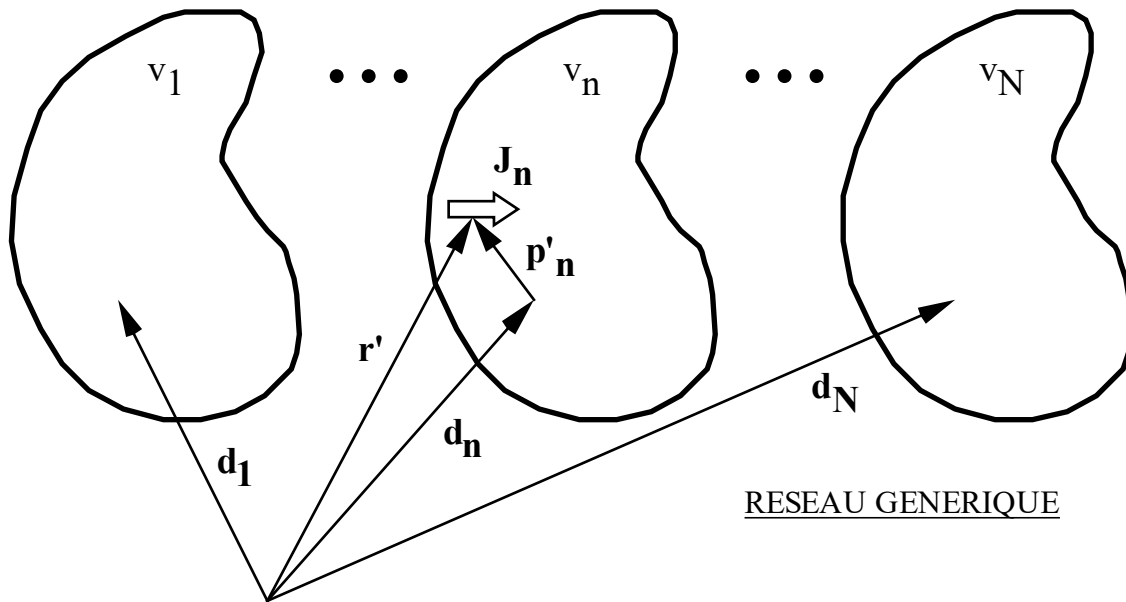
7.12 Réseaux

On peut définir un réseau d'antennes comme un ensemble d'antennes identiques possédant la même orientation. Chaque antenne du réseau est appelée un "élément du réseau"; tous les éléments d'un réseau doivent pouvoir être obtenus par translation dans l'espace d'un élément quelconque.

Le terme "réseau", utilisé au sens strict, exclut donc les groupements d'antennes où les éléments seraient identiques mais posséderaient des orientations différentes dans l'espace.

7.12.1 Courants dans le réseau : couplage mutuel

Considérons un réseau de N antennes. Ces antennes sont, bien sûr, identiques et repérées par les vecteurs de position $\mathbf{d}_n$ ($n=1,2,\dots,N$) d'un point caractéristique de chaque élément.



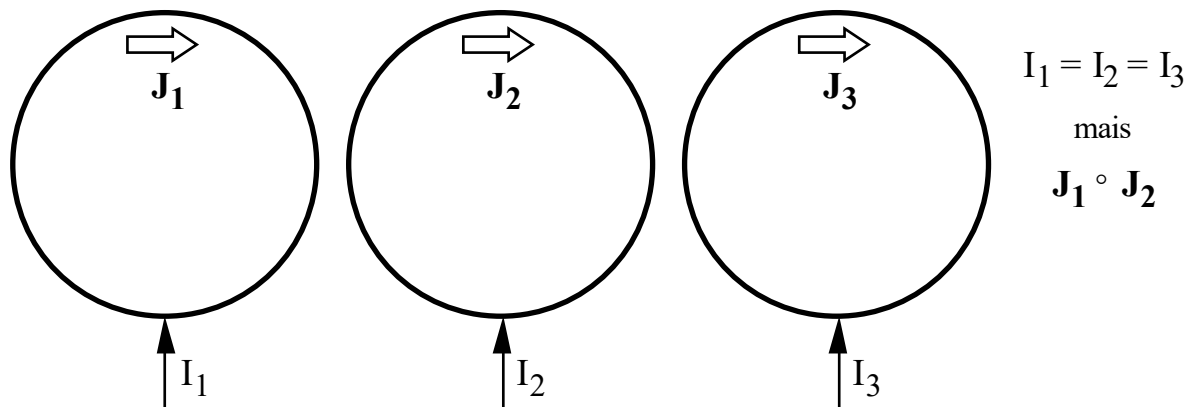
Un point quelconque de l'antenne # n est repéré par le vecteur $\mathbf{r}'$. On introduit pour chaque antenne des vecteurs de position "locaux" $\mathbf{p}'_n$, dont les composantes correspondent aux coordonnées locales ayant comme origine l'extrémité de $\mathbf{d}_n$. Donc, au sein de l'antenne # n on a :

