

**Université de Droit, d'Economie et des Sciences
Aix-Marseille III
Faculté des Sciences et Techniques de Saint-Jérôme**

Thèse

Présentée pour obtenir le titre de

Docteur en Sciences

par

Gérard TAYEB

**Contribution à l'étude de la diffraction des ondes
électromagnétiques par des réseaux.
Réflexions sur les méthodes existantes et sur
leur extension aux milieux anisotropes.**

Soutenue le 5 décembre 1990 devant la commission d'examen:

MM. J.C. BOLOMEY	Rapporteur
M. CADILHAC	
J.P. LAUDE	Rapporteur
R. PETIT	
C. VASSALLO	Rapporteur

REMERCIEMENTS

Je tiens à remercier tout particulièrement :

Monsieur le Professeur R. PETIT, qui m'a accueilli au Laboratoire d'Optique Electromagnétique et qui, tout en dirigeant mes travaux et grâce à ses talents pédagogiques, m'a fait partager une partie de son vaste savoir dans le domaine de l'électromagnétisme,

Monsieur le Professeur M. CADILHAC, qui, grâce à l'étendue de ses connaissances, m'a permis "d'éclaircir" et de résoudre de nombreux problèmes théoriques délicats. Sa contribution m'a été extrêmement précieuse,

Messieurs J.C. BOLOMEY, J.P. LAUDE et C. VASSALO, qui ont accepté de faire partie du jury et consacré une partie de leur temps à l'étude de ce mémoire,

Tous les membres du Laboratoire d'Optique Electromagnétique, qui ont mis à ma disposition leur expérience, et principalement :

Monsieur D. MAYSTRE, Directeur de Recherches au CNRS, qui m'a fait profiter de ses travaux antérieurs sur les méthodes intégrales,

Monsieur le Professeur P. VINCENT, qui m'a fourni une aide précieuse dans la programmation et m'a permis de venir à bout des déboires rencontrés dans cette partie de mon travail,

Monsieur M. SAILLARD, Chargé de Recherches au CNRS, pour les discussions fructueuses que nous avons eues ensemble dans bien des domaines,

Madame M. FRIZZI, qui s'est aimablement chargée de la frappe du manuscrit et dont j'ai pu apprécier une fois de plus à cette occasion la gentillesse et la dextérité.

TABLE DES MATIERES

<u>Chapitre 0 - Introduction</u>	6
<u>Chapitre 1 - Méthodes différentielles - Polarisation H//</u>	10
1.1. La méthode différentielle classique (M.D.C.)	10
1.1.1. Introduction	10
1.1.2. Présentation de la méthode	10
1.1.3. Résolution du problème par la M.D.C.	16
1.1.4. Le cas des "slanted gratings"	17
1.1.5. Comment s'assurer de la qualité des résultats numériques obtenus	20
1.2. Une méthode différentielle "améliorée" (M.D.A.)	21
1.3. Intégration de l'impédance ou de l'admittance	24
<u>Chapitre 2 - Méthode de synthèse</u>	29
2.1. Introduction	29
2.2. Description du problème	29
2.3. Une propriété des champs diffractés	31
2.4. Les espaces $\mathcal{V}$, $\mathcal{V}_1^+$, $\mathcal{V}_2^-$	32
2.5. Choix de familles totales dans $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$	33
2.6. Recherche d'un approximant de F	36
2.7. Produit scalaire dans $\mathcal{V}$	38
2.8. Calcul des efficacités	41
2.9. Quelques résultats numériques	42
2.10. Lien avec la méthode introduite par Yasuura	45
2.11. Conclusion	48
<u>Chapitre 3 - Propagation en milieu anisotrope</u>	49
3.1. Cadre de l'étude	49
3.2. Propagation en l'absence de sources en milieu homogène illimité	49
3.2.1. Position du problème	49
3.2.2. Une première façon de procéder	50
3.2.3. Une seconde façon de procéder	53
3.2.4. Quelques propriétés des constantes de propagation des ondes planes	54

3.3. Condition d'onde sortante (C.O.S.)	56
3.3.1. Quelques définitions et remarques préliminaires	56
3.3.2. Question posée	58
3.3.3. Réponse dans le cas d'un milieu isotrope	58
3.3.4. Réponse dans le cas d'un milieu anisotrope	59
3.3.4.1. En utilisant les vecteurs de Poynting	59
3.3.4.2. En raisonnant sur les constantes de propagation de l'onde	59
3.4. Solutions élémentaires de l'équation de propagation	62
3.4.1. Introduction, notations	62
3.4.2. Détermination des transformées de Fourier des vecteurs de Green	65
3.4.3. Inversion de Fourier - Détermination des vecteurs de Green ..	66
3.5. Quelques considérations sur la polarisation	69
<u>Chapitre 4 - Empilements de couches anisotropes</u>	72
4.1. Position du problème	72
4.2. Résolution du problème	73
4.3. Vecteur de Poynting, efficacités	75
4.4. Quelques résultats	77
4.4.1. Réfraction conique (expérience de Hamilton)	77
4.4.2. Renforcement d'une efficacité extraordinaire par utilisation de multicouches	79
<u>Chapitre 5 - Etude de réseaux anisotropes par la méthode différentielle</u>	81
5.1. Introduction	81
5.2. Présentation de la méthode	83
5.3. Les "slanted gratings"	88
5.4. Quelques résultats	89
5.4.1. Matériaux à pertes	89
5.4.2. Matériaux sans pertes	93
5.5. Conclusion	95

<u>Chapitre 6 - Etude de réseaux anisotropes par la méthode intégrale</u>	96
6.1. Introduction	96
6.2. Solutions élémentaires de l'équation de propagation	98
6.3. Généralisation des formules de Kirchhoff-Helmholtz	99
6.4. Obtention des équations intégrales	102
6.5. Résolution des équations intégrales	105
6.6. Détermination du champ diffracté	107
6.7. Résultats numériques. Critères de validité. Comparaison avec la méthode différentielle	109
6.8. Profils présentant des arêtes	112
6.9. Les opérateurs de Caldéron, définitions et propriétés	113
6.10. Sur l'usage des opérateurs de Caldéron	121
<u>REFERENCES</u>	124
<u>Annexe 1.</u> On the use of the energy balance criterion as a check of validity of computations in grating theory	130
<u>Annexe 2.</u> Compléments sur la "méthode différentielle améliorée"	139
<u>Annexe 3.</u> Obtention de la matrice impédance à partir d'une base	143
<u>Annexe 4.</u> Sur quelques propriétés des champs diffractés	144
<u>Annexe 5.</u> Un choix possible de familles totales dans $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$	149
<u>Annexe 6.</u> Permittivités associées à des milieux anisotropes "passifs"	162
<u>Annexe 7.</u> Sur la condition d'ondes sortantes en milieu anisotrope	164
<u>Annexe 8.</u> Sur la résolution de $\mathcal{L} \vec{E} = \vec{S}$ au moyen de solutions élémentaires de $\mathcal{L} \vec{g}_i = \vec{e}_i \delta$	166
<u>Annexe 9.</u> Sur les "vecteurs de Green" $\vec{g}_i$	168
<u>Annexe 10.</u> Découplage des vecteurs de Poynting	174
<u>Annexe 11.</u> Quelques lemmes pour les réseaux anisotropes	176
<u>Annexe 12.</u> Expression des $[T_i]$ et de $\mathbb{S}_2$	181
<u>Annexe 13.</u> Sur le traitement numérique des noyaux des équations intégrales	185

Le but initial de ce travail était l'étude de réseaux de diffraction constitués de matériaux anisotropes en électromagnétisme. On a évidemment tenté de généraliser les différentes méthodes déjà utilisées au Laboratoire pour l'étude des réseaux isotropes. Il est apparu que le domaine d'application des méthodes "différentielles" (limitées pour des raisons numériques, comme d'ailleurs dans le cas isotrope) ne nous permettait pas de traiter l'ensemble des cas envisagés. Quant aux méthodes "intégrales", leur transposition aux milieux anisotropes ne se fait pas sans difficultés. Pour ces raisons, il nous a paru nécessaire d'envisager des améliorations des méthodes couramment pratiquées (particulièrement pour les méthodes différentielles), ainsi que de nouvelles réflexions sur le problème général posé par les réseaux. C'est pourquoi finalement la première partie de ce mémoire (chapitres 1 et 2) est consacrée aux réseaux isotropes. Les études relatives aux structures anisotropes sont reportées dans la deuxième partie.

Dans tout le mémoire, nous utilisons un repère orthonormé $(0, x, y, z)$ muni d'une base $(\vec{e}_x, \vec{e}_y, \vec{e}_z)$. Le réseau (Figure 0.1) est une structure périodique en x , située dans la bande $0 < y < a$, et dont la périodicité (notée d) est le "pas" du réseau. Dans cette bande, la structure dépend de y . Elle est par contre invariante par translation parallèlement à $\vec{e}_z$. Cette définition très générale du "réseau", qui englobe un grand nombre de structures, convient bien aux méthodes "différentielles". Par contre, la mise en oeuvre des méthodes de type "intégrales" nécessite qu'il y ait une surface de discontinuité entre les matériaux, et nous serons donc souvent amenés à étudier des réseaux décrits par une interface périodique $y = f(x)$ (Figure 0.2). Le réseau est éclairé par une onde plane monochromatique dont le vecteur d'onde est dans le plan $(\vec{e}_x, \vec{e}_y)$. L'angle d'incidence est noté θ . Il en résulte que les champs électromagnétiques ne dépendent que des coordonnées spatiales x et y . Nous étant placés en régime harmonique de pulsation ω , nous représentons les champs en utilisant leur notation complexe, compte tenu d'une dépendance temporelle en $\exp(-i\omega t)$. La longueur d'onde dans le vide est notée λ_0 ; ϵ_0 et μ_0 désignent la permittivité et la perméabilité

du vide. Le "ka" du vide est $k_0 = 2\pi/\lambda_0 = \omega\sqrt{\epsilon_0\mu_0}$. Tous les milieux étudiés ont la même perméabilité que le vide.

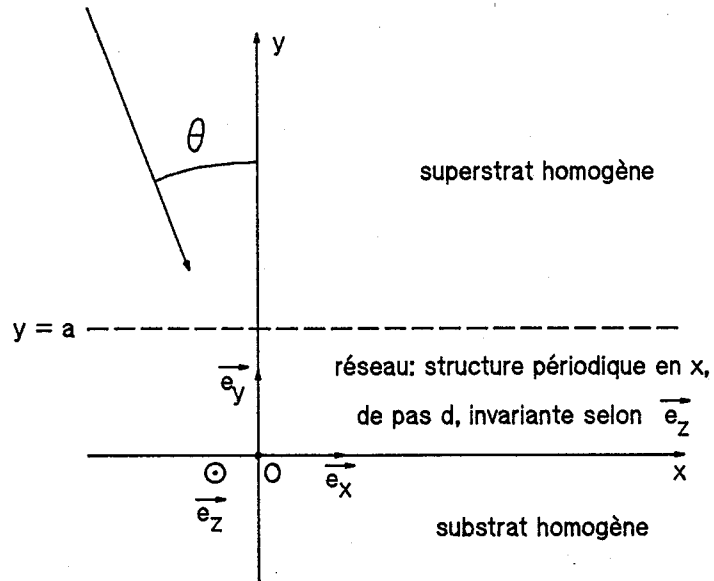


Figure 0.1

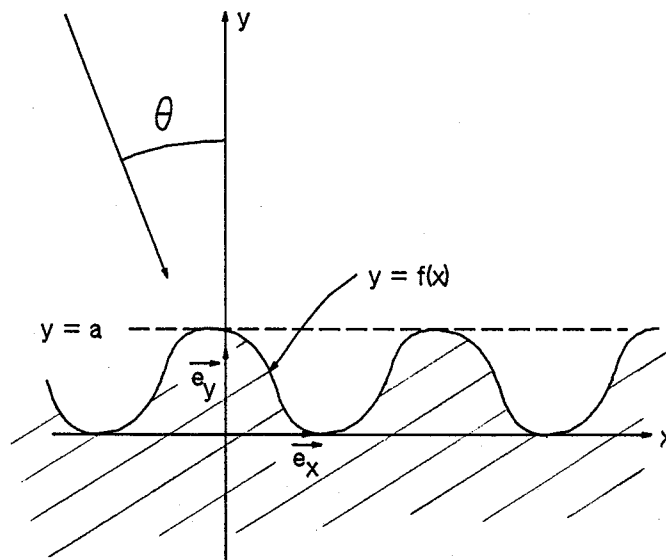


Figure 0.2

Nous serons souvent amenés à parler de la polarisation des champs. Dans toute la suite, nous parlerons de polarisation E// (ou TE) lorsque le champ électrique de l'onde incidente est parallèle à $\vec{e}_z$, et de polarisation H// (ou TM) lorsque le champ magnétique de l'onde incidente est parallèle à $\vec{e}_z$. Il est à noter que ces définitions du cas de polarisation sont uniquement relatives au champ incident. Si tous les matériaux sont isotropes, la polarisation du champ total sera la même que celle du champ incident. Par contre, en présence de matériaux anisotropes, le champ total aura en général une polarisation différente de celle du champ incident (cf § 3.5).

Nous avons limité notre expérimentation numérique au "domaine de résonance" défini par le fait que la longueur d'onde λ_0 et le pas d sont du même ordre de grandeur, ce qui signifie que le nombre d'ordres diffractés par la structure est réduit (généralement inférieur à une dizaine). L'extension des méthodes que nous proposons au cas où le nombre d'ordre diffractés serait notablement plus élevé ne pose pas a priori de problème théorique supplémentaire. Toutefois, nos programmes n'ayant pas été testés dans ce domaine, nous resterons réservés quant à leurs limites d'utilisation lorsque le pas du réseau est grand devant la longueur d'onde.

Pour satisfaire dès à présent la curiosité du Lecteur impatient, voici en quelques lignes le genre de problèmes que nous savons maintenant traiter en ce qui concerne les réseaux anisotropes [78, 79].

La méthode différentielle, très souple d'emploi, est adaptée aux cas envisagés sur la Figure 0.1, lorsque la zone $y < a$ est remplie de matériaux anisotropes quelconques. Rappelons que, puisque nous avons supposé que la perméabilité est partout égale à μ_0 , un matériau anisotrope sera caractérisé par une permittivité relative matricielle $[\epsilon]$ liant (dans les axes $\vec{e}_x$, $\vec{e}_y$ et $\vec{e}_z$ associés au réseau) le vecteur déplacement électrique $\vec{D}$ au vecteur champ électrique $\vec{E}$ selon $\vec{D} = \epsilon_0 [\epsilon] \vec{E}$. Lors de l'emploi de la méthode différentielle, les neuf éléments de la matrice $[\epsilon]$ sont a priori quelconques (sous réserve bien entendu qu'ils puissent représenter un milieu réel, avec ou sans pertes, cf annexe 6). Les limites d'utilisation de la méthode différentielle sont les mêmes que dans le cas isotrope : pour des matériaux sans pertes, on obtient de bons résultats lorsque le rapport a/λ_0 est inférieur à une valeur limite que l'on peut fixer à l'unité environ ; pour des matériaux à pertes, il semble difficile de donner une règle pratique car les limites d'emploi dépendent de nombreux paramètres, notamment la polarisation, les permittivités des matériaux ... Pour fixer les idées, nous

avons été amenés à étudier des réseaux contenant une forme anisotrope du Cobalt, et dans ce cas des résultats satisfaisants sont obtenus jusqu'à des valeurs de a/λ_0 de l'ordre de 1/10.

Pour la mise en oeuvre de la méthode intégrale, nous nous sommes jusqu'ici limités au cas où la permittivité $[\epsilon]$ est diagonale dans la base $\vec{e}_x, \vec{e}_y, \vec{e}_z$. Le programme réalisé traite le cas représenté sur la figure 0.2, la région hachurée étant homogène et anisotrope. Comme pour les réseaux isotropes, les avantages de la méthode se résument par des temps de calcul plus brefs, et la possibilité de traiter des réseaux nettement plus profonds (a/λ_0 peut facilement atteindre ou même dépasser l'unité, la limite étant imposée par les temps de calculs qui croissent rapidement avec la hauteur du réseau). En contrepartie, l'étude théorique est beaucoup plus délicate que pour la méthode différentielle, et le fait d'imposer à $[\epsilon]$ d'être diagonale restreint évidemment le champ d'application de la méthode intégrale.

1.1. LA METHODE DIFFERENTIELLE CLASSIQUE (M.D.C.)1.1.1. Introduction

Il existe de nombreuses variantes de "méthodes différentielles". Celle qui est la plus utilisée dans l'étude des réseaux a été introduite par R. Petit dès 1966 [56, 12] et a été depuis très largement développée par le Laboratoire [58] ; nous la dénommerons par la suite méthode différentielle classique (M.D.C.). Son domaine d'utilisation très large et sa souplesse d'emploi sont appréciables, et elle a permis de nombreuses études générales relatives par exemple aux réseaux [34, 47] ou au coupleur à réseau [46]. Elle s'avère particulièrement efficace dans le domaine des très courtes longueurs d'onde (UV et rayons X) [35, 48] et permet également l'étude des structures constituées de matériaux non linéaires [63, 83]. Malheureusement, des problèmes numériques apparaissent dans le cas de réseaux conducteurs profonds, et cela particulièrement dans le cas de polarisation H//. C'est la raison pour laquelle, dans ce chapitre, nous ne parlerons que de ce cas de polarisation difficile. Dans le but de mieux comprendre l'origine de ces problèmes numériques, et pour tenter d'y remédier, nous avons mené deux études présentées dans les paragraphes 1.2 ("méthode différentielle améliorée") et 1.3 (intégration de l'impédance ou de l'admittance). Pour permettre une généralisation plus facile de la M.D.C. au cas anisotrope et aussi pour préparer l'étude faite aux paragraphes 1.2 et 1.3, nous allons donner dans les lignes qui suivent une présentation de la M.D.C. légèrement différente de celle utilisée jusqu'ici [58].

1.1.2. Présentation de la méthode

Nous étudions le problème présenté sur la figure 1.1. Le problème est scalaire, autrement dit toute composante des champs peut se déduire de la composante $H_z = \vec{H} \cdot \vec{e}_z$. Le champ incident est :

$$H_z^i(x,y) = \exp(i\alpha_0 x - i\beta_0 y) , \quad (1.1)$$

$$\text{avec } \alpha_0 = k_0 n_1 \sin\theta, \quad \beta_0 = k_0 n_1 \cos\theta = \sqrt{k_0^2 n_1^2 - \alpha_0^2}. \quad (1.2)$$

Nous savons [58] que toute composante de champ $u(x,y)$ est "pseudo-périodique", ce qui signifie que :

$$u(x+d, y) = \exp(i\alpha_0 d) u(x, y), \quad (1.3)$$

d'où l'on déduit [58] qu'elle peut se mettre sous la forme d'un développement de Fourier généralisé :

$$u(x, y) = \sum_{n \in \mathbb{Z}} u_n(y) \psi_n(x), \quad (1.4)$$

$$\text{avec } \psi_n(x) = \exp(i\alpha_n x), \quad \alpha_n = \alpha_0 + nK, \quad K = 2\pi/d. \quad (1.5)$$

Les fonctions $u_n(y)$ seront appelées coefficients de Fourier généralisés de $u(x, y)$.

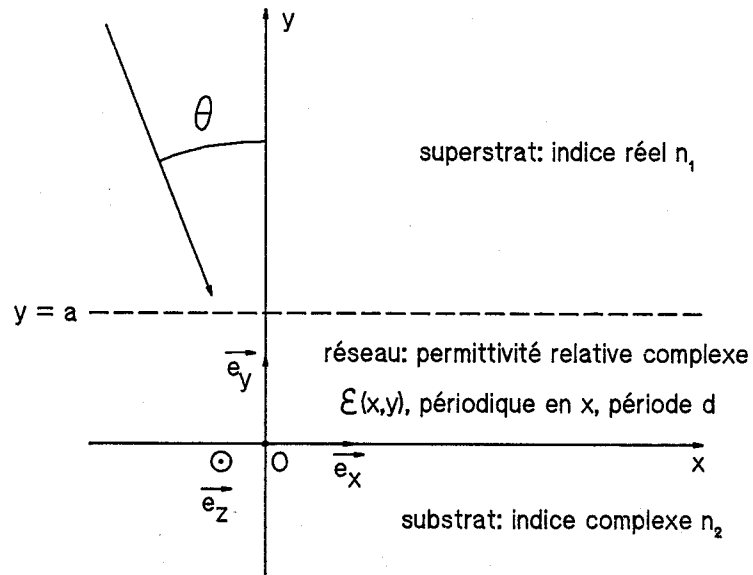


Figure 1.1

Dans la zone inhomogène $0 < y < a$, les équations de Maxwell $\text{rot } \vec{E} = i\omega \mu_0 \vec{H}$ et $\text{rot } \vec{H} = -i\omega \varepsilon_0 \varepsilon \vec{E}$, considérées au sens des distributions, conduisent aux équations :

$$\begin{cases} \partial_y H_z = -i\omega \varepsilon_0 \varepsilon E_x \\ \partial_y E_x = \partial_x \left(\frac{1}{i\omega \varepsilon_0 \varepsilon} \partial_x H_z \right) - i\omega \mu_0 H_z. \end{cases} \quad (1.6)$$

Pour simplifier l'écriture, introduisons l'impédance du vide $\eta_0 = \sqrt{\mu_0/\varepsilon_0}$, et posons pour la suite de ce chapitre :

$$u(x,y) = H_z(x,y) \quad \text{et} \quad v(x,y) = \frac{1}{\eta_0} E_x(x,y), \quad (1.7)$$

ce qui permet d'écrire (1.6) sous la forme :

$$\begin{cases} \partial_y u = -ik_0 \varepsilon v \\ \partial_y v = \partial_x \left(\frac{1}{ik_0 \varepsilon} \partial_x u \right) - ik_0 u. \end{cases} \quad (1.8)$$

$$(1.9)$$

Soient $\varepsilon_n(y)$ et $(1/\varepsilon)_n(y)$ les coefficients de Fourier des fonctions périodiques $\varepsilon(x,y)$ et $(1/\varepsilon)(x,y)$:

$$\begin{cases} \varepsilon(x,y) = \sum_{n \in \mathbb{Z}} \varepsilon_n(y) \exp(inKx), \\ (1/\varepsilon)(x,y) = \sum_{n \in \mathbb{Z}} (1/\varepsilon)_n(y) \exp(inKx). \end{cases} \quad (1.10)$$

En notant $u_n(y)$ et $v_n(y)$ les coefficients de Fourier généralisés de u et v , et après projection sur les $\psi_n(x)$, le système (1.8, 1.9) se met sous la forme :

$$\forall n, \begin{cases} \frac{du_n}{dy} = \sum_m -ik_0 \varepsilon_{n-m} v_m \\ \frac{dv_n}{dy} = \sum_m \left[-ik_0 \delta_{nm} + \frac{i\alpha_m \alpha_n}{k_0} (1/\varepsilon)_{n-m} \right] u_m \end{cases} \quad (1.11)$$

Il est commode de mettre ce système différentiel infini sous forme matricielle. Soit $F(y)$ la colonne constituée par les deux infinités de composantes $u_n(y)$ et $v_n(y)$; (1.11) s'écrit :

$$\frac{dF(y)}{dy} = A(y) F(y) , \quad (1.12)$$

où $A(y)$ est une matrice infinie connue.

Remarque : le système (1.11), vrai au sens des distributions, montre que les fonctions $u_n(y)$ et $v_n(y)$ sont continues (puisque aucune distribution singulière n'apparaît au second membre).

Dans chaque zone homogène $y > a$ et $y < 0$, u vérifie une équation de Helmholtz. Il en résulte [58] que, compte tenu des conditions de rayonnement, le développement (1.4) se met sous la forme de développements de Rayleigh :

pour $y > a$,

$$u(x,y) = \exp(i\alpha_0 x - i\beta_0 y) + \sum_{n \in \mathbb{Z}} R_n \exp(i\alpha_n x + i\beta_n y), \quad (1.13)$$

pour $y < 0$,

$$u(x,y) = \sum_{n \in \mathbb{Z}} T_n \exp(i\alpha_n x - i\beta'_n y), \quad (1.14)$$

$$\text{avec } \beta_n^2 = k_0^2 n_1^2 - \alpha_n^2, \quad \beta'_n{}^2 = k_0^2 n_2^2 - \alpha_n^2, \quad (1.15)$$

la fonction racine carrée utilisée pour déterminer β_n et β'_n est telle que :

$$\forall z \in \mathbb{C}, \quad \text{Im}(\sqrt{z}) > 0 \quad \text{ou bien} \quad \sqrt{z} > 0 \quad (1.16)$$

L'équation (1.8) nous donne immédiatement l'expression de v :

pour $y > a$,

$$v(x,y) = \frac{1}{k_0 n_1^2} \left[\beta_0 \exp(i\alpha_0 x - i\beta_0 y) - \sum_{n \in \mathbb{Z}} \beta_n R_n \exp(i\alpha_n x + i\beta_n y) \right] \quad (1.17)$$

pour $y < 0$,

$$v(x,y) = \frac{1}{k_0 n_2^2} \sum_{n \in \mathbb{Z}} \beta'_n T_n \exp(i\alpha_n x - i\beta'_n y) \quad (1.18)$$

Définissons, pour $y > a$, les champs diffractés :

$$u^d(x,y) = u(x,y) - u^i(x,y) = u(x,y) - \exp(i\alpha_0 x - i\beta_0 y), \quad (1.19)$$

$$v^d(x,y) = v(x,y) - v^i(x,y) = v(x,y) - \frac{\beta_0}{k_0 n_1^2} \exp(i\alpha_0 x - i\beta_0 y). \quad (1.20)$$

On remarquera alors que dans les deux zones considérées il existe des relations simples et indépendantes de y entre les coefficients de Fourier généralisés de la partie sortante de u et v :

$$\text{pour } y > a, \quad v_n^d(y) = - \frac{\beta_n}{k_0 n_1^2} u_n^d, \quad (1.21)$$

$$\text{pour } y < 0, \quad v_n(y) = \frac{\beta'_n}{k_0 n_2^2} u_n. \quad (1.22)$$

En d'autres termes, cela signifie que les opérateurs permettant de passer de u^d à v^d ($y > a$) et de u à v ($y < 0$) sont diagonaux dans la base des $\psi_n(x)$. Ces opérateurs sont traditionnellement appelés opérateurs impédance ; rappelons que u est une composante du champ magnétique et v une composante du champ électrique. Pour y donné, la relation entre la composante u et la composante v d'un champ diffracté peut se traduire par le fait que $F(y)$ appartient à un certain espace dépendant de y , dont on peut facilement, à l'aide de ce qui précède, mettre en évidence une base pour $y < 0$ ou pour $y > a$. Pour des raisons de commodité, nous allons supposer que chaque composante de champ est bien représentée par un nombre fini $N = 2P + 1$ (P entier positif) de composantes sur la base des $\psi_n(x)$. Cette troncature est de toute façon nécessaire pour le traitement numérique du problème. La faire dès à présent présente l'avantage de pouvoir numérotter précisément les composantes de F , ce qui est plus délicat tant que l'on conserve pour F une "double infinité" de composantes. Pour fixer les idées, plaçons nous pour $y < 0$. La colonne $F(y)$ peut alors s'écrire compte tenu de (1.14) et (1.18) :

$$\begin{aligned}
 F(y) = \begin{pmatrix} u_{-p}(y) \\ \vdots \\ u_0(y) \\ \vdots \\ u_p(y) \\ v_{-p}(y) \\ \vdots \\ v_0(y) \\ \vdots \\ v_p(y) \end{pmatrix} &= T_{-p} e^{-i\beta'_{-p} y} \underbrace{\begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \\ \chi_{-p} \\ 0 \\ \vdots \\ 0 \end{pmatrix}}_{\tau_{-p}(y)} + \dots + T_0 e^{-i\beta'_0 y} \underbrace{\begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \\ \chi_0 \\ \vdots \\ 0 \end{pmatrix}}_{\tau_0(y)} + \dots + T_p e^{-i\beta'_p y} \underbrace{\begin{pmatrix} 0 \\ \vdots \\ 0 \\ 0 \\ 1 \\ \vdots \\ 0 \\ \chi_p \end{pmatrix}}_{\tau_p(y)} \\
 &\quad (1.23)
 \end{aligned}$$

avec $\chi_k = \frac{\beta'_k}{k_0 n_2^2}$, ce qui met en évidence une base $\tau_k(y)$ de l'espace $\mathcal{E}_2^-(y)$ dans lequel $F(y)$ doit se trouver :

$$\text{pour } y < 0, \quad F(y) = \sum_k T_k \tau_k(y) . \quad (1.24)$$

En résumé, retenons que, pour $y < 0$, la colonne $F(y)$ constituée par les deux infinités de coefficients de Fourier généralisés $u_n(y)$ et $v_n(y)$ appartient à un espace $\mathcal{E}_2^-(y)$ dont on connaît explicitement une base $\tau_k(y)$. Cette appartenance traduit le fait que le champ électromagnétique vérifie dans le substrat les équations de Maxwell, la condition de pseudo-périodicité imposée par le réseau, et la condition de rayonnement. Cela justifie la notation $\mathcal{E}_2^-$ dans laquelle l'indice 2 rappelle que cet espace est relatif au milieu d'indice n_2 (substrat) et l'exposant - rappelle que l'on impose une condition de rayonnement vers les y négatifs. Remarquons enfin qu'après troncature, $F(y)$ s'exprime comme combinaison linéaire de $N = 2P + 1$ vecteurs $\tau_k(y)$, chacun de ces vecteurs possédant $2N$ composantes.

De façon analogue, pour $y > a$, désignons par $F^i(y)$ la colonne associée au champ incident $u^i(x,y)$ et $v^i(x,y)$. La colonne $F(y) - F^i(y)$ associée au champ diffracté appartient à un espace $\mathcal{E}_1^+(y)$ dont on connaît une base $\rho_k(y)$ de manière explicite. On peut écrire :

$$\text{pour } y > a, \quad F(y) - F^i(y) = \sum_k R_k \rho_k(y) . \quad (1.25)$$

L'appartenance de $F(y) - F^i(y)$ à $\mathcal{E}_1^+(y)$ traduit le fait que $F(y) - F^i(y)$ est relatif à un champ se propageant dans le milieu d'indice n_1 (superstrat) et vérifie une condition de rayonnement vers les y positifs.

1.1.3. Résolution du problème par la M.D.C.

D'après ce qui précède, nous sommes ramenés à déterminer les coefficients R_k et T_k tels que :

$$\text{pour } 0 \leq y \leq a, \quad \frac{dF(y)}{dy} = A(y) F(y), \quad (1.26)$$

$$\text{pour } y = 0, \quad F(0) = \sum_k T_k \tau_k(0), \quad (1.27)$$

$$\text{pour } y = a, \quad F(a) - F^i(a) = \sum_k R_k \rho_k(a), \quad (1.28)$$

où $A(y)$, $F^i(a)$, $\tau_k(0)$ et $\rho_k(a)$ sont explicitement connus. A ce stade, le problème apparaît comme la résolution d'un système différentiel avec conditions aux limites. Nous pouvons procéder de la manière suivante : pour tout k et pour $0 \leq y \leq a$, considérons le vecteur $\tau_k(y)$ déduit de $\tau_k(0)$ et de l'intégration numérique de $y = 0$ à $y = a$ du système :

$$\frac{d\tau_k(y)}{dy} = A(y) \tau_k(y) . \quad (1.29)$$

Il est clair que :

$$F(y) = \sum_k T_k \tau_k(y) \quad (1.30)$$

vérifie les conditions (1.26) et (1.27). Considérons l'espace $\mathcal{E}_2^-(y)$ engendré par les $\tau_k(y)$, et maintenant défini pour $y \leq a$. L'appartenance de $F(y)$ à $\mathcal{E}_2^-(y)$ traduit les conditions imposées au champ à la fois par le réseau et par la condition de rayonnement vers les y négatifs.

Le problème (1.26, 1.27, 1.28) se réduit alors au problème équivalent :

$$\begin{cases} F(a) = F^i(a) + \sum_k R_k \rho_k(a), \\ F(a) = \sum_k T_k \tau_k(a), \end{cases} \quad (1.31)$$

$$(1.32)$$

et l'on voit qu'il s'agit de déterminer l'intersection de $\mathcal{E}_2^-(a)$ et du "translaté" de $\mathcal{E}_1^+(a)$ (par translation $F^i(a)$). Nous sommes ramenés à un problème élémentaire d'algèbre linéaire. Les R_k et T_k donnent le champ "réfléchi" (ou diffracté vers le haut) et le champ "transmis". La relation (1.30) permet de connaître les valeurs du champ dans la zone du réseau.

1.1.4. Le cas des "slanted gratings"

Le terme de "slanted grating" désigne souvent dans la littérature [43] le réseau représenté sur la figure 1.2 et caractérisé par le fait que, dans la bande $0 < y < a$, la permittivité $\varepsilon(x,y)$ ne dépend que du produit scalaire $\vec{L} \cdot \vec{r}$ ($\vec{r}$ désigne le vecteur position), c'est à dire qu'elle est constante dans tout plan perpendiculaire à $\vec{L}$. Un "slanted grating" constitue donc un cas particulier du réseau décrit sur la figure 0.1, pour lequel, comme nous allons le voir, l'on peut éviter lors de la résolution numérique du problème, l'intégration des systèmes différentiels (1.29) (pour $k = -P, \dots, +P$).

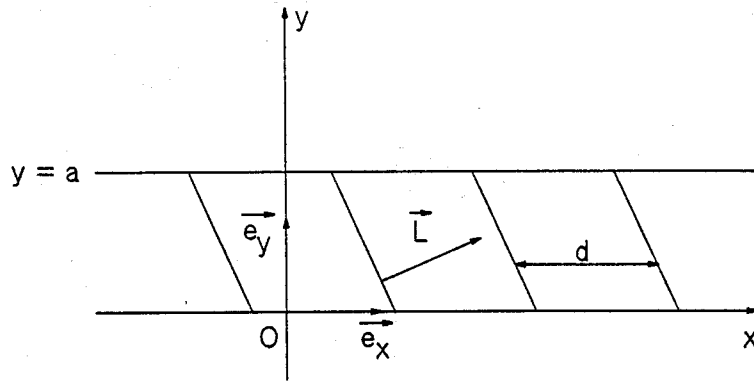


Figure 1.2. $\vec{L}$ est dans le plan $(\vec{e}_x, \vec{e}_y)$

La permittivité étant comme précédemment une fonction périodique en x (de période d), on en déduit :

$$\varepsilon(x,y) = \sum_{n \in \mathbb{Z}} \varepsilon_n(y) \exp(inKx) = f(\vec{L} \cdot \vec{r}) = f(L_x x + L_y y) \quad (1.33)$$

où f est une fonction périodique de période $(L_x d)$. On en déduit immédiatement que les coefficients de Fourier $\varepsilon_n(y)$ se mettent sous la forme :

$$\varepsilon_n(y) = c_n \exp(in\mu y), \quad (1.34)$$

$$\text{où } c_n = \frac{1}{L_x d} \int_0^{L_x d} f(t) \exp\left(-in \frac{K}{L_x} t\right) dt \quad (1.35)$$

est indépendant de y , et $\mu = K L_y / L_x$. De même, les coefficients de Fourier de $(1/\varepsilon)(x, y)$ peuvent s'écrire :

$$(1/\varepsilon)_n(y) = d_n \exp(in\mu y) \quad (1.36)$$

où d_n est indépendant de y . Le système (1.11) s'écrit alors :

$$\forall n, \begin{cases} \frac{du_n(y)}{dy} = \sum_m -ik_0 c_{n-m} \exp(i(n-m)\mu y) v_m(y) \\ \frac{dv_n(y)}{dy} = \sum_m \left[-ik_0 \delta_{nm} + \frac{i\alpha_m \alpha_n}{k_0} d_{n-m} \exp(i(n-m)\mu y) \right] u_m(y) \end{cases} \quad (1.37)$$

En effectuant le changement de variables :

$$\tilde{u}_n(y) = u_n(y) \exp(-in\mu y) \quad (1.38)$$

$$\tilde{v}_n(y) = v_n(y) \exp(-in\mu y),$$

le système (1.37) prend la forme :

$$\forall n, \begin{cases} \frac{d\tilde{u}_n}{dy} = -in\mu \tilde{u}_n + \sum_m -ik_0 c_{n-m} \tilde{v}_m \\ \frac{d\tilde{v}_n}{dy} = -in\mu \tilde{v}_n + \sum_m \left[-ik_0 \delta_{nm} + \frac{i\alpha_m \alpha_n}{k_0} d_{n-m} \right] \tilde{u}_m \end{cases} \quad (1.39)$$

En appelant $\tilde{F}(y)$ la colonne constituée par les deux infinités de composantes $\tilde{u}_n(y)$ et $\tilde{v}_n(y)$, (1.39) s'écrit :

$$\frac{d\tilde{F}(y)}{dy} = \tilde{A} \tilde{F}(y), \quad (1.40)$$

où $\tilde{A}$ est connue et est indépendante de y. Se reportant à (1.23), on constate que (1.24) peut s'écrire :

$$\text{pour } y < 0, \quad \tilde{F}(y) = \sum_k T_k \tilde{\tau}_k(y), \quad (1.41)$$

$$\text{avec } \tilde{\tau}_k(y) = \exp(-ik\mu y) \tau_k(y). \quad (1.42)$$

De même, (1.25) se met sous la forme :

$$\tilde{F}(y) - \tilde{F}^i(y) = \sum_k R_k \tilde{\rho}_k(y), \quad (1.43)$$

$$\text{avec } \tilde{\rho}_k(y) = \exp(-ik\mu y) \rho_k(y). \quad (1.44)$$

Il apparaît donc que dans le cas des "slanted gratings", le changement de variable élémentaire (1.38) permet de poser le problème (1.26 à 1.28) sous la forme équivalente :

$$\left\{ \begin{array}{ll} \text{pour } 0 \leq y \leq a, & \frac{d\tilde{F}(y)}{dy} = \tilde{A} \tilde{F}(y) \\ \text{pour } y = 0, & \tilde{F}(0) = \sum_k T_k \tilde{\tau}_k(0) \\ \text{pour } y = a, & \tilde{F}(a) - \tilde{F}^i(a) = \sum_k R_k \tilde{\rho}_k(a) \end{array} \right. \quad (1.45)$$

La résolution de ce problème se conduit de la même manière que dans le cas d'un réseau quelconque, à la différence que $\tilde{A}$ ne dépend pas de y, et que par conséquent la recherche de $\tilde{\tau}_k(a)$ connaissant $\tilde{\tau}_k(0)$ et vérifiant pour $0 \leq y \leq a$:

$$\frac{d\tilde{\tau}_k(y)}{dy} = \tilde{A} \tilde{\tau}_k(y) \quad (1.46)$$

se ramène à la détermination des valeurs propres et des vecteurs propres de $\tilde{A}$, évitant ainsi une intégration numérique plus gourmande en temps de calcul.

1.1.5. Comment s'assurer de la qualité des résultats numériques obtenus ?

Après avoir écrit un programme mettant en application les considérations précédentes, se pose le problème de savoir quelle crédibilité l'on peut accorder aux résultats numériques qu'il fournit. Rappelons que nous avons tronqué les développements en série infinis du type (1.4), en conservant un nombre fini de termes N . Nous obtenons par suite des résultats approchés, et l'erreur d'approximation (écart entre ces résultats et la solution exacte) dépend à la fois de cette valeur N et de la précision de l'intégration numérique de (1.29). Cette intégration est réalisée au moyen de l'algorithme d'Adams-Moulton du 5^{ème} ordre avec prédiction et correction, le procédé étant initialisé par la méthode de Runge-Kutta d'ordre 4 [23, 24]. La convergence de la solution en fonction du nombre de pas d'intégration J entre $y = 0$ et $y = a$ s'obtient facilement. Il est plus intéressant d'observer la convergence de la solution en fonction de l'ordre de troncature N , qui est, lors de l'utilisation des méthodes différentielles, le seul test possible de la validité des résultats numériques obtenus. On montre en effet (voir Annexe 1) que, dans le cas de réseaux constitués de matériaux sans pertes, le critère de conservation de l'énergie qui traduit le fait que toute l'énergie incidente se retrouve dans les ordres diffractés, est automatiquement vérifié par la M.D.C., quel que soit N . Il en est de même du critère de réciprocité (Fig. 1.3). Par conséquent, ces deux critères ne permettent pas de s'assurer de la qualité de la solution numérique. Tout au plus permettent-ils de renforcer notre conviction que les logiciels que nous avons utilisés ne présentent pas d'erreur de programmation ...

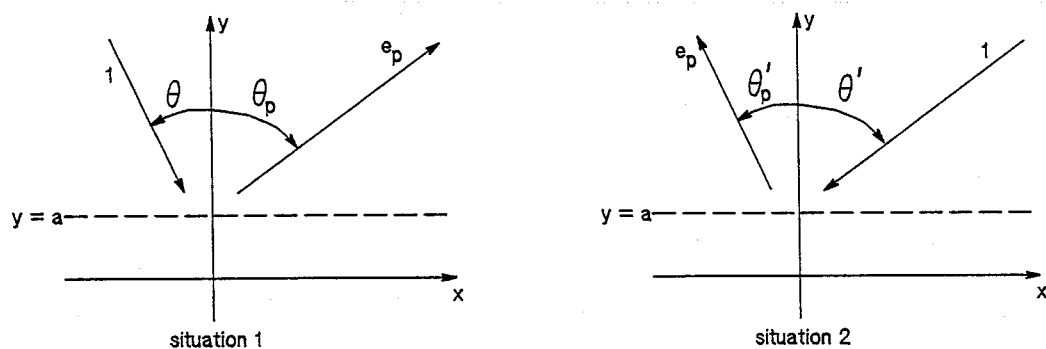


Figure 1.3. Le critère de réciprocité : Dans la situation 1, le réseau est éclairé sous une incidence θ par une onde plane de puissance unité. On récupère dans l'ordre p (angle θ_p) une puissance e_p égale à l'efficacité dans cet ordre. Dans la situation 2, le même réseau est éclairé sous l'incidence $\theta' = -\theta_p$ par une onde plane de puissance unité. On récupère dans l'ordre p (angle $\theta'_p = -\theta$), une puissance égale à e_p .

1.2. UNE METHODE DIFFERENTIELLE "AMELIOREE" (M.D.A.)

L'expérience numérique montre que, lorsque la hauteur du réseau croît, la convergence de la MDC en fonction de l'ordre de troncature N est de plus en plus lente. On observe également que lorsque N dépasse une valeur limite (qui est généralement de l'ordre d'une quarantaine), la MDC donne des résultats aberrants causés par des problèmes numériques. On atteint donc les limites de la MDC lorsque la convergence de la solution n'a pas été obtenue de façon satisfaisante avant que ces problèmes numériques ne surviennent. Pour fixer les idées, envisageons le cas concret d'un réseau sinusoïdal métallique en polarisation $H//$ (Figure 1.4).

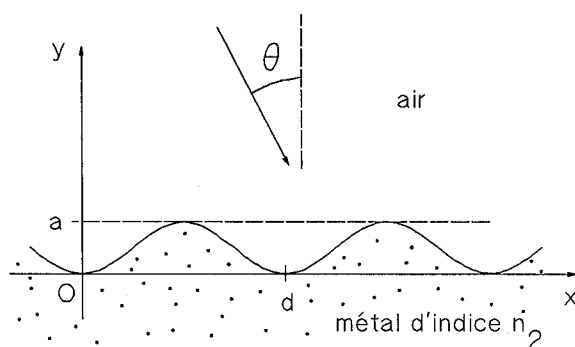


Figure 1.4. Réseau sinusoïdal de pas d et de hauteur a éclairé sous l'incidence θ ; polarisation $H//$.

Les figures 1.5, 1.6 et 1.7 montrent les résultats numériques obtenus pour des réseaux en aluminium et en or de différentes hauteurs (l'efficacité dans le $n^{\text{ième}}$ ordre est notée e_n). On constate que les valeurs maximales de N que l'on peut utiliser avec la MDC ne permettent pas toujours d'obtenir une précision suffisante sur les résultats.

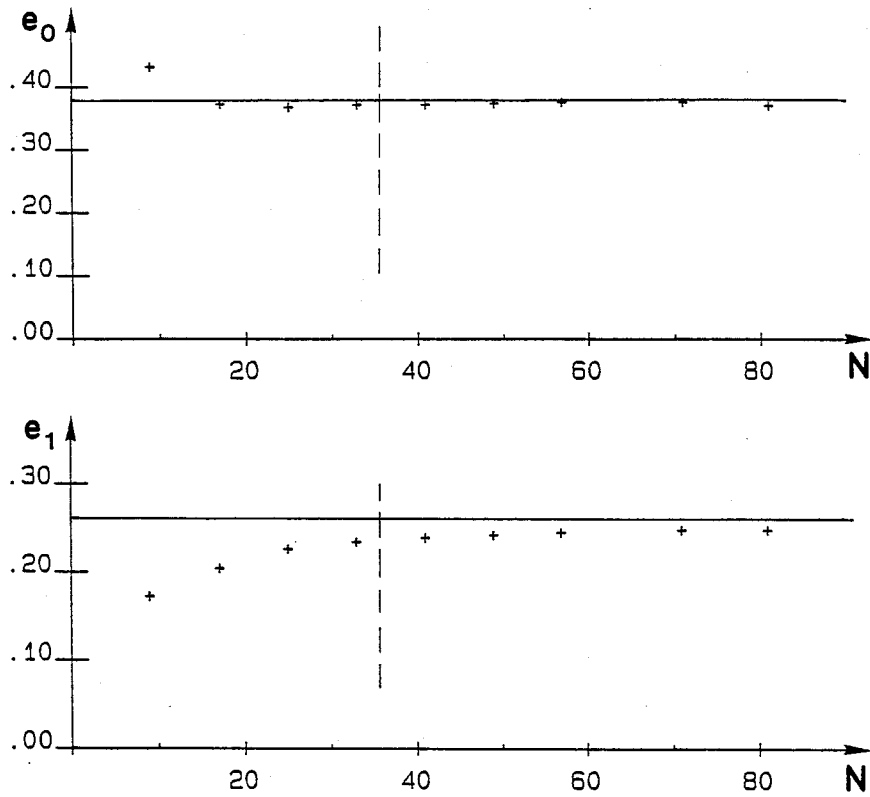


Figure 1.5. Efficacités en fonction de l'ordre de troncature N pour un réseau en aluminium, $d = 0.8333 \mu\text{m}$, $a = 0.1 \mu\text{m}$, $\lambda_0 = 0.6 \mu\text{m}$, $n_2 = 1.3 + i 7.1$, $\theta = 0$, polarisation $H//$. La droite horizontale correspond au résultat donné par la méthode intégrale, que nous considérons comme "exact". La droite verticale pointillée montre la valeur de N à partir de laquelle des difficultés numériques apparaissent lors de l'utilisation de la MDC. Au delà de cette valeur, les résultats sont obtenus par la MDA qui est l'objet de ce paragraphe.

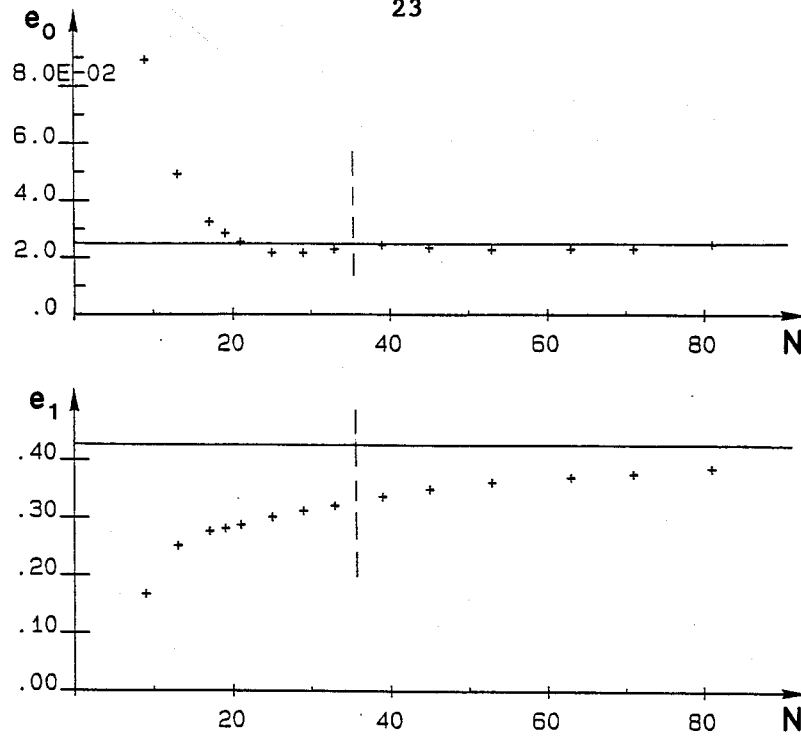


Figure 1.6. Idem que pour la fig.1.5 ; réseau en aluminium, $a = 0.2 \mu\text{m}$

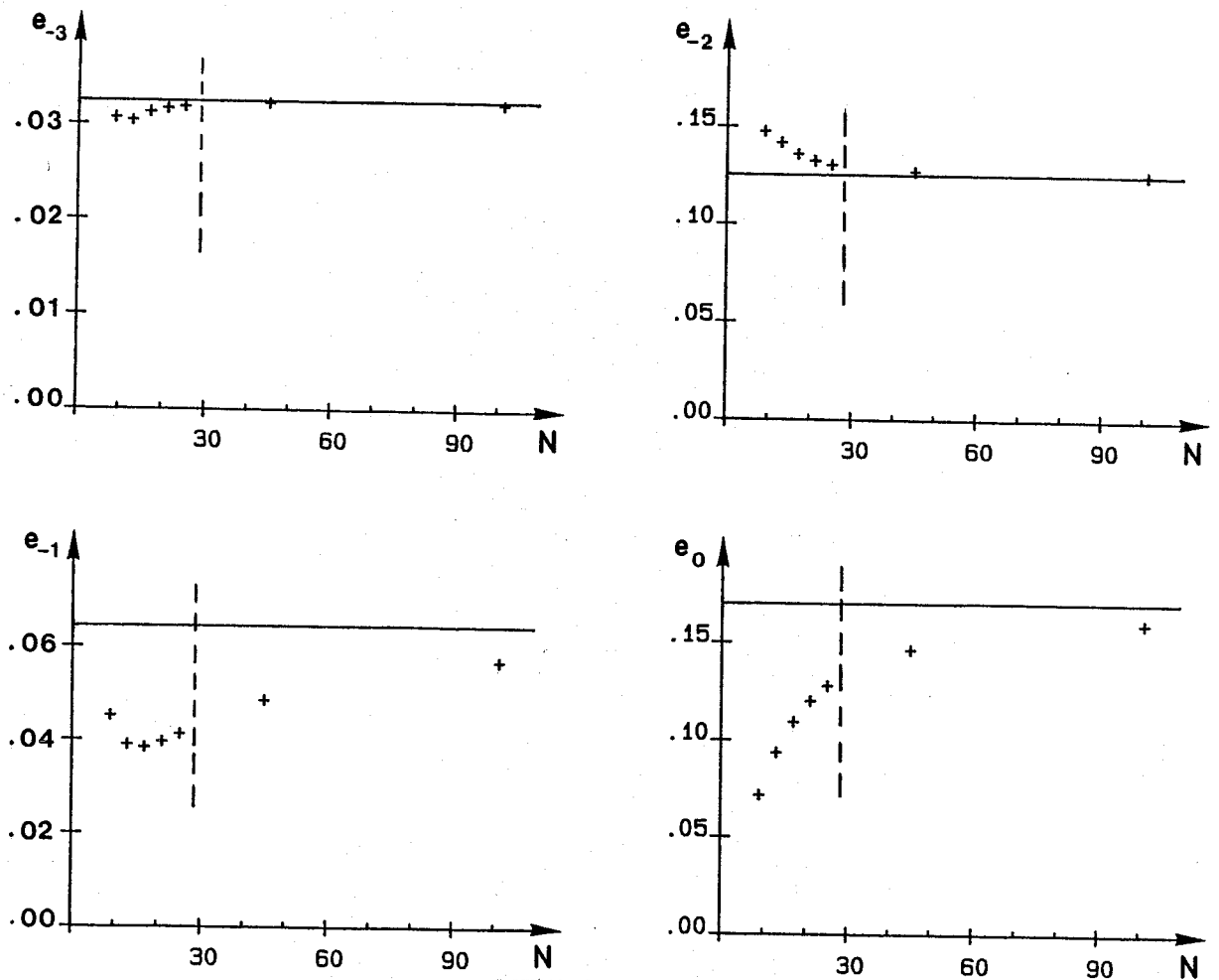


Figure 1.7. Comme pour la fig.1.5, réseau en or, $d = 1.1 \mu\text{m}$, $a = 0.22 \mu\text{m}$, $\lambda = 0.55 \mu\text{m}$, $n_2 = 0.32 + i 2.32$, $\theta = 38^\circ$ (résultats publiés dans [49]).

Pour tenter de venir à bout des difficultés numériques rencontrées avec la MDC, nous proposons une MDA qui en diffère uniquement par le traitement numérique du problème posé par les équations (1.26) à (1.28). Plus précisément, le problème consiste à déterminer une base de l'espace $\mathcal{E}_2^-(y)$ pour $0 \leq y \leq a$. Dans la MDC, cette base est obtenue par l'intégration de (1.29). Dans la MDA, à partir de la base orthogonale $\tau_k(0)$ de $\mathcal{E}_2^-(0)$, nous formons une base orthonormée. Nous dirons alors que la base de $\mathcal{E}_2^-(0)$ ainsi formée a un "degré d'orthogonalité" maximal. Si on intègre sans précaution particulière cette base orthonormée de $y = 0$ à $y = a$, on constate que les vecteurs obtenus ont un "degré d'orthogonalité" qui se dégrade au cours de la progression en y . L'idée de la MDA [72] est de tester le "degré d'orthogonalité" des vecteurs au cours de l'intégration. Lorsque ce "degré" est jugé trop faible, on stoppe l'intégration et on réorthonormalise la base grâce au procédé de Schmidt, avant de redémarrer l'intégration. Le procédé utilisé, ainsi que le critère retenu pour tester ce "degré d'orthogonalité", sont détaillés en Annexe 2. L'avantage de cette méthode est de conserver tout au long de l'intégration entre $y = 0$ et $y = a$ une bonne description de l'espace $\mathcal{E}_2^-(y)$. On s'affranchit alors des problèmes numériques rencontrés lors de l'utilisation de la MDC, ce qui permet d'utiliser, si nécessaire, de grandes valeurs de l'ordre de troncature N . En se reportant aux figures 1.5 à 1.7, on notera que les résultats donnés par la MDA convergent vers ceux donnés par la méthode intégrale. Bien évidemment, les temps de calcul deviennent rapidement très élevés au fur et à mesure que N croît, ce qui rend l'utilisation de la MDA malcommode. Nous pensons toutefois que cette étude est intéressante dans le sens où elle nous a permis de mieux comprendre l'origine des difficultés rencontrées lors de l'utilisation de la MDC. Nous sommes maintenant persuadés (et l'étude suivante par la méthode de l'intégration de la matrice impédance le confirme) que ces difficultés proviennent essentiellement de la convergence lente des développements du champ en séries de Fourier généralisées (1.4) en fonction de l'ordre de troncature N de ces séries.

1.3. INTEGRATION DE L'IMPEDANCE OU DE L'ADMITTANCE

Reconsidérons le problème exposé en 1.1.2 en introduisant un opérateur impédance $Z(y)$ défini de la façon suivante :

$$v(x,y) = Z(y) u(x,y), \quad (1.47)$$

avec toujours : $u(x,y) = H_z(x,y)$ et $v(x,y) = \frac{1}{\eta_0} E_x(x,y)$.