$$\mathbf{r}' = \mathbf{d}_n + \mathbf{p}'_n.$$

La densité de courant dans l'antenne # n s'écrit alors $\mathbf{J}(\mathbf{d}_n + \mathbf{p}'_n)$ ou, dans une notation "locale", $\mathbf{J}_n(\mathbf{p}'_n)$.

La question qu'on se pose tout de suite quant à la nature de ces densités de courants est la suivante: Puisque tous les éléments d'un réseau sont identiques, peut-on affirmer par exemple que "si tous les éléments sont excités de façon identique, on obtient la même distribution de courant dans chacun de ces éléments" ?

La réponse précise est négative. Considérons par exemple trois sphères métalliques alignées. On introduit dans trois points équivalents de ces antennes (par exemple leurs trois "pôles Sud") des courants identiques en amplitude et en phase.



On comprend aisément que la densité de courant dans le pôle Nord de la sphère du centre n'est sûrement pas la même que dans le pôle Nord de la sphère de gauche ou de droite. C'est une affaire de position relative: l'environnement géométrique n'est pas le même pour une antenne placée au centre d'un réseau et pour une antenne à l'extrémité du réseau. L'influence d'un élément du réseau sur un autre dépend de leur position relative.

Cette influence entre éléments est connu sous le nom générique de **"couplage mutuel"**. Le couplage mutuel a été longtemps le cauchemar des ingénieurs, car il est difficile à inclure dans des modèles théoriques simples pour prédire son importance. Il est vrai que dans beaucoup de situations pratiques on constate "a posteriori" que le couplage mutuel est un effet de deuxième ordre. Dans ce chapitre, on négligera systématiquement l'existence du couplage mutuel entre éléments. Pour en tenir compte, il faudrait des modèles bien plus compliqués, qui font largement appel à l'ordinateur et où les résultats purement analytiques sont presque inexistants.

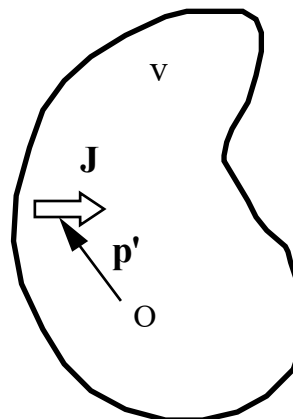
Hypothèse "couplage mutuel nul"

Soit $\mathbf{J}(\mathbf{p}')$ la distribution de courant existant dans un élément du réseau quand on le considère isolé dans l'espace et excité par un courant unité.

Si maintenant chaque élément du réseau est excité par un courant complexe I_n , on peut en première approximation écrire:

$$\mathbf{J}(\mathbf{d}_n + \mathbf{p}'_n) = \mathbf{J}_n(\mathbf{p}'_n) = I_n \mathbf{J}(\mathbf{p}') \quad \text{où} \quad \mathbf{J}_n(\mathbf{p}'_n) / \mathbf{J}_m(\mathbf{p}'_m) = I_n / I_m$$

ÉLÉMENT ISOLÉ
DE RÉFÉRENCE



ce qui peut s'exprimer en affirmant que, en absence de couplage, "le rapport entre les densités de courant dans deux éléments quelconques est égal à celui existant entre les respectifs courants d'excitation".

7.12.2 Le facteur de réseau

On démontre que le diagramme de rayonnement d'une antenne est lié à la densité de courant sur cette dernière par une grandeur appelée l'intégrale vectorielle, définie par

$$\mathbf{f}(\theta, \varphi) = \int_v dv' \mathbf{J}(\mathbf{r}') e^{j\mathbf{k}_e \cdot \mathbf{r}'}$$

où $\mathbf{e}_r(\theta, \varphi)$ est le vecteur unité dans la direction de rayonnement (θ, φ) considérée. A partir de cette intégrale, les champs lointains rayonnés sont obtenus par les relations

$$E_\theta(r, \theta, \varphi) = -\frac{jZ_c}{2\lambda} \frac{e^{jkr}}{r} \mathbf{e}_\theta \cdot \mathbf{f}(\theta, \varphi)$$

$$E_\varphi(r, \theta, \varphi) = -\frac{jZ_c}{2\lambda} \frac{e^{jkr}}{r} \mathbf{e}_\varphi \cdot \mathbf{f}(\theta, \varphi)$$

Quand on néglige le couplage mutuel, on peut calculer aisément le champ rayonné par le réseau dans un point lointain $\mathbf{r}$ à partir de l'intégrale vectorielle associée $\mathbf{f}$.

Puisque le volume du réseau est formé par l'union des volumes v_n de chaque élément, on a pour l'intégrale $\mathbf{f}$ (théorème de superposition) :

$$\begin{aligned} \mathbf{f}(\theta, \varphi) &= \sum_n \int_{v_n} dv' \mathbf{J}(\mathbf{r}') e^{j\mathbf{k}_e \cdot \mathbf{r}'} \\ &= \sum_n e^{j\mathbf{k}_e \cdot \mathbf{d}_n} \int_{v_n} dv' \mathbf{J}(\mathbf{d}_n + \mathbf{p}'_n) e^{j\mathbf{k}_e \cdot \mathbf{p}'_n} \\ &= \left[\int_{v_e} dv' \mathbf{J}(\mathbf{r}') e^{j\mathbf{k}_e \cdot \mathbf{p}'} \right] \left[\sum_n I_n e^{j\mathbf{k}_e \cdot \mathbf{d}_n} \right] \end{aligned}$$

Où on a introduit les coordonnées locales et on se à la densité de courant d'un élément isolé avec excitation unité.

v_e est le volume de l'élément de référence et l'intégrale a pu être extraite de l'intérieur de la somme.

On remarque que $\mathbf{f}$ est le produit des deux quantités bien distinctes, et on écrit symboliquement

$$\mathbf{f}(\theta, \varphi) = \underbrace{\left[\int_{v_e} dv' \mathbf{J}(\mathbf{r}') e^{j\mathbf{k}_e \cdot \mathbf{p}'} \right]}_{\mathbf{f}_e(\theta, \varphi)} \left[\sum_n I_n e^{j\mathbf{k}_e \cdot \mathbf{d}_n} \right] = \mathbf{f}_e(\theta, \varphi) AF(\theta, \varphi)$$

où $AF(\theta, \varphi) = \sum_n I_n e^{j\mathbf{k}_e \cdot \mathbf{d}_n}$

est une quantité appelée *facteur du réseau* (Array Factor).

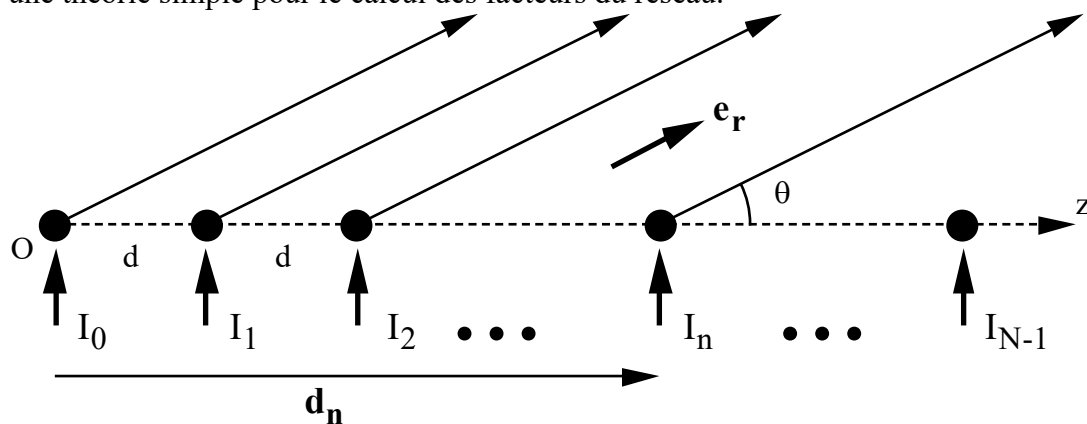
Comme la dépendance angulaire des champs rayonnés est directement donnée par l'intégrale $\mathbf{f}$, on peut affirmer que : "Le diagramme de rayonnement d'un réseau est le produit du diagramme de rayonnement d'un élément isolé et du facteur du réseau".