Il est à noter qu'avec nos notations, l'opérateur $Z(y)$ est sans dimension ; nous l'appelons par habitude opérateur impédance car il permet d'obtenir une composante du champ électrique à partir d'une composante du champ magnétique. Dans la base de Fourier généralisée $\psi_n(x)$, $u(x,y)$ et $v(x,y)$ sont représentés conformément à (1.4) par leurs composantes $u_n(y)$ et $v_n(y)$. Définissons les colonnes $U(y)$ et $V(y)$ constituées par les infinités de composantes $u_n(y)$ et $v_n(y)$. Nous pouvons alors représenter l'opérateur $Z(y)$ dans la base $\psi_n(x)$ par une matrice $Z(y)$ telle que :

$$V(y) = Z(y) U(y). \quad (1.48)$$

Le système d'équations (1.11) peut s'écrire sous la forme :

$$\begin{cases} \frac{dU(y)}{dy} = A_2(y) V(y), \\ \frac{dV(y)}{dy} = A_1(y) U(y), \end{cases} \quad (1.49)$$

$$(1.50)$$

où les matrices $A_1(y)$ et $A_2(y)$ sont connues, et dégènèrent en des matrices diagonales et indépendantes de y dans les zones homogènes $y < 0$ et $y > a$.

Pour $y \leq 0$, l'équation (1.22) montre que la matrice $Z(y)$ ne dépend pas de y et donne l'expression de $Z(0)$ telle que :

$$V(0) = Z(0) U(0). \quad (1.51)$$

Pour $y \geq a$, notons U^i et V^i les colonnes associées au champ incident, U^d et V^d les colonnes associées au champ diffracté :

$$U = U^i + U^d, \quad V = V^i + V^d. \quad (1.52)$$

L'équation (1.21) permet de lier $V^d(a)$ et $U^d(a)$ par une matrice connue (et diagonale) Z^d :

$$V^d(a) = Z^d U^d(a) \quad (1.53)$$

que l'on peut mettre sous la forme équivalente :

$$V(a) - V^i(a) = Z^d [U(a) - U^i(a)], \text{ ou bien}$$

$$[Z(a) - Z^d] U(a) = V^i(a) - Z^d U^i(a). \quad (1.54)$$

Dans la zone du réseau $0 \leq y \leq a$, les équations (1.48) à (1.50) donnent facilement :

$$\frac{dZ(y)}{dy} = A_1(y) - Z(y) A_2(y) Z(y). \quad (1.55)$$

La résolution du problème se fait de la façon suivante. Connaissant $Z(0)$, on intègre (1.55) pour obtenir $Z(y)$ pour $0 \leq y \leq a$. Connaissant alors $Z(a)$, la résolution du système (1.54) dans lequel $V^i(a)$ et $U^i(a)$ sont connus, donne le champ en dessus du réseau $U(a)$ et $V(a) = Z(a) U(a)$. Si l'on s'intéresse au champ transmis, il est nécessaire de procéder ensuite comme suit : on écrit (1.49) sous la forme $\frac{dU(y)}{dy} = A_2(y) Z(y) U(y)$ et, $Z(y)$ et $U(a)$ étant connus, on intègre ce système différentiel de $y = a$ à $y = 0$ pour déterminer $U(0)$ et $V(0) = Z(0) U(0)$.

En résumé, la méthode consiste à caractériser l'espace $\mathcal{E}_2^-(y)$ défini plus haut (§ 1.1.2) par un opérateur impédance $Z(y)$. Cet opérateur traduit une condition nécessaire et suffisante "entre les deux parties d'une colonne $F(y)$ " pour que cette colonne appartienne à $\mathcal{E}_2^-(y)$. Notons que cette façon de poser le problème est a priori différente de celle utilisée pour les MDC ou MDA dans lesquelles $\mathcal{E}_2^-(y)$ est caractérisé par une base. D'ailleurs, dès lors que les développements du champ en séries de Fourier généralisées (du type (1.4)) sont tronquées en gardant N termes, on peut remarquer que les MDC et MDA conduisent à l'intégration de N vecteurs à $2N$ composantes, alors que la méthode de "l'impédance" nécessite pour déterminer le champ réfléchi l'intégration des N^2 éléments de la matrice Z , et pour déterminer le champ transmis l'intégration des N composantes de U .

Du point de vue numérique, considérons le cas évoqué plus haut d'un réseau sinusoïdal en aluminium (Fig.1.4) avec les données : $d = 0.8333 \mu\text{m}$, $a = 0.1 \mu\text{m}$, $\lambda_0 = 0.6 \mu\text{m}$, $n_2 = 1.3 + i 7.1$, $\theta = 0$, polarisation H//. Comme pour les MDC et MDA, l'intégration numérique est réalisée au moyen de l'algorithme d'Adams-Moulton (5^{ème} ordre, avec prédiction et correction), le procédé étant initialisé par la méthode de Runge-Kutta d'ordre 4. Nous avons constaté que l'intégration est plus délicate qu'avec les MDC et MDA, et nécessite un plus grand nombre de pas d'intégration. Supposons par

exemple que l'on conserve $N = 33$ coefficients dans les développements des champs. Avec un nombre de pas d'intégration $J = 100$, la méthode d'intégration de la matrice impédance (M.I.M.I.) donne des résultats erronés dus à la divergence du procédé d'intégration numérique. Il faut resserrer les pas d'intégration (par exemple en prenant dans ce cas $J = 200$), et on obtient alors des résultats identiques à ceux des MDC et MDA. Pour fixer les idées, signalons que, pour le même problème, 50 pas d'intégration sont suffisants pour les MDC et MDA.

Intégration de la matrice admittance.

Nous avons pensé que les problèmes numériques rencontrés ici pouvaient être la conséquence du fait que l'opérateur $\mathbb{Z}$ n'est pas borné. Cela se constate facilement, du moins dans la zone homogène $y < 0$ dans laquelle l'expression de Z s'obtient de façon explicite : $Z y$ est diagonale, et les éléments diagonaux, proportionnels à β'_n (eq.(1.22)), ont un module qui croît comme $\ln|n|$ (1.15) pour de grandes valeurs de n . C'est la raison pour laquelle nous avons aussi essayé d'utiliser l'opérateur admittance, inverse de l'impédance, défini par :

$$u(x,y) = \Psi(y) v(x,y). \quad (1.56)$$

Dans la base des $\psi_n(x)$, il est représenté par une matrice $Y(y)$ telle que :

$$U(y) = Y(y) V(y). \quad (1.57)$$

Un raisonnement analogue à celui présenté plus haut nous ramène à l'intégration de la matrice admittance qui vérifie :

$$\frac{dY(y)}{dy} = A_2(y) - Y(y) A_1(y) Y(y), \quad (1.58)$$

connaissant $Y(0)$. Malheureusement, l'expérience numérique montre que la méthode d'intégration de la matrice admittance conduit à des instabilités numériques (lors de l'intégration de (1.58)) plus importantes que la MIMI. Il faut donc utiliser un nombre de pas d'intégration encore plus important que pour la MIMI, ce qui augmente les temps de calculs. Nous n'avons donc trouvé aucun avantage dans cette étude à raisonner sur l'admittance plutôt que sur l'impédance.

En conclusion, et après le verdict de l'étude numérique, ces méthodes

d'intégration des matrices impédance ou admittance ne semblent pas être d'un grand intérêt pratique, puisqu'elles conduisent aux mêmes résultats que les autres méthodes différentielles au bout de temps de calcul généralement plus longs. Signalons au passage que les MDC et MDA permettent elles aussi, s'il en est besoin, de déterminer la matrice impédance $Z(y)$ à partir d'une base de $\mathcal{E}_2^-(y)$ (Annexe 3).

Néanmoins, ces études présentent à notre avis un intérêt certain. Nous avons vérifié que, pour une valeur N donnée (N est, rappelons le, le nombre de termes conservés pour le traitement numérique dans les séries de Fourier généralisées (1.4) qui représentent les champs), toutes les méthodes différentielles que nous avons employées donnent des résultats identiques dans la mesure où l'intégration numérique "se passe bien". Les difficultés rencontrées sont donc dues au fait que, dans certains cas "difficiles" (réseaux profonds, matériaux très conducteurs, polarisation H// en particulier), ces séries de Fourier convergent lentement en fonction de N , ce qui entraîne des temps de calcul prohibitifs.

Signalons enfin qu'une étude très récente menée dans notre Laboratoire par M. Nevière [45] corrobore nos conclusions. Lors de l'utilisation de la M.D.C., on peut en effet s'affranchir des problèmes numériques grâce à l'emploi de la quadruple précision (sur les calculateurs usuels qui ne connaissent la quadruple précision que pour des nombres réels, il faut alors séparer les parties réelles et imaginaires de tous les nombres complexes utilisés !). Signalons que pour notre part, les calculs présentés ici ont été menés en double précision. Cela montre bien la nature essentiellement numérique des difficultés rencontrées.

2.1. INTRODUCTION

L'idée de la "méthode de synthèse" est de représenter le champ diffracté par un champ approché (l'approximation pouvant être réalisée avec une précision quelconque définie par avance, pourvu que "l'on y mette le prix"). Ce champ approché est créé par des sources fictives convenablement choisies, d'où le nom de la méthode. Pour faire un jeu de mots, nous pourrions dire aussi qu'il s'agit d'une méthode de synthèse en un tout autre sens ; il s'avère en effet qu'elle englobe un certain nombre de méthodes jusqu'ici proposées pour l'étude des réseaux, notamment les "improved point matching methods" introduites par W.C. Meecham [41] développées par K. Yasuura et ses collaborateurs [26] et qui ont suscité plusieurs études [42, 87, 25, 50]. Cette idée a aussi été utilisée récemment par A. Boag, Y. Leviatan et A. Boag [7] et même étendue par ces auteurs au cas de structures bidimensionnelles [8]. Ses fondements théoriques s'appuient sur des considérations d'analyse fonctionnelle récemment exposées par R. Petit et M. Cadilhac [61]. Le lecteur intéressé par les démonstrations de certains résultats que nous utiliserons ici pourra s'y reporter.

2.2. DESCRIPTION DU PROBLEME

Nous nous proposons d'étudier le réseau décrit sur la figure 2.1. La surface $y = f(x)$ (f est périodique de période d , le pas du réseau) sépare les domaines Ω^+ ($y > f(x)$) et Ω^- ($y < f(x)$). Nous nous restreignons aux cas où $f(x)$ est indéfiniment dérivable (réseau sinusoïdal par exemple). Cette hypothèse nous est utile d'un point de vue théorique pour démontrer certains résultats dont nous aurons besoin. Mais sans aucun doute, d'un point de vue pratique, l'expérimentation numérique pourra être tentée pour d'autres types de réseaux (réseaux échelonnées par exemple, cf. [7]). Les domaines Ω^+ et Ω^- sont respectivement remplis d'un diélectrique d'indice n_1 réel (permittivité $\epsilon_1 = n_1^2$; $k_1^2 = k_0^2 n_1^2$) et d'un diélectrique à pertes d'indice n_2 complexe (permittivité $\epsilon_2 = n_2^2$; $k_2^2 = k_0^2 n_2^2$). Nous notons $\mathcal{P}$ le

profil du réseau constitué par la courbe $y = f(x)$.

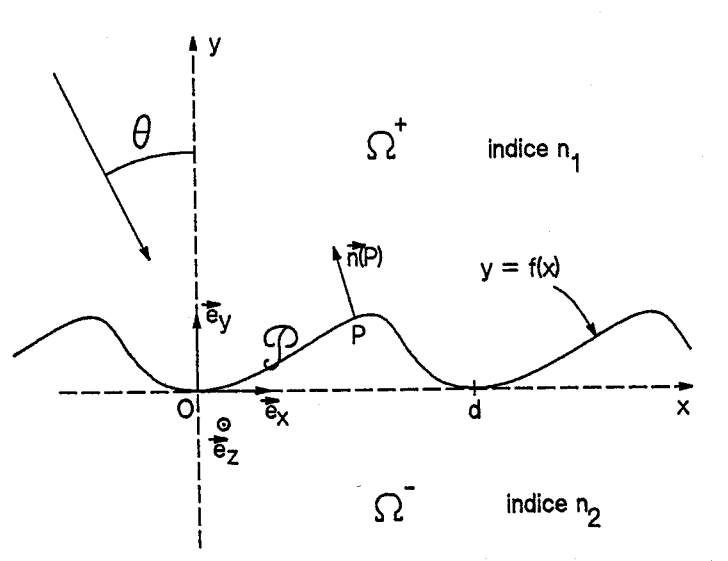


Figure 2.1. La normale en P au profil $\vec{n}(P)$ est dirigée vers Ω^+ .

Le réseau est éclairé sous l'incidence θ par une onde plane. Nous utilisons la notation suivante qui permet d'étudier simultanément les deux cas de polarisation E// et H// : dans le cas E//, le champ électrique complexe $\vec{E}$ est parallèle aux sillons du réseau et nous posons $\vec{E}_z(x, y) = u(x, y) \vec{e}_z$; dans le cas H//, c'est $\vec{H}$ qui est parallèle aux sillons et nous posons $\vec{H}_z(x, y) = u(x, y) \vec{e}_z$. Le champ incident est donc, avec les notations déjà utilisées au chapitre précédent ($\alpha_0 = k_1 \sin \theta$, $\beta_0 = \sqrt{k_1^2 - \alpha_0^2}$) :

$$u^i(x, y) = \exp(i\alpha_0 x - i\beta_0 y) . \quad (2.1)$$

Dans le domaine Ω^+ , nous définissons le champ diffracté u^d par :

$$u^d(x, y) = u(x, y) - u^i(x, y) . \quad (2.2)$$

Le problème consiste à déterminer la fonction u vérifiant :

$$a) \quad \Delta u + k_1^2 u = 0 \quad \text{dans } \Omega^+ , \quad (2.3)$$

$$b) \quad \Delta u + k_2^2 u = 0 \quad \text{dans } \Omega^- \quad (2.4)$$

c) le champ diffracté vérifie une condition de rayonnement (ou C.O.S. pour condition d'onde sortante), c'est-à-dire :

$$u^d = u - u^i \text{ vérifie une C.O.S. quand } y \rightarrow +\infty \quad (2.5)$$

$$u \text{ vérifie une C.O.S. quand } y \rightarrow -\infty \quad (2.6)$$

d) et les conditions de continuité sur $\mathcal{P}$:

$$u^+ = u^- , \quad (2.7)$$

$$Du^+ = Du^- \quad \text{dans le cas E//,} \quad (2.8)$$

$$\frac{1}{\varepsilon_1} Du^+ = \frac{1}{\varepsilon_2} Du^- \quad \text{dans le cas H//.} \quad (2.9)$$

On a adopté les notations suivantes : u^+ (resp. u^-) désigne la limite de u quand le point d'observation tend vers le profil $\mathcal{P}$ tout en restant dans Ω^+ (resp. Ω^-). Soit P un point de $\mathcal{P}$. La fonction Du^+ est définie sur $\mathcal{P}$ par :

$$Du^+(P) = \lim_{\substack{M \rightarrow P \\ M \in \Omega^+}} \vec{n}(P) \cdot \overrightarrow{\text{grad}}_M u(M) . \quad (2.10)$$

La fonction Du^- est définie de même en remplaçant Ω^+ par Ω^- , en conservant la même définition de la normale (toujours orientée vers Ω^+).

2.3. UNE PROPRIÉTÉ DES CHAMPS DIFFRACTES

Nous appelons "champ diffracté vers le haut" une fonction $\varphi(x,y)$ vérifiant les propriétés :

a) $\varphi(x,y)$ est pseudo-périodique en x , de période d et de coefficient de pseudo-périodicité a (ce qui signifie qu'il existe un nombre complexe a de module 1 tel que $\varphi(x+d, y) = a \varphi(x,y)$).

b) φ vérifie l'équation de Helmholtz $\Delta \varphi + k^2 \varphi = 0$ dans Ω^+ .

c) φ vérifie une C.O.S. quand $y \rightarrow +\infty$.

Nous rappelons la propriété suivante bien connue (elle constitue le point de départ des méthodes intégrales) : $\varphi(x,y)$ est parfaitement déterminée

dans Ω^+ dès lors que l'on connaît la valeur de φ^+ et de $D\varphi^+$ sur $\mathcal{P}$ (les fonctions φ^+ et $D\varphi^+$ ne sont d'ailleurs pas indépendantes. On peut même montrer que la connaissance d'une seule d'entre elles suffit pour déterminer $\varphi(x,y)$, cf. annexe 4).

Bien entendu, $u^d(x,y)$, qui est pseudo-périodique avec le coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$ imposé par la structure et l'onde incidente, qui vérifie dans Ω^+ l'équation $\Delta u^d + k_1^2 u^d = 0$ et une C.O.S. pour $y \rightarrow +\infty$, est un "champ diffracté vers le haut". De même, $u(x,y)$, qui est pseudo-périodique de la même façon, qui vérifie dans Ω^- l'équation $\Delta u + k_2^2 u = 0$ et une C.O.S. pour $y \rightarrow -\infty$, est un "champ diffracté vers le bas". Par conséquent, le problème posé au paragraphe précédent sera résolu quand u^{d+} et Du^{d+} (qui donnent le champ réfléchi) ainsi que u^- et Du^- (qui donnent le champ transmis) seront connues. Il apparaît donc qu'on peut caractériser les champs diffractés u^d et u par les colonnes $\begin{vmatrix} u^{d+} \\ Du^{d+} \end{vmatrix}$ et $\begin{vmatrix} u^- \\ Du^- \end{vmatrix}$ définies sur $\mathcal{P}$. Il est commode pour la suite de poser :

$$\text{dans le cas E// : } F^i = \begin{vmatrix} u^i \\ Du^i \end{vmatrix}, \quad F^d = \begin{vmatrix} u^{d+} \\ Du^{d+} \end{vmatrix}, \quad F = \begin{vmatrix} u^- \\ Du^- \end{vmatrix}, \quad (2.11)$$

$$\text{dans le cas H// : } F^i = \begin{vmatrix} u^i \\ \frac{1}{\varepsilon_1} Du^i \end{vmatrix}, \quad F^d = \begin{vmatrix} u^{d+} \\ \frac{1}{\varepsilon_1} Du^{d+} \end{vmatrix}, \quad F = \begin{vmatrix} u^- \\ \frac{1}{\varepsilon_2} Du^- \end{vmatrix}, \quad (2.12)$$

2.4. LES ESPACES $\mathcal{V}$, $\mathcal{V}_1^+$, $\mathcal{V}_2^-$

Soit $\mathcal{V}$ l'espace des colonnes $\begin{vmatrix} u \\ v \end{vmatrix}$, où u et v sont des fonctions définies sur $\mathcal{P}$, pseudo-périodiques (coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$). Pour être précis, les considérations théoriques précisées dans [61, 13] montrent qu'il est judicieux (pour des raisons d'existence et d'unicité relatives au problème posé) de choisir u dans l'espace de Sobolev $H_{\mathcal{P}}^1$ (espace des fonctions de carré sommable sur une période de $\mathcal{P}$ ainsi que leur dérivée première au sens des distributions par rapport à l'abscisse curviligne sur $\mathcal{P}$) et v dans $H_{\mathcal{P}}^0$ (qui est aussi $L_{\mathcal{P}}^2$: l'espace des fonctions de carré sommable sur une période de $\mathcal{P}$).

Nous définissons aussi l'espace $\mathcal{V}_1^+$, sous espace de $\mathcal{V}$ constitué par les

colonnes $\begin{vmatrix} u \\ Du \end{vmatrix}$ (cas E//) ou $\begin{vmatrix} u \\ \frac{1}{\varepsilon_1} Du \end{vmatrix}$ (cas H//), associées à des champs "diffractés vers le haut" dans le milieu d'indice n_1 . De même l'espace $\mathcal{V}_2^-$ est le sous espace de $\mathcal{V}$ constitué par les colonnes $\begin{vmatrix} u \\ Du \end{vmatrix}$ (cas E//) ou $\begin{vmatrix} u \\ \frac{1}{\varepsilon_2} Du \end{vmatrix}$ (cas H//) associées à des champs "diffractés vers le bas" dans le milieu d'indice n_2 . Les colonnes F^d et F définies au paragraphe précédent appartiennent donc respectivement à $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$.

En utilisant ces espaces, le problème posé au paragraphe 2.2 (équations (2.2) à (2.9)) se réduit à la recherche de F vérifiant :

$$F - F^i \in \mathcal{V}_1^+, \quad (2.13)$$

$$F \in \mathcal{V}_2^-. \quad (2.14)$$

On est ainsi ramené à la recherche de l'intersection des espaces $\mathcal{V}_2^-$ et d'un translaté de $\mathcal{V}_1^+$ (soit $\mathcal{V}_1^+ + F^i$). D'après les théorèmes d'existence et d'unicité, cette intersection se réduit à une colonne qui est l'inconnue F . Dans la méthode de synthèse, nous allons tâcher de représenter ces espaces par des familles totales (les espaces $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$ étant de dimension infinie, il est préférable d'éviter le terme de "bases"). Pour des raisons de commodité qui apparaîtront ensuite, nous réécrivons (2.13 - 2.14) sous la forme équivalente suivante : on cherche F vérifiant :

$$F^d \in \mathcal{V}_1^+, \quad (2.15)$$

$$F \in \mathcal{V}_2^-, \quad (2.16)$$

$$F^i + F^d - F = 0. \quad (2.17)$$

2.5. CHOIX DE FAMILLES TOTALES DANS $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$.

Notre but ici n'est pas de faire une étude exhaustive du choix de familles totales de ces espaces, dont il existe une "infinie" variété. Le lecteur intéressé pourra pour cela se reporter à la référence [61] qui traite de ce sujet d'un point de vue théorique. D'un point de vue pratique, nous avons réalisé deux études numériques utilisant deux choix légèrement différents de familles totales. Nous avons décidé de ne décrire que l'une

d'entre elles, qui est sans doute celle la plus simple à mettre en oeuvre.

Les espaces $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$ étant de même nature, nous porterons tout d'abord notre attention sur $\mathcal{V}_1^+$. Considérons la famille de distributions pseudo-périodiques $S_n(x,y)$ (il apparaîtra dans un instant que ces distributions pourront être considérées comme des "sources" de champ diffracté, ce qui justifie l'emploi de la lettre S) :

$$\begin{aligned} S_n(x,y) &= \sum_{m \in \mathbb{Z}} \frac{1}{d} \delta(y - y_{0n}) \exp[i\alpha_m(x - x_{0n})] \\ &= \sum_{m \in \mathbb{Z}} \delta(y - y_{0n}) \delta(x - x_{0n} - md) \exp[i\alpha_0 md] , \end{aligned} \quad (2.18)$$

où le point de coordonnées (x_{0n}, y_{0n}) est placé de telle sorte que $0 \leq x_{0n} < d$ et $y_{0n} < f(x_{0n})$ (figure 2.2).

Figure 2.2. Le caractère * représente la position du point de coordonnées (x_{0n}, y_{0n}) .

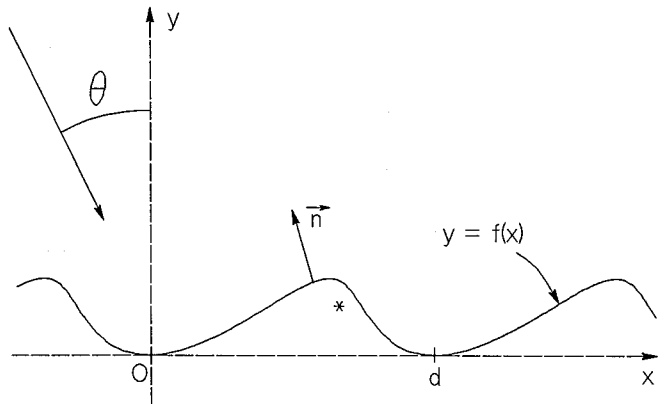
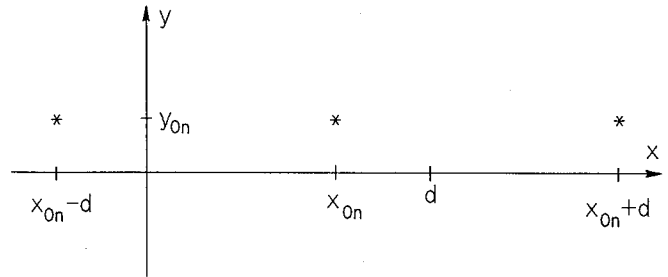


Figure 2.3. Localisation des fils (parallèles à $\vec{e}_z$), pouvant être interprétés comme le support des courants de la distribution $(i/\omega \mu_0) S_n \vec{e}_z$.



Soit $g_n(x,y)$ la solution vérifiant une C.O.S. quand $y \rightarrow \pm\infty$ de :

$$\Delta g_n + k_1^2 g_n = S_n . \quad (2.19)$$

Cette solution est bien connue en théorie des réseaux [58] et vaut :

$$g_n(x, y) = \frac{1}{2id} \sum_{m \in \mathbb{Z}} \frac{1}{\beta_m} \exp[i\alpha_m(x - x_{0n}) + i\beta_m|y - y_{0n}|] . \quad (2.20)$$

D'un point de vue plus concret, se rappelant de l'équation vérifiée au sens des distributions en régime harmonique par un champ électrique dans un milieu homogène d'indice n_1 et en présence de courants $\vec{J}$ ($\Delta \vec{E} + k_1^2 \vec{E} = -i\omega \mu_0 \vec{J}$), on notera que g_n peut être interprétée comme la composante selon $\vec{e}_z$ du champ électrique rayonné par le réseau de fils G_n représenté sur la figure 2.3, parcourus par des courants convenablement déphasés de façon à vérifier la condition de pseudo-périodicité imposée par le réseau étudié et le champ incident utilisé (ces courants étant représentés par la distribution $(i/\omega \mu_0) S_n \vec{e}_z$). La fonction g_n , qui vérifie l'équation de Helmholtz dans le complémentaire du support de ces fils, y est indéfiniment dérivable et y possède un gradient. On peut donc aisément calculer la valeur de sa dérivée normale Dg_n sur $\mathcal{P}$:

$$Dg_n(x, f(x)) = \frac{1}{2d \sqrt{1 + f'(x)^2}} \sum_{m \in \mathbb{Z}} \left[-f'(x) \frac{\alpha_m}{\beta_m} + \text{sign}(f(x) - y_{0n}) \right] \times \exp[i\alpha_m(x - x_{0n}) + i\beta_m|f(x) - y_{0n}|] . \quad (2.21)$$

Soit maintenant la colonne $\Phi_{1,n}^+$ définie sur $\mathcal{P}$ par :

$$\Phi_{1,n}^+ = \begin{vmatrix} \varphi_n \\ \psi_n \end{vmatrix} \quad \text{en polarisation E//} \quad (2.22)$$

$$\Phi_{1,n}^+ = \begin{vmatrix} \varphi_n \\ \frac{1}{\varepsilon_1} \psi_n \end{vmatrix} \quad \text{en polarisation H//,}$$

$$\text{avec } \varphi_n(x) = g_n(x, f(x)), \quad (2.23)$$

$$\text{et } \psi_n(x) = Dg_n(x, f(x)). \quad (2.24)$$

La colonne $\Phi_{1,n}^+$ représente ainsi les valeurs sur $\mathcal{P}$ du champ et de sa dérivée normale (au coefficient $1/\varepsilon_1$ près dans le cas H//) rayonné par le réseau de fils G_n . Elle appartient à $\mathcal{V}_1^+$. Du point de vue numérique, les $\Phi_{1,n}^+$ s'obtiennent en sommant les séries (2.20, 2.21). En vue d'accélérer la rapidité de convergence de ces séries, il est conseillé d'effectuer des

développements asymptotiques du facteur $\exp(i\alpha_m(x - x_{0n}) + i\beta_m|f(x) - y_{0n}|)$, ainsi que des regroupements de termes. Ces "astuces de calcul" sont analogues à celles précisées dans l'annexe 13.

Considérons une famille de distributions $S_n(x,y)$ ou, de façon équivalente, une famille de points (x_{0n}, y_{0n}) . On montre (annexe 5, corollaire 3) que si ces points sont distribués de façon dense sur une courbe $\mathcal{P}^-$ indéfiniment dérivable située au dessous du profil $\mathcal{P}$ (figure 2.4), la famille $\Phi_{1,n}^+$ associée aux S_n constitue une famille totale de $\mathcal{V}_1^+$. La démonstration de cette propriété (annexe 5), qui n'intéressera peut être guère le physicien, a été un travail difficile que nous n'aurions pu mener à bien sans l'aide de Mr. M. CADILHAC [10].

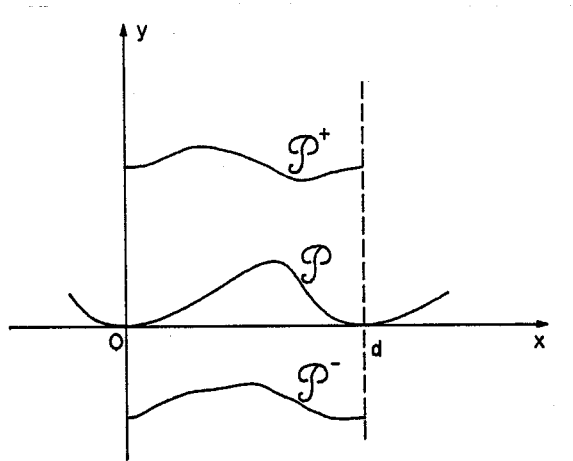


Figure 2.4. $\mathcal{P}^-$ et $\mathcal{P}^+$ sont les restrictions à l'intervalle période $[0,d]$ des courbes $y = f^-(x)$ et $y = f^+(x)$; f^- et f^+ sont indéfiniment dérivables.

Pour obtenir une famille totale de $\mathcal{V}_2^-$, nous procédons de la même manière, en remplaçant dans ce qui précède n_1 par n_2 (donc k_1 par k_2 , ε_1 par ε_2 , β_m par β_m' défini en (1.15)) et en choisissant les points (x_{0n}, y_{0n}) sur une courbe $\mathcal{P}^+$ maintenant située au dessus de $\mathcal{P}$ (figure 2.4). On détermine ainsi une famille totale de $\mathcal{V}_2^-$ que nous notons $\Phi_{2,n}^-$.

2.6. RECHERCHE D'UN APPROXIMANT DE F

On a vu (§ 2.4) que F est l'intersection de $\mathcal{V}_2^-$ et de $\mathcal{V}_1^+ + F^i$. Pour des raisons évidentes liées au calcul numérique, on est amené à considérer les espaces $\mathcal{V}_{1,N}^+$ et $\mathcal{V}_{2,N}^-$ de dimension finie N , engendrés par N colonnes

$\Phi_{1,n}^+$ et N colonnes $\Phi_{2,n}^-$. En général, les espaces $\mathcal{V}_{2,N}^-$ et $\mathcal{V}_{1,N}^+ + F^i$ n'ont pas d'intersection, et on recherche alors les coefficients $C_n^+(N)$ et $C_n^-(N)$ qui minimisent la distance entre ces espaces en cherchant le minimum de :

$$\Delta_N = \left\| F^i + \sum_{n=1}^N C_n^+(N) \Phi_{1,n}^+ - \sum_{n=1}^N C_n^-(N) \Phi_{2,n}^- \right\|_{\mathcal{V}}, \quad (2.25)$$

où $\| \cdot \|_{\mathcal{V}}$ représente la norme définie dans l'espace $\mathcal{V}$, norme que nous précisons au paragraphe suivant.

Posons :

$$F_N^d = \sum_{n=1}^N C_n^+(N) \Phi_{1,n}^+, \quad (2.26)$$

$$F_N = \sum_{n=1}^N C_n^-(N) \Phi_{2,n}^-, \quad (2.27)$$

ce qui permet de réécrire (2.25) sous la forme :

$$\Delta_N = \left\| F^i + F_N^d - F_N \right\|_{\mathcal{V}}. \quad (2.25')$$

Les coefficients $C_n^+(N)$ et $C_n^-(N)$ sont obtenus en résolvant un système "d'équations normales" (voir par exemple [21], p. 251 et suivantes). Ce système est lui-même obtenu en écrivant que pour que Δ_N soit minimal, $F^i + F_N^d - F_N$ doit être orthogonal (au sens de $\mathcal{V}$, cf. § 2.7) à tous les $\Phi_{1,n}^+$ et $\Phi_{2,n}^-$, $n = 1, N$. En d'autres termes, $F_N - F_N^d$ n'est autre que la projection orthogonale de F^i sur l'espace engendré par les $2N$ vecteurs $\Phi_{1,n}^+$ et $\Phi_{2,n}^-$ (figure 2.5).

En résumé, $F_N^d \in \mathcal{V}_1^+$, $F_N \in \mathcal{V}_2^-$, et $F^i + F_N^d - F_N \xrightarrow{N \rightarrow \infty} 0$, ce qui montre (équations (2.15), (2.16), (2.17)) que F_N^d constitue une valeur approchée de F^d , que F_N constitue une valeur approchée de F et que ces valeurs approchées tendent vers les valeurs exactes (au sens de la norme de $\mathcal{V}$) lorsque $N \rightarrow \infty$. Nous disposons de plus d'un critère (valeur de Δ_N) qui permet de contrôler la qualité de ces approximations, ce qui est précieux pour le traitement numérique du problème.

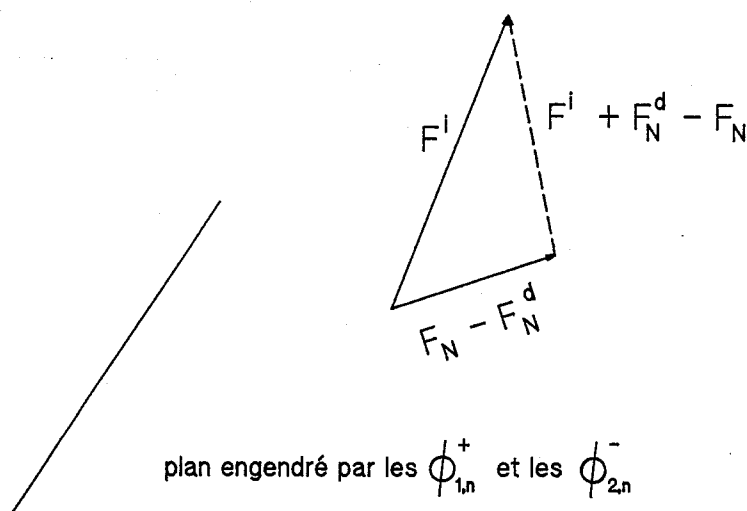


Figure 2.5. On suppose pour illustrer la projection que $N = 1$; l'espace engendré par les $\phi_{1,n}^+$ et les $\phi_{2,n}^-$ est alors réduit à un plan. Les coefficients C_n^+ et C_n^- sont obtenus en projetant orthogonalement F^i sur cet espace, ce qui minimise $\Delta_N = \|F^i + F_N^d - F_N\|_{\mathcal{V}}$.

Nous terminerons ce paragraphe par la remarque suivante : le coefficient $C_n^+(N)$ peut être interprété comme la contribution du réseau de fils G_n au champ diffracté dans Ω^+ . Au cours de l'étude numérique, il pourrait s'avérer que certains des coefficients $C_n^+(N)$ aient une valeur sensiblement supérieure aux autres ; il serait alors sans doute utile de placer d'autres "fils source" au voisinage du point (x_{0n}, y_{0n}) . Nous n'avons pas rencontré ce cas dans nos études numériques, car nous nous sommes limités à l'étude de réseaux sinusoïdaux. Mais il serait sans doute utile de procéder ainsi pour des réseaux présentant des arêtes. Cela rejoint les conclusions de Boag et coauteurs [7] qui, lors de l'étude de réseaux échellette, ont été amenés à resserrer le nombre de leurs "sources" au voisinage des arêtes du réseau.

2.7. PRODUIT SCALAIRE DANS $\mathcal{V}$

Soient $F = \begin{pmatrix} u \\ v \end{pmatrix}$ et $F' = \begin{pmatrix} u' \\ v' \end{pmatrix}$ deux colonnes appartenant à $\mathcal{V}$. Il résulte de la définition de $\mathcal{V}$ (§ 2.4) que le produit scalaire dans $\mathcal{V}$ se déduit des produits scalaires dans $H_{\mathcal{D}}^1$ et $H_{\mathcal{D}}^0$ par :

$$(F|F')_{\mathcal{V}} = (u|u')_{H_{\mathcal{P}}^1} + (v|v')_{H_{\mathcal{P}}^0}, \quad (2.28)$$

et que bien entendu la norme dans $\mathcal{V}$ est définie par :

$$\|F\|_{\mathcal{V}}^2 = (F|F)_{\mathcal{V}}. \quad (2.29)$$

La définition du produit scalaire dans $H_{\mathcal{P}}^0 = L_{\mathcal{P}}^2$ est :

$$(v|v')_{H_{\mathcal{P}}^0} = \int_{\mathcal{P}} v \overline{v'} \, d\ell, \quad (2.30)$$

l'intégrale curviligne étant réalisée sur une période de $\mathcal{P}$.

Le produit scalaire attaché à $H_{\mathcal{P}}^1$ est :

$$(u|u')_{H_{\mathcal{P}}^1} = \int_{\mathcal{P}} u \overline{u'} \, d\ell + \int_{\mathcal{P}} \frac{du}{d\ell} \frac{d \overline{u'}}{d\ell} \, d\ell, \quad (2.31)$$

où $\frac{du}{d\ell}$ désigne la dérivée de u par rapport à l'abscisse curviligne.

Pour simplifier l'écriture, nous adoptons les notations suivantes :

$$(F|F')_{\mathcal{V}} = P_1 + P_2 + P_3, \quad (2.32)$$

$$\|F\|_{\mathcal{V}}^2 = N_1 + N_2 + N_3, \quad (2.33)$$

avec :

$$\begin{aligned} P_1 &= \int_{\mathcal{P}} u \overline{u'} \, d\ell, & P_2 &= \int_{\mathcal{P}} \frac{du}{d\ell} \frac{d \overline{u'}}{d\ell} \, d\ell, & P_3 &= \int_{\mathcal{P}} v \overline{v'} \, d\ell, \\ N_1 &= \int_{\mathcal{P}} |u|^2 \, d\ell, & N_2 &= \int_{\mathcal{P}} \left| \frac{du}{d\ell} \right|^2 \, d\ell, & N_3 &= \int_{\mathcal{P}} |v|^2 \, d\ell. \end{aligned}$$

Les études numériques ont montré qu'il peut être utile d'utiliser un produit scalaire légèrement différent de la forme :

$$(F|F') = P_1 + C_2 P_2 + C_3 P_3, \quad (2.34)$$

où C_2 et C_3 sont des constantes réelles positives qui peuvent être

considérées comme des "poids" affectés aux différentes parties de ce produit scalaire. Nous avons bien entendu, au cours de nos essais, utilisé les poids $C_2 = 1$ et $C_3 = 1$ qui sont ceux déduits des normes H_ϕ^1 (pour u) et H_ϕ^0 (pour v).

Une autre idée est de déterminer les poids C_2 et C_3 au moyen des valeurs moyennes des normes N_1 , N_2 et N_3 des différents vecteurs $\Phi_{1,n}^+$ et $\Phi_{2,n}^-$; nous définissons des poids "moyennés" C_{2m} et C_{3m} par :

$$C_{2m} = \frac{\sum N_1}{\sum N_2}, \quad C_{3m} = \frac{\sum N_1}{\sum N_3} \quad (2.35)$$

où $\sum N_1$ représente la sommation sur tous les vecteurs $\Phi_{1,n}^+$ et $\Phi_{2,n}^-$ des normes N_1 associées à ces vecteurs (idem pour $\sum N_2$ et $\sum N_3$). Les poids C_{2m} et C_{3m} permettent ainsi "d'équilibrer" dans le produit scalaire (2.34) l'importance de P_1 , P_2 et P_3 .

Les conclusions de nos études numériques sont les suivantes : dans le but d'obtenir une convergence rapide de la solution vers la solution exacte lorsque N croît :

* le coefficient C_2 semble en fait n'avoir que peu d'importance ; les résultats obtenus dans les cas étudiés jusqu'ici ont été sensiblement les mêmes avec $C_2 = 0$, $C_2 = 1$, $C_2 = C_{2m}$. On peut donc essayer pour l'étude numérique de choisir $C_2 = 0$, ce qui, en cas de succès, améliore la rapidité des calculs,

* il est souvent préférable d'utiliser $C_3 = C_{3m}$ plutôt que $C_3 = 1$; bien entendu, la valeur $C_3 = 0$ conduit à des résultats erronés (si $C_3 = 0$, les conditions de continuité (2.8) et (2.9) ne sont pas prises en compte).

Il est à noter que ces différents choix de C_2 et C_3 n'influencent que sur la rapidité de convergence de la solution, mais que dans tous les cas, les solutions (nous nous sommes, en fait, intéressés principalement aux efficacités diffractées dans les différents ordres) convergent vers les mêmes valeurs.

Signalons enfin que les intégrales P_1 , P_2 et P_3 sont évaluées au moyen de sommes de Riemann ; le nombre de points de discrétisation ND devra être suffisant pour ne pas influencer sur les résultats finaux.

2.8. CALCUL DES EFFICACITES

Ici encore, l'obtention des champs diffractés dans Ω^+ (champ réfléchi) et dans Ω^- (champ transmis) sont analogues. Nous commencerons par nous intéresser au champ réfléchi $u^d(x,y)$ défini par (2.2).

A ce stade de l'étude, nous avons obtenu par l'intermédiaire des coefficients $C_n^+(N)$ une valeur approchée u_N^d de la valeur sur $\mathcal{P}$ de u^d , sous la forme d'une combinaison linéaire (de coefficients $C_n^+(N)$) des champs rayonnés par les réseaux de fils G_n situés dans Ω^- :

$$u_N^d(x, f(x)) = \sum_{n=1}^N C_n^+(N) g_n(x, f(x)) = \sum_{n=1}^N C_n^+(N) \varphi_n(x) . \quad (2.36)$$

Nous savons (§ 2.5 et annexe 5) que les φ_n constituent une famille totale dans $L^2_{\mathcal{P}}$, et que (§ 2.6) $u_N^d(x, f(x)) \xrightarrow{N \rightarrow \infty} u^d(x, f(x))$. Il en résulte (Annexe 4, propriété 3) que pour $y > f(x)$:

$$u^d(x,y) = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n^+(N) g_n(x,y) , \quad (2.37)$$

ce qui montre que :

$$u_N^d(x,y) = \sum_{n=1}^N C_n^+(N) g_n(x,y) \quad (2.38)$$

constitue un approximant du champ diffracté dans Ω^+ . L'équation (2.20) permet d'écrire pour $y > f(x)$:

$$g_n(x,y) = \frac{1}{2id} \sum_{m=-\infty}^{+\infty} \frac{1}{\beta_m} \exp[i\alpha_m(x - x_{0n}) + i\beta_m(y - y_{0n})] . \quad (2.39)$$

On recherche les "coefficients de Rayleigh" R_m du champ diffracté définis pour $y > \max(f(x))$ par :

$$u^d(x,y) = \sum_{m=-\infty}^{+\infty} R_m \exp[i\alpha_m x + i\beta_m y] . \quad (2.40)$$

Utilisant (2.37), (2.39) et (2.40), il apparaît facilement que :

$$R_m = \frac{1}{2id\beta_m} \left\{ \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n^+(N) \exp(-i\alpha_m x_{0n} - i\beta_m y_{0n}) \right\}. \quad (2.41)$$

En pratique, on utilise un nombre N fini de "sources" et on obtient une valeur approchée $R_m(N)$ qui tend donc vers R_m quand $N \rightarrow \infty$. Le calcul des coefficients de Rayleigh du champ transmis s'effectue de la même manière, et il est ensuite aisé de déterminer les efficacités dans les ordres réfléchis et transmis propagatifs.

2.9. QUELQUES RESULTATS NUMERIQUES

Nos essais numériques n'ont pour l'instant porté que sur le cas de réseaux sinusoïdaux. Pour des raisons de commodité, nous avons placé les "sources" sur des profils $\mathcal{P}^-$ et $\mathcal{P}^+$ (§ 2.5) qui sont simplement obtenus par une translation de $\mathcal{P}$ (fig. 2.6). On a placé N "sources" sur $\mathcal{P}^-$ et N "sources" sur $\mathcal{P}^+$, disposées de façon à former une partition régulière en x sur $[0, d]$. La région Ω^+ est vide, et la région Ω^- est remplie d'un diélectrique d'indice complexe n_2 . Nous avons bien entendu fait des essais avec des diélectriques sans pertes (cf. § 2.10), qui se traitent sans difficulté, mais nous ne reporterons ici que les résultats relatifs à un substrat en aluminium, d'indice $n_2 = 1.3 + i 7.1$ pour la longueur d'onde $\lambda_0 = 0.6 \mu\text{m}$ que nous avons choisie. Le pas du réseau est $d = 0.8333 \mu\text{m}$; nous nous plaçons en incidence normale ($\theta = 0$). Avec ces données, on observe trois ordres réfléchis d'efficacités e_0 et $e_1 = e_{-1}$ (incidence normale).

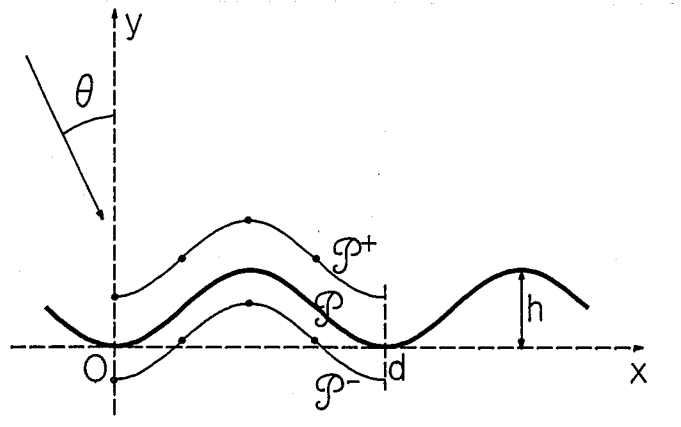


Figure 2.6. $\mathcal{P} : y = f(x)$; $\mathcal{P}^- : y = f(x) - e^-$; $\mathcal{P}^+ : y = f(x) + e^+$. Les points représentent la position des sources dans le cas où $N = 4$.

Des valeurs de e^- et e^+ (qui déterminent la position de $\mathcal{P}^-$ et $\mathcal{P}^+$) dépend la rapidité de convergence des résultats en fonction de N . Pour représenter le champ rayonné dans un diélectrique sans pertes (choix de e^-), il est apparu qu'une valeur convenable est $e^- \simeq \frac{\lambda_1}{10}$, λ_1 étant la longueur d'onde dans le diélectrique ; dans tous les résultats qui suivent, nous avons pris $e^- = 0.05 \mu\text{m}$. Pour représenter le champ rayonné dans un diélectrique à pertes (choix de e^+), il apparaît évident que l'on doit aussi tenir compte de la profondeur de pénétration dans le matériau ; avec nos données numériques, cette "épaisseur de peau" est $\lambda_0/(2\pi \text{Im}(n_2)) = 0.6/(2\pi \times 7.1) \simeq 0.013 \mu\text{m}$. Il est apparu que $e^+ = 0.05 \mu\text{m}$ constitue un "bon choix". C'est cette valeur qui sera utilisée dans tous les résultats qui suivent.

En ce qui concerne le calcul des produits scalaires (§ 2.7), nous retenons pour ce qui suit $C_2 = 0$ et $C_3 = C_{3m}$. Le nombre de points de discrétisation utilisé pour le calcul des sommes de Riemann ND est précisé dans les légendes des tableaux.

Les résultats sont comparés à ceux de la méthode intégrale pour laquelle le paramètre de convergence le plus significatif est le nombre de points (que l'on notera aussi N) utilisé pour "discrétiser" l'équation intégrale sur le profil. Les temps de calcul T donnés sont ceux obtenus sur le calculateur HP 9000-835 du Laboratoire, pour traiter les deux cas de polarisation.

Pour un réseau de hauteur $h = 0.1 \mu\text{m}$ ($h/d \simeq 1/8$; cas aussi traité dans le chapitre 1 par les méthodes différentielles), on obtient les efficacités reportées dans le tableau 1.

	N	E//		H//		T
		e_0	$e_1 = e_{-1}$	e_0	$e_1 = e_{-1}$	
Méthode de synthèse	10	.6152	.1511	.3647	.2336	9"
	15	.6078	.1495	.3802	.2589	21"
	20	.6084	.1497	.3796	.2612	37"
	30	.6088	.1498	.3791	.2609	1'28"
	50	.6087	.1497	.3792	.2609	4'16"
Méthode intégrale	15	.5809	.1425	.4158	.2909	4"
	20	.5947	.1461	.3908	.2707	6"
	30	.6040	.1485	.3822	.2636	12"
	50	.6077	.1495	.3798	.2615	33"
	80	.6086	.1497	.3792	.2610	1'28"
	120	.6087	.1497	.3792	.2610	3'33"

Tableau 1 : $h = 0.1 \mu\text{m}$; $ND = N$.

Dans le cas d'un réseau de hauteur $h = 0.8 \mu\text{m}$ ($h/d \simeq 1$), les résultats sont reportés dans le tableau 2.

	N	E//			H//			T
		e_0	$e_1 = e_{-1}$	Δ	e_0	$e_1 = e_{-1}$	Δ	
Méthode de synthèse	50	.119	.267	1.6E-1	.197	.121	1.9E-1	2'37"
	80	.170	.310	5.3E-2	.347	.160	6.2E-2	7'40"
	100	.179	.315	1.5E-2	.373	.165	2.7E-2	13'
	120	.180	.316	5.1E-3	.373	.165	1.2E-2	20'9"
	150	.180	.316	1.2E-3	.371	.165	4.3E-3	35'8"
Méthode intégrale	50	.150	.306		.438	.167		15"
	80	.171	.313		.384	.165		42"
	120	.177	.315		.374	.165		1'50"
	150	.179	.315		.373	.165		3'10"
	200	.180	.316		.372	.165		6'45"

Tableau 2 : $h = 0.8 \mu\text{m}$; $ND = 1.5 N$. "L'erreur relative" de

l'approximation Δ est définie par : $\Delta = \frac{\Delta_N}{\|F^i\|_Y}$ (cf. § 2.6).

On notera que les résultats convergent de façon tout à fait satisfaisante quand N croît. Les temps de calcul associés à la méthode de synthèse sont généralement supérieurs à ceux nécessités par la méthode intégrale, mais restent néanmoins "raisonnables".

2.10. LIEN AVEC LA METHODE INTRODUITE PAR YASUURA

Supposons ici que $\mathcal{P}^-$ soit un segment horizontal $0 \leq x \leq d$, $y = y_0^-$, avec $y_0^- < 0$ (fig. 2.4). Il apparaît alors que $g_n(x, y)$ prend, pour $y > y_0^-$, la forme (cf. eq.(2.20)) :

$$\begin{aligned} g_n(x, y) &= \frac{1}{2id} \sum_{m=-\infty}^{+\infty} \frac{1}{\beta_m} \exp\left[-i\alpha_m x_{0n} - i\beta_m y_0^-\right] \exp[i\alpha_m x + i\beta_m y] \\ &= \sum_{m=-\infty}^{+\infty} C_{n,m} \exp[i\alpha_m x + i\beta_m y] , \end{aligned} \quad (2.42)$$

$$\text{avec } C_{n,m} = \frac{1}{2id\beta_m} \exp\left[-i\alpha_m x_{0n} - i\beta_m y_0^-\right] .$$

On cherche le champ diffracté $u^d(x, y)$ dans Ω^+ sous la forme (eq.(2.38)) :

$$\begin{aligned} u_N^d(x, y) &= \sum_{n=1}^N C_n^+(N) g_n(x, y) \\ &= \sum_{n=1}^N \sum_{m=-\infty}^{+\infty} C_n^+(N) C_{n,m} \exp[i\alpha_m x + i\beta_m y] . \end{aligned} \quad (2.43)$$

Dans la "méthode de Yasuura" [26], on cherche $u_N^d(x, y)$ sous la forme (avec N impair) :

$$u_N^d(x, y) = \sum_{m=-(N-1)/2}^{(N-1)/2} Y_m^+(N) \exp[i\alpha_m x + i\beta_m y] . \quad (2.44)$$

Pour les besoins du calcul numérique, on garde dans la sommation en m (eq. 2.43) un nombre fini M de termes (M impair). On constate que si $N = M$, les deux expressions (2.43) et (2.44) sont strictement équivalentes, puisque $u_N^d(x, y)$ est recherché comme combinaison linéaire des mêmes fonctions, avec :

$$Y_m^+(N) = \sum_{n=1}^N C_n^+(N) C_{n,m} . \quad (2.45)$$

Il apparaît ainsi que si $\mathcal{P}^-$ est une horizontale située sous $\mathcal{P}$ (et de façon analogue si $\mathcal{P}^+$ est une horizontale située en dessus de $\mathcal{P}$), et pourvu que dans la partie numérique on ait choisi $M = N$, la méthode de synthèse se réduit à la "méthode de Yasuura".

Pour illustrer notre propos, considérons un réseau sinusoïdal diélectrique ($n_1 = 1$; $n_2 = 1.5$), de hauteur h , de pas $d = 0.8333 \mu\text{m}$, éclairé en incidence normale ($\theta = 0$) par une onde plane de longueur d'onde $\lambda_0 = 0.6 \mu\text{m}$. Nous avons reporté dans le tableau 3 ($h = 0.2 \mu\text{m}$) et le tableau 4 ($h = 0.4 \mu\text{m}$) les résultats obtenus dans les trois cas suivants :

- Cas 1 : les profils $\mathcal{P}^-$ et $\mathcal{P}^+$ sont obtenus par translation en y du profil $\mathcal{P}$ (fig. 2.6 ; $e^- = e^+ = 0.05 \mu\text{m}$) ; M est pris suffisant pour ne pas influencer sur les résultats.
- Cas 2 : les profils $\mathcal{P}^-$ et $\mathcal{P}^+$ sont des horizontales (respectivement $y = -0.05 \mu\text{m}$ et $y = h + 0.05 \mu\text{m}$) ; M est pris suffisant pour ne pas influencer sur les résultats.
- Cas 3 : $\mathcal{P}^-$ et $\mathcal{P}^+$ sont choisis comme dans le cas 2 ; mais $M = N$ de façon à obtenir des résultats identiques à ceux obtenus par la méthode de Yasuura.

Le produit scalaire utilisé pour ces calculs (cf. § 2.7) est tel que $C_2 = 0$ et $C_3 = 1$ (produit scalaire proposé à l'origine par Yasuura). On donne dans les tableaux les efficacités obtenues dans le cas de polarisation H// ($e_{r,0}$ est l'efficacité réfléchie dans l'ordre 0, $e_{r,1}$ l'efficacité réfléchie dans les ordres +1 et -1, $e_{t,0}$ l'efficacité transmise dans l'ordre 0, $e_{t,1}$ l'efficacité transmise dans les ordres +1 et -1, et $e_{t,2}$ l'efficacité transmise dans les ordres +2 et -2). On a aussi noté (bilan) la somme de toutes les efficacités, qui doit évidemment être égale à 1 puisque le réseau est sans pertes.

	N	$e_{r,0}$	$e_{r,1}$	$e_{t,0}$	$e_{t,1}$	$e_{t,2}$	bilan	T
Cas 1	15	.0007	.0122	.8659	.0525	.0013	.9987	9"
	21	.0007	.0124	.8676	.0524	.0010	.99987	17"
Cas 2	15	.0007	.0124	.8678	.0524	.0009	.99992	5"
	21	.0007	.0124	.8677	.0524	.0010	.999995	10"
Cas 3	15	.0007	.0124	.8678	.0524	.0009	.99991	2"
	21	.0007	.0124	.8677	.0524	.0010	.999993	6"

Tableau 3 : $h = 0.2 \mu\text{m}$.