Le facteur du réseau traduit l'effet de la position relative et de l'excitation des éléments. Si $|\mathbf{f}_e(\theta, \varphi)|=1$, alors $|\mathbf{f}| = AF$. Le facteur du réseau est donc le diagramme de rayonnement qu'on obtiendrait si tous les éléments du réseau étaient des sources isotropes. En pratique on construit souvent des réseaux avec des éléments dont le rayonnement a une dépendance angulaire peu marquée (quasi-isotropes). La forme du diagramme de rayonnement est alors contrôlée essentiellement par le facteur du réseau $AF(\theta, \varphi)$.

Le facteur du réseau n'est pas influencé par la nature des antennes et chaque élément peut être assimilé à un point en ce qui concerne le calcul du facteur du réseau. On doit finalement remarquer que le facteur du réseau est une quantité scalaire complexe. Il ne comporte donc aucune information sur la polarisation des champs rayonnés, mais agit sur leur amplitude et sur leur phase.

7.12.3 Les réseaux linéaires à éléments équidistants

On appelle réseau linéaire celui dont les éléments sont alignés le long d'une ligne droite. De plus, les éléments sont fréquemment séparés par un écart constant d . On peut alors développer une théorie simple pour le calcul des facteurs du réseau.



Traditionnellement, on fait coïncider l'axe des coordonnées z avec l'axe du réseau et on numérote les éléments depuis $n=0$ jusqu'à $n=N-1$ (un total de N éléments). Le premier élément définit l'origine de coordonnées et on a alors :

$$\mathbf{d}_n = nd \mathbf{e}_z \quad \text{et} \quad \mathbf{e}_r \cdot \mathbf{d}_n = nd \cos \theta.$$

Le facteur du réseau possède alors une symétrie de révolution autour de l'axe z et dépend donc de l'angle sphérique θ mais pas de φ . On l'écrit :

$$AF(\theta, \varphi) = \sum_n I_n e^{j\mathbf{k} \cdot \mathbf{d}_n}$$

Les courants d'excitation I_n sont des phaseurs complexes possédant une amplitude A_n et une phase α_n . On a alors $I_n = A_n e^{j\alpha_n}$ et on peut écrire :

$$AF(\theta, \varphi) = \sum_{n=0}^{n=N-1} A_n e^{(jnkd \cos \theta + \alpha_n)}$$

Ces expressions seront à la base de tous les calculs successifs dans la théorie des réseaux linéaires équidistants.

Réseaux à déphasage linéaire

Une excitation assez courante en pratique consiste en une distribution de courants $A_n \exp(j\alpha_n)$ à déphasage linéaire ($\alpha_n = n\alpha$), incluant le cas équiphase $\alpha=0$. En effet, il est relativement aisé d'introduire un déphasage en contrôlant simplement la longueur des lignes de transmission amenant la puissance aux antennes. De plus, on verra que les diagrammes de rayonnement résultants présentent des caractéristiques intéressantes et facilement modifiables.

Dans ce cas, on a pour le facteur du réseau:

$$AF(\theta, \varphi) = \sum_{n=0}^{n=N-1} A_n e^{(jn(kd \cos \theta + \alpha))} = \sum_{n=0}^{n=N-1} A_n e^{(jn\psi)} \quad \text{avec} \quad \psi = kd \cos \theta + \alpha$$

La variable auxiliaire ψ joue un rôle très important dans la théorie des réseaux.

Elle inclut notamment l'effet de:

- la fréquence (k)
- la distance entre éléments (d)
- l'angle de pointage(θ)
- le déphasage entre excitations successives (α)

Le facteur du réseau est une fonction périodique dans la variable ψ , qui a les dimensions d'un angle mais qui ne doit pas être confondue avec l'angle géométrique θ . On peut dire qu'à l'ensemble des directions physiques possibles $\theta[-\pi, \pi]$ correspond une plage de valeurs de ψ $[-kd + \alpha, kd + \alpha]$ qui constituent la région dite "*visible*" de ψ . Si $kd < \pi$, la plage de valeurs visibles de ψ ne remplit pas une période 2π . On peut alors envisager des valeurs de ψ à l'extérieur de la région visible qui correspondent à des valeurs imaginaires de l'angle géométrique θ et on parle donc de "*région invisible*". Ces concepts de région visible et invisible en ψ sont très utiles dans la synthèse de réseaux linéaires.