	N	$e_{r,0}$	$e_{r,1}$	$e_{t,0}$	$e_{t,1}$	$e_{t,2}$	bilan	T
Cas 1	21	.0053	.0003	.5962	.1943	.0022	.995	13"
	31	.0051	.0003	.6020	.1939	.0024	1.0003	29"
	41	.0051	.0003	.6017	.1939	.0024	1.00004	54"
Cas 2	21	.0052	.0002	.6105	.1991	.0010	1.016	8"
	31	.0066	.0001	.6128	.2020	.0020	1.02	19"
	41	.0131	.0017	.6006	.2035	.0019	1.03	37"
Cas 3	21	.0046	.0000	.6032	.1956	.0009	1.0006	6"
	31	.0939	.1244	.6028	.2047	.0023	1.43	20"
	41	-	-	-	-	-	-	-

Tableau 4 : $h = 0.4 \mu\text{m}$. Dans le cas 3, la valeur $N = 41$ conduit à des résultats aberrants.

On remarquera que les trois cas donnent des résultats satisfaisants pour la hauteur $h = 0.2 \mu\text{m}$. Par contre, pour $h = 0.4 \mu\text{m}$, seul le premier cas conduit à des résultats qui convergent vers la solution quand N croît. Il est facile de donner une interprétation à ces constatations. Dans les deuxième et troisième cas, lorsque la hauteur du réseau est grande, on tente de "réaliser" un champ sur $\mathcal{P}$ au moyen de "sources" qui en sont trop éloignées (pour employer des termes plus exacts, le "champ à réaliser" est F^i et les "sources", localisées sur $\mathcal{P}^-$ et $\mathcal{P}^+$, engendrent sur $\mathcal{P}$ des "champs" $\Phi_{1,n}^+$ et $\Phi_{2,n}^-$; cf. § 2.6). On se persuadera sans doute facilement que cet "éloignement" est préjudiciable à une bonne représentation du champ sur $\mathcal{P}$.

(cela paraît évident dans le cas de matériaux conducteurs ; pour des matériaux sans pertes, on peut songer que cet éloignement dégrade la " finesse " de la représentation).

En résumé, nous pouvons dire que la " méthode de synthèse " peut être considérée comme un développement de la méthode proposée par Yasuura et ses collaborateurs [26], permettant de traiter des réseaux significativement plus profonds.

2.11. CONCLUSION

La " méthode de synthèse " nous paraît intéressante pour plusieurs raisons :

* Elle s'appuie sur une idée de base qui satisfait le physicien : on essaie de représenter le champ diffracté au dessus du réseau (respectivement en dessous) comme le champ rayonné par une collection de sources fictives situées au dessous du réseau (respectivement au dessus) et rayonnant dans l'espace libre rempli du matériau d'indice n_1 (respectivement n_2). Ces deux champs diffractés sont " raccordés " sur le profil $\mathcal{P}$ du réseau par une méthode de minimisation au sens des moindres carrés.

* Les résultats numériques obtenus au stade de notre étude sont encourageants, et l'étude de réseaux métalliques relativement profonds est possible.

* De nombreux développements de la méthode sont envisageables, notamment en ce qui concerne le choix des familles totales dans $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$ [61].

* Notre présentation est soutenue par des considérations théoriques qui permettent de justifier la convergence de la solution proposée.

3.1. CADRE DE L'ETUDE

Comme annoncé au chapitre 0, nous n'étudierons que des régimes harmoniques, et utilisons pour représenter les champs la notation complexe, moyennant une dépendance temporelle en $\exp(-i\omega t)$. Nous supposons que la perméabilité des matériaux est partout égale à celle du vide μ_0 . Les matériaux anisotropes seront caractérisés par une permittivité relative qui, dans les axes $\vec{e}_x, \vec{e}_y, \vec{e}_z$ choisis pour l'étude, est décrite par une matrice $[\epsilon]$ liant les vecteurs $\vec{D}$ et $\vec{E}$ selon $\vec{D} = \epsilon_0 [\epsilon] \vec{E}$. Les équations de Maxwell harmoniques s'écrivent ainsi, au sens des distributions [59] et en l'absence de sources :

$$\text{rot } \vec{E} = i\omega \mu_0 \vec{H}, \quad (3.1)$$

$$\text{rot } \vec{H} = -i\omega \epsilon_0 [\epsilon] \vec{E}, \quad (3.2)$$

$$\text{div } [\epsilon] \vec{E} = 0, \quad (3.3)$$

$$\text{div } \vec{H} = 0. \quad (3.4)$$

Dans toute l'étude, les éléments de la matrice $[\epsilon]$ seront notés de la manière suivante :

$$[\epsilon] = \begin{bmatrix} \epsilon_{xx} & \epsilon_{xy} & \epsilon_{xz} \\ \epsilon_{yx} & \epsilon_{yy} & \epsilon_{yz} \\ \epsilon_{zx} & \epsilon_{zy} & \epsilon_{zz} \end{bmatrix}. \quad (3.5)$$

3.2. PROPAGATION EN L'ABSENCE DE SOURCES EN MILIEU HOMOGENE ILLIMITE3.2.1. Position du problème

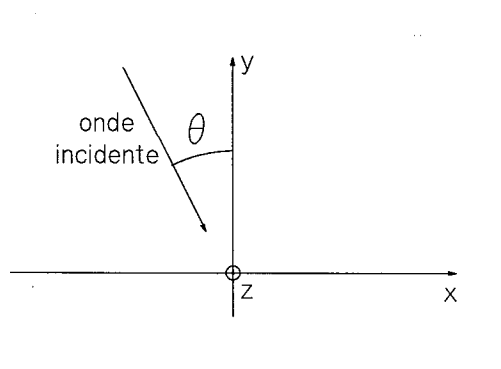
Notre objectif est ici de rechercher les ondes planes susceptibles de

se propager dans le milieu anisotrope, afin de préparer l'étude de structures telles que des empilements de couches, ou des réseaux anisotropes.

Dans le cas d'empilement de couches, si celles-ci sont disposées de telle sorte que leurs interfaces soient perpendiculaires à $\vec{e}_y$, on ne perd aucune généralité en supposant que l'onde plane incidente a un vecteur d'onde situé dans le plan $\vec{e}_x, \vec{e}_y$ (Figure 3.1). Notant α la composante du vecteur d'onde incident sur l'axe des x, les conditions de continuité imposent que toute composante de champ u est de la forme :

$$u(x,y) = \exp(i\alpha x) U(y). \quad (3.6)$$

Figure 3.1



Dans le cas de réseaux, l'axe des z est choisi parallèle aux sillons (Figures 0.1 et 0.2), et nous nous restreignons au cas où le plan d'incidence est le plan $\vec{e}_x, \vec{e}_y$. Toute composante de champ u est alors pseudo-périodique et s'écrit (eq.(1.4)) :

$$u(x,y) = \sum_{n \in \mathbb{Z}} \exp(i\alpha_n x) U_n(y). \quad (3.7)$$

Dans un cas comme dans l'autre, il apparaît que l'on est amené à rechercher, pour certaines valeurs de α réel donné, l'éventuelle propagation d'ondes de la forme $\exp(i\alpha x) U(y)$ dans le milieu anisotrope considéré.

3.2.2. Une première façon de procéder

Chacune des six composantes des champs $\vec{E}(x,y)$ et $\vec{H}(x,y)$ sera notée de façon analogue à :

$$E_x(x,y) = \exp(i\alpha x) \mathcal{E}_x(y). \quad (3.8)$$

Posons :

$$k_0 = \omega \sqrt{\epsilon_0 \mu_0}, \quad \eta_0 = \sqrt{\mu_0 / \epsilon_0}, \quad (3.9)$$

$$\begin{aligned} \epsilon_a &= \frac{1}{\epsilon_{yy}}, \quad \epsilon_b = \frac{\epsilon_{yx}}{\epsilon_{yy}}, \quad \epsilon_c = \frac{\epsilon_{yz}}{\epsilon_{yy}}, \quad \epsilon_d = \frac{\epsilon_{xy}}{\epsilon_{yy}}, \quad \epsilon_e = \frac{\epsilon_{zy}}{\epsilon_{yy}}, \\ \epsilon_f &= \epsilon_{xx} - \frac{\epsilon_{xy} \epsilon_{yx}}{\epsilon_{yy}}, \quad \epsilon_g = \epsilon_{xz} - \frac{\epsilon_{xy} \epsilon_{yz}}{\epsilon_{yy}}, \\ \epsilon_h &= \epsilon_{zx} - \frac{\epsilon_{zy} \epsilon_{yx}}{\epsilon_{yy}}, \quad \epsilon_i = \epsilon_{zz} - \frac{\epsilon_{zy} \epsilon_{yz}}{\epsilon_{yy}}. \end{aligned} \quad (3.10)$$

En développant (3.1) et (3.2), et après quelques manipulations élémentaires, on obtient :

$$\frac{d\mathcal{E}_x}{dy} = -i\alpha \epsilon_b \mathcal{E}_x - i\alpha \epsilon_c \mathcal{E}_z + i\eta_0 \left(\frac{\alpha^2}{k_0} \epsilon_a - k_0 \right) \mathcal{H}_z, \quad (3.11)$$

$$\frac{d\mathcal{E}_z}{dy} = ik_0 \eta_0 \mathcal{H}_x, \quad (3.12)$$

$$\frac{d\mathcal{H}_x}{dy} = i \frac{k_0}{\eta_0} \epsilon_h \mathcal{E}_x + \frac{i}{\eta_0} \left(k_0 \epsilon_i - \frac{\alpha^2}{k_0} \right) \mathcal{E}_z + i\alpha \epsilon_e \mathcal{H}_z, \quad (3.13)$$

$$\frac{d\mathcal{H}_z}{dy} = -i \frac{k_0}{\eta_0} \epsilon_f \mathcal{E}_x - i \frac{k_0}{\eta_0} \epsilon_g \mathcal{E}_z - i\alpha \epsilon_d \mathcal{H}_z, \quad (3.14)$$

$$\mathcal{E}_y = -\epsilon_b \mathcal{E}_x - \epsilon_c \mathcal{E}_z + \frac{\eta_0}{k_0} \epsilon_a \alpha \mathcal{H}_z, \quad (3.15)$$

$$\mathcal{H}_y = -\frac{\alpha}{k_0 \eta_0} \mathcal{E}_z. \quad (3.16)$$

Il apparaît donc que $\mathcal{E}_y$ et $\mathcal{H}_y$ peuvent simplement se déduire des autres composantes du champ (eq. (3.15) et (3.16)). Les quatre premières équations ((3.11) à (3.14)) constituent un système différentiel portant sur $\mathcal{E}_x$, $\mathcal{E}_z$, $\mathcal{H}_x$ et $\mathcal{H}_z$ qui peut s'écrire sous forme matricielle :

$$\frac{dF}{dy} = A F, \quad (3.17)$$

$$\text{en posant } F(y) = \begin{pmatrix} \mathcal{E}_x(y) \\ \mathcal{E}_z(y) \\ \mathcal{H}_x(y) \\ \mathcal{H}_z(y) \end{pmatrix}. \quad (3.18)$$

Soient λ_j et F_j ($j = 1, 2, 3, 4$) les valeurs propres et les vecteurs propres de A. La solution de notre problème est alors :

$$F(y) = \sum_{j=1}^4 C_j e^{\lambda_j y} F_j, \quad (3.19)$$

où les C_j sont quatre constantes arbitraires. Les composantes $\mathcal{E}_y$ et $\mathcal{H}_y$ s'en déduisent à l'aide de (3.15) et (3.16), de sorte que finalement :

$$\vec{E}(x,y) = \sum_{j=1}^4 C_j e^{i\alpha x + i\beta_j y} \vec{E}_j, \quad (3.20)$$

$$\vec{H}(x,y) = \sum_{j=1}^4 C_j e^{i\alpha x + i\beta_j y} \vec{H}_j, \quad (3.21)$$

en posant pour des raisons de commodité :

$$\lambda_j = i \beta_j. \quad (3.22)$$

Cela montre que, pour α donné, quatre ondes planes sont susceptibles de se propager dans le milieu anisotrope. Les constantes de propagation selon y de ces ondes sont liées simplement par (3.22) aux valeurs propres de A, et les vecteurs $\vec{E}_j$ et $\vec{H}_j$ (relatifs à la polarisation) sont déterminés par les vecteurs propres de A et par les relations (3.15) et (3.16).

Cette façon de procéder nous sera très utile par la suite, pour l'étude des réseaux anisotropes par la méthode différentielle (et aussi pour l'étude d'empilement de couches anisotropes). Toutefois, si l'on s'intéresse plus particulièrement aux constantes de propagation β_j des ondes planes dans le but d'en déterminer certaines propriétés, il semble préférable de procéder différemment. Cela est l'objet du paragraphe qui suit.

3.2.3. Une seconde façon de procéder

Au vu de ce qui précède, les solutions du problème posé au paragraphe 3.2.1 sont des ondes planes que l'on peut directement rechercher (cf VASSALLO, [82]) sous cette forme :

$$\vec{E}(x,y) = \exp(i\alpha x + i\beta_j y) \vec{E}_j. \quad (3.23)$$

Il suffit donc d'obtenir les constantes de propagation β_j et les polarisations $\vec{E}_j$ qui leur sont associées. Le champ magnétique sera ensuite obtenu trivialement par (3.1).

De (3.1) et (3.2), il vient :

$$\text{rot rot } \vec{E} - k_0^2 [\epsilon] \vec{E} = 0, \quad (3.24)$$

qui s'écrit aussi, compte tenu de (3.23) et en posant $\vec{k}_j = \alpha \vec{e}_x + \beta_j \vec{e}_y$:

$$i \vec{k}_j \wedge (i \vec{k}_j \wedge \vec{E}_j) - k_0^2 [\epsilon] \vec{E}_j = 0, \quad \text{ou bien} \quad (3.25)$$

$$\begin{bmatrix} (k_{xx}^2 - \beta_j^2) & (k_{xy}^2 + \alpha\beta_j) & k_{xz}^2 \\ (k_{yx}^2 + \alpha\beta_j) & (k_{yy}^2 - \alpha^2) & k_{yz}^2 \\ k_{zx}^2 & k_{zy}^2 & k_{zz}^2 - \alpha^2 - \beta_j^2 \end{bmatrix} \vec{E}_j = 0, \quad (3.26)$$

où l'on a posé :

$$k_{pq}^2 = k_0^2 \epsilon_{pq}. \quad (3.27)$$

Le système linéaire homogène (3.26) n'a de solutions non nulles que si son déterminant est nul :

$$\Delta = \begin{vmatrix} (k_{xx}^2 - \beta_j^2) & (k_{xy}^2 + \alpha\beta_j) & k_{xz}^2 \\ (k_{yx}^2 + \alpha\beta_j) & (k_{yy}^2 - \alpha^2) & k_{yz}^2 \\ k_{zx}^2 & k_{zy}^2 & k_{zz}^2 - \alpha^2 - \beta_j^2 \end{vmatrix} = 0. \quad (3.28)$$

On obtient ainsi une équation (équation de dispersion) du quatrième degré dont les racines sont les quatre constantes de propagation β_j recherchées. Le vecteur $\vec{E}_j$ associé à chacune de ces racines se déduit facilement de (3.26).

3.2.4. Quelques propriétés des constantes de propagation des ondes planes.

On pourra tout d'abord remarquer (eq.(3.25)) que si une onde plane de vecteur d'onde $\vec{k}_j$ est susceptible d'exister dans un milieu anisotrope, l'onde plane de vecteur $-\vec{k}_j$ possède la même propriété, avec le même vecteur $\vec{E}_j$. Il en résulte que, si $\alpha = 0$, les solutions β_j sont deux à deux opposées (cela apparaît clairement dans (3.28) qui devient alors une équation bicarrée en β_j).

Dans toute la suite (cf 3.2.1), nous nous intéresserons aux cas où $\vec{k}_j = \alpha \vec{e}_x + \beta_j \vec{e}_y$, avec α réel.

Cas d'un milieu absorbant

Dans ces conditions, l'onde est nécessairement amortie, et β_j a donc une partie imaginaire non nulle. Si dans l'équation de dispersion (3.28) on fait tendre α vers 0, les quatre solutions β_j tendent (sans traverser l'axe réel, Figure 3.2) vers des valeurs deux à deux opposées (remarque précédente). Il en résulte que, pour un milieu absorbant, deux des β_j sont telles que $\text{Im}(\beta_j) > 0$ et les deux autres sont telles que $\text{Im}(\beta_j) < 0$.

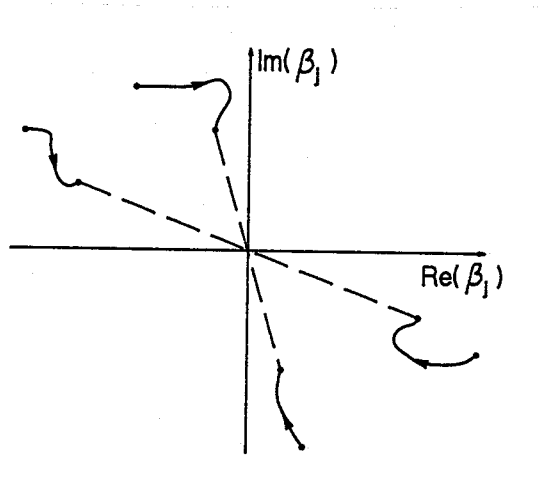


Figure 3.2. Les flèches représentent la trajectoire des β_j quand $\alpha \rightarrow 0$.

Cas d'un milieu sans pertes

Dans un milieu anisotrope sans pertes, on montre (Annexe 6, ou référence [82]) que la permittivité $[\epsilon]$ est hermitienne. La matrice $[\epsilon]$ ne dépend donc dans ce cas que de neuf paramètres réels. Si l'on développe le déter-

minant figurant dans l'équation de dispersion (3.28), on constate que cette équation se met sous la forme polynomiale suivante :

$$\mathcal{P}(\alpha, \beta_j, k_0) = \sum_{mn} P_{mn} \beta_j^m \alpha^n = 0 \quad (3.29)$$

avec pour seuls P_{mn} non nuls :

$$P_{04} = \varepsilon_{xx}$$

$$P_{40} = \varepsilon_{yy}$$

$$P_{13} = P_{31} = 2 \operatorname{Re}(\varepsilon_{xy})$$

$$P_{22} = \varepsilon_{xx} + \varepsilon_{yy}$$

$$P_{11} = 2 k_0^2 [-\varepsilon_{zz} \operatorname{Re}(\varepsilon_{xy}) + \operatorname{Re}(\varepsilon_{zx} \varepsilon_{yz})]$$

$$P_{02} = k_0^2 \left[-\varepsilon_{xx} (\varepsilon_{yy} + \varepsilon_{zz}) + |\varepsilon_{xy}|^2 + |\varepsilon_{xz}|^2 \right]$$

$$P_{20} = k_0^2 \left[-\varepsilon_{xx} (\varepsilon_{yy} + \varepsilon_{zz}) + |\varepsilon_{xy}|^2 + |\varepsilon_{yz}|^2 \right]$$

$$P_{00} = k_0^4 \left[\varepsilon_{xx} \varepsilon_{yy} \varepsilon_{zz} - \varepsilon_{xx} |\varepsilon_{yz}|^2 - \varepsilon_{yy} |\varepsilon_{xz}|^2 - \varepsilon_{zz} |\varepsilon_{xy}|^2 + 2 \operatorname{Re}(\varepsilon_{xy} \varepsilon_{yz} \varepsilon_{zx}) \right]$$

Tous les P_{mn} étant réels, les éventuelles racines non réelles β_j sont donc (pour α réel) deux à deux complexes conjuguées. C'est à dire que suivant les situations on pourra avoir :

- 4 β_j réelles,
- 2 β_j réelles et 2 β_j complexes conjuguées,
- 4 β_j deux à deux complexes conjuguées.

Remarque : l'équation (3.29) dépend de huit paramètres (les P_{mn} , moins un puisque l'équation est homogène). Nous avons vu qu'un milieu sans pertes est caractérisé par neuf paramètres réels (la donnée de sa permittivité hermitienne). Il en résulte qu'il existe donc (théoriquement) des classes de milieux anisotropes sans pertes différents, mais ayant les mêmes constantes de propagation.

Il est intéressant de tracer, pour un milieu donné, les courbes $\mathcal{P}(\alpha, \beta, k_0) = 0$, pour α et β réels. Un exemple est donné sur la figure 3.3. Ces courbes montrent clairement les vecteurs d'ondes réels (donc associés à des ondes planes qui se propagent dans le milieu sans atténuation) susceptibles d'exister dans le milieu considéré. Il est peut être utile de rappeler ici que, dans le cas d'un milieu isotrope, ces courbes se réduisent à un cercle car (en notant n l'indice du milieu) l'équation $\mathcal{P}(\alpha, \beta, k_0) = 0$ se

réduit dans ce cas à $\alpha^2 + \beta^2 = k_0^2 n^2$.

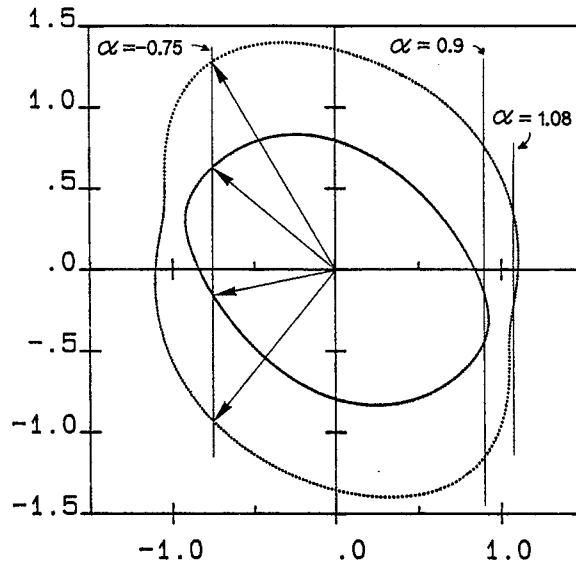


Figure 3.3. β en fonction de α pour $k_0 = 1$ et $[\epsilon] = \begin{bmatrix} 1.5 & 0.6 & 0.6 \\ 0.6 & 1.1 & 0 \\ 0 & 0.6 & 1.3 \end{bmatrix}$.

Pour $\alpha = -0.75$, on a tracé les vecteurs d'onde associés aux 4 ondes planes susceptibles de se propager. Pour $\alpha = 0.9$, on remarquera que 3 des vecteurs d'onde "possibles" sont dirigés vers le bas et qu'un seul est dirigé vers le haut. Pour $\alpha = 1.08$, la figure montre les deux β_j réels "possibles" ; les deux autres sont complexes conjugués. Pour $\alpha > 1.11$ environ, les quatre β_j sont complexes.

Que le milieu soit absorbant ou non, et toujours pour α réel, nous pourrions retenir le résultat suivant : les solutions β_j complexes (de partie imaginaire non nulle!) sont en nombre pair. La moitié d'entre elles vérifie $\text{Im}(\beta_j) > 0$ et l'autre moitié $\text{Im}(\beta_j) < 0$. Si le milieu est sans pertes, les β_j sont deux à deux conjuguées.

3.3. CONDITION D'ONDE SORTANTE (C.O.S.)

3.3.1. Quelques définitions et remarques préliminaires

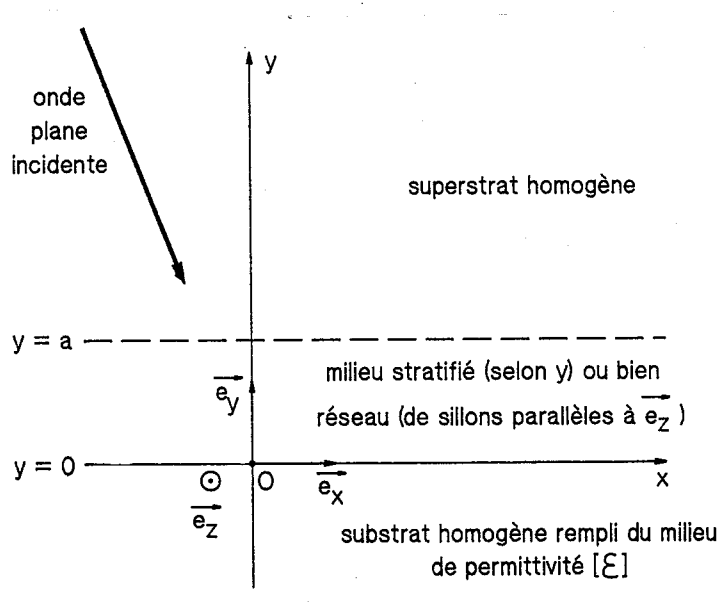
Soit $\vec{k} = \alpha \vec{e}_x + \beta \vec{e}_y$ le vecteur d'onde associé à une onde plane se propageant dans le milieu anisotrope (le fait d'avoir choisi $\vec{k}$ dans le plan $\vec{e}_x, \vec{e}_y$ ne restreint pas la généralité de l'étude. Il correspond simplement à un choix commode des axes pour un milieu anisotrope quelconque ; ce choix s'avère de plus bien adapté au problème des couches et des réseaux tel

qu'il est posé au § 3.2.1). Dans les situations qui nous préoccupent, α est imposé par l'onde incidente et la structure (cf § 3.2.1), et α est réel.

Nous adoptons le langage suivant : une onde plane sera dite propagative (ou "réelle") si son vecteur d'onde est réel. Dans ce cas, on distinguera les ondes planes propagatives montantes ou descendantes, selon qu'elles transportent de l'énergie vers les y positifs ou vers les y négatifs (lien avec le vecteur de Poynting). Dans le cas où β possède une partie imaginaire non nulle, nous dirons que l'onde plane est amortie. Nous écrirons β sous la forme $\beta = \beta' + i\beta''$ (β' et β'' réels). La dépendance en x et y de l'onde sera de la forme $\exp(-\beta''y)\exp(i\alpha x + i\beta'y)$; le terme "amortie" apparaît ainsi comme un abus de langage (en effet, si β' et β'' sont de même signe, il vaudrait peut être mieux dire que l'onde est "amplifiée") et signifie simplement que l'amplitude de l'onde varie exponentiellement en fonction de y .

Nous dirons que l'onde plane vérifie une C.O.S. vers le bas (ou quand $y \rightarrow -\infty$) si elle est propagative descendante, ou bien si elle est amortie avec $\beta'' < 0$. Dans le cas contraire, on dira qu'elle vérifie une C.O.S. vers le haut (ou quand $y \rightarrow +\infty$).

Figure 3.4



Nous serons amenés à considérer le problème suivant (Fig. 3.4). Un milieu périodique en x et invariant par translation selon z repose sur un substrat homogène anisotrope. Il est éclairé par une onde plane (propagative descendante) dont le vecteur d'onde est dans le plan $\vec{e}_x, \vec{e}_y$. Dans le

cas où la zone $0 < y < a$ est constituée d'un milieu stratifié selon y ("problème de couches"), les champs sont de la forme (3.6), et nous avons vu qu'il existe quatre ondes "possibles" dans le substrat. Dans le cas de réseaux, les champs sont pseudo-périodiques (de la forme (3.7)) et il y a alors "quatre infinités" d'ondes planes "possibles" dans le substrat (quatre dans chaque ordre du réseau). Ces ondes planes "possibles" dans le substrat ne sont pas toutes acceptables. Compte tenu des conditions d'éclairement de la structure, seules celles vérifiant une C.O.S. vers le bas seront finalement retenues. Pour simplifier l'exposé, nous nous placerons dans le cas des couches, mais le cas des réseaux se traite de façon similaire (pour l'étude de la C.O.S., chaque ordre d'un réseau donne un problème analogue à un problème de couches). Nous sommes alors conduits à nous poser la question suivante ...

3.3.2. Question posée

Le nombre α étant un réel donné, quelles sont, parmi les quatre ondes planes de vecteurs d'ondes $\vec{k}_j = \alpha \vec{e}_x + \beta_j \vec{e}_y$ susceptibles d'exister dans un milieu anisotrope infini (les quatre valeurs β_j "possibles" sont les solutions de l'équation de dispersion (3.28)), celles vérifiant une C.O.S. vers le bas et celles vérifiant une C.O.S. vers le haut.

Comme premier élément de réponse, signalons d'ores et déjà que, comme on peut s'y attendre (et à l'exception de cas singuliers pour lesquels certains des β_j sont nuls), deux de ces ondes vérifient une C.O.S. vers le bas et les deux autres une C.O.S. vers le haut (cf. Annexe 7).

3.3.3. Réponse dans le cas d'un milieu isotrope

Si l'indice d'un tel milieu est n , l'équation (3.27) devient $k_{pq}^2 = \delta_{pq} k_0^2 n^2$, et l'équation (3.28) s'écrit :

$$(k^2 - \alpha^2 - \beta_j^2)^2 = 0, \quad (k^2 = k_0^2 n^2) \quad (3.30)$$

qui admet pour solutions : $\beta_j^2 = k^2 - \alpha^2$; les quatre solutions dégénèrent en deux solutions opposées ; à chacune de ces deux solutions correspondent deux polarisations $\vec{E}_j$ qui sont associées aux deux cas fondamentaux de polarisation $E//$ et $H//$ (démonstration immédiate en considérant l'équation (3.26)). Etant donné que dans un milieu isotrope le vecteur de Poynting associé à une onde plane propagative est colinéaire au vecteur d'onde, la réponse est immédiate : dans le cas de milieux isotropes, les ondes planes

vérifiant la C.O.S. vers le bas sont celles qui sont associées à des β_j tels que $\beta_j < 0$ ou bien $\text{Im}(\beta_j) < 0$.

3.3.4. Réponse dans le cas d'un milieu anisotrope

On sait que, pour une onde plane propagative dans un milieu anisotrope, la direction de propagation de l'énergie (direction du vecteur de Poynting) est en général différente de celle du vecteur d'onde. Par exemple, il peut arriver (Fig.3.3) que l'équation de dispersion (3.28) admette quatre racines réelles, dont une est positive et trois sont négatives (ou vice-versa), ce qui correspond à un vecteur d'onde dirigé vers le haut et trois vers le bas (ou vice-versa). Mais on peut vérifier que, dans ce cas, on a en réalité deux ondes propagatives montantes et deux ondes propagatives descendantes.

3.3.4.1. En utilisant les vecteurs de Poynting

Supposons connues toutes les caractéristiques d'une onde plane se propageant dans le milieu anisotrope (§ 3.2.2 ou 3.2.3), c'est à dire son vecteur d'onde $\vec{k}_j = \alpha \vec{e}_x + \beta_j \vec{e}_y$ et la polarisation de ses champs $\vec{E}_j$ et $\vec{H}_j$. Nous avons vu que les ondes amorties ne posent aucun problème (la C.O.S. est liée au signe de $\text{Im}(\beta_j)$, de sorte que l'amplitude de l'onde reste bornée). Si l'onde est propagative, on détermine le vecteur de Poynting qui lui est associé. Le signe de sa composante selon $\vec{e}_y$ permet de savoir si l'onde vérifie une C.O.S. vers le haut ou le bas. Cette façon de procéder pour exprimer la C.O.S. a l'avantage de pouvoir être mise en oeuvre de façon simple sur ordinateur.

3.3.4.2. En raisonnant sur les constantes de propagation de l'onde

Pour les mêmes raisons qu'au paragraphe précédent, il suffit de s'intéresser aux ondes propagatives, ce qui sous entend donc que le milieu est sans pertes.

Considérons (problème initial) une onde plane dont la phase s'exprime en notation complexe (pulsation $\omega = c k_0 = k_0 / \sqrt{\epsilon_0 \mu_0}$) par $\exp(i\alpha x + i\beta(k_0)y)$, avec α et $\beta(k_0)$ réels. Nous avons vu (§ 3.2.4) que α et $\beta(k_0)$ vérifient :

$$\mathcal{P}(\alpha, \beta(k_0), k_0) = 0, \quad (3.31)$$

où le polynôme $\mathcal{P}$ dépend de la permittivité $[\varepsilon]$ du milieu. Pour des raisons de commodité, nous employons ici des notations qui rappellent que β dépend de k_0 . Il sera commode aussi d'avoir à l'esprit les courbes $\mathcal{P}(\alpha, \beta(k_0), k_0) = 0$ qui sont du type de celles représentées sur la figure 3.3 (§ 3.2.4). Nous représentons ici (Fig.3.5) une portion de cette courbe au voisinage du point $(\alpha, \beta(k_0))$.

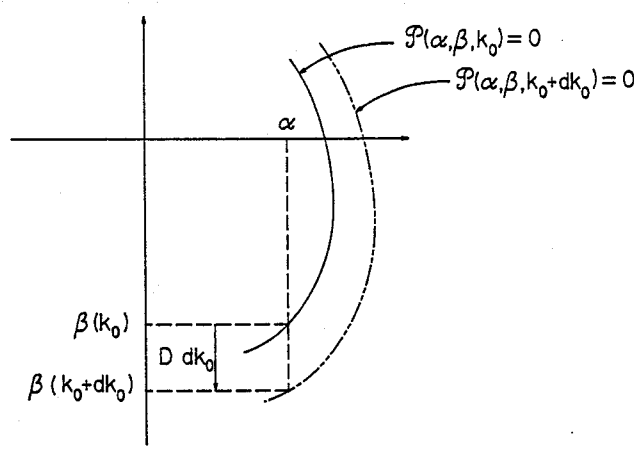


Figure 3.5. Nous avons supposé ici dk_0 réel positif.

Nous avons vu que, pour les ondes amorties, l'expression de la C.O.S. est particulièrement simple. Introduisons des pertes fictives au milieu considéré de façon à transformer cette onde plane propagative en une onde plane amortie. Considérons pour cela un problème fictif associé à des grandeurs $\tilde{\omega} = \omega + i\delta$, α et $\beta(\tilde{k}_0)$ solutions de (3.31). Nous choisissons :

$$\tilde{\omega} = \omega + i\delta, \quad \text{avec } d\omega = i\delta, \quad \delta > 0, \quad (3.32)$$

de sorte que :

$$\tilde{k}_0 = k_0 + dk_0, \quad dk_0 = i\delta/c. \quad (3.33)$$

Supposant δ "petit", ce problème fictif est donc "voisin" du problème initial, et tend vers ce dernier lorsque $\delta \rightarrow 0$. Puisque $\mathcal{P}(\alpha, \beta(\tilde{k}_0), \tilde{k}_0) = 0$, nous pouvons écrire au premier ordre :

$$\beta(\tilde{k}_0) = \beta(k_0 + dk_0) = \beta(k_0) + D dk_0. \quad (3.34)$$

La constante D peut être déterminée en supposant dk_0 réel. L'équation

$\mathcal{P}(\alpha, \beta, k_0)$ peut aussi s'exprimer sous la forme $Q(\frac{\alpha}{k_0}, \frac{\beta}{k_0}) = 0$ (vérification immédiate de cette propriété plus ou moins évidente (...) en se reportant au § 3.2.4), de sorte que les courbes $\mathcal{P}(\alpha, \beta, k_0 + dk_0)$ se déduisent simplement par homothétie des courbes $\mathcal{P}(\alpha, \beta, k_0) = 0$ (Fig.3.5). La valeur de D sera donc notée ensuite $D = \left(\frac{\partial \beta}{\partial k_0} \right)_\alpha$, étant entendu que α , β et k_0 sont liés par (3.31). Finalement, le problème fictif a pour solution une onde plane de la forme :

$$\exp \left(- \frac{\delta}{c} \left(\frac{\partial \beta}{\partial k_0} \right)_\alpha y \right) \exp (i\alpha x + i\beta(k_0)y). \quad (3.35)$$

Selon que $\left(\frac{\partial \beta}{\partial k_0} \right)_\alpha$ est négatif ou positif, cette onde amortie vérifie une C.O.S. vers le bas ou vers le haut.

Nous postulons que l'onde plane propagative associée au problème initial vérifie le même type de C.O.S. que l'onde amortie dont elle est la limite. Ces considérations ont sans aucun doute un lien avec le "principe de l'absorption limite" des mathématiciens. Nous rappelant que l'équation (3.31) est déduite de l'équation :

$$\text{rot rot } \vec{E} - k_0^2 [\epsilon] \vec{E} = 0, \quad (3.36)$$

nous constatons en effet que le fait d'ajouter à k_0 une partie purement imaginaire dk_0 est équivalent à ajouter à la permittivité $[\epsilon]$ initialement hermitienne (pas de pertes) la quantité $d[\epsilon] = - \frac{2}{k_0} [\epsilon] dk_0$; $[\epsilon] + d[\epsilon]$ n'étant pas hermitienne, ce milieu "fictif" possède donc de faibles pertes.

Pour nous résumer, et afin de savoir si une onde plane propagative vérifie une C.O.S. vers le haut ou vers le bas, il suffit d'examiner la courbe $\mathcal{P}(\alpha, \beta, k_0) = 0$ en se rappelant que la courbe $\mathcal{P}(\alpha, \beta, k_0 + dk_0)$ s'en déduit par homothétie. La nature de la C.O.S. (vers le haut ou vers le bas) est directement liée au signe de $\left(\frac{\partial \beta}{\partial k_0} \right)_\alpha$ (positif ou négatif). Un exemple est donné sur la figure 3.6 (identique à la figure 3.3). On peut vérifier que les conclusions obtenues sont les mêmes qu'en utilisant les considérations du § 3.3.4.1. Cette façon de déterminer la nature de la C.O.S. est plus " parlante " que celle s'appuyant sur les vecteurs de Poynting. Par contre, elle

a l'inconvénient de ne pas pouvoir se traduire aisément sur ordinateur. Les deux méthodes (utilisation des vecteurs de Poynting et "principe de l'absorption limite") nous paraissent ainsi complémentaires.

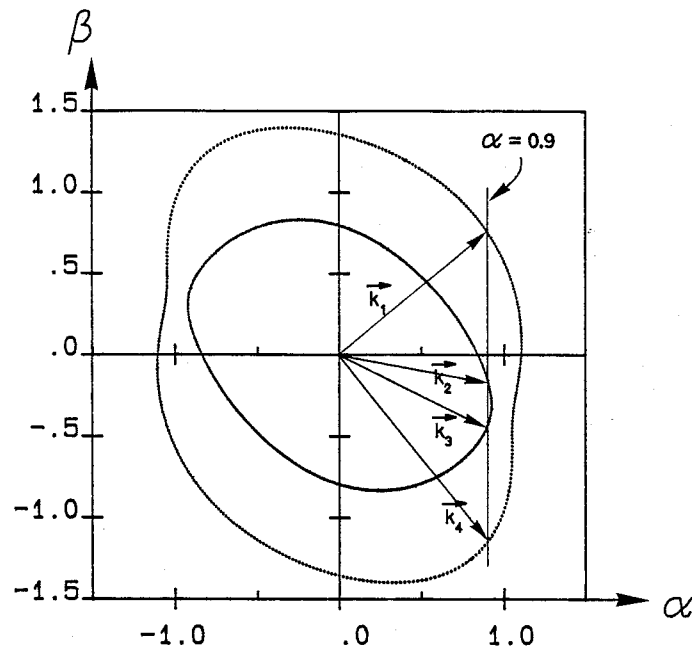


Figure 3.6. "Principe de l'absorption limite". $\left(\frac{\partial \beta}{\partial k_0}\right)_\alpha$ est positif pour $\vec{k}_1$ et $\vec{k}_2$, et négatif pour $\vec{k}_3$ et $\vec{k}_4$. Par conséquent, les ondes planes associées à $\vec{k}_1$ et $\vec{k}_2$ vérifient une COS vers le haut, et celles associées à $\vec{k}_3$ et $\vec{k}_4$ une COS vers le bas.

3.4. SOLUTIONS ELEMENTAIRES DE L'EQUATION DE PROPAGATION

3.4.1. Introduction, notations

Dans l'étude des problèmes de rayonnement en milieu anisotrope, on est amené à résoudre l'équation :

$$\text{rot rot } \vec{E} - k_0^2 [\epsilon] \vec{E} = \vec{S}, \quad \text{avec C.O.S.} \quad (3.37)$$

Cette équation appelle certains commentaires :

- $\vec{E}(x,y,z)$ et $\vec{S}(x,y,z)$ sont des distributions vectorielles qui représentent le champ électrique complexe et ses sources.
- La notation "avec C.O.S." rappelle que pour assurer l'unicité de la

solution, il est nécessaire d'imposer à $\vec{E}$ une condition d'onde sortante qui sera précisée en temps voulu.

- Elle remplace l'équation de Helmholtz (avec second membre) lorsque $\text{grad div } \vec{E}$ n'est pas nul.
- Une fois le champ électrique déterminé, le champ magnétique s'en déduit aisément à l'aide de $\text{rot } \vec{E} = i\omega \mu_0 \vec{H}$.

Pour simplifier l'écriture, nous poserons dans cette partie :

$$[A] = k_0^2 [\varepsilon] = \begin{bmatrix} A_{11} & A_{12} & A_{13} \\ A_{21} & A_{22} & A_{23} \\ A_{31} & A_{32} & A_{33} \end{bmatrix}. \quad (3.38)$$

Notre but est ici de résoudre (3.37), c'est à dire d'inverser l'opérateur $\text{rot rot} - [A]$ (la plupart des auteurs anglo-saxons diraient sans doute qu'il s'agit de déterminer les "Green's dyadics" associées à un problème anisotrope [15, 44, 81]). Le problème est loin d'être trivial et a été le sujet de nombreuses publications dans le cas isotrope (voir par exemple [80, 73]) comme dans le cas anisotrope [17, 36, 84, 74, 30, ...]. S'il est possible, pour un milieu anisotrope caractérisé par une matrice $[\varepsilon]$ quelconque, de résoudre formellement ce problème, il est en revanche beaucoup plus délicat d'en obtenir une solution explicite directement utilisable pour une étude numérique. Nous avons néanmoins obtenu une telle solution explicite en imposant certaines restrictions à $[\varepsilon]$, notamment en nous restreignant à l'étude de milieux pour lesquels $[\varepsilon]$ est diagonale (dans le trièdre $\vec{e}_x, \vec{e}_y, \vec{e}_z$ choisi pour décrire la structure). Le cas général, qui fait intervenir des transformées de Fourier au sens des distributions de fonctions de plusieurs variables, apparaît extrêmement difficile à résoudre explicitement.

Nous envisageons successivement le cas où le champ $\vec{E}$ dépend des trois coordonnées x, y, z (application à la diffraction par un obstacle de forme quelconque), puis des deux coordonnées x et y (problème à deux dimensions), et enfin le cas des réseaux dans lequel on utilise la pseudo-périodicité de $\vec{E}$ pour simplifier l'étude.

Pour alléger l'écriture, posons $\vec{e}_1 = \vec{e}_x, \vec{e}_2 = \vec{e}_y$, et $\vec{e}_3 = \vec{e}_z$. Considérons les trois distributions vectorielles $\vec{g}_i(x, y, z)$ (ou "vecteurs de Green") solutions de :

$$i = 1, 2, 3, \quad \text{rot rot } \vec{g}_i - [A] \vec{g}_i = \vec{e}_i \delta, \quad \text{avec C.O.S.} \quad (3.39)$$

La distribution δ représente ici l'unité de convolution adaptée au problème considéré. Plus précisément, dans l'étude du problème à trois dimensions où $\vec{E}$ et $\vec{S}$ sont des éléments de $\mathcal{D}'$ (distributions vectorielles définies sur l'espace $\mathcal{D}$ des fonctions de x, y, z indéfiniment dérivables à support borné [69]), δ est la classique distribution de Dirac $\delta(x, y, z)$. Dans l'étude du problème à deux dimensions, δ sera aussi la distribution de Dirac $\delta(x, y)$. Dans le cas des réseaux, $\vec{E}$ et $\vec{S}$ sont des distributions pseudo-périodiques éléments de l'espace $\mathcal{R}'$ (cet espace est noté $\mathcal{R}'_{\alpha}$ dans l'annexe 5). On trouvera dans [58] tous les renseignements utiles sur la définition de $\mathcal{R}'$. Rappelons brièvement qu'il s'agit des distributions vectorielles définies sur l'espace $\mathcal{R}$ des fonctions de x et y pseudo-périodiques en x , de période d (le pas du réseau), de coefficient de pseudo-périodicité $\exp(-ik_0 n_1 \sin \theta d)$ (notons que le réseau et l'onde incidente imposent au champ électromagnétique la pseudo-périodicité inverse), indéfiniment dérivables, et à support borné en y . On trouvera également dans [58] la définition de la convolution dans $\mathcal{R}'$, ainsi que l'expression de l'unité de convolution dans $\mathcal{R}'$ qui est la distribution :

$$\delta_{\mathcal{R}}(x, y) = \sum_{n \in \mathbb{Z}} \frac{1}{d} \delta(y) \exp(i\alpha_n x) \quad (3.40)$$

que nous avons déjà rencontrée au § 2.5. Dans le cas d'un problème de réseaux, la distribution δ figurant dans (3.39) sera donc cette unité de convolution $\delta_{\mathcal{R}}(x, y)$.

La détermination des trois vecteurs de Green $\vec{g}_i$ constitue le problème fondamental. On peut en effet montrer (Annexe 8) que si l'on note $[g]$ la matrice dont l'élément générique g_{ij} est la $j^{\text{ième}}$ composante du vecteur $\vec{g}_i$, et ${}^t[g]$ sa transposée, la solution de (3.37) n'est autre que :

$$\vec{E} = {}^t[g] * \vec{S}, \quad (3.41)$$

notation qui signifie que les composantes E_i de $\vec{E}$ se déduisent des composantes S_i de $\vec{S}$ par :

$$E_i = \sum_{j=1}^3 g_{ji} * S_j. \quad (3.42)$$

3.4.2. Détermination des transformées de Fourier des vecteurs de Green

Supposant que $\vec{E}$, $\vec{S}$ et $\vec{g}_i$ sont des distributions tempérées [69] (cela est le cas dans les problèmes rencontrés en pratique), désignons par σ_1 , σ_2 et σ_3 les variables conjuguées des coordonnées x , y et z dans la transformation de Fourier, et par $\vec{\sigma}$ le vecteur de composantes $(\sigma_1, \sigma_2, \sigma_3)$. Il s'agit de la transformation de Fourier au sens des distributions, définie de sorte que la transformée $\hat{f}$ d'une fonction sommable soit donnée par :

$$\hat{f}(\vec{\sigma}) = \iiint f(\vec{r}) \exp(-2i\pi \vec{\sigma} \cdot \vec{r}) d^3 \vec{r},$$

où $\vec{r}$ désigne le vecteur de composantes (x, y, z) .

L'équation $\text{rot rot } \vec{g}_i - [A] \vec{g}_i = \vec{e}_i \delta$ se transforme par Fourier selon :

$$-4\pi^2 \vec{\sigma} \wedge (\vec{\sigma} \wedge \widehat{\vec{g}_i}) - [A] \widehat{\vec{g}_i} = \vec{e}_i. \quad (3.43)$$

La résolution de cette équation donne, après quelques calculs (Annexe 9), l'expression générale de la transformée de Fourier des vecteurs de Green $\vec{g}_i(\vec{r})$:

$$\widehat{\vec{g}_i}(\vec{\sigma}) = [M] \vec{e}_i + 4\pi^2 \frac{\vec{\sigma} \cdot [M] \vec{e}_i}{1 - 4\pi^2 \vec{\sigma} \cdot [M] \vec{\sigma}} [M] \vec{\sigma}, \quad (3.44)$$

où la matrice $[M]$ vaut, notant $[I]$ la matrice identité et posant $\sigma^2 = \sigma_1^2 + \sigma_2^2 + \sigma_3^2$:

$$[M] = (4\pi^2 \sigma^2 [I] - [A])^{-1}, \quad (3.45)$$

On constate que l'expression des $\widehat{\vec{g}_i}(\vec{\sigma})$ est très lourde. Pour des matrices $[\varepsilon]$ (ou $[A]$) quelconques, il paraît très délicat (en tout état de cause, nous y avons au moins provisoirement renoncé) d'obtenir les $\vec{g}_i(\vec{r})$ sous forme explicite. Pour simplifier l'étude, nous avons seulement considéré des milieux pour lesquels $[A]$ est diagonal :

$$[A] = \left(\frac{2\pi}{\lambda_0} \right)^2 \begin{bmatrix} \varepsilon_{xx} & 0 & 0 \\ 0 & \varepsilon_{yy} & 0 \\ 0 & 0 & \varepsilon_{zz} \end{bmatrix} = \begin{bmatrix} A_{11} & 0 & 0 \\ 0 & A_{22} & 0 \\ 0 & 0 & A_{33} \end{bmatrix}. \quad (3.46)$$

La matrice $[A]$ peut être complexe, et c'est seulement dans le cas où elle est réelle que les milieux étudiés sont les "milieux biaxes" des cours traditionnels. Alors, les axes propres du milieu sont dirigés par $\vec{e}_1$, $\vec{e}_2$ et $\vec{e}_3$. Revenant au cas où $[A]$ est complexe, posons $\rho_i = 2\pi \sigma_i$, $\rho^2 = \rho_1^2 + \rho_2^2 + \rho_3^2$. En portant (3.46) dans (3.45) et (3.44), on obtient, après quelques astuces de calcul inspirées des cours traditionnels (Annexe 9), les composantes des transformées de Fourier des vecteurs de Green :

$$\hat{g}_{ij}(\rho_1, \rho_2, \rho_3) = \frac{\delta_{ij}}{\rho^2 - A_{ii}} - \frac{\rho^2}{\sum_k \frac{A_{kk} \rho_k^2}{\rho^2 - A_{kk}}} \frac{\rho_i \rho_j}{(\rho^2 - A_{ii})(\rho^2 - A_{jj})} \quad (3.47)$$

Le cas bidimensionnel (où les champs ne dépendent que de x et y) est obtenu immédiatement à partir de (3.47) en y faisant $\rho_3 = 0$, $\rho^2 = \rho_1^2 + \rho_2^2$; on obtient alors :

$$\hat{g}_{ij}(\rho_1, \rho_2) = \frac{\delta_{ij}}{\rho^2 - A_{ii}} - \frac{(\rho^2 - A_{11})(\rho^2 - A_{22}) \rho_i \rho_j}{(A_{11}\rho_1^2 + A_{22}\rho_2^2 - A_{11}A_{22})(\rho^2 - A_{ii})(\rho^2 - A_{jj})} \quad (3.48)$$

On pourra d'ores et déjà noter que $\hat{g}_{13}$, $\hat{g}_{23}$, $\hat{g}_{31}$ et $\hat{g}_{32}$ sont nuls, et par conséquent g_{13} , g_{23} , g_{31} et g_{32} le sont aussi. Ainsi, quand $[\varepsilon]$ est diagonale et dans le cas "bidimensionnel" où les champs sont indépendants de z , des sources "polarisées selon $\vec{e}_3 = \vec{e}_z$ " créent des champs électriques qui ont la même polarisation (éq. (3.42)). Nous reviendrons plus en détail sur les questions de polarisation au § 3.5.

3.4.3. Inversion de Fourier - Détermination des vecteurs de Green

Nous avons vu que l'inversion de Fourier des équations (3.44) est délicate. Il s'agit dans le cas général de trouver les inverses de Fourier de fonctions de trois variables d'expression très compliquées pouvant présenter des singularités. Il va sans dire que cela ne se fera pas sans rencontrer quelques distributions singulières (parties finies, ...). Même dans

le cas simplifié d'un problème à deux dimensions et pour $[\varepsilon]$ diagonale (expression (3.48)), il faudra être familier avec la transformée de Fourier à deux variables. La finalité de notre étude étant la détermination des vecteurs de Green associés à un problème de réseaux, nous n'effectuerons l'inversion de Fourier que dans ce cas. Nous allons voir que l'utilisation de la pseudo-périodicité des champs permet alors de se ramener à des inversions de Fourier de fonctions d'une seule variable.

Dans l'étude d'un problème de réseaux, les champs et la source $\vec{S}$ (eq. (3.37)) sont des distributions pseudo-périodiques appartenant à $\mathcal{R}'$. Toute distribution T de $\mathcal{R}'$ peut s'écrire [58] :

$$T(x, y) = \sum_{n \in \mathbb{Z}} T_n(y) \exp(i\alpha_n x) \quad , \quad \text{avec } T_n(y) \in \mathcal{D}' \quad (3.49)$$

et $\alpha_n = \alpha_0 + n 2\pi/d$. Rappelons que d est le pas du réseau, et que α_0 est déduit des conditions d'éclairement du réseau et de la périodicité d de la structure. Nous écrivons $\vec{E}$, $\vec{S}$ et $\vec{g}_i$ sous la forme :

$$\vec{E}(x, y) = \sum_{n \in \mathbb{Z}} \vec{E}_n(y) \exp(i\alpha_n x) = \sum_{i=1}^3 E_i \vec{e}_i \quad , \quad (3.50)$$

$$\vec{S}(x, y) = \sum_{n \in \mathbb{Z}} \vec{S}_n(y) \exp(i\alpha_n x) = \sum_{i=1}^3 S_i \vec{e}_i \quad , \quad (3.51)$$

$$\vec{g}_i(x, y) = \sum_{n \in \mathbb{Z}} \vec{g}_{in}(y) \exp(i\alpha_n x) = \sum_{j=1}^3 g_{ij} \vec{e}_j \quad . \quad (3.52)$$

L'unité de convolution est (§ 3.4.1) :

$$\delta_{\mathcal{R}}(x, y) = \sum_{n \in \mathbb{Z}} \frac{1}{d} \delta(y) \exp(i\alpha_n x) \quad . \quad (3.53)$$

D'après la définition de la convolution dans $\mathcal{R}'$ [58], l'équation (3.42) s'écrit :

$$E_{in}(y) = d \sum_{j=1}^3 g_{jin}(y) * S_{jn}(y) \quad , \quad (3.54)$$

où $g_{ijn}(y)$ est le $n^{\text{ième}}$ coefficient de Fourier de la $j^{\text{ième}}$ composante de

$\vec{g}_i$. Les vecteurs de Green $\vec{g}_i$ sont solution de (3.39), qui s'écrit ici :

$$\text{rot rot } \vec{g}_i(x,y) - [A] \vec{g}_i(x,y) = \vec{e}_i \delta_{\mathcal{R}}(x,y), \quad \text{avec C.O.S.} \quad (3.55)$$

En y remplaçant $\vec{g}_i$ et $\delta_{\mathcal{R}}$ par les expressions (3.52) et (3.53), en notant $\widehat{\vec{g}_{in}(\sigma_2)}$ la transformée de Fourier de $\vec{g}_{in}(y)$ et en posant $\vec{\sigma}_n = \begin{vmatrix} \alpha_n/2\pi \\ \sigma_2 \\ 0 \end{vmatrix}$, on obtient par transformation de Fourier de (3.55) :

$$\begin{aligned} \sum_{n \in \mathbb{Z}} -4\pi^2 \vec{\sigma}_n \wedge (\vec{\sigma}_n \wedge \widehat{\vec{g}_{in}}) \delta\left(\sigma_1 - \frac{\alpha_n}{2\pi}\right) - [A] \sum_{n \in \mathbb{Z}} \widehat{\vec{g}_{in}} \delta\left(\sigma_1 - \frac{\alpha_n}{2\pi}\right) = \\ = \vec{e}_i \sum_{n \in \mathbb{Z}} \frac{1}{d} \delta\left(\sigma_1 - \frac{\alpha_n}{2\pi}\right), \end{aligned}$$

c'est à dire que pour tout n :

$$-4\pi^2 \vec{\sigma}_n \wedge (\vec{\sigma}_n \wedge \widehat{\vec{g}_{in}}) - [A] \widehat{\vec{g}_{in}} = \frac{1}{d} \vec{e}_i. \quad (3.56)$$

Cela montre que les $\widehat{\vec{g}_{in}(\sigma_2)}$ s'obtiennent exactement de la même manière que les $\widehat{\vec{g}_i}(\sigma_1, \sigma_2)$ (§ 3.4.2), en remplaçant σ_1 par $\frac{\alpha_n}{2\pi}$ (ou bien ρ_1 par α_n). Dans le cas où $[\varepsilon]$ est diagonal, il apparaît, en comparant (3.56) et (3.43), que les $\widehat{\vec{g}_{ijn}(\rho_2)}$ s'obtiennent à partir de l'expression (3.48) par :

$$d \widehat{\vec{g}_{ijn}(\rho_2)} = \widehat{\vec{g}_{ij}(\alpha_n, \rho_2)}. \quad (3.57)$$

On en cherchera les inverses de Fourier pour obtenir finalement les $\vec{g}_{ijn}(y)$ (moyennant une certaine C.O.S. qui en assure l'unicité). Ce qui résoud le problème (cf. eq. (3.54)), qui se résume ainsi :

Lorsque $[A]$ est diagonal, les solutions de (3.55), exprimées selon (3.52), sont (cf. Annexe 9) :

$$\vec{g}_{11}(x,y) = \sum_{n \in \mathbb{Z}} \frac{i\beta_n}{2d A_{11}} \exp(i\alpha_n x + i\beta_n |y|), \quad (3.58)$$

$$\vec{g}_{22}(x,y) = \sum_{n \in \mathbb{Z}} \frac{1}{d A_{22}} \left[\frac{i A_{11} \alpha_n^2}{2 A_{22} \beta_n} - \delta(y) \right] \exp(i\alpha_n x + i\beta_n |y|), \quad (3.59)$$

$$g_{33}(x,y) = \sum_{n \in \mathbb{Z}} \frac{i}{2d \beta'_n} \exp(i\alpha_n x + i\beta'_n |y|), \quad (3.60)$$

$$g_{12}(x,y) = g_{21}(x,y) = \sum_{n \in \mathbb{Z}} \frac{\alpha_n}{2id A_{22}} \operatorname{sgn}(y) \exp(i\alpha_n x + i\beta_n |y|), \quad (3.61)$$

$$g_{13}(x,y) = g_{23}(x,y) = g_{31}(x,y) = g_{32}(x,y) = 0, \quad (3.62)$$

où $\operatorname{sgn}(y)$ représente la fonction égale à +1 si $y > 0$ et à -1 si $y < 0$, et où les coefficients β_n et β'_n sont définis par :

$$\beta_n = \sqrt{\frac{A_{11}}{A_{22}} (A_{22} - \alpha_n^2)}, \quad \beta'_n = \sqrt{A_{33} - \alpha_n^2}, \quad (3.63)$$

les racines (éventuellement complexes) étant déterminées de la façon suivante :

$$z \in \mathbb{C}, \quad \sqrt{z} > 0 \quad \text{ou bien} \quad \operatorname{Im}(\sqrt{z}) > 0. \quad (3.64)$$

3.5. QUELQUES CONSIDERATIONS SUR LA POLARISATION

Considérons une structure invariante par translation selon $\vec{e}_z$, décrite par sa permittivité relative $[\varepsilon(x,y)]$ (la perméabilité est supposée être partout égale à μ_0). Supposons que cette structure soit éclairée par un champ incident indépendant de z , de sorte que le champ total ne dépende que des seules coordonnées x et y . On notera que le cas d'un empilement de couches ou le cas d'un réseau (tel qu'il est décrit au chapitre 0) sont des cas particuliers de ce type de structure. Les équations de Maxwell en régime harmonique (on vérifiera immédiatement que les résultats de ce paragraphe s'appliquent si nécessaire aux régimes non harmoniques) s'écrivent au sens des distributions :

$$\operatorname{rot} \vec{E} = i\omega \mu_0 \vec{H} \quad \text{et} \quad \operatorname{rot} \vec{H} = -i\omega \varepsilon_0 [\varepsilon] \vec{E} \quad (3.65)$$

qui donnent en projection sur les axes :

$$\frac{\partial E_z}{\partial y} = i\omega \mu_0 H_x , \quad (3.66)$$

$$-\frac{\partial E_z}{\partial x} = i\omega \mu_0 H_y , \quad (3.67)$$

$$\frac{\partial E_y}{\partial x} - \frac{\partial E_x}{\partial y} = i\omega \mu_0 H_z , \quad (3.68)$$

$$\frac{\partial H_z}{\partial y} = -i\omega \varepsilon_0 (\varepsilon_{xx} E_x + \varepsilon_{xy} E_y + \varepsilon_{xz} E_z) , \quad (3.69)$$

$$-\frac{\partial H_z}{\partial x} = -i\omega \varepsilon_0 (\varepsilon_{yx} E_x + \varepsilon_{yy} E_y + \varepsilon_{yz} E_z) , \quad (3.70)$$

$$\frac{\partial H_y}{\partial x} - \frac{\partial H_x}{\partial y} = -i\omega \varepsilon_0 (\varepsilon_{zx} E_x + \varepsilon_{zy} E_y + \varepsilon_{zz} E_z) . \quad (3.71)$$

Lorsque $[\varepsilon(x,y)]$ est de la forme :

$$[\varepsilon(x,y)] = \begin{bmatrix} \varepsilon_{xx} & \varepsilon_{xy} & 0 \\ \varepsilon_{yx} & \varepsilon_{yy} & 0 \\ 0 & 0 & \varepsilon_{zz} \end{bmatrix} , \quad (3.72)$$

nous constatons que les six équations (3.66 à 3.71) se découpent en deux systèmes de trois équations faisant intervenir, soit les composantes E_z , H_x et H_y (équations (3.66), (3.67) et (3.71)), soit les composantes H_z , E_x et E_y (équations (3.68), (3.69) et (3.70)). Le champ incident $\vec{E}^i$, $\vec{H}^i$ peut se scinder en deux parties. La première que nous appellerons partie TE (ou E//) regroupe les composantes E_z^i , H_x^i et H_y^i , et la seconde (partie TM ou H//) regroupe les composantes H_z^i , E_x^i et E_y^i . On remarquera que les problèmes associés à un champ incident de polarisation E// et à un champ incident de polarisation H// sont indépendants. En d'autres termes, lorsque la permittivité est de la forme (3.72), un champ incident polarisé E// (resp. H//) engendrera un champ total de même polarisation. L'équation (3.71) montre de plus que le cas de polarisation E// se traite de la même façon que pour des milieux isotropes de permittivité ε_{zz} (ce qui signifie qu'il pourra être résolu à l'aide de programmes conçus pour des milieux isotropes).

Lorsque $[\epsilon(x,y)]$ est de la forme :

$$[\epsilon(x,y)] = \begin{bmatrix} \epsilon_1 & 0 & 0 \\ 0 & \epsilon_1 & 0 \\ 0 & 0 & \epsilon_{zz} \end{bmatrix}, \quad (3.73)$$

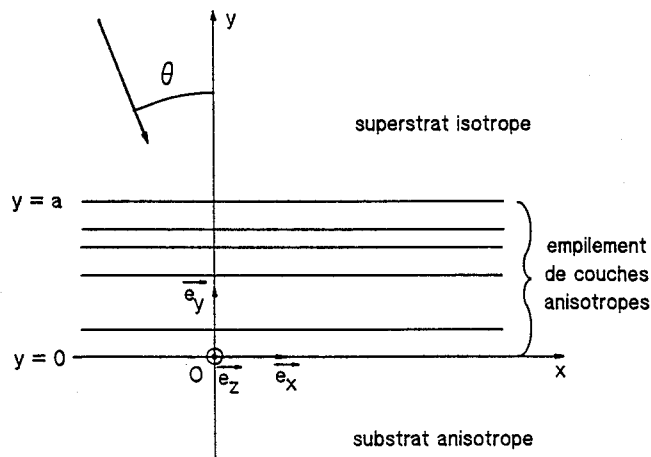
ce qui est le cas de milieux uniaxes disposés de telle sorte que leurs axes optiques soient parallèles à $\vec{e}_z$, les équations (3.69) et (3.70) montrent que le cas de polarisation H// se traite de la même façon que pour des milieux isotropes de permittivité ϵ_1 (là aussi, on pourra utiliser des programmes conçus pour des milieux isotropes pour le résoudre).

Lorsque la permittivité n'est pas de la forme (3.72), les deux cas fondamentaux de polarisation ne se découplent pas. Un champ incident polarisé E// ou H// donnera naissance à un champ total de polarisation quelconque.

4.1. POSITION DU PROBLEME

L'étude de ces structures est évidemment une étape indispensable pour qui a l'ambition de s'attaquer au cas plus difficile des réseaux. D'ailleurs, nous nous réfèrerons souvent à ce chapitre lors de l'étude des réseaux par la méthode différentielle (chapitre 5).

Nous nous intéressons ici à un empilement de couches anisotropes homogènes (Fig.4.1) situées dans la zone $0 < y < a$. Le substrat et chacune des couches sont constitués de milieux anisotropes quelconques (avec ou sans pertes ; il n'est fait aucune restriction sur leurs permittivités $[\epsilon]$). L'ensemble est éclairé par une onde plane incidente dont le vecteur d'onde est supposé être situé dans le plan $(\vec{e}_x, \vec{e}_y)$, ce qui ne restreint en rien la généralité de l'étude.

Figure 4.1

Ce genre de problème a été étudié par de nombreux auteurs [4, 89, 33, 18, 71, 85, 88, 2]. Ceux-ci ont souvent envisagé des cas particuliers pour les matrices $[\epsilon]$, dans le but de pouvoir exprimer analytiquement les "modes de propagation" dans les couches anisotropes (c'est-à-dire les vecteurs d'onde et les polarisations des ondes). Pour notre part, et utilisant les considérations développées dans le chapitre 3 (§ 3.2 et 3.3), nous avons

préférez confier cette tâche à l'ordinateur, ce qui permet de présenter la solution du problème sans lourdeur excessive. La seule restriction que nous avons faite consiste à supposer que la perméabilité est partout égale à celle du vide. Il semble que cela permette de traiter de nombreuses situations [54]. Toutefois, si cela s'avérait nécessaire, l'introduction d'une perméabilité matricielle serait possible.

4.2. RESOLUTION DU PROBLEME

Nous avons vu (§ 3.2.1) que toute composante de champ $u(x,y)$ est de la forme :

$$u(x,y) = \exp(i\alpha x) U(y), \quad (4.1)$$

avec, si n_1 est l'indice (réel) du superstrat et θ l'angle d'incidence :

$$\alpha = k_0 n_1 \sin\theta. \quad (4.2)$$

Nous avons vu aussi (§ 3.2.2) que le champ total est caractérisé par un vecteur $F(y)$ à quatre composantes vérifiant dans chaque région homogène une équation (dédue des équations de Maxwell) s'écrivant :

$$\frac{dF(y)}{dy} = A F(y). \quad (4.3)$$

On en déduit l'expression de $F(y)$ dans chaque région homogène à l'aide des valeurs propres $\lambda_j = i\beta_j$ de A et des vecteurs propres F_j associés (nous déterminons numériquement les λ_j et F_j au moyen des sous programmes classiques disponibles dans toute bonne bibliothèque) :

$$F(y) = \sum_{j=1}^4 C_j e^{i\beta_j y} F_j \quad (4.4)$$

où les C_j sont quatre constantes complexes. En d'autres termes, dans chaque région homogène, le champ est décrit par la superposition de quatre ondes planes d'amplitudes C_j . Le champ en $y = a$ s'écrira sous la forme :

$$F(a) = \sum_{j=1}^4 C_j^h e^{i\beta_j^h a} F_j^h. \quad (4.5)$$

Le superstrat (milieu du haut) étant isotrope, les constantes β_j^h dégénèrent en deux valeurs opposées (§ 3.3.3), soit $\beta_1^h = \beta_2^h = -\beta^h$ et $\beta_3^h = \beta_4^h = \beta^h$, avec $\beta^h > 0$. Il en résulte que C_1^h et C_2^h représentent le champ incident, alors que C_3^h et C_4^h représentent le champ réfléchi.

Dans le substrat (milieu du bas), la condition d'onde sortante entraîne la nullité de deux des constantes C_j . Pour exprimer cette condition d'onde sortante, on procèdera de préférence comme expliqué au § 3.3.4.1, en raisonnant sur les vecteurs de Poynting. On en déduit l'expression du champ en $y = 0$, dépendant de deux constantes C_j^b représentant le champ transmis :

$$F(0) = \sum_{j=1}^2 C_j^b F_j^b . \quad (4.6)$$

Compte tenu de la continuité de $F(y)$ aux interfaces séparant les couches (F représente en effet des composantes tangentielles de champs) et de l'expression (4.4) qui traduit la dépendance de F en y dans une couche, il est facile de déduire les quatre constantes C_j^h des deux constantes C_j^b , c'est à dire de déduire $F(a)$ de $F(0)$, en "remontant" dans la structure de $y = 0$ à $y = a$.

La linéarité du problème entraîne que les constantes C_j^h sont liées aux constantes C_j^b par une matrice de "transfert" $[T]$:

$$\begin{bmatrix} C_1^h \\ C_2^h \\ C_3^h \\ C_4^h \end{bmatrix} = \begin{bmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \\ T_{31} & T_{32} \\ T_{41} & T_{42} \end{bmatrix} \begin{bmatrix} C_1^b \\ C_2^b \end{bmatrix} . \quad (4.7)$$

Cette matrice $[T]$ est construite très facilement en utilisant par exemple la "méthode de tir". Cela consiste à effectuer deux "tirs" de la façon suivante : le premier tir est réalisé en se donnant arbitrairement le champ dans le substrat en supposant $C_1^b = 1$ et $C_2^b = 0$. De cette valeur particulière $\tilde{F}(0)$, on déduit $\tilde{F}(a)$ et les constantes $\tilde{C}_j^h$ associées, qui ne sont autres que la première colonne de $[T]$. La deuxième colonne de $[T]$ est construite au moyen d'un second tir pour lequel on choisit $C_1^b = 0$ et $C_2^b = 1$.

La matrice de transfert étant déterminée, il est plus commode d'écrire

(4.7) sous la forme :

$$\begin{bmatrix} C_1^h \\ C_2^h \end{bmatrix} = \begin{bmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{bmatrix} \begin{bmatrix} C_1^b \\ C_2^b \end{bmatrix} \quad (4.8)$$

$$\begin{bmatrix} C_3^h \\ C_4^h \end{bmatrix} = \begin{bmatrix} T_{31} & T_{32} \\ T_{41} & T_{42} \end{bmatrix} \begin{bmatrix} C_1^b \\ C_2^b \end{bmatrix} \quad (4.9)$$

Le problème est alors résolu car, le champ incident étant connu, on connaît C_1^h et C_2^h . L'équation (4.8) permet d'en déduire le champ transmis (C_1^b et C_2^b), et l'équation (4.9) donne le champ réfléchi (C_3^h et C_4^h).

On pourra noter que cette méthode de résolution s'applique aussi dans le cas où le substrat est infiniment conducteur. Dans ce cas, les conditions aux limites imposent qu'en $y = 0$ les composantes tangentielles du champ électrique (c'est à dire les deux premières composantes de $F(0)$) s'annulent. Il en résulte que les quatre constantes C_j représentant le champ dans la couche en contact avec le substrat sont liées. Il ne reste dans cette couche que deux constantes C_j indépendantes, et nous sommes ramenés au cas précédent, à la seule différence que C_1^b et C_2^b doivent être remplacées par deux constantes indépendantes représentant le champ dans la couche la plus "profonde".

4.3. VECTEUR DE POYNTING, EFFICACITES

Nous avons vu que les solutions de notre problème peuvent s'écrire dans chaque zone homogène sous la forme (4.4) d'une combinaison linéaire de quatre ondes planes, chacune de ces ondes planes étant caractérisée par sa constante de propagation β_j et sa polarisation F_j :

$$F(y) = \sum_{j=1}^4 C_j \exp(i\beta_j y) F_j . \quad (4.10)$$

Chacune de ces quatre ondes est une solution particulière des équations de Maxwell, et l'on peut calculer le vecteur de Poynting complexe $\vec{\Phi}_j$ qui lui est associé. On peut aussi calculer le vecteur de Poynting complexe $\vec{\Phi}$ associé au champ total. Le problème est de savoir si l'on a entre les composantes selon y de ces vecteurs la relation :

$$\operatorname{Re}(\mathcal{P}_y) \stackrel{?}{=} \sum_{j=1}^4 |C_j|^2 \operatorname{Re}(\mathcal{P}_{jy}). \quad (4.11)$$

Envisageons tout d'abord le cas d'un milieu isotrope d'indice n réel. Nous avons vu au § 3.3.3 que les quatre ondes planes peuvent être classées comme suit :

j	β_j	Polarisation associée à F_j
1	$-\sqrt{k_0^2 n^2 - \alpha^2}$	E//
2	$-\sqrt{k_0^2 n^2 - \alpha^2}$	H//
3	$+\sqrt{k_0^2 n^2 - \alpha^2}$	E//
4	$+\sqrt{k_0^2 n^2 - \alpha^2}$	H//

(4.12)

On montre que dans ce cas, la réponse à la question (4.11) est positive (Annexe 10). Il en résulte que l'on pourra parler de l' "énergie transportée" par chacune des quatre ondes, et ainsi définir (dans le cas où les β_j sont réels) des efficacités (réfléchies si le milieu considéré est le superstrat ; transmises si le milieu est le substrat, auquel cas les constantes C_3 et C_4 sont nulles par application de la C.O.S.). Plus précisément, si l'onde incidente est polarisée E//, c'est à dire si $C_2 = 0$ dans le superstrat, on pourra définir une efficacité ordinaire réfléchie associée à l'énergie transportée par l'onde réfléchie polarisée E// ($j = 3$) et une efficacité extraordinaire réfléchie associée à l'énergie transportée par l'onde réfléchie polarisée H// ($j = 4$). Si le substrat est isotrope, sans pertes, et qu'il ne donne pas lieu à la réflexion totale (si les β_j y sont réels), on pourra aussi définir une efficacité ordinaire transmise (onde $j = 1$ dans le substrat) et une efficacité extraordinaire transmise (onde $j = 2$ dans le substrat). On utilisera des définitions analogues si l'onde incidente est polarisée H// ; dans ce cas, les efficacités ordinaires seront liées aux ondes polarisées H// et les efficacités extraordinaires à celles polarisées E//. L'égalité (4.11) étant vraie, on est assuré que la somme des efficacités ordinaire et extraordinaire transmises est bien associée à l'énergie totale transmise (même chose pour les efficacités réfléchies).

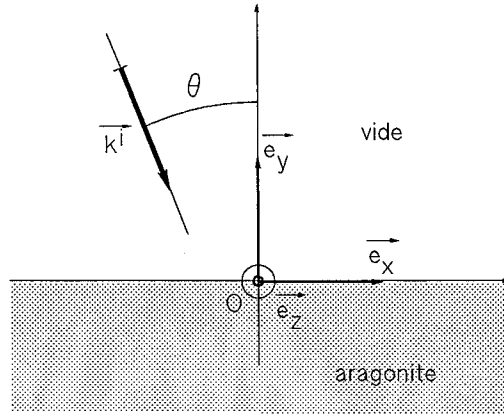
Dans le cas d'un milieu anisotrope, la réponse à la question (4.11) est en général négative. Lorsque le substrat est anisotrope, on s'est contenté de

calculer le vecteur de Poynting associé au champ total dans le substrat.

4.4. QUELQUES RESULTATS

4.4.1. Réfraction conique (expérience de Hamilton)

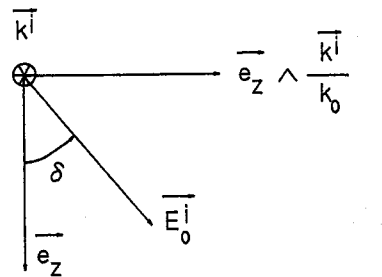
Figure 4.2



Nous avons étudié ce phénomène curieux [6, 9, 28, 1] dans le double but de mieux le comprendre et de tester dans des circonstances simples le programme qui a été réalisé. Une onde plane illumine sous l'incidence $\theta = 14.52^\circ$ (Figure 4.2) une interface plane séparant le vide et un cristal d'aragonite de permittivité diagonale, avec $\epsilon_{xx} = 2.843$, $\epsilon_{yy} = 2.341$, $\epsilon_{zz} = 2.829$. Le champ incident est décrit comme une combinaison linéaire d'ondes polarisées E// et H// (Figure 4.3) :

$$\begin{aligned} \vec{E}^i(x,y) &= \left\{ \cos\delta \vec{e}_z + \sin\delta \vec{e}_z \wedge \frac{\vec{k}^i}{k_0} \right\} \exp(i \vec{k}^i \cdot \vec{r}), \\ &= \vec{E}_0^i \exp(i \vec{k}^i \cdot \vec{r}). \end{aligned} \quad (4.13)$$

Figure 4.3



Ainsi, lorsque δ varie de $-\pi/2$ à $+\pi/2$, $\vec{E}_0^i$ tourne de 180° autour du vecteur d'onde $\vec{k}^i$. Pour chaque valeur de δ , nous calculons le vecteur de Poynting $\vec{\mathcal{P}}$ (qui dans ces conditions est réel) associé au champ total dans l'aragonite. Le point P défini par $\vec{OP} = \vec{\mathcal{P}}(\delta)$ se situe sur un cône. La figure 4.4 montre

l'intersection de ce cône avec le plan $y = -1$. On dit traditionnellement que dans ces conditions particulières, "un rayon donne une infinité de rayons répartis sur un cône". Du point de vue de l'électromagnétisme, nous dirions plutôt que, l'incidence étant fixée à cette valeur remarquable, à chaque polarisation rectiligne de l'onde incidente correspond un vecteur de Poynting bien déterminé pour l'onde réfractée, et que les directions de ces différents vecteurs de Poynting forment un cône.

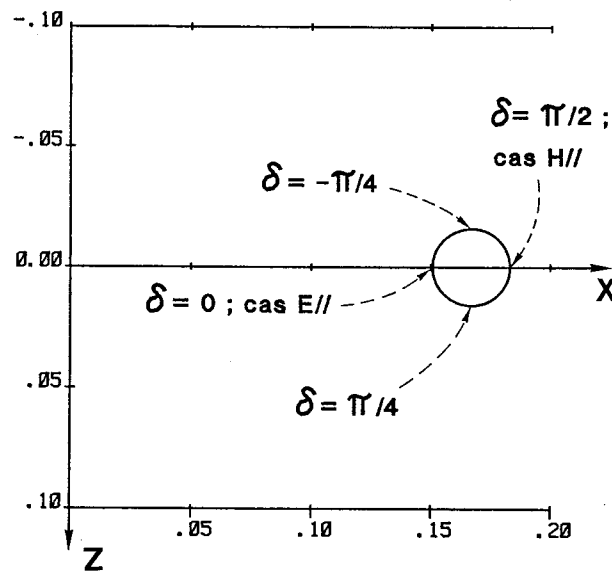


Figure 4.4. Intersection du cône et du plan $y = -1$.

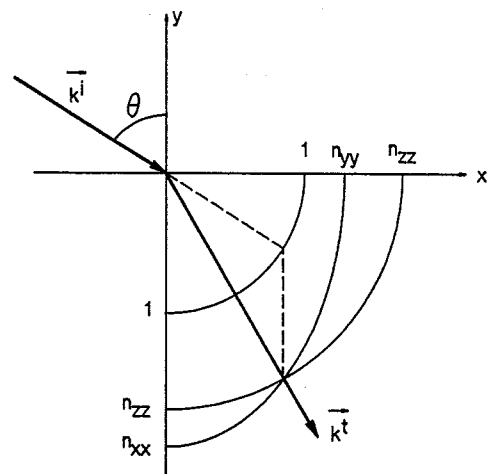


Figure 4.5. Obtention de la valeur de θ produisant le phénomène de réfraction conique à partir des surfaces des indices

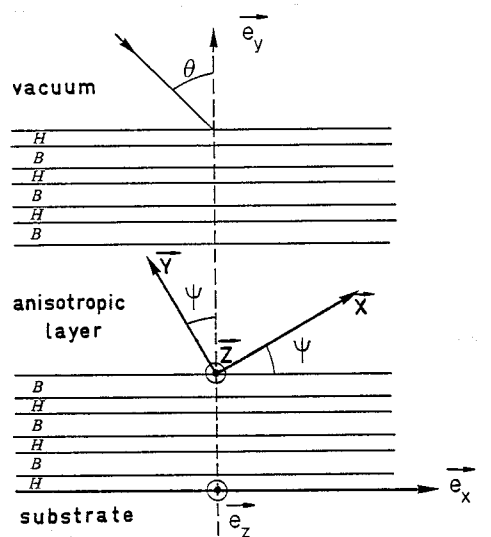
Remarque : la valeur $\theta = 14.52^\circ$ qui donne lieu à ce phénomène n'est bien entendu pas quelconque. C'est une valeur particulière pour laquelle les deux ondes vérifiant la C.O.S. dans l'aragonite ont même vecteur d'onde $\vec{k}^t$; en posant $n_{ii} = \sqrt{\epsilon_{ii}}$, elle est obtenue au moyen de la construction de la figure 4.5, déduite de considérations classiques concernant les surfaces

des indices. Ces considérations sont détaillées dans les cours traditionnels [6, 9] et c'est pourquoi nous ne les rappelons pas ici.

4.4.2. Renforcement d'une efficacité extraordinaire par utilisation de multicouches.

Suite à une suggestion de E. Pelletier [53], nous avons vérifié numériquement qu'il est possible d'augmenter considérablement l'efficacité extraordinaire réfléchie par une couche diélectrique anisotrope en l'enserrant entre deux empilements de lames minces. Nous considérons tout d'abord la situation représentée sur la figure 4.6. Les lettres H et B représentent des couches isotropes, d'indices respectifs 2.3 et 1.3, et d'épaisseurs respectives $e_H = 0.065 \mu\text{m}$ et $e_B = 0.115 \mu\text{m}$ (couches $\lambda/4$ pour la longueur d'onde $0.6 \mu\text{m}$). Le substrat est isotrope (indice 1.517). L'angle d'incidence θ a été fixé arbitrairement à $\theta = 22^\circ$. Les vecteurs $\vec{X}$, $\vec{Y}$ et $\vec{Z}$ représentent les axes principaux de la couche anisotrope, dont la permittivité relative se diagonalise dans ces axes, avec les valeurs propres $\epsilon_{XX} = 4.35$, $\epsilon_{YY} = 4.85$, $\epsilon_{ZZ} = 4.59$. Ces valeurs correspondent à une couche de TiO_2 , dont l'anisotropie est due à des conditions d'évaporation [20]. L'angle ψ est égal à 18.9° , et l'épaisseur de la couche anisotrope vaut $0.495 \mu\text{m}$. Nous traitons la situation obtenue en faisant tourner la structure représentée sur la figure 4.6 d'un angle $\delta = (\vec{e}_z, \vec{Z})$ autour de $\vec{e}_y$, le plan d'incidence restant toujours le plan $(\vec{e}_x, \vec{e}_y)$.

Figure 4.6
Représentation de la structure pour $\delta = 0$. Les axes $\vec{e}_x$, $\vec{e}_y$ et $\vec{e}_z$ sont liés à l'onde incidente. Les axes $\vec{X}$, $\vec{Y}$ et $\vec{Z}$ sont liés à la structure.

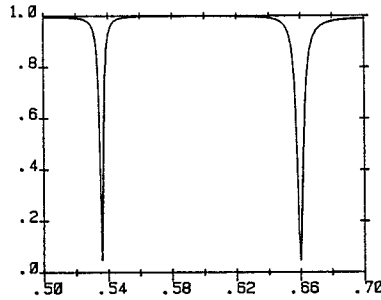


On pourra noter que si $\delta = 0$, et quel que soit ψ , les efficacités extraordinaires sont nulles. En effet, pour $\delta = 0$, la permittivité de la

couche anisotrope se met dans les axes $\vec{e}_x$, $\vec{e}_y$ et $\vec{e}_z$ sous la forme (3.72), et nous avons déjà noté (§ 3.5) que si l'onde incidente est polarisée E// (resp. H//), le champ total est partout polarisé de la même façon.

La figure 4.7 montre pour $\delta = 0$ l'efficacité ordinaire réfléchie en polarisation E// (le champ électrique incident vibre parallèlement à $\vec{e}_z$) en fonction de la longueur d'onde. On observe un minimum de réflexion pour $\lambda_0 = 0.66 \mu\text{m}$.

Figure 4.7



La figure 4.8 montre (toujours en polarisation E//) et pour la longueur d'onde $\lambda_0 = 0.66 \mu\text{m}$, les efficacités ordinaires et extraordinaires réfléchies en fonction de δ . Nous constatons que l'efficacité extraordinaire peut dépasser 10 %, alors que si les couches H et B sont supprimées, nous avons pu observer que les efficacités extraordinaires ne dépassent jamais (pour cette même couche anisotrope, dans le même domaine de longueurs d'onde, et quel que soit δ) la valeur de $2 \cdot 10^{-3}$.

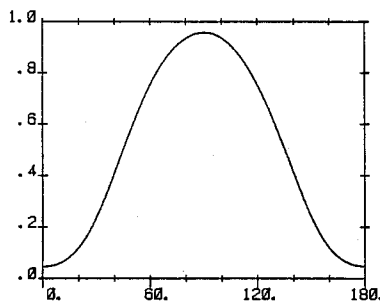


Figure 4.8.a

*Efficacité ordinaire
réfléchie en fonction de δ*

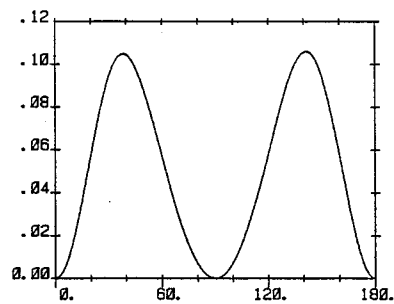


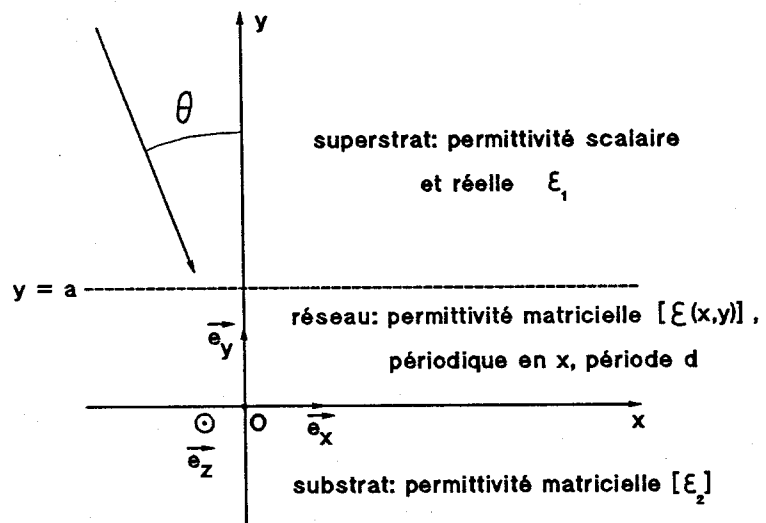
Figure 4.8.

*Efficacité extraordinaire
réfléchie en fonction de δ*

N.B. : Les résultats contenus dans ce chapitre, bien qu'ils datent en fait de plusieurs années, n'ont jamais été publiés. Leur obtention constituait essentiellement pour nous un entraînement en vue de l'étude des réseaux.

5.1. INTRODUCTION

Comme nous l'avons déjà signalé dans le cas isotrope, la méthode différentielle possède l'avantage de pouvoir être mise en oeuvre de façon relativement simple. Elle permet de traiter des structures de géométries très diverses ; toutes les situations décrites par la figure 5.1 entrent dans le cadre de cette méthode. Aucune hypothèse n'est faite sur les permittivités des milieux anisotropes. Les limites de la méthode sont les mêmes que dans le cas de milieux isotropes, à savoir une convergence de la solution (en fonction de l'ordre de troncature des séries) qui peut être assez lente (particulièrement pour des réseaux conducteurs de grande hauteur). Nous verrons tout de même au § 5.4 que, malgré cette contrainte, un grand nombre de situations pratiques intéressantes sont accessibles par cette méthode.

Figure 5.1

Nous utiliserons une adaptation au cas anisotrope de la "méthode différentielle classique" (§ 1.1). Le passage du cas isotrope au cas anisotrope se fait sans difficulté majeure ; par rapport au cas H// traité au § 1.1, cela se traduit par l'accroissement de fonctions inconnues (qui passe de 2 à 4) et, bien entendu, par un système différentiel d'expression plus lourde. C'est en réalité par cette étude que nous avons commencé notre

travail sur les réseaux anisotropes [75, 62], dans le cadre d'un contrat avec une société anglaise [51] ; on s'intéressait alors à la structure représentée sur la figure 5.2, qu'il était question d'utiliser dans le but de réaliser des mémoires optiques. Les études portant sur l'amélioration de la méthode différentielle (§ 1.2 et 1.3) ont été réalisées par la suite pour le cas isotrope, et nous n'avons pas jugé nécessaire, compte tenu des résultats du chapitre 1, de les mettre en oeuvre dans le cadre des réseaux anisotropes, jugeant que l'accroissement des possibilités de la méthode (en utilisant par exemple la "méthode différentielle améliorée") n'était pas suffisamment important en regard de l'accroissement du temps de calcul qui en résultait. Rien n'empêche toutefois de mettre en oeuvre la "méthode différentielle améliorée" pour l'étude des réseaux anisotropes si cela était jugé nécessaire.

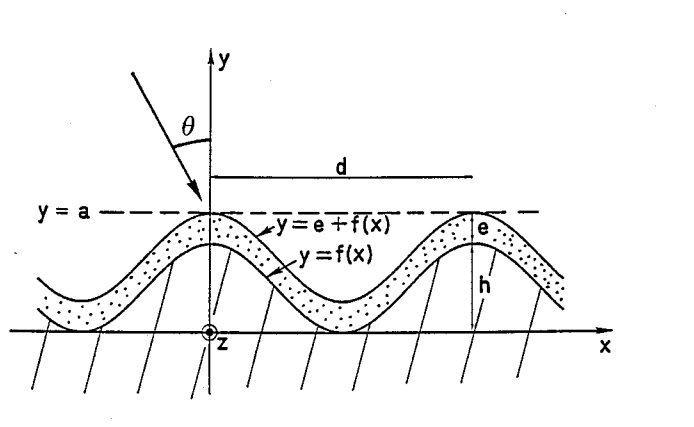


Figure 5.2. Un réseau sinusoïdal (éventuellement anisotrope) recouvert d'une couche anisotrope

On ne trouve dans la littérature que peu d'études relatives aux réseaux anisotropes. A notre connaissance, seuls une équipe de chercheurs Japonais [65, 66, 67], ainsi que Glytsis et Gaylord [22], ont publié à ce sujet. Leur approche est différente dans le sens qu'ils se sont uniquement intéressés à des "slanted gratings" qui, comme nous l'avons expliqué au § 1.1.4, constituent un cas particulier de l'étude que nous réalisons ici. Leurs études supposent de plus nombreuses restrictions sur la nature des milieux anisotropes : substrat isotrope dans [65], uniaxe dans [22] ; dans [22], il est supposé également que la zone modulée constituant le "slanted grating" est biaxe, autrement dit que la permittivité est diagonale dans les axes du réseau.

5.2. PRESENTATION DE LA METHODE

Comme au § 1.1.2, chaque composante de champ $u(x,y)$ est pseudo-périodique et peut s'écrire sous la forme :

$$u(x,y) = \sum_{n \in \mathbb{Z}} u_n(y) \psi_n(x) , \quad (5.1)$$

avec $\psi_n(x) = \exp(i\alpha_n x)$, $\alpha_n = k_0 n_1 \sin\theta + nK$, $K = 2\pi/d$; les $u_n(y)$ sont les coefficients de Fourier généralisés de $u(x,y)$.

Nous écrivons les équations de Maxwell comme il a été fait au § 3.2.2. Les quantités définies par (3.10), liées à la permittivité du milieu, sont maintenant des fonctions périodiques en x , qui peuvent s'écrire sous la forme de séries de Fourier. Par exemple, pour ε_a , nous écrirons :

$$\varepsilon_a(x,y) = \sum_{n \in \mathbb{Z}} (\varepsilon_a)_n(y) \exp(inKx). \quad (5.2)$$

Par la suite, et pour alléger l'écriture, nous ne mentionnerons plus la dépendance en y des coefficients de Fourier tels que $(\varepsilon_a)_n(y)$ que nous noterons simplement $(\varepsilon_a)_n$. En projetant les équations de Maxwell sur la base des $\psi_n(x)$, on obtient un système différentiel portant sur les coefficients de Fourier des champs :

$$\begin{aligned} \frac{dE_{xn}}{dy} = & -i\alpha_n \sum_m \{ (\varepsilon_b)_{n-m} E_{xm} + (\varepsilon_c)_{n-m} E_{zm} \} \\ & + i \frac{\eta_0}{k_0} \alpha_n \sum_m (\varepsilon_a)_{n-m} \alpha_m H_{zm} - ik_0 \eta_0 H_{zn}, \end{aligned} \quad (5.3)$$

$$\frac{dE_{zn}}{dy} = ik_0 \eta_0 H_{xn}, \quad (5.4)$$

$$\begin{aligned} \frac{dH_{xn}}{dy} = & i \frac{k_0}{\eta_0} \sum_m \{ (\varepsilon_h)_{n-m} E_{xm} + (\varepsilon_i)_{n-m} E_{zm} \} - i \frac{\alpha_n^2}{k_0 \eta_0} E_{zn} \\ & + i \sum_m (\varepsilon_e)_{n-m} \alpha_m H_{zm}, \end{aligned} \quad (5.5)$$

$$\begin{aligned} \frac{dH_{zn}}{dy} = & -i \frac{k_0}{\eta_0} \sum_m \{ (\varepsilon_f)_{n-m} E_{xm} + (\varepsilon_g)_{n-m} E_{zm} \} \\ & - i \sum_m (\varepsilon_d)_{n-m} \alpha_m H_{zm} \end{aligned} \quad (5.6)$$

et les composantes E_y et H_y peuvent encore se déduire des autres composan-

tes de champ par :

$$E_y = -\varepsilon_b E_x - \varepsilon_c E_z - i \frac{\eta_0}{k_0} \varepsilon_a \frac{\partial H_z}{\partial x}, \quad (5.7)$$

$$H_y = \frac{i}{k_0 \eta_0} \frac{\partial E_z}{\partial x}. \quad (5.8)$$

On notera que ces équations se réduisent :

- aux équations (3.11) à (3.16) dans un problème de "couches anisotropes" (c'est-à-dire dans une zone homogène, si les champs s'expriment de façon analogue à (3.8)),
- aux équations (1.11) dans le cas d'un réseau isotrope et en polarisation H//.

Pour simplifier cet exposé, nous supposons dès à présent que les composantes des champs sont "bien" représentées par un nombre fini $N = 2P + 1$ (P entier positif) de composantes sur la base des $\psi_n(x)$; en d'autres termes, les développements du type (5.1) sont tronqués en faisant "courir" n de $-P$ à $+P$. Cette troncature sera de toute manière nécessaire pour le traitement numérique.

Pour toute valeur de y , le champ est ainsi caractérisé par une colonne $F_N(y)$ possédant $4N$ composantes qui sont les $E_{xn}(y)$, $E_{zn}(y)$, $H_{xn}(y)$ et $H_{zn}(y)$. On écrira dorénavant le système (5.3) à (5.6) sous la forme matricielle :

$$\frac{dF_N(y)}{dy} = A(y) F_N(y) \quad (5.9)$$

où $A(y)$ est une matrice $4N \times 4N$ entièrement connue.

Pour $y > a$, chacune des composantes de champ vérifie l'équation de Helmholtz. Leurs coefficients de Fourier s'écrivent :

$$E_{xn}(y) = E_{xn}^- \exp(-i\beta_n y) + E_{xn}^+ \exp(i\beta_n y) \quad (5.10)$$

$$E_{zn}(y) = E_{zn}^- \exp(-i\beta_n y) + E_{zn}^+ \exp(i\beta_n y) \quad (5.11)$$

$$H_{xn}(y) = H_{xn}^- \exp(-i\beta_n y) + H_{xn}^+ \exp(i\beta_n y) \quad (5.12)$$

$$H_{zn}(y) = H_{zn}^- \exp(-i\beta_n y) + H_{zn}^+ \exp(i\beta_n y) \quad (5.13)$$

où β_n est toujours défini par $\beta_n^2 = k_0^2 \varepsilon_1 - \alpha_n^2$, la détermination de β_n étant imposée par (1.16). Les équations de Maxwell imposent de plus des relations entre les composantes de champ ; en effet, les équations (5.3) à (5.6) s'écrivent pour $y > a$:

$$\frac{dE_{xn}}{dy} = - \frac{i\eta_0 \beta_n^2}{k_0 \varepsilon_1} H_{zn} \quad (5.14)$$

$$\frac{dE_{zn}}{dy} = ik_0 \eta_0 H_{xn} \quad (5.15)$$

$$\frac{dH_{xn}}{dy} = \frac{i\beta_n^2}{k_0 \eta_0} E_{zn} \quad (5.16)$$

$$\frac{dH_{zn}}{dy} = - \frac{ik_0 \varepsilon_1}{\eta_0} E_{xn} \quad (5.17)$$

Parmi ces quatre équations, seules deux d'entre elles sont indépendantes compte tenu de (5.10) à (5.13). Retenons les équations (5.15) et (5.17) qui entraînent :

$$H_{xn}^\mp = \mp \frac{\beta_n}{k_0 \eta_0} E_{zn}^\mp \quad (5.18)$$

$$E_{xn}^\mp = \pm \frac{\eta_0 \beta_n}{k_0 \varepsilon_1} H_{zn}^\mp \quad (5.19)$$

Cela montre qu'après troncature, le champ dans la zone $y > a$ est déterminé par les 4N constantes E_{zn}^- , E_{zn}^+ , H_{zn}^- , H_{zn}^+ . En écrivant (5.10) à (5.13) pour $y = a$, il vient :

$$E_{xn}(a) = \frac{\eta_0 \beta_n}{k_0 \varepsilon_1} \left(H_{zn}^- e^{-i\beta_n a} - H_{zn}^+ e^{i\beta_n a} \right) \quad (5.20)$$

$$E_{zn}(a) = E_{zn}^- e^{-i\beta_n a} + E_{zn}^+ e^{i\beta_n a} \quad (5.21)$$

$$H_{xn}(a) = - \frac{\beta_n}{k_0 \eta_0} \left(E_{zn}^- e^{-i\beta_n a} - E_{zn}^+ e^{i\beta_n a} \right) \quad (5.22)$$

$$H_{zn}(a) = H_{zn}^- e^{-i\beta_n a} + H_{zn}^+ e^{i\beta_n a} \quad (5.23)$$

Ce système s'inverse immédiatement en :

$$E_{zn}^- = \frac{1}{2} e^{i\beta_n a} \left[E_{zn}(a) - \frac{k_0 \eta_0}{\beta_n} H_{xn}(a) \right] \quad (5.24)$$

$$H_{zn}^- = \frac{1}{2} e^{i\beta_n a} \left[H_{zn}(a) + \frac{k_0 \varepsilon_1}{\eta_0 \beta_n} E_{xn}(a) \right] \quad (5.25)$$

$$E_{zn}^+ = \frac{1}{2} e^{-i\beta_n a} \left[E_{zn}(a) + \frac{k_0 \eta_0}{\beta_n} H_{xn}(a) \right] \quad (5.26)$$

$$H_{zn}^+ = \frac{1}{2} e^{-i\beta_n a} \left[H_{zn}(a) - \frac{k_0 \varepsilon_1}{\eta_0 \beta_n} E_{xn}(a) \right] \quad (5.27)$$

Ces 4N équations permettent ainsi d'obtenir de façon élémentaire (équations (5.10) à (5.13)) le champ dans la zone $y > a$ dès que l'on connaît ses coefficients de Fourier en $y = a$.

pour $y < 0$, la matrice A figurant dans (5.9) ne dépend pas de y (zone homogène). Bien évidemment, si le substrat était isotrope, on pourrait procéder comme on vient de le faire (pour $y > a$) : on obtiendrait ainsi le champ pour $y < 0$ sous la forme de développements de Rayleigh dont les coefficients seraient directement liés aux coefficients de Fourier du champ en $y = 0$.

Dans le cas d'un substrat anisotrope, nous utilisons une procédure analogue à celle décrite dans le cas des couches (§ 4.2). La solution de

$$\frac{dF_N(y)}{dy} = A F_N(y) \quad (5.28)$$

s'écrit en fonction des valeurs propres $\lambda_j = i\beta_j$ de A et des vecteurs propres F_j associés (les λ_j et F_j sont déterminés numériquement) sous la forme :

$$F_N(y) = \sum_{j=1}^{4N} C_j e^{i\beta_j y} F_j \quad (5.29)$$

Le champ s'exprime donc pour $y < 0$ sous la forme d'une combinaison linéaire de 4N ondes planes d'amplitudes C_j . Ces 4N ondes planes ne sont pas toutes acceptables pour représenter le champ dans le substrat. Seule la moitié d'entre elles vérifie une condition d'onde sortante vers le bas. Elles sont

sélectionnées en utilisant les résultats du § 3.3 (pour les ondes propagatives, on raisonne de préférence sur les vecteurs de Poynting, § 3.3.4.1).

Remarque : dans la zone homogène $y < 0$, seul le coefficient de Fourier central des permittivités est non nul : $(\epsilon_a)_{n,m} = \epsilon_a \delta_{nm}$, et idem pour $\epsilon_b, \dots, \epsilon_i$. En comparant les équations (5.3) à (5.6) aux équations (3.11) à (3.15), on constate que l'on a affaire à N "problèmes de couches" indépendants, chacun de ces "problèmes de couches" étant obtenu en remplaçant α par α_n dans les équations (3.11) à (3.15). Rappelons que dans un "problème de couche" (§ 3.3), deux et deux seulement des quatre ondes planes vérifient la C.O.S. Il en résulte que dans le problème de réseaux (tronqué à l'ordre N), $2N$ et $2N$ seulement des $4N$ ondes planes vérifient la C.O.S.

En appelant $\mathcal{J}$ l'ensemble des indices j associés aux ondes vérifiant la C.O.S. (de cardinal égal à $2N$), le champ s'écrit donc pour $y < 0$:

$$F_N(y) = \sum_{j \in \mathcal{J}} C_j e^{i\beta_j y} F_j, \quad (5.30)$$

d'où :

$$F_N(0) = \sum_{j \in \mathcal{J}} C_j F_j. \quad (5.31)$$

Nous pouvons résumer ce qui précède de la façon suivante :

* pour $y > a$, le champ s'écrit en fonction de $4N$ constantes : E_{zn}^- et H_{zn}^- (associées au champ incident), E_{zn}^+ et H_{zn}^+ (associées au champ "réfléchi"). Ces $4N$ constantes s'obtiennent trivialement à partir de $F_N(a)$: équations (5.24) à (5.27) ;

$$* \text{ pour } 0 < y < a, \quad \frac{dF_N(y)}{dy} = A(y) F_N(y) \quad (5.32)$$

* pour $y < 0$, le champ s'écrit en fonction de $2N$ constantes C_n , et

$$F_N(0) = \sum_{j \in \mathcal{J}} C_j F_j. \quad (5.33)$$

Ce problème est résolu comme au § 4.2 : la linéarité du problème entraîne qu'il existe deux matrices de "transfert" T_1 et T_2 (de dimension $2N \times 2N$) telles que (on note E_z^- la colonne à N éléments contenant les E_{zn}^- , idem pour H_z^- , E_z^+ et H_z^+ ; on note C la colonne à $2N$ éléments contenant les C_n) :

$$\begin{bmatrix} E_z^- \\ H_z^- \end{bmatrix} = T_1 \begin{bmatrix} C \end{bmatrix} \quad (5.34)$$

$$\begin{bmatrix} E_z^+ \\ H_z^+ \end{bmatrix} = T_2 \begin{bmatrix} C \end{bmatrix}. \quad (5.35)$$

On effectue 2N "tirs" en donnant à C des valeurs particulières (pour le $i^{\text{ème}}$ tir, on prend $C_i = 1$ et pour $n \neq i$, $C_n = 0$). Pour chaque tir, C étant fixé, on en déduit $F_N(0)$ (éq.(5.33)), puis $F_N(a)$ par intégration numérique de (5.32) (on utilise le même algorithme qu'au § 1.1.5), et enfin E_z^- , H_z^- , E_z^+ et H_z^+ . Le $i^{\text{ème}}$ tir permet ainsi de construire les $i^{\text{èmes}}$ colonnes de T_1 et T_2 . Une fois les 2N tirs effectués, T_1 et T_2 sont déterminées. On donne alors à E_z^- et H_z^- les valeurs imposées par le champ incident du problème considéré, on en déduit le champ "transmis" C par inversion de (5.34), puis le champ "réfléchi" par (5.35).

Remarque : cette manière de décrire la résolution du problème est probablement la plus concise, et le programme que nous avons réalisé suit cette procédure. Il aurait été aussi possible de faire une présentation légèrement différente (comme nous l'avons fait au § 1.1 pour le cas isotrope) en termes d'espaces vectoriels, caractérisés par des bases. Les relations (5.18) et (5.19) traduisent en effet (pour $y > a$) l'appartenance du champ diffracté (lié aux 4N constantes $E_{x_n}^+$, $E_{z_n}^+$, $H_{x_n}^+$, $H_{z_n}^+$) à un espace vectoriel $\mathcal{E}_1^+(y)$ de dimension 2N, dont on peut déterminer explicitement une base. La relation (5.30) traduit (pour $y < 0$) l'appartenance du champ à un espace vectoriel $\mathcal{E}_2^-(y)$ dont on connaît une base. Le problème se présente ainsi sous une forme analogue à celle décrite par les équations (1.26) à (1.28). Le raisonnement du § 1.1.3 pourrait alors être reconduit ici.

5.3. LES "SLANTED GRATINGS"

A notre connaissance, toutes les publications étrangères au Laboratoire relatives aux réseaux anisotropes traitent uniquement de ce type de réseaux [65, 22, 66, 67] qui se caractérisent par le fait que la permittivité $[\epsilon]$ ne dépend (dans la zone modulée $0 < y < a$) que du produit scalaire

$\vec{L} \cdot \vec{r}$ (Figure 1.2) ; elle est donc constante dans tout plan perpendiculaire à $\vec{L}$. Nous n'en reprendrons pas une étude détaillée, la méthode exposée au § 1.1.4 pouvant facilement être adaptée au cas anisotrope. Rappelons que moyennant un changement de fonctions inconnues approprié, la problème se met sous la même forme que dans le cas plus général que nous avons envisagé au § 5.2, la simplification venant du fait que, dans l'équation (5.32), $A(y)$ est remplacée par une matrice $\tilde{A}$ indépendante de y . Le principe de résolution reste identique, mais l'intégration numérique entre $y = 0$ et $y = a$ est remplacée par une recherche de valeurs propres et de vecteurs propres de $\tilde{A}$, en tous points semblable à celle que nous avons exposée dans le cas des couches anisotropes au § 4.2 (la seule différence étant la taille de la matrice : 4×4 pour les couches, $4N \times 4N$ pour les réseaux, c'est-à-dire que l'on exprime le champ dans la zone $0 < y < a$ comme combinaison linéaire de $4N$ ondes planes). Nous n'avons pas réalisé pour l'instant de logiciel mettant en application ces considérations. Nous ne sommes pas convaincus du fait que l'on puisse espérer pouvoir traiter ainsi (et sans approximation) des réseaux de profondeur importante. En effet, si l'on retient une valeur N élevée en vue de bien représenter les composantes de champ en séries de Fourier, on ne manquera pas d'obtenir, parmi les $4N$ ondes planes envisagées plus haut, un grand nombre d'ondes évanescentes et antiévanescentes (c'est-à-dire ayant des constantes de propagation selon y présentant une partie imaginaire) qui risquent fort de conduire à des problèmes numériques si le réseau est profond.

5.4. QUELQUES RESULTATS

5.4.1. Matériaux à pertes

Dans le cadre d'un contrat avec une société privée (P.A. Technology, Cambridge) qui s'intéressait à la réalisation de mémoires optiques, nous avons été amenés à considérer le cas d'un réseau sinusoïdal isotrope recouvert d'une couche anisotrope (Fig. 5.2). La couche anisotrope était constituée d'un matériau gyrotrope (ou gyromagnétique) : c'est le cas de certains métaux ou composés métalliques qui, en présence d'un champ magnétique statique, acquièrent une anisotropie. Nous avons entre autres étudié une couche de Cobalt qui, dans les conditions des expériences envisagées (champ magnétique statique uniforme colinéaire à $\vec{e}_y$) possède une permittivité :

$$[\varepsilon] = \varepsilon_c \begin{bmatrix} 1 & 0 & Q \\ 0 & 1 & 0 \\ -Q & 0 & 1 \end{bmatrix},$$

avec $\varepsilon_c = -8.19 + i 16.38$ et $Q = 0.0069 + i 0.0268$. On notera que cette permittivité est très faiblement anisotrope, et on s'attend donc à avoir des effets anisotropes (tels que le changement de polarisation) peu marqués. Pour illustrer cela, envisageons tout d'abord le cas d'une couche de ce matériau déposée sur un substrat plan, éclairée par une onde plane (Figure 5.3). On s'intéresse principalement à l'onde réfléchie, et aux changements de polarisation de cette onde par rapport à l'onde incidente.

Figure 5.3

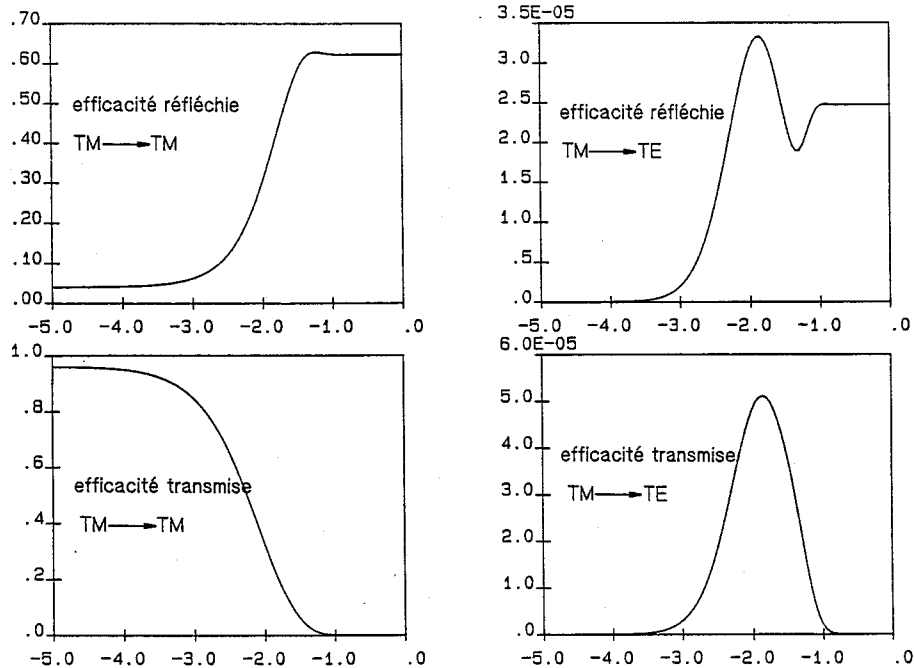
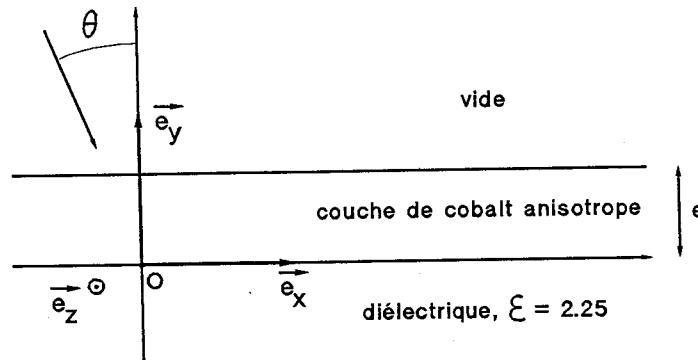


Figure 5.4. Efficacités en fonction de $\log_{10}(e)$ pour la structure de la figure 5.3 ; $\theta = 0$, $\lambda_0 = 0.6328 \mu\text{m}$; e est exprimé en μm .