Réseaux équiampitude à déphasage linéaire

Si toutes les amplitudes des excitations sont identiques ($A_n = 1$) le facteur du réseau est une progression géométrique facilement sommable:

$$AF(\theta, \varphi) = \sum_{n=0}^{n=N-1} A_n e^{(jn(kd \cos \theta + \alpha))} = \sum_{n=0}^{n=N-1} A_n e^{(jn\psi)} = \sum_{n=0}^{n=N-1} e^{(jn\psi)} = \frac{e^{jN\psi} - 1}{e^{j\psi} - 1}$$

$$AF(\theta, \varphi) = \frac{e^{jN\psi} - 1}{e^{j\psi} - 1} = \frac{\sin\left(N\frac{\psi}{2}\right)}{\sin\left(\frac{\psi}{2}\right)} e^{j\left(N-1\right)\frac{\psi}{2}}$$

En pratique, on ne s'intéresse souvent qu'à l'amplitude des champs rayonnés et non pas à la phase. On prend alors la norme du facteur et on trouve:

$$|AF(\theta, \varphi)| = \left| \frac{\sin\left(N \frac{\psi}{2}\right)}{\sin\left(\frac{\psi}{2}\right)} \right|$$

Aussi, ce qui importe souvent est le niveau relatif des champs quand la direction θ change. On utilise alors un facteur de réseau normalisé NAF (normalized array factor) dont la valeur ne dépasse pas l'unité. On a $NAF = |AF| / \max|AF|$ et dans notre cas d'un réseau équiampplitude à déphasage linéaire :

$$NAF(\theta, \varphi) = \left| \frac{\sin\left(N \frac{\psi}{2}\right)}{N \sin\left(\frac{\psi}{2}\right)} \right|$$

dans l'écriture courante on se passe souvent des barres de module, mais il est sous-entendu que un "NAF" ne comporte que des valeurs réelles positives

On voit facilement dans le NAF que le faisceau principal de l'antenne est compris entre les directions de rayonnement nul $N\psi/2 = \pm\pi$. Donc la largeur du faisceau est $4\pi/N$, en termes de variable ψ , ce qui veut dire qu'un faisceau mince demande un nombre d'éléments élevé. Pour N grand, le lobe latéral atteint son maximum approximativement pour $\psi = 3\pi/2$. Cette valeur maximale est alors $1/(N \sin 3\pi/2N)$. Le niveau du lobe secondaire tend alors vers la valeur limite $2/3\pi$, soit environ -13.3 dB, lorsque N augmente indéfiniment.

Ces données sont facilement transformables en termes d'angle géométrique θ .

7.12.4 Réseaux à pointage variable

Il est évident qu'un réseau équiampplitude à déphasage linéaire émet un rayonnement maximum pour $\psi=0$, c'est à dire dans la direction $\theta_{\max}$ solution de:

$$kd \cos \theta_{\max} + \alpha = 0 \quad (11)$$

Ceci implique qu'on peut "pointer" le réseau vers n'importe quelle direction en introduisant un déphasage entre éléments $\alpha = -kd \cos \theta_{\max}$. Ceci est le principe même des antennes à pointage électronique variable où il est possible de dépasser le faisceau sans faire tourner mécaniquement l'antenne. Mais, attention!, contrairement au pointage purement mécanique le pointage électronique déforme le diagramme de rayonnement et par exemple le niveau des lobes secondaires peut monter fortement quand on force l'antenne à pointer dans certaines directions par contrôle du déphasage.

Si les éléments du réseau sont alimentés avec la même amplitude et la même phase ($\alpha=0$), le réseau émet un rayonnement maximum pour $\theta=\pi/2$ et ceci pour toute valeur de kd . On a affaire

à un réseau rayonnant essentiellement perpendiculairement à son axe. Ce cas est connu sous le nom "broadside" dans la littérature.

Si l'on souhaite un rayonnement max. dans l'axe du réseau (configuration "endfire") il faut alors $\theta_{\max} = 0, \pi$ et donc $\alpha = \pm kd$.

On remarquera finalement que la condition rayonnement maximum n'est pas seulement liée à la valeur $\psi=0$ mais en réalité aussi à tous les multiples $\psi=2n\pi$. Si $kd > \pi$ ($d > \lambda/2$), il se peut qu'un ou plusieurs de ces valeurs multiples tombe dans la région visible de ψ . La conséquence est que ces réseaux où l'espacement entre éléments est supérieur à la demi-longueur d'onde peuvent montrer plusieurs directions $\theta_{\max}$ de rayonnement maximum. On parle alors souvent de lobes d'ambiguïté (anglais : grating lobes).