La figure 5.4 montre, dans le cas de polarisation TM (champ magnétique de

l'onde incidente colinéaire à $\vec{e}_z$) et pour $\theta = 0$, les efficacités réfléchies et transmises par cette structure. La notation $TM \rightarrow TM$ indique qu'il s'agit de l'efficacité associée à la partie de l'onde ayant la même polarisation que l'onde incidente, et $TM \rightarrow TE$ indique qu'il s'agit de l'efficacité associée à la fraction de l'onde polarisée TE. Bien entendu, pour des épaisseurs e très faibles, la couche ne joue aucun rôle, alors que pour des épaisseurs importantes (à partir de $\lambda_0/4$ environ) le champ ne traverse pas la couche et le substrat n'a aucune influence. Entre ces valeurs, on constate un maximum de changement de polarisation pour $\log_{10}(e) \simeq -1.9$ c'est-à-dire $e \simeq 0.013 \mu\text{m} \approx \lambda_0/50$. La dépolarisation de l'onde reste néanmoins extrêmement faible.

Le but de l'étude que nous avons menée était de savoir si l'on pouvait accroître ces changements de polarisation en remplaçant les interfaces planes par des réseaux. La figure 5.5 montre l'évolution de l'efficacité réfléchie $TM \rightarrow TE$ lorsque l'on déforme progressivement la structure de la figure 5.3 pour obtenir le réseau sinusoïdal représenté sur la figure 5.2. On constate que cette efficacité décroît quand la hauteur h du réseau augmente. Pour évaluer la précision de nos calculs, nous les avons effectués pour deux valeurs différentes de l'ordre de troncature N ($N = 11$ et $N = 19$). La convergence n'étant pas assurée de façon parfaite, nous avons également comparé les efficacités $TE \rightarrow TE$ obtenues au moyen du programme différentiel anisotrope à celles données par un programme intégral isotrope pour lequel on a remplacé la couche anisotrope par une couche isotrope (en prenant $Q = 0$) ; cela est possible compte tenu de la très faible valeur de l'anisotropie Q . Les courbes obtenues (Fig. 5.6) montrent que la précision du programme différentiel peut être jugée satisfaisante tant que la hauteur h reste inférieure à $0.1 \mu\text{m}$ environ.

Figure 5.5. Efficacité $TM \rightarrow TE$ réfléchie dans l'ordre 0 en fonction de h (en μm) pour le réseau de la fig. 5.2. $\theta = 0$, $e = 0.013 \mu\text{m}$, pas du réseau : $d = 0.3 \mu\text{m}$, couche de cobalt anisotrope, indice du substrat : $n = 1.5$, $\lambda_0 = 0.6328 \mu\text{m}$. En tireté : $N = 11$; en trait mixte : $N = 19$.

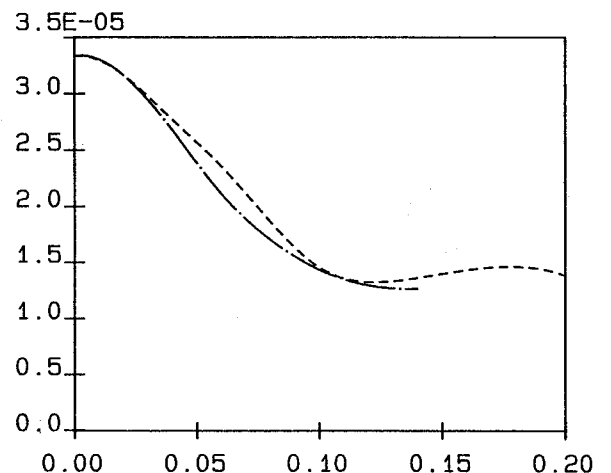
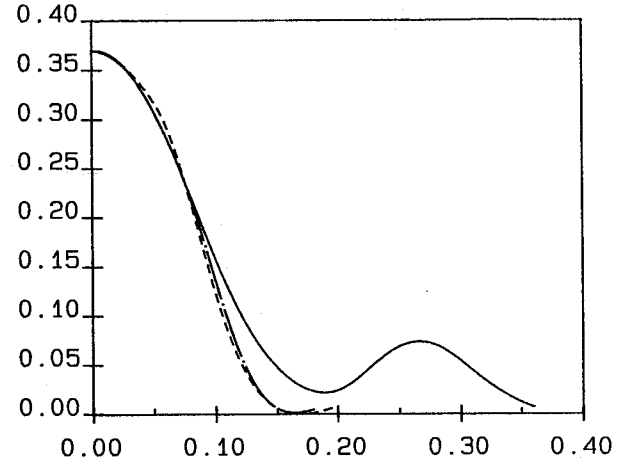


Figure 5.6. Efficacité $TM \rightarrow TM$ réfléchie dans l'ordre 0 en fonction de h (en μm). Mêmes valeurs numériques que pour la fig.5.5. En tireté : $N = 11$, en trait mixte : $N = 19$, en trait continu : résultats du programme intégral pour lequel on a pris $Q = 0$.



Les résultats précédents montrent qu'il n'est pas possible d'accroître les changements de polarisation en remplaçant les interfaces planes par des réseaux (pour les matériaux envisagés ici tout au moins). Nous savons par ailleurs qu'un réseau isotrope conducteur peut, dans le cas de polarisation TM, donner lieu au phénomène d'absorption totale. S'il n'est pas possible d'augmenter les efficacités $TM \rightarrow TE$, on peut envisager néanmoins de rendre très faible, voire nulle, l'efficacité $TM \rightarrow TM$, de façon à augmenter l'ellipticité de l'onde réfléchie. Nous avons recherché les paramètres (d et h) qui donnent le phénomène d'absorption totale pour l'onde $TM \rightarrow TM$ réfléchie dans l'ordre 0, pour l'incidence $\theta = 0^\circ$ et pour la valeur $e = 0.013 \mu m$. On obtient : $d \simeq 0.62 \mu m$ et $h \simeq 0.124 \mu m$.

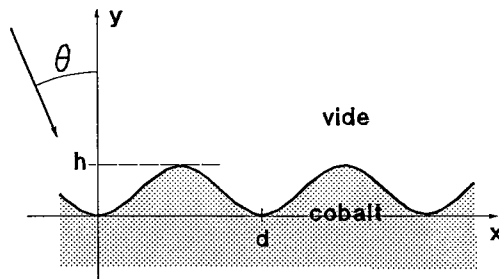


Figure 5.7. Réseau sinusoïdal. Le substrat est constitué de cobalt de permittivité $[\epsilon] = \epsilon_c \begin{bmatrix} 1 & 0 & Q \\ 0 & 1 & 0 \\ -Q & 0 & 1 \end{bmatrix}$, $\epsilon_c = -8.19 + i 16.38$, $Q = 0.0069 + i 0.0268$. $h = 0.1 \mu m$, $\lambda_0 = 0.6328 \mu m$, $d = 0.6 \mu m$, $\theta = 30^\circ$. Il existe donc deux ordres diffractés (ordres -1 et 0).

Pour illustrer le cas d'un substrat anisotrope à pertes, envisageons le réseau représenté sur la figure 5.7. Compte tenu de la très faible ani-

sotropie du matériau, la dépolarisation de l'onde est extrêmement faible. On peut ici aussi comparer les résultats obtenus à ceux donnés par un programme intégral isotrope (tableau 5.8). On constate un accord excellent dans le cas d'une onde incidente polarisée TE, et une convergence assez lente de la méthode différentielle en fonction de l'ordre de troncature N pour une onde incidente polarisée TM (cette convergence serait bien plus rapide pour des réseaux de moindre hauteur).

Onde incidente TE :

	Ordre -1		Ordre 0	
	TE $\rightarrow$ TE	TE $\rightarrow$ TM	TE $\rightarrow$ TE	TE $\rightarrow$ TM
N = 13	.1048	3.7E-6	.5427	15E-6
N = 17	.1048	4.0E-6	.5430	14E-6
N = 21	.1049	4.3E-6	.5431	14E-6
N = 25	.1049	4.4E-6	.5431	14E-6
N = 29	.1049	4.5E-6	.5432	14E-6
Méthode intégrale	.1050		.5430	

Onde incidente TM :

	Ordre -1		Ordre 0	
	TM $\rightarrow$ TM	TM $\rightarrow$ TE	TM $\rightarrow$ TM	TM $\rightarrow$ TE
N = 13	.207	3.9E-6	.282	15E-6
N = 17	.216	4.2E-6	.287	14E-6
N = 21	.223	4.4E-6	.290	14E-6
N = 25	.227	4.5E-6	.292	14E-6
N = 29	.230	4.7E-6	.294	14E-6
Méthode intégrale	.251		.306	

Tableau 5.8. Efficacités réfléchies par le réseau de la figure 5.7. Résultats de la méthode différentielle pour différentes valeurs de N comprises entre 13 et 29 (80 pas d'intégration sont utilisés pour l'intégration de (5.32)). On a aussi reporté les résultats donnés par une méthode intégrale dans le cas du substrat isotrope obtenu en prenant $Q = 0$.

5.4.2. Matériaux sans pertes

Considérons le réseau représenté sur la figure 5.9. La couche anisotrope est constituée d'un dépôt de TiO_2 dont les caractéristiques ont été extraites de [20], et qui a déjà été étudiée au § 4.4.2 dans le cadre des

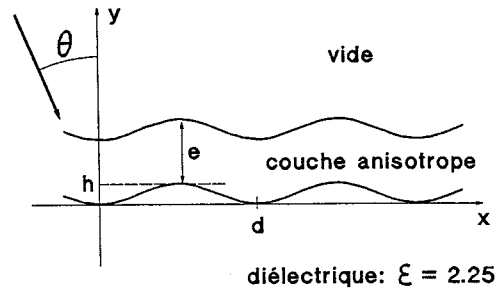


Figure 5.9. Réseau sinusoïdal ($d = 0.8 \mu\text{m}$, $h = 0.1 \mu\text{m}$) recouvert d'une couche anisotrope d'épaisseur $e = 0.495 \mu\text{m}$.

empilements de couches. La permittivité de ce matériau se diagonalise avec les valeurs propres $\epsilon_{xx} = 4.35$, $\epsilon_{yy} = 4.85$, $\epsilon_{zz} = 4.59$. Après deux rotations définies au § 4.4.2 (d'un angle $\psi = 18.6^\circ$ autour de $\vec{e}_z$, puis d'un angle $\delta = 40^\circ$ autour de $\vec{e}_y$), la permittivité de cette couche devient dans la base $(\vec{e}_x, \vec{e}_y, \vec{e}_z)$:

$$[\epsilon] = \begin{bmatrix} 4.4800 & -0.1174 & 0.0923 \\ -0.1174 & 4.7975 & 0.0985 \\ 0.0923 & 0.0985 & 4.5125 \end{bmatrix}.$$

Le tableau 5.10, extrait de [62], montre les efficacités obtenues dans les deux cas suivants : tout d'abord en retenant $n = 15$ termes dans les développements des champs en séries de Fourier généralisées (5.1) et en utilisant $J = 150$ pas pour l'intégration de (5.32) entre $y = 0$ et $y = a$, puis pour $N = 9$ et $J = 80$. L'accord entre les deux calculs peut être jugé satisfaisant. On notera que la hauteur de la région modulée du réseau est ici relativement importante, voisine de la longueur d'onde de la lumière : $(e + h)/\lambda_0 = 0.85$.

	Order	R e f l e c t e d			T r a n s m i t t e d		
		Angle (degrees)	TM $\rightarrow$ TM	TM $\rightarrow$ TE	Angle	TM $\rightarrow$ TM	TM $\rightarrow$ TE
N=15 150 steps	- 2				-66.48	0.5719 E-2	0.1656 E-3
	- 1	-30.03	0.2305 E-1	0.5374 E-4	-19.49	0.8689 E-2	0.1063 E-3
	0	22.00	0.4364 E-1	0.4088 E-3	14.46	0.8490	0.1137 E-1
	1				56.42	0.5614 E-1	0.1610 E-2
N=9 80 steps	- 2				-66.48	0.5548 E-2	0.1626 E-3
	- 1	-30.03	0.2301 E-1	0.5155 E-4	-19.49	0.8508 E-2	0.1055 E-3
	0	22.00	0.4361 E-1	0.4040 E-3	14.46	0.8502	0.1137 E-1
	1				56.42	0.5557 E-1	0.1591 E-2

Tableau 5.10. Efficacités réfléchies et transmises par le réseau de la figure 5.9 pour $\lambda_0 = 0.7 \mu\text{m}$, $\theta = 22^\circ$, et un champ incident polarisé TM.

5.5. CONCLUSION

Les études numériques montrent ainsi que le domaine d'utilisation de la méthode différentielle est sensiblement identique pour les matériaux anisotropes et les matériaux isotropes. Dans le domaine de résonance (λ_0/d voisin de l'unité), on peut espérer de bons résultats pour des matériaux sans pertes tant que la hauteur de la région modulée (notée a sur la figure 5.1) reste inférieure à la longueur d'onde. Pour des matériaux fortement conducteurs, on ne pourra espérer traiter avec une précision convenable (de l'ordre de 1 % sur les efficacités) que des réseaux pour lesquels le rapport a/λ_0 reste inférieur à 1/10 environ ; les principales difficultés (convergence lente de la solution en fonction de l'ordre de troncature N) apparaissent principalement lorsque la lumière incidente est polarisée TM.

Nonobstant ces restrictions sur la hauteur du réseau, il faut souligner le large éventail des structures pouvant être étudiées par cette méthode. En effet, tous les réseaux pouvant être décrits comme schématisé sur la figure 5.1 sont susceptibles d'être traités par la méthode différentielle, sans aucune contrainte sur les permittivités. On peut par exemple envisager l'étude de structures dans lesquelles la permittivité est continûment et périodiquement modulée (sous l'action de contraintes extérieures par exemple).

6.1. INTRODUCTION

Nous avons vu précédemment (Chapitres 1 et 5) que la méthode différentielle ne suffit pas à nos besoins, principalement dans le cas de réseaux conducteurs ou de hauteur importante. Les diverses études réalisées au Laboratoire dans le cas de réseaux isotropes ont montré que les méthodes "intégrales" sont généralement plus performantes. Il nous a ainsi paru indispensable de tenter de généraliser ces dernières au cas de structures anisotropes. Dans le cas isotrope, de nombreux articles ont été publiés au sujet des méthodes intégrales. Le lecteur intéressé pourra trouver les principales références relatives à ce problème dans l'article de R. Petit [57] qui dresse un état des connaissances sur le sujet (en 1975), dans "Integrals methods" par D. Maystre ([58], Chapitre 3, 1980) ou bien dans "Rigorous vector theories of diffraction gratings" par D. Maystre ([38], 1984). Le propos de ce chapitre est de détailler nos résultats annoncés dans [76, 77, 78, 79]. Nous utilisons un repère orthonormé $0, \vec{e}_1, \vec{e}_2, \vec{e}_3$. Le réseau est la surface d'équation $y = f(x)$, invariante par translation suivant $\vec{e}_3$ (Figure 6.1). La fonction périodique f , de période d , décrit le profil $\mathcal{P}$ de sa section droite. Nous la supposons de classe C^2 . La perméabilité est partout égale à μ_0 . Le domaine Ω^+ ($y > f(x)$) est rempli d'un diélectrique isotrope d'indice n_1 réel (permittivité $\epsilon_1 = n_1^2$) et le domaine Ω^- ($y < f(x)$) par un milieu anisotrope de permittivité $[\epsilon_2] = \begin{bmatrix} \epsilon_{xx} & 0 & 0 \\ 0 & \epsilon_{yy} & 0 \\ 0 & 0 & \epsilon_{zz} \end{bmatrix}$

éventuellement complexe, mais nécessairement diagonale dans la base $\vec{e}_1, \vec{e}_2, \vec{e}_3$ "attachée" au réseau. La raison pour laquelle nous supposons la matrice $[\epsilon_2]$ diagonale résulte uniquement du fait que nous n'avons pas encore pu expliciter les "vecteurs de Green" du problème dans le cas où $[\epsilon_2]$ est quelconque (chapitre 3). La structure est éclairée sous l'incidence θ par une onde plane dont le vecteur d'onde est situé dans le plan $(\vec{e}_1, \vec{e}_2)$.

Nous notons $\vec{n}(M)$ ou $\vec{n}(x)$ le vecteur unitaire normal au profil au point $M(x, f(x))$ et dirigé vers Ω^+ , et $\vec{t}(x)$ le vecteur unitaire tangent en M au profil et dirigé dans le sens des x croissants (Figure 6.1).

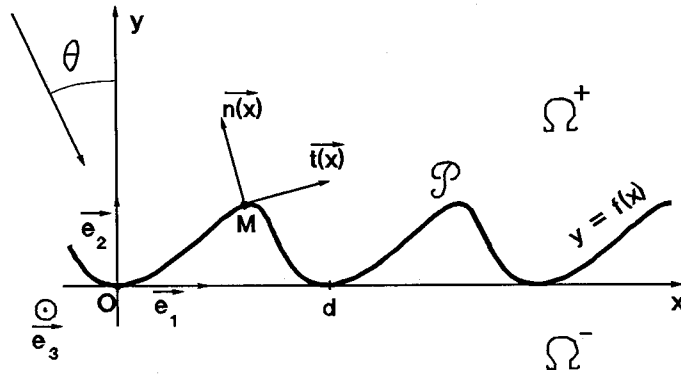


Figure 6.1. Ω^+ est rempli d'un milieu isotrope de permittivité ϵ_1 et Ω^- d'un milieu anisotrope de permittivité $[\epsilon_2]$.

La méthode de résolution que nous utilisons est une adaptation au cas anisotrope d'une procédure devenue classique. Les étapes essentielles sont les suivantes :

- détermination des solutions élémentaires des équations de propagation dans l'espace libre rempli des milieux de permittivité $[\epsilon_2]$ ou ϵ_1 ,
- choix de fonctions inconnues définies sur $\mathcal{P}$; expression du champ diffracté dans chaque domaine Ω^+ et Ω^- à partir de ces fonctions inconnues (formules de Kirchhoff-Helmholtz généralisées),
- obtention d'équations intégrales pour les fonctions inconnues après un passage à la limite convenable (faire tendre un point P situé dans Ω^+ ou Ω^- vers un point M du profil),
- résolution numérique des équations intégrales obtenues,
- détermination du champ diffracté.

Compte tenu de la nature vectorielle du problème, nous sommes amenés à considérer quatre fonctions inconnues (les composantes tangentielles du champ sur $\mathcal{P}$). Nous obtenons (étape c) un système de quatre équations intégrales couplées. Nous ne nous sommes pas attachés à essayer de réduire le nombre de fonctions inconnues (et conséquemment le nombre d'équations intégrales) comme l'ont proposé D. Maystre et quelques autres auteurs [86, 29] dans le cas isotrope. Cela ne paraît pas être nécessaire compte tenu de la capacité et de la puissance des moyens informatiques mis à notre disposition actuellement. Il est néanmoins possible de réduire le nombre de fonctions inconnues de moitié (§ 6.10, ou [76]).

Nous avons vu précédemment (§ 3.5) que, compte tenu de la forme diagonale de $[\epsilon_2]$, les deux cas fondamentaux de polarisation se découplent. Cela entraîne un découplage des fonctions inconnues, et un découplage du système à résoudre en deux systèmes indépendants, chacun d'eux constituant un sous-système de deux équations intégrales couplées. Toujours d'après le § 3.5, le cas TE peut se résoudre éventuellement au moyen de programmes conçus pour des milieux isotropes, en remplaçant le milieu situé dans Ω^- par un milieu isotrope de permittivité scalaire ϵ_{zz} . Pour la clarté et la généralité de l'exposé, il ne nous a pas paru nécessaire de prendre en compte cette remarque dès le début de l'étude. Dans les étapes a, b et c, nous considérerons donc le problème vectoriel.

Dans le traitement numérique du problème (étape d), deux procédés déjà concurremment utilisés au Laboratoire en milieu isotrope peuvent être envisagés. Dans la "méthode des séries de Fourier" (Fourier Series Method, ou FSM) [60, 55, 39, 40], on représente les fonctions inconnues (pseudo-périodiques) par leurs développements en séries de Fourier généralisées, et on cherche à déterminer les coefficients de Fourier de ces fonctions inconnues. Dans la PMM (Point Matching Method), on s'intéresse à la valeur des fonctions inconnues en certains points de discrétisation localisés sur le profil. Nous avons retenu pour cette étude la FSM qui, bien qu'un peu plus difficile à mettre en oeuvre que la PMM, donne entière satisfaction.

6.2. SOLUTIONS ELEMENTAIRES DES EQUATIONS DE PROPAGATION

Comme auparavant, nous posons :

$$\alpha_0 = k_0 n_1 \sin \theta, \quad \alpha_n = \alpha_0 + n 2\pi/d. \quad (6.1)$$

Notre formulation nécessite la connaissance des solutions élémentaires de l'équation de propagation dans l'espace libre rempli des milieux de permittivité $[\epsilon_2]$ et ϵ_1 . Considérons d'abord le cas du milieu anisotrope de permittivité $[\epsilon_2]$ et, pour $i = 1, 2, 3$, les solutions $\vec{g}_i(x, y)$ ("vecteurs de Green") vérifiant une condition d'onde sortante pour $y \rightarrow \pm \infty$ de l'équation :

$$\text{rot rot } \vec{g}_i - k_0^2 {}^t[\epsilon_2] \vec{g}_i = \vec{e}_i \sum_{n \in \mathbb{Z}} \frac{1}{d} \delta(y) \exp(-i\alpha_n x) = \vec{e}_i \delta_{\mathcal{R}}(-x, y) \quad (6.2)$$

où $\delta_{\mathcal{R}}(x, y)$ est l'unité de convolution de $\mathcal{R}'$ (§ 3.4.1).

Pour des raisons qui apparaîtront ensuite, on notera que nous considérons ici la transposée de la matrice de permittivité du matériau. Cela peut paraître superflu compte tenu du fait que nous nous limitons à des matrices $[\varepsilon_2]$ diagonales ; mais ce serait nécessaire si elles ne l'étaient pas.

La recherche des $\vec{g}_i$ a été détaillée au § 3.4.3. Leurs composantes sur $\vec{e}_1, \vec{e}_2, \vec{e}_3$ sont données par les équations (3.58) à (3.62) dans lesquelles il suffit de remplacer α_n par $-\alpha_n$. Les solutions élémentaires de l'équation de propagation dans le milieu isotrope de permittivité ε_1 s'en déduisent immédiatement en remplaçant $[\varepsilon_2]$ par $\varepsilon_1 [I]$ ($[I]$: matrice identité).

6.3. GENERALISATION DES FORMULES DE KIRCHHOFF-HELMHOLTZ

Dans un premier temps, voyons comment il est possible d'exprimer le champ en tout point de Ω^- (rempli du milieu de permittivité $[\varepsilon_2]$) en fonction de la valeur de certaines de ses composantes sur $\mathcal{P}$. Considérons deux points $P(x, y)$ et $P'(x', y')$ (x appartient à l'intervalle période $[0, d[$, figure 6.2) et posons $\vec{G}_i(P', P) = \vec{g}_i(x' - x, y' - y)$. De (6.2), nous déduisons :

$$\text{rot}_{P'} \text{rot}_{P'} \vec{G}_i(P', P) - k_0^2 {}^t[\varepsilon_2] \vec{G}_i(P', P) = \vec{e}_i \delta_{\mathcal{R}}(-(x' - x), y' - y) \quad (6.3)$$

Soit $\vec{E}$ un champ électrique vérifiant l'équation de propagation dans Ω^- :

$$\text{rot rot } \vec{E}(P') - k_0^2 [\varepsilon_2] \vec{E}(P') = 0 \quad \text{dans } \Omega^-, \quad (6.4)$$

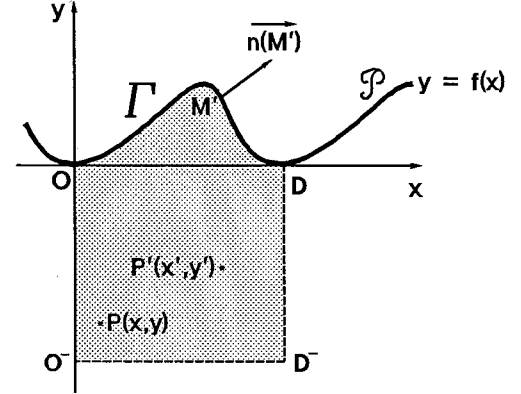
et une condition d'ondes sortantes pour $y \rightarrow -\infty$. Multiplions scalairement (6.3) par $\vec{E}(P')$ et (6.4) par $\vec{G}_i(P', P)$ et retranchons les deux équations obtenues. On vérifie immédiatement que la quantité $\vec{E}(P') \cdot \left\{ {}^t[\varepsilon_2] \vec{G}_i(P', P) \right\} - \vec{G}_i(P', P) \cdot \left\{ [\varepsilon_2] \vec{E}(P') \right\}$ est nulle (c'est ici que se justifie le choix de la matrice ${}^t[\varepsilon_2]$ introduite dans (6.2)). On utilise l'identité :

$$\begin{aligned} \vec{E}(P') \cdot \text{rot}_{P'} \text{rot}_{P'} \vec{G}_i(P', P) - \vec{G}_i(P', P) \cdot \text{rot rot } \vec{E}(P') &= \\ &= \text{div}_{P'} \left\{ \vec{G}_i(P', P) \wedge \text{rot } \vec{E}(P') - \vec{E}(P') \wedge \text{rot}_{P'} \vec{G}_i(P', P) \right\} \end{aligned}$$

pour obtenir :

$$\begin{aligned} \operatorname{div}_{P'} \left\{ \vec{G}_i(P', P) \wedge \operatorname{rot} \vec{E}(P') - \vec{E}(P') \wedge \operatorname{rot}_{P'} \vec{G}_i(P', P) \right\} = \\ = \vec{e}_i \cdot \vec{E}(x', y') \sum_{n \in \mathbb{Z}} \frac{1}{d} \delta(y - y') \exp(i\alpha_n(x - x')) \end{aligned} \quad (6.5)$$

Figure 6.2. Γ est l'arc OD de $\mathcal{P}$



En remplaçant, dans le second membre de (6.5), $\sum_{n \in \mathbb{Z}} \frac{1}{d} \exp(i\alpha_n(x - x'))$ par $\sum_{n \in \mathbb{Z}} \exp(i\alpha_0 nd) \delta(x - x' - nd)$, puis en intégrant à l'intérieur du domaine grisé (figure 6.2), on obtient :

$$\begin{aligned} \int_{\Gamma \cup DD^- \cup O^-O} \vec{n}(M') \cdot \left\{ \vec{G}_i(M', P) \wedge \operatorname{rot} \vec{E}(M'_-) - \vec{E}(M'_-) \wedge \operatorname{rot}_{M'} \vec{G}_i(M', P) \right\} dM' = \\ = \begin{cases} \vec{e}_i \cdot \vec{E}(P) & \text{si } y < f(x) \\ 0 & \text{si } y > f(x) \end{cases} \end{aligned} \quad (6.6)$$

Dans cette expression, la normale $\vec{n}(M')$ est orientée vers l'extérieur du domaine grisé et $\operatorname{rot} \vec{E}(M'_-)$ désigne la limite de $\operatorname{rot} \vec{E}(Q)$ lorsque Q tend vers M' en restant dans Ω^- (idem pour $\vec{E}(M'_-)$). Les contributions des côtés OO^- et DD^- sont opposées et disparaissent. Les conditions d'ondes sortantes vérifiées par $\vec{G}_i(M', P)$ et $\vec{E}(M')$ entraînent la nullité de la contribution du segment O^-D^- (Annexe 11, lemme 3). D'où :

$$\begin{aligned} - \int_{\Gamma} \left\{ \vec{G}_i(M', P) \cdot (\vec{n}(M') \wedge \operatorname{rot} \vec{E}(M'_-)) + \operatorname{rot}_{M'} \vec{G}_i(M', P) \cdot (\vec{n}(M') \wedge \vec{E}(M'_-)) \right\} dM' = \\ = \begin{cases} \vec{e}_i \cdot \vec{E}(P) & \text{si } y < f(x) \\ 0 & \text{si } y > f(x) \end{cases} \end{aligned} \quad (6.7)$$

En remplaçant $\text{rot } \vec{E}$ par $i\omega \mu_0 \vec{H}$, on obtient ainsi les composantes de $\vec{E}(P)$ en fonction des composantes tangentielles de $\vec{E}$ et $\vec{H}$ sur Γ sous la forme :

$$\begin{aligned} - \int_{\Gamma} [T_1(M', P)] (\vec{n}(M') \wedge \vec{H}(M')) dM' - \int_{\Gamma} [T_2(M', P)] (\vec{n}(M') \wedge \vec{E}(M')) dM' = \\ = \begin{cases} \vec{E}(P) & \text{si } y < f(x) \\ 0 & \text{si } y > f(x) \end{cases} \quad (6.8) \end{aligned}$$

On en déduit $\vec{H}(P)$ en prenant le rotationnel par rapport à P de (6.8) ; l'égalité obtenue se met sous la forme :

$$\begin{aligned} - \int_{\Gamma} [T_3(M', P)] (\vec{n}(M') \wedge \vec{H}(M')) dM' - \int_{\Gamma} [T_4(M', P)] (\vec{n}(M') \wedge \vec{E}(M')) dM' = \\ = \begin{cases} \vec{H}(P) & \text{si } y < f(x) \\ 0 & \text{si } y > f(x) \end{cases} \quad (6.9) \end{aligned}$$

Dans les relations (6.8) et (6.9), les $[T_\ell(M', P)]$ désignent des matrices 3×3 dont les éléments s'obtiennent à partir des $G_{ij}(M', P) = g_{ij}(x' - x, y' - y)$ et de leurs dérivées partielles par rapport à x, x', y et y' . On trouvera en annexe 12 l'expression des $[T_\ell(M', P)]$ dans le cas où $[\varepsilon]$ est diagonale.

En résumé, les équations (6.8) et (6.9) permettent d'obtenir en tout point P de Ω^- la valeur d'un champ pseudo-périodique $\vec{E}(P), \vec{H}(P)$ se propageant dans Ω^- supposé rempli du milieu de permittivité $[\varepsilon_2]$ et vérifiant une COS quand $y \rightarrow -\infty$, lorsque l'on connaît ses composantes tangentielles sur Γ . Ces équations s'appliquent donc en particulier au champ total dans Ω^- .

Il est également nécessaire pour la suite d'obtenir des expressions analogues (que nous appellerons (6.8') et (6.9')) pour un champ se propageant dans Ω^+ supposé rempli du milieu de permittivité ε_1 [I] et vérifiant une COS quand $y \rightarrow +\infty$. Ces équations (6.8') et (6.9') s'obtiennent à partir de (6.8) et (6.9) en y remplaçant tous les signes $-$ par des signes $+$, en permutant le sens des inégalités figurant aux deuxièmes membres, et en changeant $\varepsilon_{xx}, \varepsilon_{yy}$ et ε_{zz} en ε_1 dans les expressions des $[T_\ell(M', P)]$. Ces équations s'appliquent au champ obtenu en retranchant le champ incident du champ total dans Ω^+ (on obtient ainsi le champ "réfléchi" qui vérifie une COS quand $y \rightarrow +\infty$).

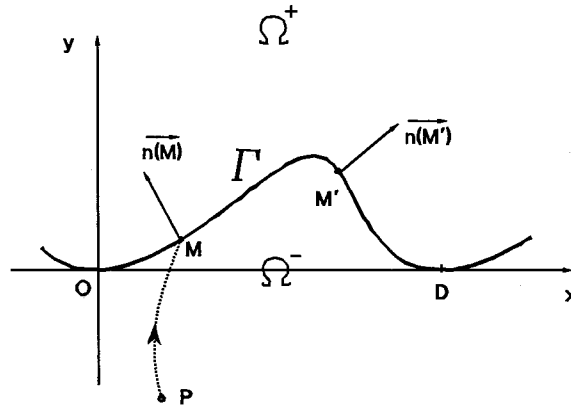
6.4. OBTENTION DES EQUATIONS INTEGRALES

Nous allons essayer dans les lignes qui suivent de montrer comment l'on peut déduire des équations intégrales des formules de Kirchhoff-Helmholtz généralisées précédemment établies. Pour réaliser cette étape, on est conduit à écrire des expressions très "encombrantes" que nous ne pouvons pas reproduire ici. Nous allons donc utiliser une formulation "imagée", nécessairement vague et imprécise, qui prétend néanmoins se rapprocher le plus possible de notre démarche. En un point Q de l'espace, nous désignons par $\vec{EH}(Q)$ un être regroupant les champs $\vec{E}(Q)$ et $\vec{H}(Q)$. Les équations (6.8) et (6.9) se mettent sous forme symbolique :

$$\int_{\Gamma} [T(M', P)] (\vec{n}(M') \wedge \vec{EH}(M')) dM' = \begin{cases} \vec{EH}(P) & \text{si } P \in \Omega^- \\ 0 & \text{si } P \in \Omega^+ \end{cases}$$

qui donne par multiplication vectorielle avec $\vec{n}(M)$ (M est un point de Γ ; cf figure 6.3) :

Figure 6.3



$$\vec{n}(M) \wedge \int_{\Gamma} [T(M', P)] (\vec{n}(M') \wedge \vec{EH}(M')) dM' = \begin{cases} \vec{n}(M) \wedge \vec{EH}(P) & \text{si } P \in \Omega^- \\ 0 & \text{si } P \in \Omega^+ \end{cases}$$

Le premier membre de cette équation est un opérateur intégral que l'on peut aussi écrire sous la forme :

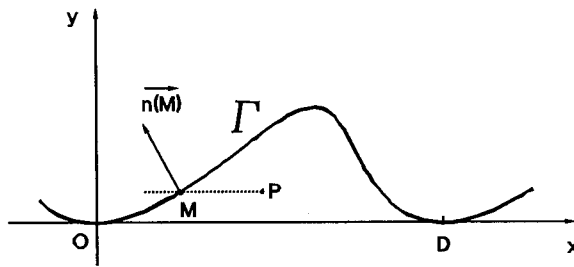
$$\int_{\Gamma} [T(M, M', P)] (\vec{n}(M') \wedge \vec{EH}(M')) dM' = \begin{cases} \vec{n}(M) \wedge \vec{EH}(P) & \text{si } P \in \Omega^- \\ 0 & \text{si } P \in \Omega^+ \end{cases} \quad (6.10)$$

Faisons maintenant "tendre le point P vers le point M". Ce passage à la limite soulève en réalité des problèmes théoriques difficiles. On pourra à ce sujet consulter l'ouvrage de Colton et Kress [16]. Nous donnerons un sens à ce passage à la limite, d'une manière plus prosaïque, à l'aide des

remarques suivantes (nous rejoignons sur ce point les études de Maystre [37], Pavageau et Bousquet [52]) :

a) Soient x_0 et $y_0 = f(x_0)$ les coordonnées de M (Figure 6.4) et x et y celles de P. On peut envisager qu'au cours du passage à la limite $P \rightarrow M$, la trajectoire de P au voisinage de M soit l'horizontale $y = y_0$. Pour $y = y_0$, le second membre de (6.10) peut s'écrire comme une fonction pseudo-périodique en x et admet un développement en série de Fourier généralisé par rapport à la variable x . Il s'agit de la série de Fourier d'une fonction discontinue (elle vaut 0 si $x < x_0$ et $\vec{n}(M) \wedge \vec{EH}(P)$ si $x > x_0$). On sait [3] qu'une telle série prend pour $x = x_0$ la demi somme de sa valeur à droite et de sa valeur à gauche, soit : $\frac{1}{2} \vec{n}(M) \wedge \vec{EH}(M)$.

Figure 6.4



b) On admettra que, lors du passage à la limite, le premier membre de (6.10) tend vers l'intégrale obtenue en y remplaçant P par M, pourvu que l'on puisse donner un sens à cette intégrale. Nous verrons par la suite que les noyaux des équations intégrales que nous retiendrons sont continus, ou bien présentent une faible singularité (logarithmique) intégrable.

Après ce passage à la limite, l'équation (6.10) devient donc :

$$\int_{\Gamma} [T(M, M', M)] (\vec{n}(M') \wedge \vec{EH}(M')) dM' = \frac{1}{2} \vec{n}(M) \wedge \vec{EH}(M) \quad (6.11)$$

Il est maintenant temps de mettre un peu "d'ordre" dans ce qui précède et d'écrire proprement l'équation (6.11). Soit $\vec{E}, \vec{H}$ un champ électromagnétique. Nous nous intéressons à ses composantes tangentielles sur $\mathcal{P}$ (d'équation $y = f(x)$) décrites par les quatre fonctions :

$$F_1(x) = \vec{e}_1 \cdot \{\vec{n}(x) \wedge \vec{E}(x, f(x))\} = \vec{e}_3 \cdot \vec{E}(x, f(x)) / \sqrt{1 + f'(x)^2} \quad (6.12)$$

$$F_2(x) = \vec{e}_3 \cdot \{\vec{n}(x) \wedge \vec{E}(x, f(x))\} = - \vec{t}(x) \cdot \vec{E}(x, f(x)) / \sqrt{1 + f'(x)^2} \quad (6.13)$$

$$F_3(x) = \vec{e}_1 \cdot \{\vec{n}(x) \wedge \vec{H}(x, f(x))\} = \vec{e}_3 \cdot \vec{H}(x, f(x)) / \sqrt{1 + f'(x)^2} \quad (6.14)$$

$$F_4(x) = \vec{e}_3 \cdot \{\vec{n}(x) \wedge \vec{H}(x, f(x))\} = - \vec{t}(x) \cdot \vec{H}(x, f(x)) / \sqrt{1 + f'(x)^2} \quad (6.15)$$

On vérifie immédiatement que les projections sur $\vec{e}_2$ de $\vec{n}(x) \wedge \vec{E}(x, f(x))$ et $\vec{n}(x) \wedge \vec{H}(x, f(x))$ ne sont autres que $f'(x) F_1(x)$ et $f'(x) F_3(x)$. Les quatre fonctions ci-dessus caractérisent donc parfaitement les composantes tangentielles de $\vec{E}$ et $\vec{H}$ sur $\mathcal{P}$. Nous les regroupons sous la forme d'une matrice colonne $F(x)$ à quatre éléments. Nous désignons par $F(x)$ la colonne associée par (6.12 - 6.15) au champ total et par $F^{inc}(x)$ celle associée au champ incident.

L'équation (6.11), déduite de (6.8) et (6.9), s'applique au champ total dans Ω^- . En utilisant la définition de F précédente, et en notant que les composantes tangentielles du champ total sont continues sur $\mathcal{P}$, nous écrirons (6.11) sous la forme :

$$\mathbb{S}_2 F = - F \quad (6.16)$$

où $\mathbb{S}_2$ est un opérateur intégral matriciel attaché au milieu de permittivité $[\epsilon_2]$. L'expression de $\mathbb{S}_2$ est explicitement donnée en annexe 12. De même, à partir de (6.8') et (6.9') qui s'appliquent au champ réfléchi dans Ω^+ , on déduit une équation qui se met sous la forme :

$$\mathbb{S}_1 (F - F^{inc}) = F - F^{inc} , \quad (6.17)$$

$\mathbb{S}_1$ étant un opérateur attaché au milieu de permittivité ϵ_1 , qui se déduit de $\mathbb{S}_2$ en y remplaçant partout ϵ_{xx} , ϵ_{yy} et ϵ_{zz} par ϵ_1 .

Le champ incident peut également être considéré comme un champ se propageant dans Ω^- supposé rempli du milieu de permittivité ϵ_1 , et vérifie une COS quand $y \rightarrow -\infty$. A ce titre, ses composantes tangentielles sur $\mathcal{P}$ vérifient l'équation déduite de (6.16) en remplaçant la permittivité $[\epsilon_2]$ par $\epsilon_1[I]$, soit :

$$\mathbb{S}_1 F^{inc} = - F^{inc} , \quad (6.18)$$

ce qui permet de réécrire (6.17) sous la forme :

$$\mathbb{S}_1 F = F - 2 F^{inc} . \quad (6.19)$$

Les équations (6.16) et (6.19) constituent un système de 8 équations

intégrales couplées pour les quatre fonctions inconnues F_i . On est naturellement amené à se demander si ces équations sont indépendantes. Pour fixer les idées, numérotions ces équations de la façon suivante :

$\mathbb{S}_2 F = - F$	$\longleftrightarrow$	équation 1 équation 2 équation 3 équation 4
$\mathbb{S}_1 F = F - 2 F^{inc}$	$\longleftrightarrow$	équation 5 équation 6 équation 7 équation 8

Envisageons par exemple le système $\mathbb{S}_2 F = - F$. Une étude attentive de ce qui précède montre que, puisque (6.9) se déduit de (6.8) en en prenant le rotationnel, les équations 3 et 4 sont conséquence des équations 1 et 2. Il en est de même pour le système $\mathbb{S}_1 (F - F^{inc}) = F - F^{inc}$ et les équations 7 et 8 sont conséquence des équations 5 et 6. On pourrait donc se contenter pour résoudre notre problème de retenir les équations 1, 2, 5 et 6. Malheureusement, les équations 2, 4, 6 et 8 contiennent des noyaux très singuliers (singularités du type parties finies, cf. l'expression de $\mathbb{S}_2$ en annexe 12). Nous montrerons plus loin (§ 6.9, propriété 11) en utilisant le formalisme des opérateurs de Calderon, que les équations 2 et 4 peuvent se déduire des équations 1 et 3, et que de même les équations 6 et 8 se déduisent des équations 5 et 7. Nous avons choisi de ne conserver pour la résolution du problème que les équations 1, 3, 5 et 7, ce qui évite de rencontrer les noyaux les plus singuliers et se révèle plus agréable pour le traitement numérique.

6.5. RESOLUTION DES EQUATIONS INTEGRALES

Il résulte du fait que nous avons choisi $[\epsilon_2]$ diagonale que les deux cas de polarisation E// et H// (champ électrique ou magnétique incident parallèle aux sillons du réseau) donnent naissance à un champ total qui conserve la même polarisation (§ 3.5). Par conséquent, dans le cas E//, F_2 et F_3 sont nuls, et dans le cas H//, F_1 et F_4 sont nuls. Les quatre équations 1, 3, 5 et 7 se découpent en deux systèmes indépendants : le cas E// se résoud au moyen des équations 1 et 5 et le cas H// se résoud à l'aide des équations 3 et 7. Par ailleurs, nous avons déjà noté au § 3.5 que le cas E// est identique au problème isotrope obtenu en remplaçant $[\epsilon_2]$ par la permittivité scalaire ϵ_{zz} et peut ainsi se résoudre à l'aide de programmes

conçus pour les milieux isotropes. Nous ne nous intéresserons donc dans la suite qu'au cas de polarisation H//. Les deux équations intégrales couplées à résoudre (équations 3 et 7) s'écrivent (cf. annexe 12) :

$$\begin{aligned} i = 1, 2 ; \quad \sum_{q=2,3} \int_0^d N_{3q}^i(x, x') F_q(x') \sqrt{1 + f'(x')^2} dx' = \\ = \lambda_i F_3(x) + \mu_i F_3^{inc}(x) \end{aligned} \quad (6.20)$$

où l'indice i est relatif au milieu (N_{3q}^1 fait intervenir la permittivité ε_1 et N_{3q}^2 la permittivité $[\varepsilon_2]$) et où $\lambda_1 = 1$, $\lambda_2 = -1$, $\mu_1 = -2$, $\mu_2 = 0$. En posant :

$$\rho(x) = \sqrt{1 + f'(x)^2} , \quad (6.21)$$

nous pouvons réécrire (6.20) sous la forme :

$$\begin{aligned} \sum_{q=2,3} \int_0^d \left\{ \frac{\rho(x) N_{3q}^i(x, x')}{\exp(i\alpha_0(x-x'))} \right\} \left\{ \frac{F_q(x') \rho(x')}{\exp(i\alpha_0 x')} \right\} dx' = \\ = \lambda_i \left\{ \frac{F_3(x) \rho(x)}{\exp(i\alpha_0 x)} \right\} + \mu_i \left\{ \frac{F_3^{inc}(x) \rho(x)}{\exp(i\alpha_0 x)} \right\} \end{aligned} \quad (6.22)$$

où nous faisons apparaître entre accolades des fonctions (ou distributions) périodiques (de période d) par rapport à x et x' . Nous pouvons les représenter par leurs séries de Fourier (une double série pour les distributions des deux variables x et x'). Les coefficients de Fourier de $\left\{ \exp(-i\alpha_0 x) F_3^{inc}(x) \rho(x) \right\}$ sont calculés par une transformation de Fourier discrète. Cela consiste à évaluer les intégrales sur $[0, d]$ qui donnent ces coefficients de Fourier par la "méthode des rectangles" ; cette méthode s'avère en effet particulièrement bien adaptée à l'intégration de fonctions périodiques continues (cf. D. Maystre, dans [58]). Nous avons déjà noté (annexe 12) que les noyaux N_{32}^i présentent une singularité. En retranchant à ces noyaux une fonction adéquate des variables x et x' présentant le même type de singularité (cette fonction est choisie de façon à pouvoir calculer analytiquement ses coefficients de Fourier, cf annexe 13), on obtient une fonction continue dont les coefficients de Fourier se calculent par transformée de Fourier discrète. Les noyaux N_{33}^i étant continus, leurs coefficients de Fourier s'obtiennent aussi par transformée de Fourier discrète. Les calculs relatifs à l'extraction de la singularité des $N_{32}^i(x, x')$ et quelques astuces destinées à accélérer la convergence des séries donnant

les noyaux sont détaillés dans l'annexe 13.

En remplaçant dans (6.22) les accolades par leurs séries de Fourier, l'intégration se fait de façon élémentaire, et on est conduit à un système linéaire infini dont les inconnues sont les coefficients de Fourier des $\{\exp(-i\alpha_0 x) F_q(x) \rho(x)\}$ pour $q = 2, 3$. Pour les besoins du calcul numérique, ce système est tronqué en ne retenant dans chacune des séries de Fourier (pour les champs et les noyaux) qu'un nombre fini de termes. La résolution de ce système permet d'obtenir les composantes tangentielles du champ total sur le profil $\mathcal{P}$ du réseau.

6.6. DETERMINATION DU CHAMP DIFFRACTE

Le champ diffracté (c'est-à-dire les champs réfléchis et transmis) se déduit des composantes tangentielles du champ total sur $\mathcal{P}$ grâce aux formules de Kirchhoff-Helmholtz. Les calculs se mènent sans difficulté mais conduisent à des expressions encombrantes. Nous nous contenterons ici d'en donner les principales étapes.

Attachons nous tout d'abord au champ transmis. Dans le cas de polarisation $H//$ qui nous préoccupe ici, il suffit de déterminer le champ magnétique qui ne possède qu'une composante (H_z). Le champ électrique s'en déduit par :

$$\text{rot } \vec{H} = \overrightarrow{\text{grad}} H_z \wedge \vec{e}_3 = -i \omega \epsilon_0 [\epsilon_2] \vec{E} . \quad (6.23)$$

De l'équation (6.9), nous pouvons déduire $H_z(P) = H_z(x, y)$ en tout point P situé dans Ω^- . Si nous supposons de plus que $y < \min(f(x))$, les valeurs absolues figurant dans les séries donnant $[T_3(M', P)]$ et $[T_4(M', P)]$ disparaissent (cf annexe 12). Nous obtenons H_z sous la forme :

$$y < \min(f(x)), \quad H_z(x, y) = \sum_{n \in \mathbb{Z}} H_{zn}^t \exp(i\alpha_n x - i\beta_{2,n} y) \quad (6.24)$$

où $\beta_{2,n} = \sqrt{\frac{\epsilon_{xx}}{\epsilon_{yy}} (k_0^2 \epsilon_{yy} - \alpha_n^2)}$ avec la détermination $\beta_{2,n} > 0$ ou $\text{Im}(\beta_{2,n}) > 0$,

et où les H_{zn}^t sont obtenus par intégration sur le profil $\mathcal{P}$ d'expressions périodiques faisant intervenir les composantes tangentielles du champ total :

$$H_{zn}^t = \frac{1}{2d} \int_0^d e^{-i\alpha_n x + i\beta_{2,n} f(x)} \left[\left(1 + \frac{\varepsilon_{xx}}{\varepsilon_{yy}} \frac{\alpha_n}{\beta_{2,n}} f'(x) \right) F_3(x) - \frac{\varepsilon_{xx} k_0}{\eta_0 \beta_{2,n}} F_2(x) \right] \times \rho(x) dx \quad (6.25)$$

Le champ électrique se déduit de (6.23). On peut alors calculer le vecteur de Poynting associé au champ transmis et en déduire, dans le cas où le substrat est sans pertes, les efficacités e_n^t transmises dans les différents ordres propagatifs ; on obtient :

$$e_n^t = |H_{zn}^t|^2 \frac{\varepsilon_1 \beta_{2,n}}{\varepsilon_{xx} \beta_{1,0}} \quad (6.26)$$

où $\beta_{1,0} = \sqrt{k_0^2 \varepsilon_1 - \alpha_0^2} = k_0 n_1 \cos \theta$. Toujours dans le cas d'un substrat sans pertes, on obtient pour l'onde plane associée au $n^{\text{ième}}$ ordre propagatif un vecteur d'onde réel égal à $(\alpha_n \vec{e}_1 - \beta_{2,n} \vec{e}_2)$, tandis que le vecteur de Poynting est colinéaire à $\left(\frac{\alpha_n}{\varepsilon_{yy}} \vec{e}_1 - \frac{\beta_{2,n}}{\varepsilon_{xx}} \vec{e}_2 \right)$. On notera que les ordres transmis se propagent dans des directions différentes de celles que l'on obtiendrait dans le cas de polarisation E// ; dans ce cas en effet (qui est rappelons-le identique à un problème isotrope) vecteur d'onde et vecteur de Poynting sont tous deux colinéaires à $(\alpha_n \vec{e}_1 - \sqrt{k_0^2 \varepsilon_{zz} - \alpha_n^2} \vec{e}_2)$.

Pour déterminer le champ réfléchi, on peut procéder de façon analogue. Mais il est préférable d'utiliser les propriétés établies pour des réseaux isotropes. Rappelons [59, p.306] que si u et u' sont des fonctions pseudo-périodiques de coefficients de pseudo-périodicité respectifs $\exp(i\alpha_0 d)$ et $\exp(-i\alpha_0 d)$, vérifiant l'équation de Helmholtz dans Ω^+ (avec $k_1 = k_0 \sqrt{\varepsilon_1}$) et pouvant s'écrire pour $y > \max(f(x))$ sous la forme :

$$u(x,y) = \exp(i\alpha_0 x - i\beta_{1,0} y) + \sum_{m \in \mathbb{Z}} B_m \exp(i\alpha_m x + i\beta_{1,m} y)$$

$$u'(x,y) = \exp(-i\alpha_n x - i\beta_{1,n} y)$$

où n est choisi de sorte que $\beta_{1,n} = \sqrt{k_0^2 \varepsilon_1 - \alpha_n^2}$ soit positif, alors on a (Du désigne la dérivée normale de u sur le profil) :

$$\int_0^d [u(x, f(x)) Du'(x, f(x)) - u'(x, f(x)) Du(x, f(x))] \rho(x) dx = -2id\beta_{1,n} B_n$$

En appliquant ce résultat à $u(x, y) = H_z(x, y)$ qui s'écrit pour $y > \max(f(x))$:

$$H_z(x, y) = \exp(i\alpha_0 x - i\beta_{1,0} y) + \sum_{m \in \mathbb{Z}} H_{zm}^r \exp(i\alpha_m x + i\beta_{1,m} y),$$

et en remarquant qu'alors $Du(x, f(x)) = i\omega \varepsilon_0 \varepsilon_1 F_2(x)$ et que $u(x, f(x)) = \rho(x) F_3(x)$, on obtient facilement :

$$H_{zn}^r = \frac{1}{2d} \int_0^d e^{-i\alpha_n x - i\beta_{1,n} f(x)} \left[\left(1 - \frac{\alpha_n}{\beta_{1,n}} f'(x) \right) F_3(x) + \frac{\varepsilon_1 k_0}{\eta_0 \beta_{1,n}} F_2(x) \right] \times \rho(x) dx \quad (6.27)$$

D'où l'on déduit l'efficacité réfléchie dans le $n^{\text{ième}}$ ordre propagatif :

$$e_n^r = \left| H_{zn}^r \right|^2 \frac{\beta_{1,n}}{\beta_{1,0}}. \quad (6.28)$$

6.7. RESULTATS NUMERIQUES. CRITERES DE VALIDITE. COMPARAISON AVEC LA METHODE DIFFERENTIELLE.

Les résultats donnés dans ce paragraphe se rapportent à des matériaux sans pertes. Ce choix résulte de deux principales raisons :

- le critère de conservation de l'énergie (qui n'est pas automatiquement vérifié pour les méthodes intégrales, contrairement aux méthodes différentielles) permet de contrôler la validité des résultats,
- la méthode différentielle permettant de traiter des réseaux de hauteur plus importante dans le cas de matériaux sans pertes, une confrontation plus significative entre les deux méthodes est possible.

Mais des essais numériques ont été également conduits pour des matériaux conducteurs. Cela nous a paru intéressant dans la mesure où la méthode intégrale que nous proposons diffère de celle (proposée par D. Maystre) couramment utilisée au Laboratoire pour les réseaux isotropes. Nous avons par exemple choisi pour inconnues les coefficients de Fourier généralisés du champ tangentiel sur le profil alors que D. Maystre choisit pour inconnues les valeurs du champ en différents points de discrétisation sur le profil. La confrontation entre les deux méthodes pour des réseaux isotropes

très conducteurs (substrat en aluminium par exemple) a montré que les deux façons de procéder sont sensiblement équivalentes.

Les résultats présentés ci-dessous concernent le cas d'un substrat ayant pour permittivités $\epsilon_{xx} = 6.31$, $\epsilon_{yy} = 6.81$, $\epsilon_{zz} = 7.34$. Ces permittivités correspondent d'après [27] à la lithargite, qui a la propriété de posséder une anisotropie importante. Mais il faut bien avouer que la réalisation de réseaux au moyen de ce matériau ne semble guère concevable. Il s'agit donc d'un exemple très académique. Le superstrat est vide ($\epsilon_1 = 1$). Le réseau est éclairé par une onde incidente (longueur d'onde $\lambda_0 = 0.6 \mu\text{m}$) de polarisation $H//$ sous l'incidence $\theta = 20^\circ$. Le pas du réseau est $d = 0.5 \mu\text{m}$. La profondeur des sillons, notée h , est la différence $\max(f(x)) - \min(f(x))$. On obtient dans ces conditions deux ordres réfléchis : les ordres -1 et 0 qui se propagent sous les angles -59.09° et 20° (ces angles sont obtenus par la classique "formule des réseaux") et quatre ordres transmis : ordres -2 , -1 , 0 et 1 . Les directions des vecteurs d'onde et des vecteurs de Poynting de ces ondes transmises sont obtenues comme indiqué au § 6.6. On trouve respectivement les angles -53.11° , -19.88° , 7.82° et 37.27° pour les vecteurs d'onde et -50.99° , -18.53° , 7.25° et 35.19° pour les vecteurs de Poynting.

Nous notons par la suite :

- NSOM le nombre de termes retenus pour la sommation des séries donnant les noyaux (la sommation est effectuée de $-NSOM$ à $+NSOM$),
- NF le nombre total de coefficients de Fourier retenus pour les développements des champs et des noyaux (accolades de l'équation (6.22)) en séries de Fourier,
- ND le nombre total de points de discrétisation (par période) utilisés pour les intégrations sur le profil. Ce nombre de points de discrétisation est utilisé d'une part pour obtenir les transformées de Fourier discrètes des accolades de (6.22) liées aux noyaux et au champ incident, et d'autre part pour le calcul des H_{zn}^t et H_{zn}^r (équations (6.25) et (6.27)).

L'expérience numérique montre (tableau 6.1) une rapide convergence de la solution lorsque NSOM, NF et ND croissent. Il est intéressant de comparer les résultats à ceux obtenus par la méthode différentielle (cf. Chapitre 5). Nous les avons fait figurer dans ce tableau où N désigne le nombre de coefficients de Fourier retenus pour les développements des champs et J le nombre de pas d'intégration utilisés (notations du § 5.4.2). On observe que la convergence de la méthode différentielle est moins rapide

que celle de la méthode intégrale. Nous avons également reporté dans ce tableau les temps de calcul T_1 (sur le HP 9000-835 du Laboratoire) et T_2 (sur IBM 3090 "vectorisé", ce qui explique que T_1 et T_2 ne sont pas proportionnels).

	ordre	méthode différentielle			méthode intégrale	
		N=15,J=90	N=21,J=150	N=41,J=150	NF=7,ND=15 NSOM=20	NF=11,ND=21 NSOM=30
efficacités réfléchies	-1	0.0695	0.0695	0.0695	0.0695	0.0695
	0	0.0730	0.0731	0.0733	0.0736	0.0735
efficacités transmises	-2	0.0015	0.0016	0.0018	0.0019	0.0019
	-1	0.0573	0.0573	0.0572	0.0571	0.0571
	0	0.6357	0.6337	0.6314	0.6288	0.6291
	1	0.1630	0.1648	0.1668	0.1691	0.1689
somme efficacités		1.0000	1.0000	1.0000	1.0000	1.0000
T1		9 min	aie!	houla!	9 s	23 s
T2		3 s	12 s	61 s	1 s	3 s

Tableau 6.1. Réseau sinusoïdal, $h = 0.1 \mu\text{m}$.

	ordre	méthode différentielle		méthode intégrale		
		N=15,J=90	N=21,J=200	NF=7,ND=15 NSOM=10	NF=11,ND=25 NSOM=20	NF=25,ND=45 NSOM=30
efficacités réfléchies	-1	0.0914	0.0913	0.0914	0.0923	0.0923
	0	0.0002	0.0002	0.0003	0.0003	0.0003
efficacités transmises	-2	0.0055	0.0061	0.0082	0.0078	0.0078
	-1	0.2512	0.2521	0.2608	0.2530	0.2530
	0	0.2436	0.2369	0.2185	0.2230	0.2231
	1	0.4081	0.4130	0.4275	0.4238	0.4235
somme efficacités		1.0000	0.9996	1.0067	1.0002	1.0000
T1		9 min	aie!	7 s	29 s	5 min
T2		3 s	16 s	1 s	3 s	11 s

Tableau 6.2. Réseau sinusoïdal, $h = 0.2 \mu\text{m}$.

Pour une hauteur $h = 0.2 \mu\text{m}$ (tableau 6.2), la méthode différentielle commence à "peiner". On voit en effet que pour $N = 21$ et $J = 200$, le bilan énergétique n'est pas strictement égal à l'unité, ce qui montre que des problèmes numériques commencent à poindre (rappelons que la conservation de l'énergie est théoriquement automatiquement vérifiée lors de l'emploi des méthodes différentielles, cf. annexe 1). On constate toutefois que la méthode différentielle donne une précision qui pourra être jugée suffisante pour d'éventuelles confrontations avec l'expérience.

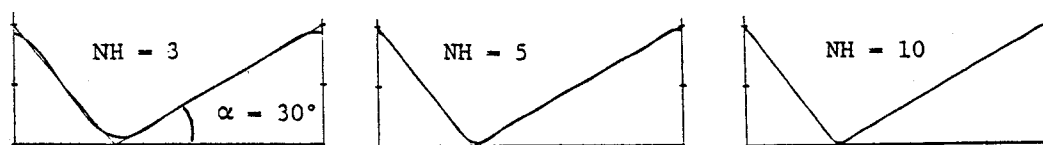
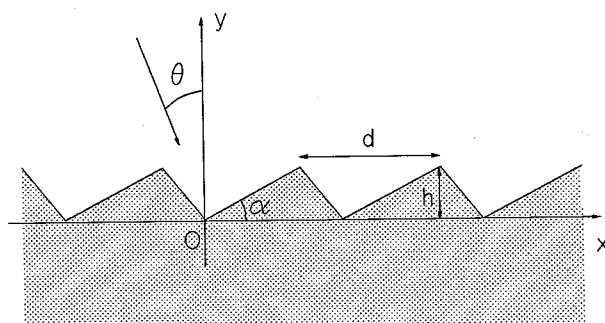
Les tableaux 6.1 et 6.2 montrent que le bilan énergétique permet (dans le cas de milieux sans pertes) d'estimer l'ordre de grandeur de la précision des résultats fournis par la méthode intégrale. Au niveau des temps de calcul, la méthode intégrale est bien entendu beaucoup plus performante.

Par contre, la méthode différentielle, rappelons le, permet de traiter des réseaux de nature plus complexe (réseau recouvert par une couche, réseau à modulation d'indice, sans restriction sur les permittivités, cf Chapitre 5).

6.8. PROFILS PRESENTANT DES ARETES

Nous avons supposé au début de l'étude que le profil est de classe C^2 (l'étude numérique des noyaux $N_{33}^i(x, x')$ nécessite la connaissance de $f''(x)$). Il est tout de même possible d'étudier des profils qui ne vérifient pas cette hypothèse en remplaçant la fonction $f(x)$ par la fonction obtenue en sommant les NH premiers termes de sa série de Fourier (NH représente ici le nombre d' "harmoniques" : $NH = 3$ signifie que l'on retient 7 termes dans la série de Fourier). Pour traiter par exemple le cas du réseau de la figure 6.5, nous avons retenu successivement $NH = 3, 5$ puis 10 harmoniques. Les efficacités obtenues sont reportées sur la figure 6.6. Nous avons aussi représenté sur cette figure la fonction reconstruite à partir de sa série de Fourier tronquée. Les résultats obtenus convergent rapidement lorsque NH croît.

Figure 6.5. Profil échelonne ;
 $h = 0.2 \mu m, d = 0.5 \mu m, \lambda_0 = 0.6 \mu m,$
 $\alpha = 30^\circ, \theta = 20^\circ$; permittivités du
 substrat : $\epsilon_{xx} = 6.31, \epsilon_{yy} = 6.81,$
 $\epsilon_{zz} = 7.34$. On étudie le cas de
 polarisation $H//$.



ordre	efficacités réfléchies	efficacités transmises	efficacités réfléchies	efficacités transmises	efficacités réfléchies	efficacités transmises
-2		0.0119		0.0126		0.0128
-1	0.1226	0.1339	0.1167	0.1363	0.1166	0.1365
0	0.0126	0.4782	0.0142	0.4766	0.0143	0.4758
1		0.2408		0.2436		0.2439

Figure 6.6. Calculs effectués avec $NF = 25, ND = 45, NSOM = 30$.

6.9. LES OPERATEURS DE CALDERON. DEFINITIONS ET PROPRIETES

Ce paragraphe résultant d'une longue réflexion théorique n'a pas donné lieu jusqu'ici à des applications numériques. Il devrait néanmoins être très utile dans le cas où nous déciderions de poursuivre plus avant l'étude entreprise.

Dans toute la suite, toutes les fonctions considérées sont supposées être pseudo-périodiques en x , de période d (le pas du réseau), et de coefficient de pseudo-périodicité fixé.

Considérons le problème suivant. Se donnant quatre fonctions définies sur $\mathcal{P}$ (le profil du réseau d'équation $y = f(x)$) :

$$F(x) = \begin{pmatrix} F_1(x) \\ F_2(x) \\ F_3(x) \\ F_4(x) \end{pmatrix}, \quad (6.29)$$

nous définissons les "courants" :

$$\vec{j}_e \stackrel{\text{def}}{=} \begin{pmatrix} F_1 \\ f'(x) F_1 \\ F_2 \end{pmatrix}, \quad \vec{j}_m \stackrel{\text{def}}{=} \begin{pmatrix} F_3 \\ f'(x) F_3 \\ F_4 \end{pmatrix}, \quad (6.30)$$

et recherchons les champs $\vec{E}(x,y)$ et $\vec{H}(x,y)$ vérifiant au sens des distributions :

$$\text{rot } \vec{E} - i \omega \mu_0 \vec{H} = \vec{j}_e \delta_{\mathcal{P}}, \quad (6.31)$$

$$\text{rot } \vec{H} + i \omega \epsilon_0 [\epsilon] \vec{E} = \vec{j}_m \delta_{\mathcal{P}}, \quad (6.32)$$

ainsi qu'une condition d'onde sortante quand $y \rightarrow \pm\infty$. Dans (6.31) et (6.32), la distribution $\delta_{\mathcal{P}}$ représente la distribution de Dirac sur le profil $\mathcal{P}$ (cf annexe 5, § 1, ou bien [58], appendice A). En explicitant les rotationnels (avec les notations usuelles de la théorie des distributions [68], $\vec{n}$ est la normale au profil dirigée vers Ω^+ , $\sigma(\vec{n} \wedge \vec{E})$ représente le saut de $\vec{n} \wedge \vec{E}$ lorsqu'on traverse $\mathcal{P}$ dans le sens de $\vec{n}$) :

$$\text{rot } \vec{E} = (\text{rot } \vec{E}) + \sigma (\vec{n} \wedge \vec{E}) \delta_{\mathcal{P}},$$

il apparaît que le problème considéré peut aussi s'énoncer de la façon suivante : trouver les champs $\vec{E}$ et $\vec{H}$ solutions des équations de propagation dans un milieu anisotrope de permittivité $[\varepsilon]$ remplissant les domaines Ω^+ et Ω^- , vérifiant une C.O.S. pour $y \rightarrow \pm \infty$, en supposant connus les sauts de leurs composantes tangentielles sur $\mathcal{P}$ (ces sauts étant caractérisés par $\vec{j}_e = \sigma(\vec{n} \wedge \vec{E})$ et $\vec{j}_m = \sigma(\vec{n} \wedge \vec{H})$, c'est-à-dire par la colonne F).

Un autre énoncé du même problème est constitué par les conditions (6.33) à (6.35) :

$$\text{rot rot } \vec{E} - k_0^2 [\varepsilon] \vec{E} = \text{rot } (\vec{j}_e \delta \varphi) + i \omega \mu_0 \vec{j}_m \delta \varphi \quad (6.33)$$

(directement déduite de (6.31) et (6.32),

$$\vec{H} = \frac{1}{i \omega \mu_0} \text{rot } \vec{E} \quad \text{dans } \Omega^+ \text{ et } \Omega^-, \quad (6.34)$$

$$\vec{E} \text{ et } \vec{H} \text{ vérifient une C.O.S. pour } y \rightarrow \pm \infty. \quad (6.35)$$

Nous avons vu (Annexe 8) que la solution de ce problème peut s'exprimer lorsque l'on connaît les vecteurs de Green $\vec{g}_i$ vérifiant une C.O.S et solutions de :

$$\text{rot rot } \vec{g}_i - k_0^2 [\varepsilon] \vec{g}_i = \vec{e}_i \delta_{\mathcal{R}}. \quad (6.36)$$

La solution est (annexe 8) :

$$\vec{E} = {}^t[g] * \left(\text{rot } (\vec{j}_e \delta \varphi) + i \omega \mu_0 \vec{j}_m \delta \varphi \right), \quad (6.37)$$

de laquelle $\vec{H}$ se déduit par (6.34).

On connaît alors $\vec{E}$ et $\vec{H}$ dans Ω^+ et Ω^- . Définissons F^+ et F^- à partir des limites des composantes tangentielles de $\vec{E}$ et $\vec{H}$ sur $\mathcal{P}$ (+ (resp. -) désigne une limite en venant de Ω^+ (resp. Ω^-)) : la colonne F^+ (resp. F^-) est déduite de $\vec{n} \wedge \vec{E}^+$ et $\vec{n} \wedge \vec{H}^+$ (resp. $-\vec{n} \wedge \vec{E}^-$ et $-\vec{n} \wedge \vec{H}^-$) conformément à (6.12)-(6.15).

Par le procédé suggéré ci-dessus, on associe à toute colonne F deux colonnes F^+ et F^- . Les opérateurs de Caldéron (associés au milieu anisotrope, au profil $\mathcal{P}$ et à la pseudo-périodicité considérés) sont les opérateurs P^+ et P^- définis par :

$$F^+ = P^+ F, \quad (6.38)$$

$$F^- = P^- F. \quad (6.39)$$

Remarque : Les opérateurs P^+ et P^- sont des opérateurs intégraux dont l'expression peut être obtenue en explicitant la convolution de ${}^t[g]$ et de $\text{rot}[\sigma(\vec{n} \wedge \vec{E}) \delta_{\mathcal{P}}] + i\omega \mu_0 \sigma(\vec{n} \wedge \vec{H}) \delta_{\mathcal{P}}$, le calcul de $\vec{H} = \frac{1}{i\omega \mu_0} \text{rot } \vec{E}$, puis en recherchant les limites des champs obtenus lorsque le point d'observation tend vers un point du profil par valeur supérieure (obtention de P^+) ou inférieure (obtention de P^-). Ces opérations sont très semblables à celles décrites aux § 6.3 et 6.4. Nous expliquerons un peu plus loin comment l'on peut obtenir la forme explicite de P^+ et P^- en utilisant les résultats du § 6.4. Il est important de noter que ces opérateurs peuvent être obtenus dès lors que l'on connaît les vecteurs de Green $\vec{g}_i$ associés au milieu anisotrope considéré et à la structure étudiée (profil $\mathcal{P}$, pseudo-périodicité des champs).

Pour nous résumer, considérons les énoncés suivants :

Problème 1 : Résoudre les équations de Maxwell dans Ω^+ et Ω^- remplis du même milieu anisotrope, en imposant sur $\mathcal{P}$ aux champs $\vec{E}$ et $\vec{H}$ le saut de leurs composantes tangentielles, et une C.O.S. quand $y \rightarrow \pm\infty$.

$$\text{Formulation 1 : } \left\{ \begin{array}{l} \vec{E} = {}^t[g] * \left[\text{rot}(\sigma(\vec{n} \wedge \vec{E}) \delta_{\mathcal{P}}) + i\omega \mu_0 \sigma(\vec{n} \wedge \vec{H}) \delta_{\mathcal{P}} \right], \\ \vec{H} = \frac{1}{i\omega \mu_0} \text{rot } \vec{E} \quad \text{dans } \Omega^+ \text{ et } \Omega^-. \end{array} \right.$$

Associations aux sauts des composantes tangentielles de $\vec{E}$ et $\vec{H}$ la colonne F définie sur $\mathcal{P}$ par :

$$F = \begin{cases} F_1 = \vec{e}_1 \cdot \sigma(\vec{n} \wedge \vec{E}) \\ F_2 = \vec{e}_3 \cdot \sigma(\vec{n} \wedge \vec{E}) \\ F_3 = \vec{e}_1 \cdot \sigma(\vec{n} \wedge \vec{H}) \\ F_4 = \vec{e}_3 \cdot \sigma(\vec{n} \wedge \vec{H}) \end{cases} \quad (6.40)$$

et aux valeurs au bord des composantes tangentielles sur $\mathcal{P}$ de $\vec{E}$ et $\vec{H}$ les colonnes :

$$F^+ = \begin{cases} F_1^+ = \vec{e}_1 \cdot (\vec{n} \wedge \vec{E}^+) \\ F_2^+ = \vec{e}_3 \cdot (\vec{n} \wedge \vec{E}^+) \\ F_3^+ = \vec{e}_1 \cdot (\vec{n} \wedge \vec{H}^+) \\ F_4^+ = \vec{e}_3 \cdot (\vec{n} \wedge \vec{H}^+) \end{cases} \quad (6.41)$$

$$F^- = \begin{cases} F_1^- = -\vec{e}_1 \cdot (\vec{n} \wedge \vec{E}^-) \\ F_2^- = -\vec{e}_3 \cdot (\vec{n} \wedge \vec{E}^-) \\ F_3^- = -\vec{e}_1 \cdot (\vec{n} \wedge \vec{H}^-) \\ F_4^- = -\vec{e}_3 \cdot (\vec{n} \wedge \vec{H}^-) \end{cases} \quad (6.42)$$

Nous pouvons considérer la :

$$\text{Formulation 2 : } \begin{cases} F^+ = P^+ F \\ F^- = P^- F \end{cases}$$

et énoncer la :

Propriété 1 : Les formulations 1 et 2 sont deux formulations équivalentes permettant de résoudre le problème 1.

Preuve : Nous avons déjà établi un peu plus haut que la formulation 1 représente une façon possible de résoudre le problème 1. Quant à la formulation 2 : les sauts des composantes tangentielles des champs sur $\mathcal{P}$ étant donnés, F est connu, on en déduit F^+ et F^- par application de P^+ et P^- , puis le champ partout au moyen des formules de Kirchhoff-Helmholtz (§ 6.3).

Propriété 2 : $P^+ + P^- = I$ (= opérateur unité).

En effet, d'après (6.40) à (6.42), on a quel que soit F :

$$F = F^+ + F^-. \quad (6.43)$$

D'où d'après (6.38, 6.39) :

$$F = (P^+ + P^-) F.$$

Propriété 3 : Etant donné une colonne F , considérons les colonnes F^+ et F^- obtenues par action de P^+ et P^- (eq.(6.38 - 6.39)), on a :

$$P^+ F^+ = F^+, \quad P^- F^- = F^-,$$

$$P^- F^+ = 0, \quad P^+ F^- = 0.$$

Preuve : F étant donnée, considérons le champ $(\vec{E}, \vec{H})$ solution du problème 1 lorsque le saut de ses composantes tangentielles sur $\mathcal{P}$ est lié à F par (6.40). Ce champ possède sur $\mathcal{P}$ des valeurs au bord F^+ et F^- (6.41, 6.42). Considérons le champ $(\vec{\tilde{E}}, \vec{\tilde{H}})$ solution du problème 1 lorsque le saut de ses composantes tangentielles sur $\mathcal{P}$ est lié à $\tilde{F} = F^+$. La solution de ce problème est évidente et vaut $(\vec{\tilde{E}}, \vec{\tilde{H}}) = (\vec{E}, \vec{H})$ dans Ω^+ et $\vec{\tilde{E}} = \vec{\tilde{H}} = 0$ dans Ω^- , de sorte que :

$$\tilde{F}^+ = F^+ \quad \text{et} \quad \tilde{F}^- = 0. \quad (6.44)$$

Nous savons par ailleurs que (formulation 2) :

$$\tilde{F}^+ = P^+ \tilde{F} \quad \text{et} \quad \tilde{F}^- = P^- \tilde{F}. \quad (6.45)$$

Puisque $\tilde{F} = F^+$ et en utilisant (6.44), (6.45) s'écrit :

$$F^+ = P^+ F^+ \quad \text{et} \quad 0 = P^- F^+. \quad (6.46)$$

On peut montrer que $P^- F^- = F^-$ et que $P^+ F^- = 0$ en suivant une démarche analogue, ou plus simplement en utilisant (6.46), (6.43), (6.38) et (6.39).

Propriété 4 : $P^- P^+ = P^+ P^- = 0$,

$$P^- P^- = P^-,$$

$$P^+ P^+ = P^+.$$

La démonstration est immédiate en appliquant les opérateurs à une colonne F quelconque et en utilisant les propriétés précédentes. Par exemple, en utilisant successivement (6.39), la propriété 3, et à nouveau (6.39) :

$$P^- P^- F = P^- F^- = F^- = P^- F.$$

Cette propriété montre que P^+ et P^- sont des projecteurs (opérateurs idempotents).

Remarque et définitions. Nous n'avons pas précisé à quels espaces doivent appartenir les fonctions définies sur $\mathcal{P}$ telles que $\vec{j}_e$, $\vec{j}_m$ ou F ni les champs $\vec{E}$ et $\vec{H}$ pour que les problèmes considérés plus haut admettent une solution unique. Il s'agit d'un problème mathématique délicat. On pourra se reporter à la référence [13] pour avoir quelques précisions à ce sujet (il y est traité un problème analogue dans le cas isotrope). On y voit apparaître des espaces de Sobolev tels que $H^{-1/2}(\text{div}, \mathcal{P})$ et $H^{-1/2}(\text{rot}, \mathcal{P})$. Nous nous contenterons ici de noter $\mathcal{V}$ l'espace des colonnes F définies sur $\mathcal{P}$, de définir l'espace $\mathcal{V}^+$ auquel appartiennent les colonnes F^+ comme l'image de $\mathcal{V}$ par P^+ , et l'espace $\mathcal{V}^-$ auquel appartiennent les colonnes F^- comme l'image de $\mathcal{V}$ par P^- .

Propriété 5 : $\mathcal{V}$ est la somme directe de $\mathcal{V}^+$ et $\mathcal{V}^-$.

En effet, nous avons vu que toute colonne $F \in \mathcal{V}$ peut s'écrire de façon unique sous la forme $F = (P^+ + P^-)F = F^+ + F^-$ (utilisation de la propriété 2, (6.38 - 6.39)).

Propriété 6 : $\mathcal{V}^+$ (resp. $\mathcal{V}^-$) est l'ensemble des colonnes représentant (conformément à (6.41) - resp. (6.42) -) les valeurs au bord sur $\mathcal{P}$ des composantes tangentielles de champs vérifiant les équations de Maxwell dans le domaine Ω^+ (resp. Ω^-) supposé rempli du milieu de permittivité $[\epsilon]$ et vérifiant une C.O.S. dans ce domaine.

Preuve : Considérons par exemple le cas de $\mathcal{V}^+$. Pour alléger la présentation, nous dirons ici qu'une colonne est une "valeur au bord" plutôt que de dire qu'elle représente, conformément à (6.41), les valeurs au bord des composantes tangentielles de champs vérifiant les équations de Maxwell dans le domaine Ω^+ supposé rempli du milieu de permittivité $[\epsilon]$ et vérifiant une C.O.S. dans Ω^+ .

Il résulte de la définition de $\mathcal{V}^+$ et de celle de P^+ que tout élément de $\mathcal{V}^+$ est une "valeur au bord". Montrons que réciproquement, toute "valeur au bord" appartient à $\mathcal{V}^+$. Soit $(\vec{E}, \vec{H})$ un champ solution des équations de Maxwell dans Ω^+ et vérifiant la C.O.S. dans ce domaine, et soit $\tilde{F}^+$ sa valeur au bord. Soit $(\vec{E}, \vec{H})$ le champ égal à $(\vec{E}, \vec{H})$ dans Ω^+ et nul dans Ω^- ; il est solution du problème 1 lorsque l'on impose un saut de composantes

tangentielles $F = F^+ + F^-$, avec $F^- = 0$ et $F^+ = \tilde{F}^+$. Par conséquent $\tilde{F}^+ = F^+ = P^+ F$ appartient à $\mathcal{V}^+$.

Propriété 7 : La relation $P^+ F = F$ (équivalente d'après la propriété 2 à $P^- F = 0$) caractérise l'appartenance de F à $\mathcal{V}^+$, autrement dit caractérise le fait que F représente une valeur au bord sur $\mathcal{P}$ des composantes tangentielles d'un champ vérifiant la C.O.S. dans Ω^+ .

De même, la relation $P^- F = F$ (ou $P^+ F = 0$) caractérise l'appartenance de F à $\mathcal{V}^-$.

Preuve :

- Soit F une colonne appartenant à $\mathcal{V}^+$; d'après la définition de $\mathcal{V}^+$, il existe une colonne $G \in \mathcal{V}$ telle que $F = P^+ G$, d'où $P^+ F = P^+ P^+ G = P^+ G$ (propriété 4). Donc $P^+ F = F$.

- Réciproquement, toute colonne vérifiant $F = P^+ F$ appartient à $\mathcal{V}^+$ (définition de $\mathcal{V}^+$).

Définition : Il est commode d'introduire un opérateur $\mathbb{S}$, lié au milieu anisotrope, au profil $\mathcal{P}$ et à la pseudo-périodicité considérés, et défini de la façon suivante :

$$\mathbb{S} = P^+ - P^- . \quad (6.47)$$

De la propriété 2, on déduit immédiatement :

$$P^+ = \frac{1}{2} (\mathbb{I} + \mathbb{S}) ; \quad P^- = \frac{1}{2} (\mathbb{I} - \mathbb{S}) . \quad (6.48)$$

On notera que $\mathbb{S}$ est un opérateur involutif :

$$\mathbb{S} \mathbb{S} = \mathbb{I} . \quad (6.49)$$

En effet, pour tout $F \in \mathcal{V}$,

$$\begin{aligned} \mathbb{S} \mathbb{S} F &= \mathbb{S} (P^+ - P^-) F = \mathbb{S} (F^+ - F^-) \\ &= (P^+ - P^-) (F^+ - F^-) = F^+ + F^- = F. \end{aligned}$$

Propriété 8 : F étant une colonne quelconque, F^+ et F^- étant les colonnes qui lui sont associées (équations (6.38-6.39)), on a :

$$\mathbb{S} F^+ = F^+ , \quad (6.50)$$

$$\mathbb{S} F^- = - F^- . \quad (6.51)$$

La démonstration est immédiate ; par exemple :

$$\mathbb{S} F^- = (P^+ - P^-) F^- = 0 - F^- .$$

Propriété 9 : En utilisant les définitions précédentes, nous pouvons réécrire la propriété 7 sous la forme :

$$F \in \mathcal{V}^+ \Leftrightarrow P^+ F = F \Leftrightarrow P^- F = 0 \Leftrightarrow \mathbb{S} F = F ;$$

$$F \in \mathcal{V}^- \Leftrightarrow P^- F = F \Leftrightarrow P^+ F = 0 \Leftrightarrow \mathbb{S} F = - F .$$

Remarque : En se reportant aux paragraphes 6.3 et 6.4, on remarquera que l'opérateur $\mathbb{S}$ défini ici n'est autre que celui noté $\mathbb{S}_2$ dans l'équation (6.16) (le domaine Ω^- était alors rempli d'un milieu de permittivité $[\varepsilon_2]$, ce qui justifie la présence de l'indice 2) ou bien celui noté $\mathbb{S}_1$ dans l'équation (6.17) (le domaine Ω^+ était rempli d'un milieu de permittivité $\varepsilon_1[I]$). On notera que, dans les équations (6.16) et (6.17), il n'était pas fait de distinction entre les valeurs au bord "en venant de Ω^+ ou de Ω^- " car les composantes tangentielles des champs étaient continues dans le problème physique considéré. L'opérateur $\mathbb{S}$ est donc bien celui donné explicitement en annexe 12, de sorte que les opérateurs P^+ et P^- sont également connus de façon explicite.

Propriété 10 : Pour une colonne F appartenant à l'un quelconque des espaces $\mathcal{V}^+$ ou $\mathcal{V}^-$, la donnée de deux quelconques des composantes de F entraîne la détermination des deux autres.

Cette propriété provient du fait que $\mathcal{V}$, considéré comme l'ensemble de colonnes formées de quatre fonctions définies sur le profil, est un espace de dimension quatre. Puisque $\mathcal{V}^+ \oplus \mathcal{V}^- = \mathcal{V}$, $\mathcal{V}^+$ et $\mathcal{V}^-$ forment des sous espaces de $\mathcal{V}$ de dimension deux. Une colonne F de $\mathcal{V}^+$ (ou $\mathcal{V}^-$) ne contient donc que deux composantes indépendantes, dont les deux autres peuvent être déduites. Cela signifie par exemple que la donnée de F_1 et F_3 (c'est à dire la valeur sur $\mathcal{P}$ des composantes de champ E_z et H_z , cf (6.12) et (6.14)) suffit pour déterminer F_2 et F_4 , et par conséquent pour obtenir la valeur du

champ vérifiant la C.O.S. dans Ω^+ (ou Ω^-) et ayant E_z et H_z pour limites sur $\mathcal{P}$ (par application des formules de Kirchhoff-Helmholtz). On remarquera que cette propriété se réduit dans le cas particulier de milieux isotropes à la propriété 2 de l'annexe 4.

Propriété 11 : Parmi les quatre équations que traduit l'égalité $\mathbb{S} F = F$ (ou bien $\mathbb{S} F = -F$), les équations 2 et 4 sont conséquence des équations 1 et 3.

Preuve : Supposons qu'une colonne $F \in \mathcal{V}$ vérifie les équations 1 et 3 de l'égalité :

$$\mathbb{S} F = F, \quad (6.52)$$

qui s'écrit aussi (puisque $\mathbb{S} = \mathbb{I} - 2 \mathbb{P}^-$) :

$$\mathbb{P}^- F = 0. \quad (6.53)$$

Posons $G = \mathbb{P}^- F$. Si F vérifie les équations 1 et 3 de (6.52), elle vérifie aussi les équations 1 et 3 de (6.53), et par conséquent $G_1 = G_3 = 0$. Or G appartient à $\mathcal{V}^-$, ce qui entraîne que $G_2 = G_4 = 0$. (propriété 10). D'où $G = \mathbb{P}^- F = 0$, d'où l'on déduit que $F \in \mathcal{V}^+$ (propriété 9) et donc que les quatre équations de (6.52) sont vérifiées.

6.10. SUR L'USAGE DES OPERATEURS DE CALDERON

Considérons à nouveau le problème de la diffraction par le réseau de la figure 6.1. Le champ total $(\vec{E}, \vec{H})$ possède sur $\mathcal{P}$ des composantes tangentiellles continues, décrites par la colonne F (équations (6.12-6.15)) qui représente l'inconnue du problème. Les composantes tangentiellles sur $\mathcal{P}$ du champ incident sont représentées par la colonne F^{inc} . Nous notons $\mathbb{P}_1^+$, $\mathbb{P}_1^-$, $\mathbb{S}_1$, $\mathcal{V}_1^+$ et $\mathcal{V}_1^-$ les opérateurs et les espaces associés au milieu isotrope de permittivité $\varepsilon_1[I]$, et $\mathbb{P}_2^+$, $\mathbb{P}_2^-$, $\mathbb{S}_2$, $\mathcal{V}_2^+$, $\mathcal{V}_2^-$ ceux associés au milieu anisotrope de permittivité $[\varepsilon_2]$. On remarquera que $F \in \mathcal{V}_2^-$ ("valeur au bord" d'un champ vérifiant la C.O.S. dans Ω^-), que $F - F^{\text{inc}} \in \mathcal{V}_1^+$ ("valeur au bord" du champ réfléchi vérifiant la C.O.S. dans Ω^+), et que $F^{\text{inc}} \in \mathcal{V}_1^-$ ("valeur au bord" d'un champ vérifiant la C.O.S. dans Ω^- supposé rempli du milieu isotrope). Le problème de diffraction (recherche de F) se résoud en

écrivait que :

$$F - F^{\text{inc}} \in \mathcal{V}_1^+ \quad \text{et} \quad F \in \mathcal{V}_2^- . \quad (6.54)$$

Compte tenu de la propriété 9 (§ 6.9), (6.54) peut aussi s'écrire :

$$\mathbb{P}_1^- (F - F^{\text{inc}}) = 0 \quad \text{et} \quad \mathbb{P}_2^+ F = 0 , \quad (6.55)$$

ou bien :

$$\mathbb{S}_1 (F - F^{\text{inc}}) = F - F^{\text{inc}} \quad \text{et} \quad \mathbb{S}_2 F = -F , \quad (6.56)$$

ou encore, en utilisant le fait que $\mathbb{S}_1 F^{\text{inc}} = -F^{\text{inc}}$ (car $F^{\text{inc}} \in \mathcal{V}_1^-$) :

$$\mathbb{S}_1 F = F - 2 F^{\text{inc}} \quad \text{et} \quad \mathbb{S}_2 F = -F . \quad (6.57)$$

Toutes ces formulations sont strictement équivalentes. On pourra obtenir d'autres équations en réalisant des combinaisons linéaires des précédentes dans le but de faciliter le traitement numérique (on peut par exemple envisager d'éliminer les noyaux les plus singuliers figurant dans ces opérateurs intégraux au moyen de telles combinaisons linéaires).

Les conditions (6.54 à 6.57) représentent chacune huit équations intégrales couplées pour quatre fonctions inconnues. Nous avons vu à la fin du § 6.4 comment, au moyen de la propriété 11 du § 6.9, l'on peut se ramener à un système de quatre équations intégrales couplées.

Il est également possible de se ramener à un système de seulement deux équations intégrales en introduisant deux fonctions inconnues auxiliaires. L'inconnue F peut en effet s'écrire sous la forme :

$$F = \begin{vmatrix} F_1 \\ F_2 \\ F_3 \\ F_4 \end{vmatrix} = \begin{vmatrix} G_1 = F_1 \\ G_2 \\ G_3 = F_3 \\ G_4 \end{vmatrix} + \begin{vmatrix} 0 \\ J_2 = F_2 - G_2 \\ 0 \\ J_4 = F_4 - G_4 \end{vmatrix} = G + J$$

où G est choisi de telle sorte que $G \in \mathcal{V}_1^-$. Cela est possible car tout élément de $\mathcal{V}_1^-$ est entièrement déterminé par la donnée de sa première et de sa troisième composante, soit $G_1 = F_1$ et $G_3 = F_3$ (propriété 10, § 6.9). On a donc $\mathbb{P}_1^+ F = \mathbb{P}_1^+ G + \mathbb{P}_1^+ J = \mathbb{P}_1^+ J$, qui donne, par addition avec la première

égalité de (6.55) :

$$F = P_1^+ J + P_1^- F^{inc} = P_1^+ J + F^{inc} . \quad (6.58)$$

En appliquant P_2^+ à (6.58), et compte tenu de la deuxième égalité de (6.55), on obtient :

$$P_2^+ P_1^+ J + P_2^+ F^{inc} = 0 , \quad (6.59)$$

ce qui représente un système de quatre équations intégrales (non indépendantes) pour les deux inconnues J_2 et J_4 . Il suffit de retenir deux d'entre elles, obtenant ainsi deux équations intégrales couplées dont la résolution donne J_2 et J_4 . La solution du problème s'obtient alors par (6.58). On notera que cette façon de procéder constitue une généralisation au cas des réseaux anisotropes de la méthode proposée par D. Maystre [37] pour des structures isotropes.

Nous espérons avoir ainsi convaincu le lecteur de l'utilité du formalisme des opérateurs dans ce type de problèmes. Ils permettent notamment d'écrire les équations sous une forme condensée, plus apte à dégager des propriétés générales des "êtres" que nous sommes amenés à manipuler et qui facilite la mise en équations du problème en faisant apparaître plus clairement les différentes méthodes de résolution possibles.

Enfin, on pourra noter que les espaces $\mathcal{V}$, $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$ introduits dans le chapitre 2 sont bien entendu les mêmes que ceux que nous utilisons ici (il s'agissait alors de milieux isotropes, et les éléments de ces espaces ne comportaient que deux composantes au lieu de quatre). On retiendra que, dans la méthode de synthèse, il s'agissait de caractériser ces espaces par des familles totales, alors que pour l'utilisation des méthodes intégrales, nous les caractérisons par des opérateurs.

REFERENCES

- [1] V.N. Alexandroff, Les solutions des équations de Maxwell dans le cas de réfraction conique interne, J. Optics (Paris), Vol.15, n°4, 1984, p.219-228.
- [2] B. Audone, P. Uslenghi, Reflection and transmission for general multi-layered anisotropic structures. Proceedings of the 1989 URSI Symposium on electromagnetic theory, Stockholm, p.286-288, 1989.
- [3] J. Bass, Cours de Mathématiques, Tome III, Masson, 1971.
- [4] D.W. Berreman, Optics in stratified and anisotropic media : 4×4 - matrix formulation, J. Opt. Soc. Am., Vol.62, n°4, April 1972, p.502-510.
- [5] E. Oran Brigham, The Fast Fourier Transform, Prentice-Hall, Inc., Englewood Cliffs, New Jersey, 1974.
- [6] G. Bruhat, Optique, Masson (Paris) 1947.
- [7] A. Boag, Y. Leviatan, A. Boag, Analysis of diffraction from echelette gratings, using a strip-current model, J. Opt. Soc. Am. A, Vol.6, n°4, April 1989, p.543-549.
- [8] A. Boag, Y. Leviatan, A. Boag, Analysis of electromagnetic scattering from doubly periodic arrays using a patch current model, Proceedings of the 1989 URSI Symposium on electromagnetic theory, Stockholm, Août 1989, p.605-606.
- [9] Born and Wolf, Principles of Optics, Pergamon Press, 1965.
- [10] M. Cadilhac, Laboratoire d'Optique Electromagnétique, Faculté de St-Jérôme Marseille.
- [11] M. Cadilhac et R. Petit, Sur la "totalité" de certaines suites de fonctions souvent utilisées pour représenter des champs électromagnétiques. Comptes rendus du 9^{ème} Colloque "Optique Hertzienne et Diélectriques", Pise, Septembre 1987, p.78-81.
- [12] G. Cerutti-Maori, R. Petit, M. Cadilhac, Etude numérique du champ diffracté par un réseau, C.R. Acad. Sci. Paris, t.268, Avril 1969, p.1060-1063.
- [13] M. Cessenat, The use of the Calderon projectors and the capacity operators in scattering, dans Scattering in Volumes and Surfaces, édité par Nieto-Vesperinas et J.C. Dainty, North-Holland, Elsevier Science Publishers B.V., 1990.
- [14] L. Chambadal, Dictionnaire de Mathématiques, Hachette, 1981.
- [15] R.E. Collin, Field theory of guided waves, Mc Graw-Hill Book Company, 1960.

- [16] D. Colton, R. Kress, Integral equation methods in scattering theory, Wiley-Interscience, 1983.
- [17] Corti et Franceschetti, Radiation of electromagnetic waves in very general time-invariant linear media, Alta Frequenza, Vol.38, N. Speciale, 1969 (selected papers from the URSI Symposium on electromagnetic waves, Stresa, June 1968), p.9-15.
- [18] N. Damaskos, A. Maffett, P. Uslenghi, Reflection and transmission for gyroelectromagnetic biaxial layered media. J. Opt. Soc. Am. A, Vol.2, n°3, p.454-461, 1985.
- [19] R. Dautray, J.L. Lions, Analyse mathématique et calcul numérique, Masson, 1988.
- [20] F. Flory, H. Youlin, E. Pelletier, Optical anisotropy of thin film materials measured with guided waves techniques, OSA, Seattle, Oct. 1986.
- [21] F.R. Gantmacher, Théorie des matrices, tome 1, Collection universitaire de mathématiques, Dunod, Paris, 1966.
- [22] E.N. Glytsis et T.K. Gaylord, Rigorous three-dimensional coupled-wave diffraction analysis of single and cascaded anisotropic gratings, J. Opt. Soc. Am. A, Vol.4, n°11, Nov. 1987, p.2061-2080.
- [23] P. Henrici, Discrete variable methods in ordinary differential equations, John Wiley, New-York, 1962, p.192 et 249.
- [24] P. Henrici, Elements of numerical analysis, John Wiley, New York, 1964.
- [25] J.P. Hugonin, R. Petit, M. Cadilhac, Plane-wave expansions used to describe the field diffracted by a grating, J. Opt. Soc. Am., 71, 5, 1981, p.593-598.
- [26] H. Ikuno, K. Yasuura, Improved point matching method with applications to scattering from a periodic surface, IEEE Trans. Ant. and Propag., AP-21, 5, 1973, p.657-662.
- [27] A. Jenkins et E. White, Fundamentals of optics, Mc Graw-Hill Book Company, International student edition, 1981, p.554.
- [28] D.S. Jones, The Theory of Electromagnetism, Pergamon Press, 1964.
- [29] R. Kleiman, P. Martin, On single integral equations for the transmission problem of acoustics, SIAM J. Appl. Math., Vol.48, n°2, 1988, p.307-325.
- [30] A. Lakhtakia, V.V. Varadan et V.K. Varadan, Time-harmonic and time-dependent dyadic Green's functions for some uniaxial gyro-electromagnetic media, Applied Optics, Vol.28, n°6, March 1989, p.1049-1052.
- [31] J. Lavoine, Transformation de Fourier des pseudo-fonctions, Editions du CNRS, 1963.

- [32] M. Lavrentiev et B. Chabat, Méthodes de la théorie des fonctions d'une variable complexe, Editions de Moscou, 1972.
- [33] P.J. Lin-Chung and S. Teitler, 4×4 matrix formalism for optics in stratified anisotropic media, J. Opt. Soc. Am. A, Vol.1, n°7, July 1984, p.703-705.
- [34] E.G. Loewen, M. Nevière, D. Maystre, Grating efficiency theory as it applies to blazed and holographic gratings, Applied Optics, Vol.16, p.2711-2721, 1977.
- [35] E.G. Loewen and M. Nevière, Simple selection rules for VUV and XUV diffraction gratings. Applied Optics, Vol.17, n°7, p.1087-1092, 1978.
- [36] N. Marcuvitz, On the formal theory of electromagnetic fields, Alta Frequenza, Vol.38, N. Speciale, 1969 (selected papers from the URSI Symposium on electromagnetic waves, Stresa, June 1968), p.1-8.
- [37] D. Maystre, Thèse d'Etat, Université Aix-Marseille III, 1974.
- [38] D. Maystre, Rigorous vector theories of diffraction gratings, in Progress in Optics ~~XXI~~, Elsevier Science Publishers B.V., 1984.
- [39] D. Maystre, R. Petit, Sur la détermination de la répartition angulaire de l'énergie diffractée par un réseau échelette infiniment conducteur, C. R. Acad. Sci. Paris, t.271, 1970, p.400-403.
- [40] D. Maystre, R. Petit, Sur la détermination du champ diffracté par un réseau holographique, Opt. Commun., Vol.2, n°7, 1970, p.309-311.
- [41] W.C. Meecham, J. Appl. Phys., 27, 361, 1956.
- [42] R.F. Millar, The Rayleigh hypothesis and a related least squares solution to scattering problems for periodic surfaces and other scatterers, Radio Science, 8, 1973, p.785-796.
- [43] M.G. Moharan and T.K. Gaylord, Rigorous coupled wave analysis of planar-grating diffraction, J. Opt. Soc. Am., Vol.71, n°7, p.811-818, July 1981.
- [44] Morse and Feshbach, Methods of Theoretical Physics, Mc Graw-Hill Book Company, 1953.
- [45] M. Nevière, Laboratoire d'Optique Electromagnétique, Faculté de St-Jérôme, Marseille.
- [46] M. Nevière, R. Petit, M. Cadilhac, About the theory of optical grating coupler-waveguide systems, Optics Commun., Vol.8, n°2, p.113-117, 1973.
- [47] M. Nevière, D. Maystre, P. Vincent, Application du calcul des modes de propagation à l'étude théorique des anomalies des réseaux recouverts de diélectrique, J. Optics (Paris), Vol.8, n°4, p.231-242, 1977.
- [48] M. Nevière and J. Flamand, Electromagnetic theory as it applies to X-ray and XUV gratings, Nuclear Instruments and Methods 172, p.273-279, 1980.

- [49] M. Nevière and P. Vincent, Differential theory of gratings : answer to an objection on its validity for TM polarization, J. Opt. Soc. Am. A, Vol.5, n°9, Sept. 1988, p.1522-1524.
- [50] Y. Okuno and H. Ikuno, The Yasuura method, Proceedings of the 1989 URSI Symposium on the electromagnetic theory, Stockholm, Août 1989, p.619-621.
- [51] P.A. Technology, Cambridge.
- [52] J. Pavageau, J. Bousquet, Diffraction par un réseau conducteur. Nouvelle méthode de résolution, Optica Acta, Vol.17, n°6, 1970, p.469-478.
- [53] E. Pelletier, Ecole Nationale Supérieure de Physique de Marseille, Laboratoire d'Optique des Surfaces et des Couches Minces, 13397 Marseille Cedex 13.
- [54] P.S. Pershan, Magneto-optical effects, J. Appl. Phys., Vol.38, n°3, March 1967, p.1482-1490.
- [55] R. Petit, Etude numérique de la diffraction par un réseau. C.R. Acad. Sci. Paris, t.260, 1965, p.4454-4457.
- [56] R. Petit, Rev. Opt., n°8, Août 1966, p.353-370.
- [57] R. Petit, Electromagnetic grating theories : limitations and successes, Nouv. Rev. Optique, t.6, n°3, 1975, p.129-135.
- [58] R. Petit, Editor, Electromagnetic theory of gratings, Topics in Current Physics, Vol.22, Springer-Verlag, 1980.
- [59] R. Petit, Ondes électromagnétiques en Radioélectricité et en Optique, Masson, 1989.
- [60] R. Petit, M. Cadilhac, Etude théorique de la diffraction par un réseau, C.R. Acad. Sci. Paris, t.259, 1964, p.2077-2080.
- [61] R. Petit, M. Cadilhac, Electromagnetic theory of gratings : some advances and some comments on the use of the operator formalism, J. Opt. Soc. Am. A., Vol.7, n°9, September 1990, p.1666-1674.
- [62] R. Petit, G. Tayeb, About the electromagnetic theory of gratings made with anisotropic materials, Proceedings of SPIE, San Diego, Vol.815, 1987, p.11-16.
- [63] R. Reinisch, M. Nevière, Electromagnetic theory of diffraction in non-linear optics and surface-enhanced nonlinear optical effects, Phys. Rev. B, Vol.28, n°4, p.1870-1885, 1983.
- [64] K. Rokushima and J. Yamakita, Scattering and waveguiding properties of general slanted gratings consisting of anisotropic media, Proceedings of SPIE, Vol.401, April 1983, p.174-179.
- [65] K. Rokushima and J. Yamakita, Analysis of anisotropic dielectric gratings, J. Opt. Soc. Am., Vol.73, n°7, July 1983, p.901-908.

- [66] K. Rokushima and J. Yamakita, Analysis of diffraction in periodic liquid crystals : the optics of the chiral smectic C phase, J. Opt. Soc. Am. A, Vol.4, n°1, January 1987, p.27-33.
- [67] K. Rokushima, J. Yamakita, S. Mori, K. Tominaga, Unified approach to wave diffraction by space-time periodic anisotropic media, IEEE Trans. MTT, Vol.MTT-35, n°11, Nov. 1987, p.937-945.
- [68] L. Schwartz, Méthodes mathématiques pour les sciences physiques, Hermann, Paris, 1965.
- [69] Laurent Schwartz, Théorie des distributions, Hermann, Paris, 1966.
- [70] L. Schwartz, Analyse ; deuxième partie : topologie générale et analyse fonctionnelle, Hermann, 1970.
- [71] O. Schwelb, Stratified lossy anisotropic media : general characteristics, J. Opt. Soc. Am. A, Vol.3, n°2, February 1986, p.188-193.
- [72] M.R. Scott and H.A. Watts, Computational solution of linear two-point boundary value problems via orthonormalization, SIAM J. Numer. Anal., Vol.14, n°1, March 1977, p.40-70.
- [73] Shao and Zheng, A proving method of symmetries of magnetic dyadic Green's function, Proceedings of the URSI International Symposium on Electromagnetic Theory, Budapest, Août 1986, p.612-614.
- [74] G. Tayeb, Sur la résolution de l'équation $\text{rot rot } \vec{E} - [A]\vec{E} = \vec{S}$ en théorie des réseaux anisotropes, Comptes rendus du 9^{ème} colloque Optique Hertzienne et Diélectriques, Pise, Septembre 1987, p.201-204.
- [75] G. Tayeb, R. Petit, M. Cadilhac, On the theoretical and numerical study of diffraction gratings coated with anisotropic layers, Proceedings of SPIE, Montréal, August 1987, Vol.813, p.407-408.
- [76] G. Tayeb, M. Cadilhac, R. Petit, Sur l'étude théorique de réseaux de diffraction constitués de matériaux anisotropes, C.R. Acad. Sci. Paris, t.307, Série II, 1988, p.711-714.
- [77] G. Tayeb, Sur l'étude numérique de réseaux de diffraction constitués de matériaux anisotropes, C.R. Acad. Sci. Paris, t.307, Série II, 1988, p.1501-1504.
- [78] G. Tayeb, R. Petit, Researches on gratings made with anisotropic materials : how is the work progressing in our Laboratory, Proceedings of SPIE, Hambourg, Sept. 1988, Vol. 1029, p.170-177.
- [79] G. Tayeb, R. Petit, Periodic structures and anisotropic media : a comparison of the numerical results obtained by integral and differential methods, Proceedings of SPIE, Hambourg, Sept. 1988, Vol.1029, p.168-169.
- [80] J. Van Bladel, Some remarks on Green's dyadic for infinite space, IRE Trans. on Ant. and Propag., November 1961, p.563-566.

- [81] J. Van Bladel, *Electromagnetic Fields*, Mc Graw-Hill Book Company, 1964.
- [82] C. Vassallo, *Electromagnétisme classique dans la matière*, Dunod, Collection Technique et Scientifique des Télécommunications, 1980.
- [83] P. Vincent, N. Paraire, M. Nevière, A. Koster, R. Reinisch, Gratings in nonlinear optics and optical bistability, *J. Opt. Soc. Am. B*, Vol.2, p.1106-1116, 1985.
- [84] Weiglhofer and Papousek, Scalar Hertz potentials and dyadic Green's functions for gyrotropic media, *Proceedings of the URSI International Symposium on Electromagnetic Theory*, Budapest, Août 1986, p.615-617.
- [85] R.S. Weis and T.K. Gaylord, Electromagnetic transmission and reflection characteristics of anisotropic multilayered structures, *J. Opt. Soc. Am. A*, Vol.4, n°9, Sept. 1987, p.1720-1740.
- [86] A. Wirgin, Green function theory of the scattering of electromagnetic waves from a cylindrical boundary of arbitrary shape, *Optics Commun.*, Vol.7, n°1, 1973, p.65-69.
- [87] A. Wirgin, *C. R. Acad. Sci. Paris*, A 289, 259, 1979.
- [88] H. Wöhler, G. Haas, M. Fritsch, D.A. Mlynski, Faster 4×4 matrix method for uniaxial inhomogeneous media, *J. Opt. Soc. Am. A*, Vol.5, N°9, Sept. 1988, p.1554-1557.
- [89] P. Yeh, Electromagnetic propagation in birefringent layered media, *J. Opt. Soc. Am.*, Vol.69, n°5, May 1979, p.742-755.

Annexe 1.

On the use of the energy balance criterion
as a check of validity of computations in grating theory

R. Petit and G. Tayeb

Laboratoire d'Optique Electromagnétique, U.A. CNRS n° 843, Faculté des
Sciences, Centre de St-Jérôme, 13397 Marseille Cedex 13, France

ABSTRACT

We consider rigorous theories of diffraction gratings in which the electromagnetic field can be considered as infinite series. Numerical implementation of these theories needs truncation of the series. We show that in many circumstances, some properties of the exact solution (conservation of energy, reciprocity relations) are also verified by the truncated solution. Consequently these properties can not be systematically used as a check of validity of the truncated solution.

1. DEFINITION OF THE GRATING PROBLEM

We deal with time harmonic fields represented by complex vectors taking into account a time dependence in $\exp(-i\omega t)$. We use a rectangular coordinate system $Oxyz$ and denote by $\hat{e}_x$, $\hat{e}_y$, $\hat{e}_z$ the unit vectors of x , y , z axes. The permeability is equal to μ_0 everywhere. Referring to fig.1, we consider a structure composed of three regions. The superstrate ($y > a$) and the substrate ($y < 0$) are homogeneous regions of relative permittivities ϵ_1 and ϵ_s (ϵ_1 is supposed to be real). In the region $0 < y < a$, the relative permittivity ϵ is a function of x and y which is, for fixed y , periodic with respect to x (period d). Some particular cases of this periodic structure are for instance the coated grating (fig.2) and the "slanted" grating. The grating is illuminated by a plane wave and we denote the total field by $\vec{E}$ and $\vec{H}$. We assume for the sake of simplicity that the incident wave vector lies in the xy plane and that consequently the fields are z independent. We put $k_0 = \omega \sqrt{\epsilon_0 \mu_0}$. In all the paper, $\bar{Z}$ denotes the complex conjugate of Z , and δ_{nm} is the Kronecker symbol.

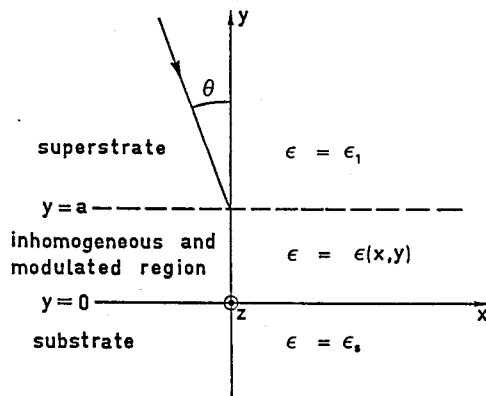


Figure 1. The structure we study.
In the region $0 < y < a$, ϵ is, for fixed y , a periodic function of x .

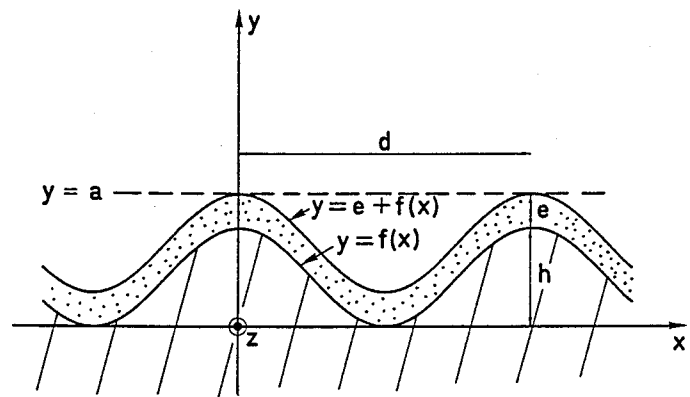


Figure 2. A particular case of fig.1 : the coated grating.

2. THE EQUATIONS VERIFIED BY THE FIELDS

We recall here the principles of the so-called differential method,² in the general case of an arbitrary polarization. We know that the incident plane wave gives birth to "pseudo-periodic" fields,² which means that any component (say for instance E_x) can be written as :

$$E_x(x, y) = \sum_{n=-\infty}^{+\infty} E_{xn}(y) \exp(i\alpha_n x), \quad (1)$$

with $\alpha_n = \alpha_0 + nK$, $\alpha_0 = \sqrt{\epsilon_1} k_0 \sin \theta$, $k_0 = \omega \sqrt{\epsilon_0 \mu_0} = 2\pi/\lambda$ and $K = 2\pi/d$.

The right member of (1) can be interpreted as a generalized Fourier series, and consequently we will say that the $E_{xn}(y)$ are the Fourier coefficients of the field $E_x(x, y)$. Noticing that $E_x(x+d, y) = \exp(i\alpha_0 d) E_x(x, y)$, we say that $\exp(i\alpha_0 d)$ is the pseudo periodicity coefficient

of the field. The Fourier coefficients are taken as the unknowns of the problem. We will subsequently denote by $|E_x|$ the infinite column matrix whose elements are the E_{xn} . This matrix is a function of y . The same notation will be used for the other field's components.

Elementary calculations detailed in appendix A show that :

- 1) The two components E_y and H_y can be deduced from the four others (eq.(A1) and (A2)) ; we will take interest only in E_x , E_z , H_x , H_z , and more precisely in the column matrices $|E_x|$, $|E_z|$, $|H_x|$ and $|H_z|$.

- 2) Maxwell equations can be written in the form :

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_x| = A(y) |H_z| , \\ \frac{d}{dy} |E_z| = B(y) |H_x| , \\ \frac{d}{dy} |H_x| = C(y) |E_z| , \\ \frac{d}{dy} |H_z| = D(y) |E_x| , \end{array} \right. \quad (2)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_x| = A(y) |H_z| , \\ \frac{d}{dy} |E_z| = B(y) |H_x| , \\ \frac{d}{dy} |H_x| = C(y) |E_z| , \\ \frac{d}{dy} |H_z| = D(y) |E_x| , \end{array} \right. \quad (3)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_x| = A(y) |H_z| , \\ \frac{d}{dy} |E_z| = B(y) |H_x| , \\ \frac{d}{dy} |H_x| = C(y) |E_z| , \\ \frac{d}{dy} |H_z| = D(y) |E_x| , \end{array} \right. \quad (4)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_x| = A(y) |H_z| , \\ \frac{d}{dy} |E_z| = B(y) |H_x| , \\ \frac{d}{dy} |H_x| = C(y) |E_z| , \\ \frac{d}{dy} |H_z| = D(y) |E_x| , \end{array} \right. \quad (5)$$

where A, B, C, D are "infinite square" matrices.

It is worth noting that the equations (2) to (5) are valid in the sense of distributions,² which means that the continuity relations on the surfaces where ϵ is discontinuous (i.e. the boundary conditions) are automatically taken into account. In general, matrices A, B, C, D are y dependent, and eq.(2) to (5) have to be integrated between $y = 0$ and $y = a$ using a numerical algorithm.

3. EXPRESSION OF THE LORENTZ RECIPROCITY THEOREM

Let us consider two solutions $(\vec{E}, \vec{H})$ and $(\vec{E}', \vec{H}')$ of the harmonic Maxwell equations in a volume V bounded by the closed surface S . If there is no current distribution in V , and if $\vec{n}$ denotes the unit normal of S , the Lorentz reciprocity theorem³ takes the form :

$$\iint_S (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{n} \, dS = 0 . \quad (6)$$

If S is the surface depicted in fig.3, then, since the fields are z independent :

$$\int_{A_1 B_1 B_2 A_2} (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{n} \, dl = 0 \quad (7)$$

Let us suppose now that $(\vec{E}, \vec{H})$ is a pseudo-periodic field which can be written as in (1), and that $(\vec{E}', \vec{H}')$ is another pseudo-periodic field whose components can be written as :

$$E'_x(x, y) = \sum_{n=-\infty}^{+\infty} E'_{xn}(y) \exp(-i\alpha_n x) . \quad (8)$$

Let us note that $(\vec{E}', \vec{H}')$ has a pseudo periodicity coefficient inverse of that of $(\vec{E}, \vec{H})$. It is also worth noting that, if $(\vec{E}, \vec{H})$ represents for instance the electromagnetic field corresponding to the propagating plane waves depicted fig.4a, $(\vec{E}', \vec{H}')$ can represent the electromagnetic field of another problem where the incident plane wave falls on the grating with an incidence $-\theta_p$. Fig. 4b illustrates this case, supposing $p = 1$.

As shown in appendix B, (7) implies that the integral :

$$I = \int_0^d (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{e}_y \, dx , \quad \text{is } y \text{ independent} \quad \left(\frac{dI(y)}{dy} = 0 \right) . \quad (9)$$

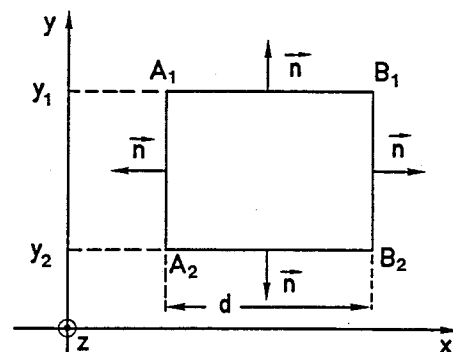


Figure 3. y_1 and y_2 can take any value ; the "walls" $(A_1 A_2)$ and $(B_1 B_2)$ are distant from the grating spacing d . S is a cylinder of unit length whose cross section is $(A_1 B_1 B_2 A_2)$.

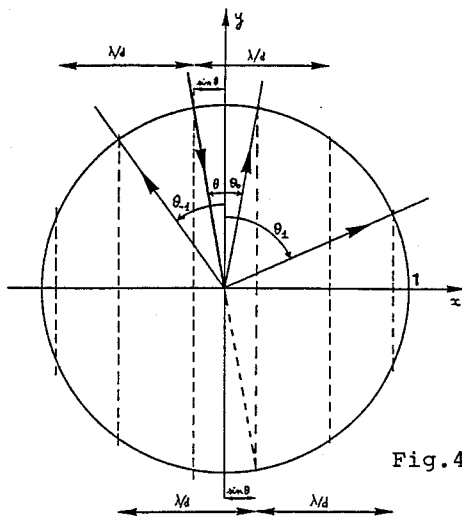


Fig.4a

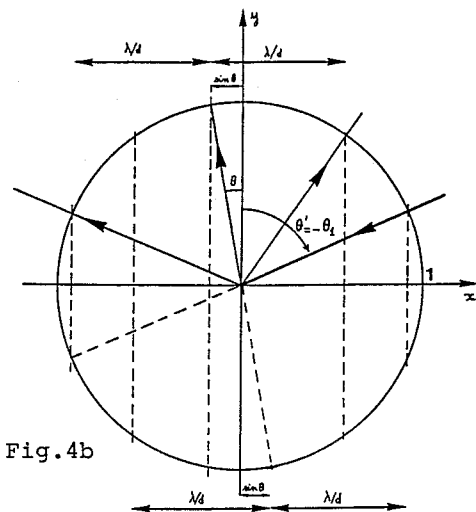


Fig.4b

Figures 4. Only the plane waves propagating in the superstrate are represented. λ is the wavelength in the superstrate. The radius of the circle is equal to unity. The incidence angles are measured anticlockwise, while the diffracted angles are measured clockwise.

We want to emphasize that the property (9), which is true for the exact solution of the problem, still holds true (see appendix B) for a truncated solution which verifies the truncated equations (2) to (5). We call truncated solution a solution obtained assuming that a Fourier series such as (1) or (8) can be replaced by a finite sum, an assumption which is obviously necessary for the numerical implementation.

4. CONSERVATION OF ENERGY

In this section, we deal with unlossy materials, consequently $\epsilon(x,y)$ is a real function. In eq.(9), $(\vec{E}', \vec{H}')$ represents an electromagnetic field solution of the harmonic Maxwell equations, which has a coefficient of pseudo periodicity inverse of that of $(\vec{E}, \vec{H})$. It is easy to show (Appendix C) that $(\vec{E}' = \vec{E}, \vec{H}' = -\vec{H})$ verifies these conditions. We make this choice in this section, and from (9), we can claim that :

$$\int_0^d (\vec{E} \wedge \vec{H} + \vec{E}' \wedge \vec{H}') \cdot \vec{e}_y dx = 2 \int_0^d \text{Re} (\vec{E} \wedge \vec{H}) \cdot \vec{e}_y dx, \quad \text{is independent of } y. \quad (10)$$

We recognize here an expression of the Poynting theorem*, which as it is well known, leads to the famous energy balance criterion**. Therefore, as a consequence of considerations developed in section 3, it appears that this criterion is verified not only by the exact solution, but also by any truncated solution verifying the equations deduced from (2) to (5) by truncature. To lay stress on this rather amazing result, consider for instance a grating problem in which the incident wave gives rise to N propagating diffracted waves in the superstrate. If the approximate numerical solution is obtained by retaining only P Fourier coefficients for the fields (P being possibly less than N, and even equal to 1), the energy balance criterion will be verified whatever P, provided that the integration of eq.(2) to (5) between $y = 0$ and $y = a$ is accurately performed. Indeed, as claimed by Nevière et al.,⁴ and at least for certain integration algorithms, it seems that the energy balance criterion is verified whatever the integration step. Anyway, in all the computations we have done using the differential method, this criterion has been verified with a precision of the order of 10^{-5} (even when using unreasonable truncatures and very large integration steps ; cf table 1.)

* Let us remark that (10) involves a spatially averaged (on the grating spacing d) Poynting vector.

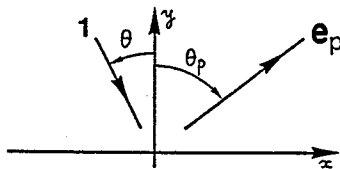
** i.e. the sum of all the diffracted efficiencies is equal to unity.

	order	reflected waves		transmitted waves		sum of efficiencies
		angle	efficiency	angle	efficiency	
$P = 1$	0	20°	0.002497	13.18°	0.997503	1.00000
$P = 11$	-3			-76.41°	0.000293	1.00000
	-2	-59.09°	0.005353	-34.89°	0.000275	
	-1	-14.95°	0.018351	-9.90°	0.061074	
	0	20°	0.005751	13.18°	0.778558	
	1	70.39°	0.020611	38.90°	0.109734	

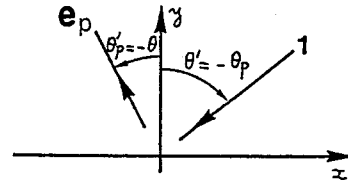
Table 1. Relative to a sinusoidal grating described fig.2, with :
 $e = 0$, $d = 1 \mu\text{m}$, $h = 0.2 \mu\text{m}$, $\lambda = 0.6 \mu\text{m}$, $\theta = 20^\circ$, superstrate is vacuum, $\epsilon = 2.25$ for the substrate. P is the number of Fourier coefficients used to represent the fields.
Polarization : $\vec{E}$ is parallel to Oz (TE case).
Integration between $y = 0$ and $y = a$ is performed using 38 steps with Runge-Kutta (4th order) algorithm.

5. RECIPROCITY

In both the superstrate ($y > a$) and the substrate ($y < 0$), each Fourier coefficient of the fields can be written as the product of a constant (usually called the Rayleigh coefficient) by a function exponentially dependent on y .



Case 1



Case 2

Figure 5. In case 1, the grating G is illuminated under the incidence θ and we turn our attention to the p^{th} diffracted wave (angle θ_p). In case 2, G is illuminated under the incidence $\theta' = -\theta_p$. The reciprocity theorem claims that the p^{th} diffracted wave in case 2 and the incident wave in case 1 propagate in opposite directions, and that the efficiencies in the p^{th} order are the same (say e_p) in both cases. The theorem holds also for a transmitted order, when dealing with lossless substrates.

		TE	TM
$P = 3$	case 1	0.1190	0.0619
	case 2	0.1179	0.0674
$P = 5$	case 1	0.1104	0.0627
	case 2	0.1119	0.0644
$P = 7$	case 1	0.109741	0.06312
	case 2	0.109727	0.06322
$P = 11$	case 1	0.109606	0.063275
	case 2	0.109685	0.063295

Table 2. Relative to a sinusoidal grating described fig.2, with : $e = 0$, $d = 1 \mu\text{m}$, $h = 0.2 \mu\text{m}$, $\lambda = 0.6 \mu\text{m}$, superstrate is vacuum, $\epsilon = 2.25$ for the substrate. The notations are those of fig.5, for the transmitted wave in the first order ($p = 1$). In case 1, $\theta = 20^\circ$. In case 2, the incident field comes from the substrate, with $\theta' = 38.904^\circ$. P is the number of Fourier coefficients used to represent the fields. We denote by TE (resp. TM) the situation where $\vec{E}$ (resp. $\vec{H}$) is parallel to Oz . The table gives the values of e_1 . Integration between $y = 0$ and $y = a$ is performed using 40 steps with Adams-Moulton (initiated by Runge-Kutta) algorithm.

For $y > a$ (resp. $y < 0$) (see figure 1), it turns out that the integral appearing in (9) can be written as a finite sum. Each term of the sum contains the Rayleigh coefficients corresponding to the only propagating waves of the incident and reflected (resp. of the transmitted) fields. Using for $(\vec{E}', \vec{H}')$ a field associated to an incident plane wave which falls on the grating with the same direction as one of the diffracted waves of $(\vec{E}, \vec{H})$ (as explained in section 2 and fig.4), it is easy to deduce the well known reciprocity theorem (fig.5). These considerations are developed in detail by D. Maystre for a simple case of polarization in Progress in Optics,⁵ and by A. Roger in ref. 6.

Let us now consider truncated fields, described in the superstrate and in the substrate by truncated Rayleigh developments. From section 3 we know that (9) still holds true. Consequently, provided that the truncated developments contain the orders represented fig.5.1 and fig.5.2, the reciprocity theorem still holds true. We may confess that our numerical computations does not perfectly agree with this prediction (cf table 2). Up to now, we have not been able to explain this discrepancy satisfactorily (it could perhaps be due to the numerical process).

6. THE SLANTED GRATING

In this section we show that the problem of a slanted grating¹ described fig.6, even when solved using an eigenvalues problem, is a particular case of the differential method, and that the results of sections 3, 4, and 5 still hold true.

In the region $0 < y < a$, the permittivity can be written as :

$$\begin{aligned}\epsilon(x,y) &= \sum_n \epsilon_n(y) \exp(inKx) = f(\vec{L} \cdot \vec{r}) = \\ &= f(L_x \cdot x + L_y \cdot y),\end{aligned}$$

where f is a periodic function of period $(L_x \cdot d)$. It is easy to show that, in these circumstances, the Fourier coefficients of ϵ take the form :

$$\epsilon_n(y) = c_n \exp(in\mu y),$$

where

$$c_n = \frac{1}{L_x \cdot d} \int_0^{L_x \cdot d} f(t) \exp(-in \frac{K}{L_x} t) dt$$

is independent of y , and $\mu = K L_y / L_x$.

For the sake of simplicity, let us suppose that the field is polarized $\vec{E} // O_z$, and put $u(x,y) = E_z(x,y) = \sum_n u_n(y) \exp(i\alpha_n x)$.

Then, system (A_3) to (A_6) is equivalent to :

$$\Delta u(x,y) + k_0^2 \epsilon(x,y) u(x,y) = 0$$

which implies, for any n :

$$u_n''(y) - \alpha_n^2 u_n(y) + k_0^2 \sum_m \epsilon_{n-m}(y) u_m(y) = 0.$$

Putting now $u_n(y) = w_n(y) \exp(in\mu y)$, we obtain :

$$w_n'' + 2in\mu w_n' - (n^2\mu^2 + \alpha_n^2) w_n + k_0^2 \sum_m c_{n-m} w_m = 0,$$

which can be written in matrix form as :

$$|w''| = R |w'| + S |w|, \quad (11)$$

where R and S are y independent infinite matrices whose elements are :

$$R_{nm} = -2in\mu \delta_{nm}; \quad S_{nm} = (n^2\mu^2 + \alpha_n^2) \delta_{nm} - k_0^2 c_{n-m}.$$

It is important to note that (11) is strictly equivalent to the system (2) to (5).

We can also write (11) in the form :

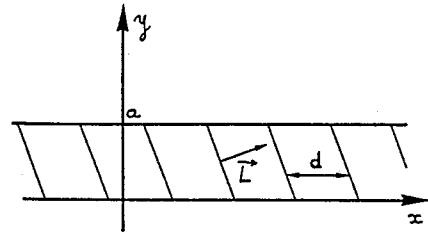


Figure 6. The slanted grating

$$\begin{pmatrix} |w'| \\ |w''| \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ S & R \end{pmatrix} \begin{pmatrix} |w| \\ |w'| \end{pmatrix} \quad (12)$$

Comparing (12) to its equivalent expression (2) to (5), it is worth noting that now the square matrices are y independent. This remark allows us to reduce the solving of (12) to an eigenvalues problem, avoiding so a numerical integration. Anyway, except for their numerical treatment, the problems are strictly equivalent. The results of sections 3, 4, and 5 are still valid, and we agree with the paper by Russell.⁷

7. A LAST EXAMPLE

Obviously, the previous considerations are reminiscent with a question which we refer to during the last SPIE Symposium. In a paper by R. Petit, J.L. Suratteau and M. Cadilhac,⁸ we had to write the continuity of two functions on a bounded interval, which is equivalent to express the vanishing of two other functions (i.e. their jumps on this interval). This was done by projecting on a convenient basis in order to get an algebraic system. We can use either the same basis or two different basis. In the first case, the energy balance is not automatically satisfied after truncature. On the other hand, when using two different basis as done by the Australian group of Sidney University,⁹ we are led to the opposite conclusion.¹⁰ In our opinion, this is an important point since the energy balance is often used as a check of validity of the numerical results obtained after truncature.

APPENDIX A

$\vec{E}$ and $\vec{H}$ verify the Maxwell equations :

$$\text{curl } \vec{E} = i\omega\mu_0 \vec{H} ; \quad \text{curl } \vec{H} = -i\omega\epsilon_0 \vec{E} .$$

Putting $Z_0 = \sqrt{\frac{\mu_0}{\epsilon_0}}$, we get :

$$\begin{cases} E_Y = -i \frac{Z_0}{k_0} \frac{1}{\epsilon} \frac{\partial H_Z}{\partial x} \end{cases} \quad (A1)$$

$$\begin{cases} H_Y = \frac{i}{k_0 Z_0} \frac{\partial E_Z}{\partial x} \end{cases} \quad (A2)$$

$$\begin{cases} \frac{\partial E_X}{\partial y} = -i \frac{Z_0}{k_0} \frac{\partial}{\partial x} \left(\frac{1}{\epsilon} \frac{\partial H_Z}{\partial x} \right) - ik_0 Z_0 H_Z \end{cases} \quad (A3)$$

$$\begin{cases} \frac{\partial E_Z}{\partial y} = i k_0 Z_0 H_X \end{cases} \quad (A4)$$

$$\begin{cases} \frac{\partial H_X}{\partial y} = \frac{i}{k_0 Z_0} \frac{\partial^2 E_Z}{\partial x^2} + i \frac{k_0}{Z_0} \epsilon E_Z \end{cases} \quad (A5)$$

$$\begin{cases} \frac{\partial H_Z}{\partial y} = -i \frac{k_0}{Z_0} \epsilon E_X . \end{cases} \quad (A6)$$

Taking into account the periodicity of ϵ and $1/\epsilon$, we write their Fourier series :

$$\epsilon(x, y) = \sum_{n=-\infty}^{+\infty} \epsilon_n(y) \exp(inKx) , \quad (A7)$$

$$\left(\frac{1}{\epsilon} \right)(x, y) = \sum_{n=-\infty}^{+\infty} \left[\frac{1}{\epsilon} \right]_n(y) \exp(inKx) . \quad (A8)$$

Replacing in (A3) to (A6) the field's components by their expressions (1), ϵ and $1/\epsilon$ by

their expressions (A7) and (A8), and with a projection on the $\exp(i\alpha_n x)$ basis, we obtain for any value of n :

$$\left\{ \begin{array}{l} \frac{dE_{xn}}{dy} = \sum_{m=-\infty}^{+\infty} \left[i \frac{Z_0}{k_0} \alpha_n \alpha_m \left[\frac{1}{\epsilon} \right]_{n-m} - i k_0 Z_0 \delta_{nm} \right] H_{zm} \end{array} \right. \quad (A9)$$

$$\left\{ \begin{array}{l} \frac{dE_{zn}}{dy} = i k_0 Z_0 H_{xn} \end{array} \right. \quad (A10)$$

$$\left\{ \begin{array}{l} \frac{dH_{xn}}{dy} = \sum_{m=-\infty}^{+\infty} \left[- \frac{i}{k_0 Z_0} \alpha_n^2 \delta_{nm} + i \frac{k_0}{Z_0} \epsilon_{n-m} \right] E_{zm} \end{array} \right. \quad (A11)$$

$$\left\{ \begin{array}{l} \frac{dH_{zn}}{dy} = \sum_{m=-\infty}^{+\infty} - i \frac{k_0}{Z_0} \epsilon_{n-m} E_{xm} \end{array} \right. \quad (A12)$$

which can be written in matrix form :

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_x| = A(y) |H_z| \end{array} \right. \quad (A13)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |E_z| = B |H_x| \end{array} \right. \quad (A14)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |H_x| = C(y) |E_z| \end{array} \right. \quad (A15)$$

$$\left\{ \begin{array}{l} \frac{d}{dy} |H_z| = D(y) |E_x| \end{array} \right. \quad (A16)$$

where A, B, C, D are infinite matrices whose elements are given by :

$$\left\{ \begin{array}{l} A_{nm}(y) = i \frac{Z_0}{k_0} \alpha_n \alpha_m \left[\frac{1}{\epsilon} \right]_{n-m} - i k_0 Z_0 \delta_{nm} \end{array} \right. \quad (A17)$$

$$\left\{ \begin{array}{l} B_{nm} = i k_0 Z_0 \delta_{nm} \end{array} \right. \quad (A18)$$

$$\left\{ \begin{array}{l} C_{nm}(y) = - \frac{i}{k_0 Z_0} \alpha_n^2 \delta_{nm} + i \frac{k_0}{Z_0} \epsilon_{n-m} \end{array} \right. \quad (A19)$$

$$\left\{ \begin{array}{l} D_{nm}(y) = - \frac{ik_0}{Z_0} \epsilon_{n-m} \end{array} \right. \quad (A20)$$

APPENDIX B

In this section we propose us to calculate the quantity

$$J = \int_{A_1 B_1 B_2 A_2} (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{n} dl \quad (\text{notations are exposed in section 3})$$

in terms of the Fourier coefficients of the fields.

Let u be a component of $(\vec{E}, \vec{H})$ and u' a component of $(\vec{E}', \vec{H}')$; eq.(1) and (8) show that:

$$u(x + d, y) = \exp(i\alpha_0 d) u(x, y), \quad (B1)$$

$$u'(x + d, y) = \exp(-i\alpha_0 d) u'(x, y), \quad (B2)$$

which means that $\vec{E} \wedge \vec{H}'$ and $\vec{E}' \wedge \vec{H}$ are d periodic with respect to x . This implies that the contributions of segments $(B_1 B_2)$ and $(A_2 A_1)$ cancel each other, and J takes the form :

$$J = \int_{A_1 B_1} (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{e}_y dl - \int_{A_2 B_2} (\vec{E} \wedge \vec{H}' - \vec{E}' \wedge \vec{H}) \cdot \vec{e}_y dl = I(y_1) - I(y_2), \quad (B3)$$

$$\text{with } I(y) = \int_0^d [\vec{E}(x, y) \wedge \vec{H}'(x, y) - \vec{E}'(x, y) \wedge \vec{H}(x, y)] \cdot \vec{e}_y dx. \quad (B4)$$

Let us now calculate $\frac{dI(y)}{dy}$; developing the vector-products, we find :

$$\frac{dI(y)}{dy} = \int_0^d \left[\left(\frac{\partial E_z}{\partial y} H'_x - \frac{\partial E'_z}{\partial y} H_x \right) + \left(E_z \frac{\partial H'_x}{\partial y} - E'_z \frac{\partial H_x}{\partial y} \right) - \left(\frac{\partial E_x}{\partial y} H'_z - \frac{\partial E'_x}{\partial y} H_z \right) - \left(E_x \frac{\partial H'_z}{\partial y} - E'_x \frac{\partial H_z}{\partial y} \right) \right] dx \quad (B5)$$

Let us study for example the quantity $\frac{dI_3(y)}{dy} = \int_0^d \left(- \frac{\partial E_x}{\partial y} H'_z + \frac{\partial E'_x}{\partial y} H_z \right) dx$.

Using eq.(1) and (8), we get :

$$\frac{dI_3(y)}{dy} = \int_0^d \sum_{n,m} \left(- \frac{dE_{xn}}{dy} H'_{zm} + \frac{dE'_{xm}}{dy} H_{zn} \right) e^{i(n-m)Kx} dx.$$

Noting that $\int_0^d e^{i(n-m)Kx} dx = d \delta_{nm}$, we obtain :

$$\frac{dI_3(y)}{dy} = d \sum_n \left(- \frac{dE_{xn}}{dy} H'_{zn} + \frac{dE'_{xn}}{dy} H_{zn} \right). \quad (B6)$$

We have shown in Appendix A that $\frac{dE_{xn}}{dy} = \sum_m A_{nm} H_{zm}$. In the same way, we find that :

$$\frac{dE'_{xn}}{dy} = \sum_m A'_{nm} H'_{zm}, \quad \text{where :}$$

$$A'_{nm} = i \frac{Z_0}{k_0} \alpha_n \alpha_m \left[\frac{1}{\epsilon} \right]_{m-n} - ik_0 Z_0 \delta_{nm}. \quad (B7)$$

Eq.(B6) becomes :

$$\begin{aligned} \frac{dI_3(y)}{dy} &= d \sum_{n,m} - A_{nm} H_{zm} H'_{zn} + A'_{nm} H'_{zm} H_{zn} \\ &= d \sum_{n=-\infty}^{+\infty} \sum_{m=-\infty}^{+\infty} (-A_{nm} + A'_{mn}) H_{zm} H'_{zn}. \end{aligned} \quad (B8)$$

It is now a trivial matter to verify that $A_{nm} = A'_{mn}$ and that consequently each term of the double series in (B8) vanishes. The same can be done for the three other terms in (B5) which also vanish. The important result is the following : the quantity $I(y)$ is independent of y even if the developments of the fields in generalized Fourier series are truncated.

APPENDIX C

a) Each component of $(\vec{E}, \vec{H})$, denoted for instance by $u(x,y)$ expresses as :

$$u(x,y) = \sum_n u_n(y) \exp(i\alpha_n x); \text{ and consequently } (\alpha_n \text{ is real}) :$$

$$\bar{u}(x,y) = \sum_n \bar{u}_n(y) \exp(-i\alpha_n x).$$

b) If $(\vec{E}, \vec{H})$ verifies the harmonic Maxwell equations : $\text{curl } \vec{E} = i\omega\mu_0 \vec{H}$;

$\text{curl } \vec{H} = -i\omega\epsilon_0 \vec{E}$, it is clear that $(\vec{E}, -\vec{H})$ verifies the same equations in the case where ϵ is real.

ACKNOWLEDGMENTS

Thanks are due to Professor M. Cadilhac for valuable discussions and advices.

The computing facilities used in this paper have been given by the Scientific Council of the "Centre de Calcul Vectoriel pour la Recherche".

REFERENCES

1. Moharam and Gaylord, "Rigorous coupled wave analysis of planar-grating diffraction", J. Opt. Soc. Am., Vol. 71, n° 7, July 1981, p. 811-818.
2. Electromagnetic Theory of Gratings, Editor R. Petit, Topics in Current Physics, vol. 22, Springer Verlag, (1980).
3. E. Collin, Field Theory of Guided Waves, Mc Graw Hill Book Company, (1960).
4. M. Nevière, G. Cerutti-Maori, M. Cadilhac, "Sur une nouvelle méthode de résolution du problème de la diffraction d'une onde plane par un réseau infiniment conducteur", Optics Commun., vol. 3, n° 1, 48-52, (1971).
5. E. Wolf, Progress in Optics XXI, Elsevier Science Publishers B.V., (1984).
6. A. Roger, "Generalized reciprocity for gratings of finite conductivity", Optica Acta, vol. 30, n° 5, 575-585, (1983).
7. P. St. J. Russel, "Power conservation and field structures in uniform dielectric gratings", J. Opt. Soc. Am. A, Vol. 1, n° 3, 293-299, (1984).
8. R. Petit, J.Y. Suratteau, M. Cadilhac, "On the numerical study of deep lamellar gratings in the resonance domain", Proceedings of SPIE, vol. 503, 160-167 (1984).
9. Botten, Craig, R. C. Mc Phedran, Adams, Andrewertha. Optica Acta, vol. 28, n° 3, 413-428, (1981) and Optica Acta, vol. 28, n° 8, 1087-1102 (1981).
10. J.Y. Suratteau, Thesis, Université d'Aix Marseille III, 14 Juin 1985.

Annexe 2.Compléments sur la "méthode différentielle améliorée".

Le problème est le suivant : déterminer une base de l'espace $\mathcal{E}_2^-(y)$ pour $0 \leq y \leq a$, connaissant une base $\tau_k(0)$ de $\mathcal{E}_2^-(0)$ et sachant que, pour $0 \leq y \leq a$, les variations des vecteurs τ_k sont données par :

$$\frac{d\tau_k(y)}{dy} = A(y) \tau_k(y) , \quad (\text{A2.1})$$

où $A(y)$ est une matrice connue. Rappelons que, pour les besoins du traitement numérique, les bases ont un nombre fini de vecteurs égal à N ($k = 1, \dots, N$) et que chacun des vecteurs $\tau_k(y)$ possède $2N$ composantes.

Remarque : dans le chapitre 1, k varie de $-P$ à $+P$. Pour des raisons de commodité, nous renumérotions ici les vecteurs τ_k de sorte que k varie de 1 à N .

1. PRELIMINAIRES

a) Comment intégrer un "vecteur" de sorte que sa "norme" reste constante ?

Soit à intégrer un système différentiel du genre :

$$\frac{d\tau(y)}{dy} = A(y) \tau(y) , \quad (\text{A2.2})$$

où $\tau(y)$ est un vecteur à $2N$ composantes et $A(y)$ une matrice connue $2N \times 2N$. Définissons le produit scalaire de deux vecteurs τ_1 et τ_2 de composantes $\tau_{1,i}$ et $\tau_{2,i}$ par :

$$\langle \tau_1, \tau_2 \rangle = \sum_{i=1}^{2N} \tau_{1,i} \overline{\tau_{2,i}} . \quad (\text{A2.3})$$

Nous désignerons par A^* la matrice adjointe de A .

Le système (A2.2) est équivalent à :

$$\frac{d\varphi(y)}{dy} = \left(A(y) - \frac{d\mu(y)}{dy} \right) \varphi(y) , \quad (A2.4)$$

$$2 \frac{d\mu(y)}{dy} = \langle \varphi(y), (A(y) + A^*(y)) \varphi(y) \rangle , \quad (A2.5)$$

$$\tau(y) = e^{\mu(y)} \varphi(y) , \quad \text{avec } \mu \text{ réel} \quad (A2.6)$$

$$\langle \varphi(y), \varphi(y) \rangle = 1 . \quad (A2.7)$$

Ce résultat se démontre aisément. Il permet de remplacer l'intégration du système à $2N$ fonctions inconnues (A2.2) par le système à $2N + 1$ fonctions inconnues (A2.4) et (A2.5). On est ainsi ramené à l'intégration d'un vecteur à $2N$ composantes $\varphi(y)$ colinéaire à $\tau(y)$, qui reste normé au cours de l'intégration, et d'un scalaire $\mu(y)$ représentant le logarithme de la norme de $\tau(y)$.

b) Caractérisation du "degré d'orthogonalité" d'une famille de vecteurs normés.

Soit φ_k une famille de N vecteurs normés ayant chacun $2N$ composantes. En utilisant le procédé d'orthogonalisation de Schmidt, on peut fabriquer à partir des φ_k une famille de vecteurs $\hat{\varphi}_k$ orthonormée. Soit P la matrice de passage des φ_k aux $\hat{\varphi}_k$ et son inverse P^{-1} la matrice de passage des $\hat{\varphi}_k$ aux φ_k . Par construction, P^{-1} est une matrice triangulaire, et son déterminant D est égal au produit de ses éléments diagonaux. Si la famille φ_k était initialement orthogonale, $P = P^{-1} = I$, et D serait égal à 1. Il apparaît également que si les φ_k étaient linéairement dépendants, au moins un des vecteurs $\hat{\varphi}_k$ serait nul, et par conséquent D serait nul. Nous définirons $|D|$ comme étant une mesure du degré d'orthonormalité de la famille φ_k , $|D|$ décroissant à partir de la valeur 1 au fur et à mesure que φ_k "s'éloigne" d'une famille orthogonale.

Remarque : cette façon de procéder nous a été dictée par le fait que le nombre de vecteurs de la famille considérée φ_k est différent du nombre de composantes de chaque vecteur. Il est facile de vérifier que si ces deux nombres étaient identiques et égaux à un entier K , la valeur de $|D|$ serait égale à la valeur absolue du déterminant de la base φ_k dans la base canonique de $\mathbb{C}^K$, et l'on retrouverait ainsi une manière plus classique de caractériser l'orthogonalité de ces vecteurs.

2. INTEGRATION DE (A2.1) DANS LA M.D.A.

Pour intégrer entre $y = 0$ et $y = a$ les N vecteurs $\tau_k(y)$ vérifiant A2.1, on peut commencer par fabriquer à partir de la base orthogonale $\tau_k(0)$ une base orthonormée $\hat{\tau}_k(0)$. On remplace (A2.1) par le système équivalent (A2.8) et (A2.9) :

$$d\varphi_k(y) = \left(A(y) - \frac{d\mu_k(y)}{dy} \right) \varphi_k(y), \quad (\text{A2.8})$$

$$\frac{d\mu_k(y)}{dy} = \left\langle \varphi_k(y), (A(y) + A^*(y)) \varphi_k(y) \right\rangle, \quad (\text{A2.9})$$

$$\text{avec } \hat{\tau}_k(y) = e^{\mu_k(y)} \varphi_k(y), \quad (\text{A2.10})$$

$$\text{et } \langle \varphi_k(y), \varphi_k(y) \rangle = 1. \quad (\text{A2.11})$$

Ce procédé a donc pour avantage de caractériser $\mathcal{E}_2^-(y)$ par une base $\varphi_k(y)$ normée, initialement orthogonale (puisque $\varphi_k(0) = \hat{\tau}_k(0)$). On utilise pour intégrer (A2.8) et (A2.9) (comme on le fait pour la M.D.C.) l'algorithme d'Adams-Moulton. On constate alors que le degré d'orthogonalité des vecteurs $\varphi_k(y)$ égal à $|D(y)|$ décroît au cours de l'intégration à partir de la valeur 1. L'expérience numérique montre même que cette décroissance est pratiquement exponentielle (la représentation de $\text{Log}|D(y)|$ en fonction de y est quasiment linéaire). Dans la M.D.A., nous testons à chaque pas d'intégration en y la valeur de $|D|$, et lorsque le degré d'orthogonalité des $\varphi_k(y)$ est jugé insuffisant, nous réorthonormalisons la base $\varphi_k(y)$ au moyen du procédé de Schmidt avant de poursuivre l'intégration. A la fin de l'intégration, nous disposons ainsi d'une base orthonormale $\varphi_k(a)$ de $\mathcal{E}_2^-(a)$, et le problème est alors réduit (cf § 1.1.3) à la résolution de :

$$F(a) = F^i(a) + \sum_k R_k \rho_k(a), \quad (\text{A2.12})$$

$$F(a) = \sum_k \hat{T}_k \varphi_k(a). \quad (\text{A2.13})$$

La relation (A2.13) traduit l'appartenance de $F(a)$ à $\mathcal{E}_2^-(a)$. Evidemment, compte tenu des orthonormalisations effectuées entre $y = 0$ et $y = a$, les $\hat{T}_k$

figurant dans (A2.13) sont différents des T_k (définis au § 1.1.3) permettant de déterminer le champ transmis. Le système (A2.12) (A2.13) permet donc de déterminer directement le champ réfléchi déduit des R_k , ce qui est généralement suffisant dans le cas de l'étude des réseaux conducteurs. Si néanmoins on s'intéresse au champ transmis, il faudra prendre soin de conserver, lors de chaque orthonormalisation au cours de l'intégration, les matrices de passage (en fait le produit de ces matrices de passage est suffisant) pour obtenir les T_k en fonction des $\hat{T}_k$.

Annexe 3Obtention de la matrice impédance à partir d'une base

Nous utilisons ici les notations du chapitre 1. On suppose connue une base de l'espace $\mathcal{E}_2^-(y)$. Cette base est formée de N vecteurs $\tau_n(y)$, chacun de ces vecteurs ayant $2N$ composantes. Un élément $F(y)$ de $\mathcal{E}_2^-(y)$ possède $2N$ composantes ; les N premières (coefficients de Fourier de $H_z(x,y)$) peuvent être considérées comme les composantes d'un vecteur $U(y)$, les N suivantes (coefficients de Fourier de $\frac{1}{\eta_0} E_x(x,y)$) comme les composantes d'un vecteur $V(y)$:

$$F = \begin{vmatrix} U \\ V \end{vmatrix}.$$

Soient P et Q les matrices $N \times N$ que nous déduisons de la base $\tau_n(y)$ et dont les éléments sont :

$$P_{i,j} = \tau_{j,i}, \quad Q_{i,j} = \tau_{j,N+i},$$

où $\tau_{j,k}$ désigne la $k^{\text{ième}}$ composante de $\tau_j(y)$. Nous savons (équation (1.24)) que $F(y)$ peut s'écrire sous la forme : $F(y) = \sum_n T_n \tau_n(y)$, c'est-à-dire :

$$U_i = \sum_n T_n \tau_{n,i} = \sum_n P_{i,n} T_n,$$

$$V_i = \sum_n T_n \tau_{n,N+i} = \sum_n Q_{i,n} T_n,$$

c'est-à-dire sous forme matricielle, en notant T la colonne des T_n :
 $U = P T$ et $V = Q T$. L'impédance Z définie par $V = Z U$ est donc :

$$Z = Q P^{-1}.$$

Annexe 4Sur quelques propriétés des champs diffractés1. PROPRIÉTÉ 1

Hypothèses (avec les notations du paragraphe 2.3) :

- a) $\varphi(x,y)$ est pseudo-périodique en x , de période d , de coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$.
- b) φ vérifie $\Delta\varphi + k^2\varphi = 0$ dans Ω^+ .
- c) φ vérifie une condition d'onde sortante quand $y \rightarrow +\infty$.

(A4.1)

Conclusion : φ se calcule aisément dans Ω^+ dès que l'on connaît la valeur de φ^+ et de $D\varphi^+$ sur $\mathcal{P}$ (formules de Kirchhoff-Helmholtz adaptées aux réseaux).

Preuve : Considérons la solution élémentaire de l'équation :

$$\Delta g(x,y) + k^2 g(x,y) = \delta_{\mathcal{R}}(x,y) , \quad (A4.2)$$

$$\text{avec } \delta_{\mathcal{R}}(x,y) = \delta(y) \sum_{n \in \mathbb{Z}} \delta(x - nd) \exp(i\alpha_0 nd) . \quad (A4.3)$$

En imposant à $g(x,y)$ une C.O.S. pour $y \rightarrow \pm\infty$, l'unique solution de (A4.2) est [58] :

$$g(x,y) = \frac{1}{2id} \sum_{n \in \mathbb{Z}} \frac{1}{\beta_n} \exp(i\alpha_n x + i\beta_n |y|) . \quad (A4.4)$$

Soient M et P les points de coordonnées (x,y) et (x_0, y_0) , (fig. A4.1). Posons :

$$G(M,P) = g(x_0 - x, y_0 - y) . \quad (A4.5)$$

On notera que $G(M,P)$, considérée comme fonction de x et y , est pseudo-périodique avec le coefficient de pseudo-périodicité $\exp(-i\alpha_0 d)$: elle

satisfait également une C.O.S. quand $y \rightarrow \pm\infty$ et vérifie (propriétés du Laplacien) :

$$\Delta_M G(M, P) + k^2 G(M, P) = \delta_{\mathcal{R}}(x_0 - x, y_0 - y) . \quad (\text{A4.6})$$

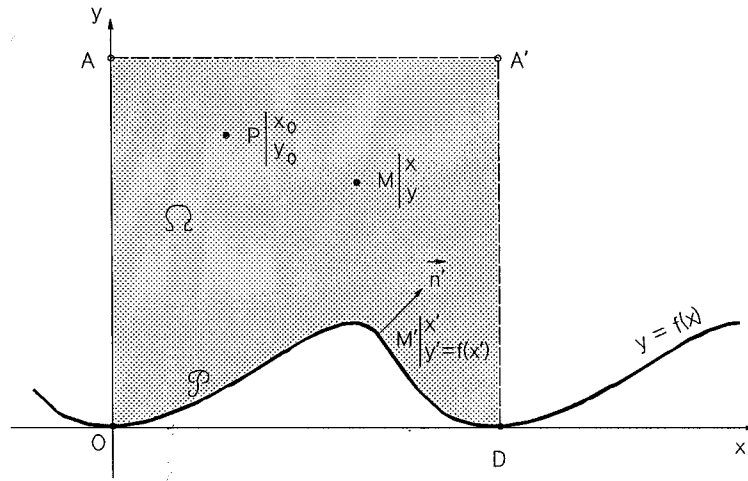


Figure A4.1. Ω^+ est le domaine $y > f(x)$, Ω^- est le domaine $y < f(x)$, Ω est le domaine grisé. Dans cette annexe, $\int_{\mathcal{P}} \xi(M') d\ell'$ désigne l'intégrale curviligne de la fonction ξ sur une période de $\mathcal{P}$; on notera que toutes les fonctions ξ considérées sont périodiques en x' .

Multipliant (A4.6) par φ , (A4.1) par G et intégrant dans le domaine grisé Ω (fig. A4.1), on obtient, après utilisation de l'identité de Green :

$$\begin{aligned} \int_{\mathcal{P} \cup (OAA'D)} \left[\varphi(M') \frac{DG(M', P)}{M'} - G(M', P) D\varphi(M') \right] d\ell' \\ = \iint_{\Omega} \varphi(M) \delta_{\mathcal{R}}(x_0 - x, y_0 - y) dx dy . \end{aligned} \quad (\text{A4.7})$$

Si l'horizontale AA' est à une ordonnée supérieure à y_0 , il est facile de constater compte tenu de (A4.3) que le second membre de (A4.7) est égal à $\varphi(P)$. A cause des pseudo-périodicités des différentes fonctions intervenant dans le premier membre, les contributions des segments OA et $A'D$ sont opposées et disparaissent. L'intégrale sur AA' figurant dans le premier membre de (A4.7) est nulle à cause des C.O.S. imposées à φ et à G . Pour le vérifier, il suffit d'écrire dans la région $y > y_0$ les fonctions φ et G sous forme de développements de Rayleigh vérifiant une condition d'onde sortante quand $y \rightarrow +\infty$ et d'explicitier l'intégrale pour en constater la nullité (cf. [59], p.305-306). L'équation (A4.7) se réduit donc à :

$$\varphi(P) = \int_{\mathcal{P}} \left[\varphi^+(M') \frac{DG(M', P)}{M'} - G(M', P) D\varphi^+(M') \right] d\ell' . \quad (\text{A4.8})$$

On constate donc qu'une fois la fonction de Green g déterminée, la connaissance de φ^+ et $D\varphi^+$ sur $\mathcal{P}$ entraîne la connaissance de φ dans Ω^+ .

2. PROPRIETE 2

En réalité, on peut montrer qu'un champ diffracté tel que φ (vérifiant a), b) et c)) est parfaitement déterminé dans Ω^+ par la valeur de φ^+ sur $\mathcal{P}$. Autrement dit, il n'est pas indispensable de connaître $D\varphi^+$; cela est lié à l'existence d'un opérateur qui donne $D\varphi^+$ (sur $\mathcal{P}$) connaissant φ^+ (sur $\mathcal{P}$), traditionnellement appelé impédance (ou admittance selon les notations et le cas de polarisation ...).

Justification : Considérons la fonction $\Gamma(M, P)$ solution de :

$$\Delta_M \Gamma(M, P) + k^2 \Gamma(M, P) = \delta_{\mathcal{R}}(x_0 - x, y_0 - y) , \quad (\text{A4.9})$$

vérifiant une C.O.S. pour $y \rightarrow \pm\infty$, et telle que si M' est un point de $\mathcal{P}$, alors $\Gamma(M', P) = 0$. Le physicien admettra sans doute l'existence et l'unicité de Γ , qui peut être interprétée (cf § 2.5) comme le champ rayonné par un réseau de "fils" disposés en $y = y_0$ et $x = x_0 + nd$ et parcourus par des courants convenablement déphasés, en présence d'un réseau infiniment conducteur situé dans Ω^- ; le champ rayonné par ces fils est polarisé $E//$ et s'annule donc sur $\mathcal{P}$.

On peut reproduire le raisonnement décrit plus haut en remplaçant $G(M, P)$ par $\Gamma(M, P)$. L'équation (A4.8) devient alors, compte tenu de la nullité de $\Gamma(M', P)$:

$$\varphi(P) = \int_{\mathcal{P}} \varphi^+(M') \frac{D\Gamma(M', P)}{M'} d\ell' , \quad (\text{A4.10})$$

ce qui montre la propriété annoncée.

3. PROPRIETE 3

Soit $u^d(x,y)$ un "champ diffracté vers le haut", c'est à dire :

- a) $u^d(x,y)$ est pseudo-périodique en x , période d , coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$.
- b) $\Delta u^d + k_1^2 u^d = 0$ dans Ω^+ ($y > f(x)$).
- c) u^d vérifie une C.O.S. quand $y \rightarrow +\infty$.

Soit $\varphi_n(x,y)$ une famille de "champs diffractés vers le haut" vérifiant les mêmes conditions a), b) et c).

Si la famille $\varphi_n(x, f(x))$ est une famille totale de $H_{\mathcal{D}}^0 = L_{\mathcal{D}}^2$ telle que :

$$u^d(x, f(x)) = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n^+(x, f(x)) , \quad (A4.11)$$

alors, pour $y > f(x)$, $\sum_{n=1}^N C_n(N) \varphi_n(x,y)$ converge uniformément vers $u^d(x,y)$ dans tout compact situé dans Ω^+ lorsque $N \rightarrow \infty$.

Justification : D'après la propriété 2, pour tout point $P \begin{vmatrix} x \\ y \end{vmatrix}$ dans Ω^+ , on peut écrire pour tout champ diffracté vers le haut u :

$$u(P) = \int_{\mathcal{P}} u^+(M') \frac{D\Gamma(M', P)}{M'} d\ell' , \quad \text{ou bien :}$$

$$u(x,y) = \int_{\mathcal{P}} u^+(x', f(x')) \gamma(x', x, y) d\ell' , \quad (A4.12)$$

avec $\gamma(x', x, y) = \frac{D\Gamma(M', P)}{M'}$;

D'après (A4.9), $\gamma(x', x, y)$ est une fonction continue de x et y dès lors que $y > f(x)$. Il en résulte que :

$$\|\gamma(x', x, y)\|_{L_{\mathcal{D}}^2} = \left[\int_{\mathcal{P}} |\gamma(x', x, y)|^2 d\ell' \right]^{1/2}$$

est également une fonction continue de x et y , et est par conséquent majorée par une constante Q sur tout compact situé dans Ω^+ . D'après (A4.12), et utilisant l'inégalité de Schwartz, on obtient :

$$\begin{aligned}
|u(x,y)| &\leq \|u^+(x', f(x'))\|_{L^2_{\mathcal{D}}} \times \|\gamma(x', x, y)\|_{L^2_{\mathcal{D}}} \\
&\leq \|u^+(x', f(x'))\|_{L^2_{\mathcal{D}}} \times Q.
\end{aligned}
\tag{A4.13}$$

Si nous prenons maintenant $u(x,y) = u^d(x,y) - \sum_{n=1}^N C_n(N) \varphi_n(x,y)$, on obtient, compte tenu de (A4.13) et (A4.11) :

$$u(x,y) \xrightarrow[N \rightarrow \infty]{} 0 \quad (\text{convergence uniforme dans le compact consid\'er\'e}),$$

$$\text{c'est-\`a-dire} \quad u^d(x,y) = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n(x,y).$$

Annexe 5

Un choix possible de familles totales dans $\mathcal{V}_1^+$ et $\mathcal{V}_2^-$.

Malgré nos efforts de rédaction, la lecture de cette annexe essentiellement mathématique reste difficile. Nous avons en effet été amenés à donner un bon nombre de définitions et à établir plusieurs résultats intermédiaires dans le but de démontrer le corollaire 3 qui conclut cette annexe.

1. DEFINITION DE QUELQUES ESPACES FONCTIONNELS

Dans cette annexe, lorsque l'on parle de fonctions périodiques ou pseudo-périodiques, la période est toujours d .

Nous allons être amenés à utiliser les espaces suivants :

Espace $\mathcal{D}$: Il s'agit de l'espace usuel des fonctions d'une variable réelle, indéfiniment dérivables, à support borné.

Topologie de $\mathcal{D}$: $\varphi_n \xrightarrow{\mathcal{D}} \varphi$ si les supports des φ_n sont contenus dans un même ensemble borné et si, quand $n \rightarrow \infty$, φ_n et ses dérivées convergent uniformément vers φ et ses dérivées.

L'espace $\mathcal{D}$ est un espace vectoriel topologique localement convexe [69]. Il n'est sans doute pas nécessaire d'approfondir cette notion car la plupart des espaces vectoriels topologiques rencontrés en analyse sont localement convexes [14].

Espace $\mathcal{D}'$: Il s'agit du dual topologique de $\mathcal{D}$. C'est l'espace usuel des distributions d'une variable.

Topologie de $\mathcal{D}'$: $T_n \xrightarrow{\mathcal{D}'} T \iff \forall \varphi \in \mathcal{D}, \quad \langle T_n, \varphi \rangle \longrightarrow \langle T, \varphi \rangle$.

L'espace $\mathcal{D}'$ est un espace vectoriel topologique localement convexe [69].

Espace $\mathcal{D}_\alpha$ ($\alpha \in \mathbb{R}$) : C'est l'ensemble des fonctions d'une variable réelle, indéfiniment dérivables, pseudo-périodiques de coefficient de pseudo-périodicité $\exp(-i\alpha d)$.

Topologie de $\mathcal{D}_\alpha$: $\varphi_n \xrightarrow{\mathcal{D}_\alpha} \varphi$ si, quand $n \rightarrow \infty$, φ_n et ses dérivées convergent uniformément vers φ et ses dérivées.

Nous admettrons que $\mathcal{D}_\alpha$ est un espace vectoriel topologique localement convexe.

Espace $\mathcal{D}'_\alpha$: C'est le dual topologique de $\mathcal{D}_\alpha$ (ensemble des fonctionnelles linéaires continues sur $\mathcal{D}_\alpha$), et on dira qu'un élément de $\mathcal{D}'_\alpha$ est une distribution de $\mathcal{D}'_\alpha$.

Topologie de $\mathcal{D}'_\alpha$: $T_n \xrightarrow{\mathcal{D}'_\alpha} T \iff \forall \varphi \in \mathcal{D}_\alpha, \quad \langle T_n, \varphi \rangle \rightarrow \langle T, \varphi \rangle$.

Nous admettrons que $\mathcal{D}'_\alpha$ est un espace vectoriel topologique localement convexe.

Exemple : une fonction ψ pseudo-périodique, de coefficient de pseudo-périodicité $\exp(i\alpha d)$, sommable sur $[0, d]$, peut être considérée comme un élément de $\mathcal{D}'_\alpha$ en posant :

$$\forall \varphi \in \mathcal{D}_\alpha, \quad \langle \psi, \varphi \rangle = \int_a^{a+d} \psi(x) \varphi(x) dx.$$

Exemple : Soit $x_0 \in [0, d[$, nous définissons la distribution T_{x_0} de $\mathcal{D}'_\alpha$ par :

$$\forall \varphi \in \mathcal{D}_\alpha, \quad \langle T_{x_0}, \varphi \rangle = \varphi(x_0).$$

T_{x_0} joue dans $\mathcal{D}'_\alpha$ le même rôle que δ_{x_0} dans $\mathcal{D}'$ (δ_{x_0} étant la distribution de Dirac notée $\delta(x - x_0)$ en physique).

Espace $\mathcal{R}_\alpha$: C'est l'ensemble des fonctions de deux variables réelles x et y , indéfiniment dérivables, pseudo-périodiques en x (coefficient de pseudo-périodicité $e^{-i\alpha d}$), à support borné en y .

Cet espace est noté $\mathcal{R}$ dans l'appendice A de [58].

Topologie de $\mathcal{R}_\alpha$: $\varphi_n \xrightarrow{\mathcal{R}_\alpha} \varphi$ si les supports des φ_n sont contenus dans un même ensemble borné en y et si, quand $n \rightarrow \infty$, φ_n et ses dérivées convergent uniformément vers φ et ses dérivées.

Espace $\mathcal{R}'_\alpha$: C'est le dual topologique de $\mathcal{R}_\alpha$ (ensemble des fonctionnelles linéaires continues sur $\mathcal{R}_\alpha$). Cet espace est noté $\mathcal{R}'$ dans l'appendice A de [58] où il est étudié plus en détail.

Topologie de $\mathcal{R}'_\alpha$: la notion de convergence est identique à celles de $\mathcal{D}'$ et $\mathcal{D}'_\alpha$.

Exemple : si $\psi(x,y)$ est pseudo-périodique en x (coefficient de pseudo-périodicité $e^{i\alpha d}$) et localement sommable, ψ peut être considérée comme un élément de $\mathcal{R}'_\alpha$ en posant :

$$\forall \varphi \in \mathcal{R}_\alpha, \quad \langle \psi, \varphi \rangle = \int_{x=a}^{a+d} \int_{y=-\infty}^{+\infty} \psi(x,y) \varphi(x,y) dx dy.$$

Exemple : Soit s une distribution de $\mathcal{D}'_\alpha$ et $\mathcal{P}$ une courbe définie par $y = f(x)$, avec f périodique et indéfiniment dérivable. On définit la distribution $s \delta_{\mathcal{P}}$ de $\mathcal{R}'_\alpha$ par :

$$\forall \varphi \in \mathcal{R}_\alpha, \quad \langle s \delta_{\mathcal{P}}, \varphi \rangle = \langle s(x), \varphi(x, f(x)) \sqrt{1 + f'(x)^2} \rangle.$$

On notera que si ψ est une fonction pseudo-périodique (coefficient de pseudo-périodicité $e^{i\alpha d}$), sommable sur $[0,d]$, elle est une distribution de $\mathcal{D}'_\alpha$ et :

$$\begin{aligned} \forall \varphi \in \mathcal{R}_\alpha, \quad \langle \psi \delta_{\mathcal{P}}, \varphi \rangle &= \langle \psi(x), \varphi(x, f(x)) \sqrt{1 + f'(x)^2} \rangle \\ &= \int_a^{a+d} \psi(x) \varphi(x, f(x)) \sqrt{1 + f'(x)^2} dx \end{aligned}$$

où $\sqrt{1 + f'(x)^2} dx$ représente l'élément d'abscisse curviligne sur $\mathcal{P}$.

2. THEOREME DE HAHN-BANACH ET COROLLAIRES

Nous nous plaçons pour ce qui suit sur le corps des complexes. Commençons par rappeler un des corollaires du théorème de Hahn-Banach ([70], page 275) :

Corollaire 1 : Soit $\mathcal{E}$ un espace vectoriel topologique localement convexe. Soit $\mathcal{F}$ un sous-espace vectoriel de $\mathcal{E}$. Soit f une forme linéaire continue sur $\mathcal{F}$. Alors il existe une forme linéaire $\tilde{f}$ continue sur $\mathcal{E}$ prolongeant f .

On en déduit le corollaire suivant :

Corollaire 2 : Soit $\mathcal{E}$ un espace vectoriel topologique localement convexe. Soit $\mathcal{F}$ un sous-espace vectoriel fermé de $\mathcal{E}$, différent de $\mathcal{E}$. Alors il existe une forme linéaire continue et non nulle sur $\mathcal{E}$, s'annulant sur $\mathcal{F}$.

Preuve : Soit $v \in \mathcal{E}$, $v \notin \mathcal{F}$, et soit $\mathcal{V} = \{\lambda v, \lambda \in \mathbb{C}\}$. Soit $\mathcal{G} = \mathcal{F} \oplus \mathcal{V}$ la somme directe de $\mathcal{F}$ et $\mathcal{V}$. Tout élément w de $\mathcal{G}$ se décompose de façon unique sous la forme $w = u + \lambda v$, avec $u \in \mathcal{F}$ et $\lambda \in \mathbb{C}$. Soit g la forme linéaire définie sur $\mathcal{G}$ par : $\forall w \in \mathcal{G}$, $g(w) = \lambda$, où λ est obtenu par la décomposition précédente.

Montrons que la forme g est continue. Soit une suite d'éléments de $\mathcal{G}$: $w_n = u_n + \lambda_n v$ tendant vers zéro. Si $g(w_n) = \lambda_n$ ne tendait pas vers zéro, pour tout $\varepsilon > 0$, on pourrait extraire de w_n une suite $w_m = u_m + \lambda_m v$ tendant vers zéro et telle que $|\lambda_m| > \varepsilon$; on aurait alors $u_m/\lambda_m \rightarrow -v$, ce qui est impossible car $\mathcal{F}$ étant fermé, $u_m/\lambda_m \in \mathcal{F}$ alors que $v \notin \mathcal{F}$.

D'après le corollaire 1, il existe une forme linéaire $\tilde{g}$ continue sur $\mathcal{E}$ prolongeant g . Cette forme linéaire $\tilde{g}$ satisfait aux conditions de l'énoncé du corollaire 2.

3. FAMILLES TOTALES DANS $\mathcal{D}_\alpha$ ET $\mathcal{D}'_\alpha$.

Théorème 1 : Soit φ_n une famille de fonctions de $\mathcal{D}_\alpha$. Les deux propriétés suivantes sont équivalentes :

Propriété 1 : $\forall T \in \mathcal{D}'_\alpha, (\forall n, \langle T, \varphi_n \rangle = 0 \implies T = 0)$.

Propriété 2 : $\forall \varphi \in \mathcal{D}_\alpha$, il existe une suite de coefficients complexes $C_n(N)$

telle que :

$$\varphi = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n.$$

Preuve : Montrons que la propriété 2 entraîne la propriété 1. Supposons la propriété 2 vérifiée. Soit $T \in \mathcal{D}'_\alpha$ et supposons que :

$$\forall n, \langle T, \varphi_n \rangle = 0. \tag{A5.1}$$

Montrons que $T = 0$. D'après la propriété 2,

$$\forall \varphi \in \mathcal{D}_\alpha, \quad \langle T, \varphi \rangle = \langle T, \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n \rangle.$$

La continuité de T permet d'écrire :

$$\begin{aligned} \langle T, \varphi \rangle &= \lim_{N \rightarrow \infty} \langle T, \sum_{n=1}^N C_n(N) \varphi_n \rangle \\ &= 0 \quad \text{d'après (A5.1),} \end{aligned} \quad \text{C.Q.F.D.}$$

Montrons que la propriété 1 entraîne la propriété 2. Soit $\mathcal{F}$ l'ensemble des combinaisons linéaires des φ_n et des limites de ces combinaisons. $\mathcal{F}$ est un sous-espace vectoriel fermé de $\mathcal{D}_\alpha$. La propriété 1 s'énonce : toute distribution $T \in \mathcal{D}'_\alpha$ (forme linéaire continue sur $\mathcal{D}_\alpha$) s'annulant sur $\mathcal{F}$ est nulle sur $\mathcal{D}_\alpha$. Puisque $\mathcal{D}_\alpha$ est un espace vectoriel topologique localement convexe, il résulte du corollaire 2 que $\mathcal{F} = \mathcal{D}_\alpha$ (en effet, si $\mathcal{F}$ était différent de $\mathcal{D}_\alpha$ le corollaire 2 serait mis en contradiction). La propriété 2 est donc vérifiée.

Définition : Nous dirons qu'une famille φ_n de fonctions de $\mathcal{D}_\alpha$ est totale dans $\mathcal{D}_\alpha$ si elle possède les propriétés 1 ou 2.

On notera l'analogie avec la définition d'une famille totale dans un espace de Hilbert ; mais ici les produits scalaires sont remplacés par des "crochets de distributions".

Théorème 2 : Soit T_n une famille de distributions de $\mathcal{D}'_\alpha$. Les deux propriétés suivantes sont équivalentes :

Propriété 3 : $\forall \varphi \in \mathcal{D}_\alpha, \quad (\forall n, \langle T_n, \varphi \rangle = 0 \implies \varphi = 0).$

Propriété 4 : $\forall T \in \mathcal{D}'_\alpha$, il existe une suite de coefficients complexes $C_n(N)$ telle que :

$$T = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) T_n.$$

Preuve : On utilisera ici le fait que le dual topologique de $\mathcal{D}'_\alpha$ est $\mathcal{D}_\alpha$. Il en résulte que pour tout $\varphi \in \mathcal{D}_\alpha$,

$$\varphi = 0 \iff \forall T \in \mathcal{D}'_\alpha, \quad \langle T, \varphi \rangle = 0. \quad (\text{A5.2})$$

Montrons que la propriété 4 entraîne la propriété 3. Supposons donc que :

$$\forall T \in \mathcal{D}'_{\alpha}, \quad T = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) T_n.$$

Soit $\varphi \in \mathcal{D}_{\alpha}$ vérifiant $\forall n, \langle T_n, \varphi \rangle = 0$. Il résulte immédiatement de la notion de convergence dans $\mathcal{D}'_{\alpha}$ que $\forall T \in \mathcal{D}'_{\alpha}, \langle T, \varphi \rangle = 0$. Nous en déduisons par (A5.2) que $\varphi = 0$.

Montrons que la propriété 3 entraîne la propriété 4. Soit $\mathcal{F}$ l'ensemble des combinaisons linéaires des T_n et des limites de ces combinaisons. $\mathcal{F}$ est un sous-espace vectoriel fermé de $\mathcal{D}'_{\alpha}$. La propriété 3 s'énonce : toute fonction $\varphi \in \mathcal{D}_{\alpha}$ (forme linéaire continue sur $\mathcal{D}'_{\alpha}$) s'annulant sur $\mathcal{F}$ est nulle sur $\mathcal{D}'_{\alpha}$. Puisque $\mathcal{D}'_{\alpha}$ est un espace vectoriel topologique localement convexe, il résulte du corollaire 2 que $\mathcal{F} = \mathcal{D}'_{\alpha}$, ce qui entraîne la propriété 4.

Définition : Nous dirons qu'une famille de distributions T_n de $\mathcal{D}'_{\alpha}$ est totale dans $\mathcal{D}'_{\alpha}$ si elle possède les propriétés 3 ou 4.

Théorème 3 : Soit X un ensemble de réels dense sur $[0, d]$. La famille de distributions T_{x_0} de $\mathcal{D}'_{\alpha}$, $x_0 \in X$, (T_{x_0} est définie plus haut au § 1) constitue une famille totale dans $\mathcal{D}'_{\alpha}$.

Preuve : Soit $\varphi \in \mathcal{D}_{\alpha}$ telle que :

$$\forall x_0 \in X, \quad \langle T_{x_0}, \varphi \rangle = 0.$$

Compte tenu de la définition de T_{x_0} , cela s'écrit aussi :

$$\forall x_0 \in X, \quad \varphi(x_0) = 0.$$

L'ensemble X étant dense, et φ étant continue, il en résulte que $\varphi = 0$. La famille T_{x_0} vérifie donc la propriété 3 ; elle est totale dans $\mathcal{D}'_{\alpha}$.

4. UN THEOREME DE RECIPROCITE

Théorème 4.

Hypothèses (il est conseillé de se reporter à la figure A5.1) :

Soit $\mathcal{P}$ une courbe définie par $y = f(x)$, avec f périodique de période d , indéfiniment dérivable.

Soit $\tilde{\mathcal{P}}$ une courbe définie par $y = \tilde{f}(x)$, avec $\tilde{f}$ périodique de période d , indéfiniment dérivable.

Nous supposons que les courbes $\mathcal{P}$ et $\tilde{\mathcal{P}}$ ne se rencontrent pas ; supposons par exemple que $\tilde{\mathcal{P}}$ soit situé "en dessous" de $\mathcal{P}$, c'est-à-dire que $\forall x, \tilde{f}(x) < f(x)$.

Soit s une distribution de $\mathcal{D}'_{-\alpha}$.

Soit $\tilde{s}$ une distribution de $\mathcal{D}'_{\alpha}$.

Soit u la solution unique dans $\mathcal{R}'_{\alpha}$ de $(\Delta + k^2)u = s \delta_{\mathcal{P}}$, vérifiant une condition d'onde sortante quand $y \rightarrow \pm\infty$.

Soit $\tilde{u}$ la solution unique dans $\mathcal{R}'_{\alpha}$ de $(\Delta + k^2)\tilde{u} = \tilde{s} \delta_{\tilde{\mathcal{P}}}$, vérifiant une condition d'onde sortante quand $y \rightarrow \pm\infty$.

Soit φ la restriction de $\tilde{u}$ à $\mathcal{P}$.

Soit $\tilde{\varphi}$ la restriction de u à $\tilde{\mathcal{P}}$.

Conclusion : $\langle s, \varphi \rangle = \langle \tilde{s}, \tilde{\varphi} \rangle$.

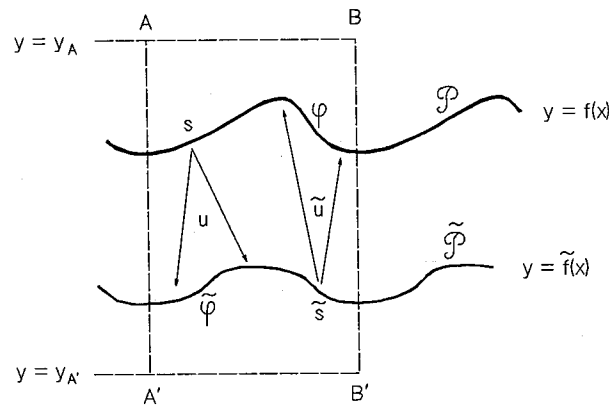


Figure A5.1, servant de "support visuel" pour le théorème 4 ;

$$AB = A'B' = d.$$

Preuve : Multiplions par $\tilde{u}$ l'équation $\Delta u + k^2 u = s \delta_{\mathcal{P}}$ (cette opération a un sens car $\tilde{u}$ est indéfiniment dérivable dans un voisinage contenant le support $\mathcal{P}$ de $s \delta_{\mathcal{P}}$).

Multiplions de même par u l'équation $\Delta \tilde{u} + k^2 \tilde{u} = \tilde{s} \delta_{\tilde{\mathcal{P}}}$.

Retranchons membre à membre, il vient :

$$\tilde{u} \Delta u - u \Delta \tilde{u} = \tilde{u} s \delta_{\mathcal{P}} - u \tilde{s} \delta_{\tilde{\mathcal{P}}}.$$

Chaque membre de cette égalité représente une distribution périodique (élément de $\mathcal{R}'_0$) dont le support se réduit aux profils $\mathcal{P}$ et $\tilde{\mathcal{P}}$. Il est donc possible de les faire agir sur une fonction test ψ de $\mathcal{R}_0$ indépendante de x et prenant la valeur 1 dans le rectangle $ABB'A'$ (Fig.A5.1) :

$$\langle \tilde{u} \Delta u - u \Delta \tilde{u}, \psi \rangle = \langle \tilde{u} s \delta_{\mathcal{P}} - u \tilde{s} \delta_{\tilde{\mathcal{P}}}, \psi \rangle.$$

Montrons que le membre de gauche est nul :

$$\begin{aligned} \langle \tilde{u} \Delta u - u \Delta \tilde{u}, \psi \rangle &= \langle \operatorname{div}(\tilde{u} \overrightarrow{\operatorname{grad} u} - u \overrightarrow{\operatorname{grad} \tilde{u}}), \psi \rangle \\ &= \sum_{i=1}^2 \left\langle \frac{\partial}{\partial x_i} \left(\tilde{u} \frac{\partial u}{\partial x_i} - u \frac{\partial \tilde{u}}{\partial x_i} \right), \psi \right\rangle \\ &= - \sum_{i=1}^2 \left\langle \tilde{u} \frac{\partial u}{\partial x_i} - u \frac{\partial \tilde{u}}{\partial x_i}, \frac{\partial \psi}{\partial x_i} \right\rangle \end{aligned}$$

soit, en revenant aux variables x et y ,

$$= - \left\langle \tilde{u} \frac{\partial u}{\partial y} - u \frac{\partial \tilde{u}}{\partial y}, \frac{\partial \psi}{\partial y} \right\rangle$$

puisque ψ est indépendante de x . En utilisant de plus le fait que ψ est constante dans le rectangle $ABB'A'$ et que $\tilde{u} \frac{\partial u}{\partial y} - u \frac{\partial \tilde{u}}{\partial y}$ est indéfiniment dérivable pour $y < y_A$, et pour $y > y_A$, il vient en posant $\mathcal{Y} =]-\infty, y_A[\cup]y_A, +\infty[$:

$$\begin{aligned} \langle \tilde{u} \Delta u - u \Delta \tilde{u}, \psi \rangle &= - \int_{x=0}^d \int_{y \in \mathcal{Y}} \left(\tilde{u} \frac{\partial u}{\partial y} - u \frac{\partial \tilde{u}}{\partial y} \right) \frac{\partial \psi}{\partial y} dx dy \\ &= - \int_{y \in \mathcal{Y}} \frac{\partial \psi}{\partial y} \left(\int_0^d \left(\tilde{u} \frac{\partial u}{\partial y} - u \frac{\partial \tilde{u}}{\partial y} \right) dx \right) dy . \end{aligned}$$

Pour $y > \max(f(x))$ et $y < \min(\tilde{f}(x))$, u et $\tilde{u}$ peuvent s'écrire sous la forme :

$$u = \sum_{n \in \mathbb{Z}} u_n \exp(i \alpha_n x + i \beta_n y)$$

$$\tilde{u} = \sum_{n \in \mathbb{Z}} \tilde{u}_n \exp(i \tilde{\alpha}_n x + i \tilde{\beta}_n y)$$

avec $K = 2\pi/d$, $\alpha_n = -\alpha + nK$, $\tilde{\alpha}_n = \alpha + nK$, $\beta_n = \sqrt{k^2 - \alpha_n^2}$, $\tilde{\beta}_n = \sqrt{k^2 - \tilde{\alpha}_n^2}$,

les racines étant déterminées par $\beta_n > 0$ ou $\text{Im}(\beta_n) > 0$, et même chose pour $\tilde{\beta}_n$.

On remarque que $\alpha_{-n} = -\tilde{\alpha}_n$ et $\beta_{-n} = \tilde{\beta}_n$, ce qui permet d'écrire u et $\tilde{u}$ de la façon suivante :

$$u = \sum_{m \in \mathbb{Z}} u_{-m} \exp(-i \tilde{\alpha}_m x + i \tilde{\beta}_m y) \quad (\text{A5.3})$$

$$\tilde{u} = \sum_{n \in \mathbb{Z}} \tilde{u}_n \exp(i \tilde{\alpha}_n x + i \tilde{\beta}_n y). \quad (\text{A5.4})$$

Il suffit de calculer l'intégrale $\int_0^d \left(\tilde{u} \frac{\partial u}{\partial y} - u \frac{\partial \tilde{u}}{\partial y} \right) dx$ au moyen des développements (A5.3, A5.4) pour constater qu'elle est nulle (cf. [59], page 306, lemme 2'),

d'où finalement $\langle \tilde{u} s \delta_{\mathcal{P}} - u \tilde{s} \delta_{\tilde{\mathcal{P}}}, \psi \rangle = 0$,

qui s'écrit aussi $\langle s \delta_{\mathcal{P}}, \tilde{u} \psi \rangle = \langle \tilde{s} \delta_{\tilde{\mathcal{P}}}, u \psi \rangle$,

ou bien $\langle s, \varphi \rangle = \langle \tilde{s}, \tilde{\varphi} \rangle$.

5. OBTENTION DE FAMILLES TOTALES DANS $\mathcal{V}_1^+$ OU $\mathcal{V}_2^-$.

Théorème 5.

Hypothèses : $\mathcal{P}$ et $\tilde{\mathcal{P}}$ sont des courbes définies comme au théorème 4.

Soit $\tilde{s}_n$ une famille totale de distributions de $\mathcal{D}'_\alpha$.

Soit $\tilde{u}_n$ la solution unique dans $\mathcal{R}'_\alpha$ de $(\Delta + k^2) \tilde{u}_n = \tilde{s}_n \delta_{\tilde{\mathcal{P}}}$, vérifiant une C.O.S. quand $y \rightarrow \pm\infty$, et soit φ_n la restriction de $\tilde{u}_n$ à $\mathcal{P}$.

Conclusion : Les φ_n forment une famille totale de $\mathcal{D}_{-\alpha}$.

Preuve : Soit s une distribution de $\mathcal{D}'_\alpha$; nous considérons la fonction $\tilde{\varphi}$, élément de $\mathcal{D}_\alpha$, définie comme au théorème 4 (voir aussi figure A5.1).

Du théorème 4, il résulte que :

$$\forall n, \langle s, \varphi_n \rangle = \langle \tilde{s}_n, \tilde{\varphi} \rangle.$$

Les $\tilde{s}_n$ formant une famille totale de distributions de $\mathcal{D}'_\alpha$, nous pouvons écrire (propriété 3) :

$$\forall n, \langle \tilde{s}_n, \tilde{\varphi} \rangle = 0 \implies \tilde{\varphi} = 0,$$

$$\Leftrightarrow \forall n, \langle s, \varphi_n \rangle = 0 \implies \tilde{\varphi} = 0.$$

Montrons que $\tilde{\varphi} = 0$ entraîne $s = 0$. La fonction $\tilde{\varphi}$ étant la restriction à $\tilde{\mathcal{P}}$ de $u(x,y)$ vérifiant l'équation de Helmholtz $\Delta u + k^2 u = 0$ dans le domaine $y < f(x)$, et une condition d'onde sortante quand $y \rightarrow -\infty$, le fait que u s'annule sur $\tilde{\mathcal{P}}$ entraîne la nullité de u dans la région $y < f(x)$ (propriété de l'équation de Helmholtz). Considérons, pour ε réel non nul, la fonction $u_\varepsilon(x) = u(x, f(x) + \varepsilon)$ qui appartient à $\mathcal{D}_\alpha$, et qui peut donc aussi être considérée comme un élément de $\mathcal{D}'_\alpha$. Du fait que u vérifie $\Delta u + k^2 u = s \delta_{\tilde{\mathcal{P}}}$ ($s \in \mathcal{D}'_\alpha$), nous déduisons [10] que u_ε tend, quand ε tend vers 0, et indépendamment du signe de ε , vers une distribution $u_0 \in \mathcal{D}'_\alpha$. Le fait que u soit nul dans la région $y < f(x)$ entraîne que $u_0 = 0$. Autrement dit, $u(x,y)$ admet, dans la région $y > f(x)$, une valeur au bord sur $\mathcal{P}$ presque partout nulle. Alléguant à nouveau le fait que $\Delta u + k^2 u = 0$ pour $y > f(x)$ et que u vérifie une C.O.S. quand $y \rightarrow +\infty$, on en déduit que $u = 0$ pour $y > f(x)$. Le fait que u soit nulle pour $y > f(x)$ et pour $y < f(x)$,

entraîne que $s = 0$.

Nous pouvons donc écrire pour tout $s \in \mathcal{D}'_{-\alpha}$:

$$\forall n, \langle s, \varphi_n \rangle = 0 \Rightarrow s = 0.$$

La famille φ_n vérifie donc la propriété 1. Elle est totale dans $\mathcal{D}_{-\alpha}$.

Théorème 6.

Soit φ_n une famille totale dans $\mathcal{D}_{-\alpha}$. Soit $\mathcal{P}$ un "profil" défini par $y = f(x)$, f étant une fonction périodique indéfiniment dérivable. Alors φ_n constitue une famille totale dans $H^1_{\mathcal{P}}$.

Preuve : On utilise ici le fait que $\mathcal{D}_{-\alpha}$ est dense dans l'ensemble des fonctions pseudo-périodiques de $H^1_{\mathcal{P}}$ (voir par exemple [19], chapitre IV, p.885 et chapitre VII, p.1313, où il est question d'espaces légèrement différents : $\mathcal{D}(\mathbb{R}^n)$ et $H^s(\mathbb{R}^n)$). On peut aussi en faire une démonstration directe en écrivant toute fonction pseudo-périodique de $H^1_{\mathcal{P}}$ sous la forme d'une série de Fourier généralisée). Si φ_n est une famille totale de $\mathcal{D}_{-\alpha}$, il en résulte que $\forall u \in H^1_{\mathcal{P}}$, il existe une suite de coefficients complexes $C_n(N)$

telle que $u = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n$, ce qui s'écrit également (§ 2.7) :

$$\int_{\mathcal{P}} \left| u - \sum_{n=1}^N C_n(N) \varphi_n \right|^2 d\ell + \int_{\mathcal{P}} \left| \frac{du}{d\ell} - \sum_{n=1}^N C_n(N) \frac{d\varphi_n}{d\ell} \right|^2 d\ell \xrightarrow{N \rightarrow \infty} 0$$

Cela montre que φ_n est totale dans $H^1_{\mathcal{P}}$.

Théorème 7.

Soit $\mathcal{P}$ un "profil" défini par $y = f(x)$, f étant une fonction périodique indéfiniment dérivable.

Soit $\tilde{u}_n$ une famille de distributions de $\mathcal{R}_{\alpha}$ définie comme au Théorème 5, et $D \varphi_n$ leurs dérivées normales sur $\mathcal{P}$.

Alors les $D \varphi_n$ constituent une famille totale dans $H^0_{\mathcal{P}}$ (c'est à dire $L^2_{\mathcal{P}}$).

En clair, les dérivées normales sur $\mathcal{P}$ de champs rayonnés par des sources

formant une famille totale de $\mathcal{D}'_\alpha$ forment elles-mêmes une famille totale de $H^0_\mathcal{P}$.

Preuve : Puisque les $\tilde{u}_n$ sont indéfiniment dérivables dans un voisinage de $\mathcal{P}$, les $D \varphi_n$ appartiennent à $\mathcal{D}'_\alpha$. Nous utilisons ci-après les opérateurs impédance ($\mathbb{Z}$) et admittance (Ψ) tels qu'ils sont définis dans [61].

Soit $v \in H^0_\mathcal{P}$; on en déduit [61] que $u = \mathbb{Z} v$ appartient à $H^1_\mathcal{P}$.

La famille φ_n définie comme au Théorème 5 est totale dans $H^1_\mathcal{P}$ (Théorème 6).

On peut donc trouver une suite de coefficients complexes $C_n(N)$ telle que :

$$u = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) \varphi_n .$$

Or on sait [61] que $\Psi u = v$, et que $\Psi \varphi_n = D \varphi_n$. On en déduit ainsi, utilisant la continuité de Ψ [61] :

$$v = \lim_{N \rightarrow \infty} \sum_{n=1}^N C_n(N) D \varphi_n ,$$

ce qui traduit le fait que les $D \varphi_n$ forment une famille totale dans $H^0_\mathcal{P}$.

On peut donc en particulier énoncer le résultat suivant que nous avons utilisé au chapitre 2, § 2.5, et qui justifie cette annexe :

Corollaire 3.

Hypothèses : Soit X un ensemble de réels dense sur $[0, d]$.

Soit $\mathcal{P}$ une courbe définie par $y = f(x)$, f périodique de période d , indéfiniment dérivable.

Soit $\tilde{\mathcal{P}}$ une courbe définie par $y = \tilde{f}(x)$, $\tilde{f}$ périodique de période d , indéfiniment dérivable, avec $\forall x, \tilde{f}(x) < f(x)$.

Pour $x_0 \in X$:

Soit $\tilde{u}_{x_0}$ la solution de $(\Delta + k_1^2) \tilde{u}_{x_0} = T_{x_0} \delta_{\tilde{\mathcal{P}}}$, vérifiant une C.O.S. quand $y \rightarrow \pm\infty$.

Soit φ_{x_0} la restriction de $\tilde{u}_{x_0}$ à $\mathcal{P}$.

Soit $D \varphi_{x_0}$ la dérivée normale de $\tilde{u}_{x_0}$ sur $\mathcal{P}$.

Conclusion : La famille de colonnes $\begin{vmatrix} \varphi_{x_0} \\ D\varphi_{x_0} \end{vmatrix}$ est totale dans $\mathcal{V}_1^+$.

Preuve : D'après le Théorème 3, T_{x_0} est une famille totale de $\mathcal{D}'_\alpha$. D'après le Théorème 5, φ_{x_0} est une famille totale de $\mathcal{D}_\alpha$, et d'après le Théorème 6, elle est aussi une famille totale de $H^1_\mathcal{P}$. D'après le Théorème 7, $D\varphi_{x_0}$ est une famille totale de $H^0_\mathcal{P}$. D'où le résultat annoncé.

Remarque : pour obtenir une famille totale dans $\mathcal{V}_2^-$, il suffit de remplacer k_1 par k_2 et $\tilde{\mathcal{P}}$ par une courbe située au-dessus de $\mathcal{P}$.

Annexe 6

Permittivités associées à des milieux anisotropes "passifs"

Comme dans le reste du mémoire, nous considérons uniquement des régimes harmoniques ; la perméabilité est supposée être partout égale à celle du vide (μ_0).

On sait que si un milieu isotrope est sans pertes, sa permittivité est un nombre réel. Plus généralement, le fait qu'un milieu isotrope soit dissipatif impose que la partie imaginaire de sa permittivité soit un nombre positif. L'objet de cette annexe est d'énoncer les conditions que doivent vérifier les permittivités des milieux anisotropes passifs (avec ou sans pertes).

Nous supposons que les champs complexes sont liés par la relation diélectrique (dans le trièdre orthonormé de référence) :

$$\vec{D} = \epsilon_0 [\epsilon] \vec{E}. \quad (\text{A6.1})$$

Des équations de Maxwell $\text{rot } \vec{E} = i\omega \mu_0 \vec{H}$ et $\text{rot } \vec{H} = -i\omega \epsilon_0 [\epsilon] \vec{E}$, nous pouvons déduire la divergence du vecteur de Poynting complexe ($\vec{P} = \frac{1}{2} \vec{E} \wedge \vec{H}$) :

$$\begin{aligned} 2 \text{div } \vec{P} &= \text{rot } \vec{E} \cdot \vec{H} - \vec{E} \cdot \text{rot } \vec{H} \\ &= i\omega \mu_0 \vec{H} \cdot \vec{H} - i\omega \epsilon_0 \vec{E} \cdot [\epsilon] \vec{E}. \end{aligned} \quad (\text{A6.2})$$

Le milieu étant passif, le flux de la partie réelle de $\vec{P}$ à travers toute surface fermée doit être négatif ou nul, ce qui peut s'écrire :

$$\text{Re} (\text{div } \vec{P}) \leq 0. \quad (\text{A6.3})$$

D'après (A6.2), cette inégalité peut s'écrire aussi :

$$\text{Im} (\vec{E} \cdot [\epsilon] \vec{E}) \leq 0. \quad (\text{A6.4})$$

En notant ε_{jk} les éléments de la matrice $[\varepsilon]$, et E_j les composantes de $\vec{E}$, (A6.4) prend la forme :

$$\sum_{j=1}^3 \sum_{k=1}^3 i (\bar{\varepsilon}_{jk} - \varepsilon_{kj}) E_j \bar{E}_k \geq 0. \quad (A6.5)$$

Posons : $h_{jk} = i (\bar{\varepsilon}_{jk} - \varepsilon_{kj})$. (A6.6)

Remarquant que $\bar{h}_{jk} = h_{kj}$, la forme :

$$h(\vec{E}, \vec{E}) \stackrel{\text{def}}{=} \sum_{j=1}^3 \sum_{k=1}^3 h_{jk} E_j \bar{E}_k \quad (A6.7)$$

est une forme hermitique [21]. Le fait que le milieu soit passif peut donc se traduire par :

$$h(\vec{E}, \vec{E}) = \sum_{j=1}^3 \sum_{k=1}^3 h_{jk} E_j \bar{E}_k \geq 0. \quad (A6.8)$$

Dans les cas de milieux anisotropes sans pertes, toutes les inégalités précédentes se réduisent à des égalités. L'équation (A6.5) montre qu'alors la permittivité du matériau est hermitique : $\bar{\varepsilon}_{jk} = \varepsilon_{kj}$.

Dans le cas de milieux anisotropes possédant des pertes, les inégalités précédentes deviennent des inégalités strictes (pour tout champ $\vec{E}$ non nul). La forme hermitique $h(\vec{E}, \vec{E})$ est alors définie positive. On démontre ([21], p.339) que cela est équivalent aux trois conditions suivantes, portant sur les mineurs principaux de la matrice $[h]$ d'éléments h_{jk} :

- a) $h_{11} > 0$,
- b) $\begin{vmatrix} h_{11} & h_{12} \\ h_{21} & h_{22} \end{vmatrix} > 0$,
- c) $\det [h] > 0$.

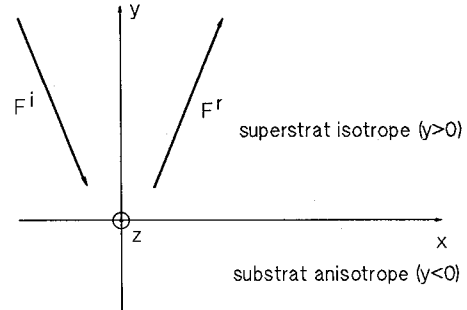
Ces conditions peuvent être testées très simplement sur ordinateur. On remarquera que la première condition ($h_{11} > 0$) est équivalente à $\text{Im}(\varepsilon_{11}) > 0$. Par permutation des vecteurs de la base utilisée, on en déduit une condition nécessaire (mais non suffisante) pour que le milieu possède des pertes : tous les éléments diagonaux de la permittivité doivent avoir une partie imaginaire positive.

Annexe 7.Sur la condition d'ondes sortantes en milieu anisotrope

Propriété : Parmi les quatre ondes planes de vecteur d'onde $\vec{k}_j = \alpha \vec{e}_x + \beta_j \vec{e}_y$ (α réel donné) susceptibles d'exister dans un milieu anisotrope homogène, deux d'entre elles vérifient une C.O.S. vers le haut et les deux autres une C.O.S. vers le bas.

Nous ne connaissons pas de démonstration vraiment rigoureuse de cette propriété. Mais le Lecteur se laissera sans doute convaincre par ce qui suit de la véracité de la propriété annoncée.

Figure A7.1. F^i et F^r représentent les champs incident et réfléchi



Supposons que le milieu anisotrope remplisse le demi-espace $y < 0$, et que le demi-espace $y > 0$ soit rempli (pour simplifier) d'un milieu isotrope (Fig.A7.1). Imaginons que cette structure soit éclairée par une onde plane provenant des y positifs, dont le vecteur d'onde, situé dans le plan (x,y) , possède sur l'axe des x une composante α . Nous avons vu (§ 3.2) que le champ électromagnétique est périodique en x et peut être représenté dans chaque zone homogène $y > 0$ et $y < 0$ par une superposition de quatre ondes planes. Nous avons montré également (§ 3.2) que le champ est parfaitement décrit dans chacun de ces demi-espaces par une colonne :

$$\text{pour } y > 0 : F^h(x,y) = e^{i\alpha x} \mathcal{F}^h(y) = \sum_{j=1}^4 c_j^h e^{i\alpha x + i\beta_j^h y} \mathcal{F}_j^h, \quad (\text{A7.1})$$

$$\text{pour } y < 0 : F^b(x, y) = e^{i\alpha x} \mathcal{F}^b(y) = \sum_{j=1}^4 C_j^b e^{i\alpha x + i\beta_j^b y} \mathcal{F}_j^b. \quad (\text{A7.2})$$

Le superstrat étant isotrope, les quatre valeurs β_j^h dégénèrent en deux valeurs réelles opposées $\pm\beta^h$ (§ 3.3.3) :

$$F^h(x, y) = e^{i\alpha x} \left[(C_1^h \mathcal{F}_1^h + C_2^h \mathcal{F}_2^h) e^{-i\beta^h y} + (C_3^h \mathcal{F}_3^h + C_4^h \mathcal{F}_4^h) e^{i\beta^h y} \right] \quad (\text{A7.3})$$

où les colonnes $\mathcal{F}_1^h$ et $\mathcal{F}_2^h$ sont associées aux deux cas fondamentaux de polarisation de l'onde incidente, et $\mathcal{F}_3^h$ et $\mathcal{F}_4^h$ sont associées aux deux cas fondamentaux de polarisation de l'onde réfléchie (§ 3.3.3). L'onde incidente étant connue, le problème consiste à déterminer les constantes C_3^h et C_4^h (onde réfléchie) et les quatre constantes C_j^b (onde transmise). Pour cela, on écrira la continuité du champ en $y = 0$, soit $F^h(x, 0) = F^b(x, 0)$, obtenant ainsi un système de quatre équations pour les six inconnues C_3^h , C_4^h et C_j^b . Admettons l'existence et l'unicité de la solution. Il est alors nécessaire que deux des constantes C_j^b soient nulles.

Ainsi, l'existence et l'unicité de la solution entraînent qu'il existe deux "catégories" distinctes d'ondes planes dans un milieu donné. Nous demandons au Lecteur d'admettre (ou de se persuader) que les deux ondes planes associées aux constantes C_j^b devant s'annuler vérifient une C.O.S. vers le haut (telle qu'elle a été définie au § 3.3.1) et que les deux autres (celles qui sont par conséquent conservées pour résoudre le problème) vérifient une C.O.S. vers le bas.

Annexe 8

Sur la résolution de $\mathcal{L} \vec{E} = \vec{S}$ au moyen de solutions
élémentaires de $\mathcal{L} \vec{g}_i = \vec{e}_i \delta$.

Propriété : Soient $\vec{E}$, $\vec{S}$ et $\vec{g}_i$ des distributions vectorielles des variables x , y et z , et $\mathcal{L}$ un opérateur qui, agissant sur une distribution vectorielle $\vec{U} = \sum_i U_i \vec{e}_i$, donne :

$$\mathcal{L} \vec{U} = \sum_{p,q} (\mathcal{L}_{pq} U_q) \vec{e}_p, \quad p \text{ et } q = 1, 2, 3 \quad (\text{A8.1})$$

où les $\mathcal{L}_{ij}$ sont des opérateurs différentiels linéaires à coefficients constants du type

$$\sum_{\ell, m, n} C_{\ell, m, n} \frac{\partial^{\ell+m+n}}{\partial_x^\ell \partial_y^m \partial_z^n}.$$

Supposons que $\vec{E}$ soit l'unique solution de :

$$\mathcal{L} \vec{E} = \vec{S}, \quad (\text{A8.2})$$

et que, δ étant l'unité de convolution, on ait pour $i = 1, 2, 3$:

$$\mathcal{L} \vec{g}_i = \vec{e}_i \delta. \quad (\text{A8.3})$$

Alors les composantes E_i de $\vec{E}$ peuvent être déduites des composantes g_{ij} de $\vec{g}_i$ et des composantes S_i de $\vec{S}$ par :

$$E_i = \sum_{j=1}^3 g_{ji} * S_j, \quad (\text{A8.4})$$

qui peut s'écrire sous forme matricielle :

$$\vec{E} = {}^t[g] * \vec{S}. \quad (\text{A8.5})$$

Preuve : Ayant supposé que $\vec{E}$, définie comme solution de (A8.2), est unique (dans le cadre du § 3.4.1, l'unicité est assurée par l'imposition d'une C.O.S.), il suffit de vérifier que l'application de $\mathcal{L}$ sur ${}^t[g] * \vec{S}$ donne effectivement $\vec{S}$; de (A8.1) et (A8.4), il vient :

$$\mathcal{L} ({}^t[g] * \vec{S}) = \sum_{pq} (\mathcal{L}_{pq} \sum_j g_{jq} * s_j) \vec{e}_p = \sum_{pqj} (\mathcal{L}_{pq} g_{jq} * s_j) \vec{e}_p .$$

Or d'après (A8.3) et (A8.1) :

$$\mathcal{L} \vec{g}_j = \sum_{p,q} (\mathcal{L}_{pq} g_{jq}) \vec{e}_p = \vec{e}_j \delta . \quad \text{D'où :}$$

$$\mathcal{L} ({}^t[g] * \vec{S}) = \sum_j \vec{e}_j \delta * s_j = \sum_j \vec{e}_j s_j = \vec{S} . \quad \text{C.Q.F.D.}$$

Annexe 9Sur les "vecteurs de Green" $\vec{g}_i$.

Nous détaillons ici quelques calculs relatifs au § 3.4.

1. OBTENTION DE L'EQUATION (3.44) A PARTIR DE (3.43)

$$-4\pi^2 \vec{\sigma} \wedge (\vec{\sigma} \wedge \hat{\vec{g}}_i) - [A] \hat{\vec{g}}_i = \vec{e}_i \quad (\text{A9.1})$$

$$\Leftrightarrow -4\pi^2 \left[(\vec{\sigma} \cdot \hat{\vec{g}}_i) \vec{\sigma} - \sigma^2 \hat{\vec{g}}_i \right] - [A] \hat{\vec{g}}_i = \vec{e}_i$$

$$\Leftrightarrow (4\pi^2 \sigma^2 - [A]) \hat{\vec{g}}_i = \vec{e}_i + 4\pi^2 (\vec{\sigma} \cdot \hat{\vec{g}}_i) \vec{\sigma} \quad (\text{A9.2})$$

Soit $[M(\vec{\sigma})]$ "l'opérateur inverse" de $4\pi^2 \sigma^2 - [A]$. Nous pouvons écrire (A9.2) sous la forme :

$$\hat{\vec{g}}_i = [M] \vec{e}_i + 4\pi^2 (\vec{\sigma} \cdot \hat{\vec{g}}_i) [M] \vec{\sigma} \quad (\text{A9.3})$$

Ces notations nécessitent quelques commentaires. Supposons pour fixer les idées que le milieu anisotrope soit sans pertes. On sait alors (Annexe 6) que $[\varepsilon]$, donc $[A]$ est hermitienne. Les valeurs propres λ_k de $[A]$ sont réelles et $[M(\vec{\sigma})]$ présentera donc des singularités pour $\sigma^2 = \|\vec{\sigma}\|^2 = \lambda_k / (4\pi^2)$. On peut s'affranchir de ce problème en ajoutant à $[A]$ une petite partie imaginaire $i[\delta]$, c'est à dire en rajoutant de faibles pertes fictives au milieu considéré, de sorte que les valeurs propres de $[\tilde{A}] = [A] + i[\delta]$ ne soient plus sur l'axe réel. L'opérateur $4\pi^2 \sigma^2 - [\tilde{A}]$ est alors inversible (nous noterons $[\tilde{M}]$ son inverse) et nous pouvons poursuivre la détermination des $\hat{\vec{g}}_i(\vec{\sigma})$. Quand ils auront été obtenus, il restera à envisager le passage à la limite $[\delta] \rightarrow [0]$, c'est à dire $[\tilde{A}] \rightarrow [A]$ et $[\tilde{M}] \rightarrow [M]$. Nous poursuivons en choisissant pour inconnue auxiliaire :

$$\psi(\vec{\sigma}) = \vec{\sigma} \cdot \hat{\vec{g}}_i(\vec{\sigma}). \quad (\text{A9.4})$$

De (A9.3), il vient :

$$\begin{aligned}\psi(\vec{\sigma}) &= \vec{\sigma} \cdot [\tilde{M}] \vec{e}_i + 4\pi^2 \psi(\vec{\sigma}) \vec{\sigma} \cdot [\tilde{M}] \vec{\sigma} \\ \Rightarrow \psi(\vec{\sigma}) (1 - 4\pi^2 \vec{\sigma} \cdot [\tilde{M}] \vec{\sigma}) &= \vec{\sigma} \cdot [\tilde{M}] \vec{e}_i.\end{aligned}$$

Ici encore nous sommes confrontés à un problème de division de distributions. Nous supposons que $[\delta]$ peut être choisie de façon à ce que $\psi(\vec{\sigma})$ puisse s'exprimer par :

$$\psi(\vec{\sigma}) = \frac{\vec{\sigma} \cdot [\tilde{M}] \vec{e}_i}{1 - 4\pi^2 \vec{\sigma} \cdot [\tilde{M}] \vec{\sigma}} \quad (\text{A9.5})$$

d'où finalement :

$$\hat{g}_i(\vec{\sigma}) = [M] \vec{e}_i + 4\pi^2 \frac{\vec{\sigma} \cdot [M] \vec{e}_i}{1 - 4\pi^2 \vec{\sigma} \cdot [M] \vec{\sigma}} [M] \vec{\sigma} \quad (\text{A9.6})$$

qu'il serait préférable d'écrire :

$$\hat{g}_i(\vec{\sigma}) = \lim_{[\delta] \rightarrow [0]} \left\{ [\tilde{M}] \vec{e}_i + 4\pi^2 \frac{\vec{\sigma} \cdot [\tilde{M}] \vec{e}_i}{1 - 4\pi^2 \vec{\sigma} \cdot [\tilde{M}] \vec{\sigma}} [\tilde{M}] \vec{\sigma} \right\}. \quad (\text{A9.7})$$

2. OBTENTION DES TRANSFORMEES DE FOURIER DES VECTEURS DE GREEN DANS LE CAS OU $[\varepsilon]$ EST DIAGONAL

Nous supposons que $[A]$ est de la forme (3.46). Si nécessaire (en particulier si $[A]$ est réelle), on ajoutera à $[A]$ une petite partie imaginaire $i[\delta]$ ($[\delta]$ diagonale) pour que les calculs que nous allons effectuer aient un sens. Pour simplifier l'écriture, nous ne le ferons pas, et ce qui suit doit être considéré comme une présentation formelle de résultats justifiables au moyen de la remarque précédente.

Nous posons $\vec{\rho} = 2\pi \vec{\sigma}$; la matrice $[M]$ (eq.3.45) est diagonale et a pour éléments $M_{ij} = \delta_{ij}/(\rho^2 - A_{ii})$. La $j^{\text{ème}}$ composante de $\hat{g}_i$ s'écrit (eq.3.44) :

$$\hat{g}_{ij} = \frac{\delta_{ij}}{\rho^2 - A_{ii}} + \frac{\vec{\rho} \cdot [M] \vec{e}_i}{1 - \vec{\rho} \cdot [M] \vec{\rho}} \sum_k M_{jk} \rho_k . \quad (A9.8)$$

Nous calculons le dénominateur $1 - \vec{\rho} \cdot [M] \vec{\rho}$ de la façon suivante :

$$1 - \vec{\rho} \cdot [M] \vec{\rho} = 1 - \sum_k \frac{\rho_k^2}{\rho^2 - A_{kk}} = \sum_k \left(\frac{\rho_k^2}{\rho^2} - \frac{\rho_k^2}{\rho^2 - A_{kk}} \right) = \sum_k \frac{-A_{kk} \rho_k^2}{\rho^2 (\rho^2 - A_{kk})} .$$

D'où :

$$\hat{g}_{ij} = \frac{\delta_{ij}}{\rho^2 - A_{ii}} - \frac{\rho^2}{\sum_k \frac{A_{kk} \rho_k^2}{\rho^2 - A_{kk}}} \frac{\rho_i \rho_j}{(\rho^2 - A_{ii})(\rho^2 - A_{jj})} . \quad (A9.9)$$

3. DETERMINATION DES VECTEURS DE GREEN POUR UN MILIEU DE PERMITTIVITE DIAGONALE. CAS DES RESEAUX

L'énoncé du problème est le suivant : trouver les $\vec{g}_i(x,y)$ pseudo-périodiques en x avec la période d , de coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$, solution de :

$$\text{rot rot } \vec{g}_i - [A] \vec{g}_i = \vec{e}_i \delta_{\mathcal{R}} , \quad i = 1, 2, 3, \quad (A9.10)$$

et vérifiant une C.O.S. pour $y \rightarrow \pm\infty$, dans le cas où $[A]$ est diagonale.

On a vu au § 3.4.3 que les $\vec{g}_i$ s'écrivent :

$$\vec{g}_i(x,y) = \sum_{j=1}^3 g_{ij}(x,y) \vec{e}_j , \quad \text{avec}$$

$$g_{ij}(x,y) = \sum_{n \in \mathbb{Z}} g_{ijn}(y) \exp(i\alpha_n x) , \quad (A9.11)$$

et $\alpha_n = \alpha_0 + n 2\pi/d$.

En appelant σ_2 la variable conjuguée de y par la transformation de Fourier

$\mathcal{F}^-$, les transformées de Fourier des $g_{ijn}(y)$ sont obtenues par (3.57), avec $\rho_1 = \alpha_n$, $\rho_2 = 2\pi\sigma_2$, $\rho^2 = \rho_1^2 + \rho_2^2$, $\rho_3 = 0$:

$$d \hat{g}_{ijn}(\rho_2) = \frac{\delta_{ij}}{\rho^2 - A_{ii}} - \frac{(\rho^2 - A_{11})(\rho^2 - A_{22}) \rho_i \rho_j}{(A_{11}\rho_1^2 + A_{22}\rho_2^2 - A_{11}A_{22})(\rho^2 - A_{ii})(\rho^2 - A_{jj})} \quad (A9.12)$$

D'où en posant $F(\sigma_2) = A_{11}\alpha_n^2 + 4\pi^2 A_{22} \sigma_2^2 - A_{11} A_{22}$:

$$\hat{g}_{11n}(\sigma_2) = \frac{1}{d} \frac{A_{22} - \alpha_n^2}{F(\sigma_2)}$$

$$\hat{g}_{22n}(\sigma_2) = \frac{1}{d} \frac{A_{11} - 4\pi^2 \sigma_2^2}{F(\sigma_2)}$$

$$\hat{g}_{33n}(\sigma_2) = \frac{1}{d} \frac{1}{\alpha_n^2 + 4\pi^2 \sigma_2^2 - A_{33}}$$

$$\hat{g}_{12n}(\sigma_2) = \hat{g}_{21n}(\sigma_2) = -\frac{1}{d} \frac{2\pi\alpha_n \sigma_2}{F(\sigma_2)}$$

$$\hat{g}_{13n}(\sigma_2) = \hat{g}_{31n}(\sigma_2) = \hat{g}_{23n}(\sigma_2) = \hat{g}_{32n}(\sigma_2) = 0.$$

Introduisons les constantes β_n et β'_n définies par :

$$\beta_n^2 = \frac{A_{11}}{A_{22}} (A_{22} - \alpha_n^2) \quad (A9.13)$$

$$\beta_n'^2 = A_{33} - \alpha_n^2, \quad (A9.14)$$

et choisies de telle sorte que :

$$\beta_n > 0 \quad \text{ou bien} \quad \text{Im}(\beta_n) > 0 \quad \text{et idem pour } \beta'_n. \quad (A9.15)$$

Nous supposons que α_n est tel que β_n ou β'_n ne soient jamais nuls, c'est à dire que nous évitons de nous placer sur une anomalie de Rayleigh [58] qui correspond au cas où l'un des ordres du réseau se propage parallèlement à $\vec{e}_x$. Les $\hat{g}_{ijn}(\sigma_2)$ s'écrivent alors :

$$\hat{g}_{11n}(\sigma_2) = - \frac{1}{d A_{22}} \frac{A_{22} - \alpha_n^2}{\beta_n^2 - 4\pi^2 \sigma_2^2} \quad (A9.16)$$

$$\hat{g}_{22n}(\sigma_2) = - \frac{1}{d A_{22}} \frac{A_{11} - 4\pi^2 \sigma_2^2}{\beta_n^2 - 4\pi^2 \sigma_2^2} \quad (A9.17)$$

$$\hat{g}_{33n}(\sigma_2) = - \frac{1}{d} \frac{1}{\beta_n'^2 - 4\pi^2 \sigma_2^2} \quad (A9.18)$$

$$\hat{g}_{12n}(\sigma_2) = \hat{g}_{21n}(\sigma_2) = \frac{1}{d A_{22}} \frac{2\pi\alpha_n \sigma_2}{\beta_n^2 - 4\pi^2 \sigma_2^2} \quad (A9.19)$$

La détermination des inverses de Fourier de ces $\hat{g}_{ijn}(\sigma_2)$ nous donne l'occasion d'illustrer les remarques que nous avons faites précédemment. En effet, il se peut que β_n^2 ou $\beta_n'^2$ soient réels positifs, et les $\hat{g}_{ijn}$ présentent alors des singularités. On ajoutera alors à [A] une petite partie imaginaire de façon à ce que les β_n^2 (ou $\beta_n'^2$) ne soient plus réels positifs. Compte tenu de la détermination (A9.15), on aura alors $\text{Im}(\beta_n)$ et $\text{Im}(\beta_n')$ positifs. On est conduit à rechercher les inverses de Fourier de fonctions de la forme $\frac{1}{\beta_n^2 - 4\pi^2 \sigma_2^2}$, $\frac{2i\pi\sigma_2}{\beta_n^2 - 4\pi^2 \sigma_2^2}$, $\frac{-4\pi^2 \sigma_2^2}{\beta_n^2 - 4\pi^2 \sigma_2^2}$. Elles s'obtiennent aisément en utilisant les propriétés des distributions et de leur transformée de Fourier, ou bien en consultant par exemple la référence [31] :

$$\frac{1}{2i\beta_n} \exp(i\beta_n |y|) \xrightarrow{\mathcal{F}^-} \frac{1}{\beta_n^2 - 4\pi^2 \sigma_2^2}$$

$$\frac{1}{2} \text{sgn}(y) \exp(i\beta_n |y|) \xrightarrow{\mathcal{F}^-} \frac{2i\pi\sigma_2}{\beta_n^2 - 4\pi^2 \sigma_2^2}$$

$$\left[\delta(y) + \frac{i\beta_n}{2} \right] \exp(i\beta_n |y|) \xrightarrow{\mathcal{F}^-} \frac{-4\pi^2 \sigma_2^2}{\beta_n^2 - 4\pi^2 \sigma_2^2}$$

D'où les $g_{ij}(x, y)$:

$$g_{11}(x, y) = \sum_{n \in \mathbb{Z}} \frac{i\beta_n}{2d A_{11}} \exp(i\alpha_n x + i\beta_n |y|)$$

$$g_{22}(x,y) = \sum_{n \in \mathbb{Z}} \frac{1}{d A_{22}} \left[\frac{i A_{11} \alpha_n^2}{2 A_{22} \beta_n} - \delta(y) \right] \exp(i\alpha_n x + i\beta_n |y|)$$

$$g_{33}(x,y) = \sum_{n \in \mathbb{Z}} \frac{i}{2d \beta'_n} \exp(i\alpha_n x + i\beta'_n |y|)$$

$$g_{12}(x,y) = g_{21}(x,y) = \sum_{n \in \mathbb{Z}} \frac{\alpha_n}{2id A_{22}} \operatorname{sgn}(y) \exp(i\alpha_n x + i\beta_n |y|)$$

$$g_{13}(x,y) = g_{23}(x,y) = g_{31}(x,y) = g_{32}(x,y) = 0.$$

On notera pour terminer que dans ces expressions finales, on pourra retirer à [A] la partie imaginaire que l'on y avait ajoutée, et que les vecteurs de Green ainsi obtenus vérifient une C.O.S. quand $y \rightarrow \pm\infty$ compte tenu de (A9.15).

Annexe 10Découplage des vecteurs de Poynting

Hypothèses : Dans un milieu isotrope sans pertes, considérons en régime harmonique deux ondes planes séparément solutions des équations de Maxwell représentées par leurs champs complexes ($\vec{k}_1$ et $\vec{k}_2$ sont supposés réels) :

Champ 1 : $\vec{E}_1(\vec{r}) = \vec{E}_{01} \exp(i \vec{k}_1 \cdot \vec{r})$; $\vec{H}_1(\vec{r}) = \vec{H}_{01} \exp(i \vec{k}_1 \cdot \vec{r})$,

Champ 2 : $\vec{E}_2(\vec{r}) = \vec{E}_{02} \exp(i \vec{k}_2 \cdot \vec{r})$; $\vec{H}_2(\vec{r}) = \vec{H}_{02} \exp(i \vec{k}_2 \cdot \vec{r})$.

Soit $\vec{E} = \vec{E}_1 + \vec{E}_2$, $\vec{H} = \vec{H}_1 + \vec{H}_2$ le champ résultant de la somme de ces deux ondes planes.

Soient $\vec{\mathcal{P}} = \vec{E} \wedge \vec{H}$, $\vec{\mathcal{P}}_1 = \vec{E}_1 \wedge \vec{H}_1$, $\vec{\mathcal{P}}_2 = \vec{E}_2 \wedge \vec{H}_2$ les vecteurs de Poynting complexes qui sont associés à ces champs (au coefficient multiplicatif 2 près).

Nous considérons une base orthonormée $\vec{e}_x, \vec{e}_y, \vec{e}_z$, et noterons par un indice x (ou y) les composantes des vecteurs sur $\vec{e}_x$ (ou $\vec{e}_y$).

Conclusion : La proposition $\text{Re}(\mathcal{P}_y) = \text{Re}(\mathcal{P}_{1y}) + \text{Re}(\mathcal{P}_{2y})$ est vraie dans les cas suivants (voir figure) :

Cas 1 : lorsque $k_{1x} = k_{2x} = k_x$ et $k_{1y} = -k_{2y} = k_y$.

Cas 2 : lorsque $\vec{k}_1 = \vec{k}_2$ et si les polarisations des champs 1 et 2 sont orthogonales.

Démonstration : Le calcul de $\vec{\mathcal{P}}$ donne :

$$\vec{\mathcal{P}} = \vec{\mathcal{P}}_1 + \vec{\mathcal{P}}_2 = \vec{E}_1 \wedge \vec{H}_2 + \vec{E}_2 \wedge \vec{H}_1.$$

Il s'agit donc de montrer que dans les cas 1 et 2,

$$\operatorname{Re}(\vec{E}_1 \wedge \vec{H}_2 + \vec{E}_2 \wedge \vec{H}_1)_y = 0.$$

Ecrivons les champs magnétiques sous la forme :

$$\vec{H}_1 = \frac{1}{\omega \mu_0} \vec{k}_1 \wedge \vec{E}_1 ; \quad \vec{H}_2 = \frac{1}{\omega \mu_0} \vec{k}_2 \wedge \vec{E}_2$$

d'où :

$$\omega \mu_0 (\vec{E}_1 \wedge \vec{H}_2 + \vec{E}_2 \wedge \vec{H}_1) = (\vec{E}_1 \cdot \vec{E}_2) \vec{k}_2 - (\vec{E}_1 \cdot \vec{k}_2) \vec{E}_2 + (\vec{E}_2 \cdot \vec{E}_1) \vec{k}_1 - (\vec{E}_2 \cdot \vec{k}_1) \vec{E}_1 \quad (\text{A10.1})$$

Cas 1 : On peut remplacer $\vec{k}_1$ par $(\vec{k}_2 + 2k_y \vec{e}_y)$ et $\vec{k}_2$ par $(\vec{k}_1 - 2k_y \vec{e}_y)$.
En tenant compte de $\vec{k}_1 \cdot \vec{E}_1 = \vec{k}_2 \cdot \vec{E}_2 = 0$, il vient :

$$\omega \mu_0 (\vec{E}_1 \wedge \vec{H}_2 + \vec{E}_2 \wedge \vec{H}_1)_y = - (\vec{E}_1 \cdot \vec{E}_2) k_y \vec{e}_y + 2k_y E_{1y} \vec{E}_{2y} + (\vec{E}_2 \cdot \vec{E}_1) k_y \vec{e}_y - 2k_y E_{2y} \vec{E}_{1y}.$$

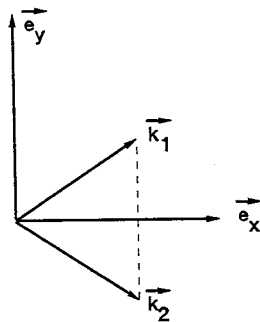
Il apparaît alors que cette quantité est imaginaire pure. C.Q.F.D.

Cas 2 : Dans ce cas, $\vec{E}_1 \cdot \vec{E}_2 = 0 = \vec{E}_2 \cdot \vec{E}_1$; $\vec{k}_1 = \vec{k}_2$,

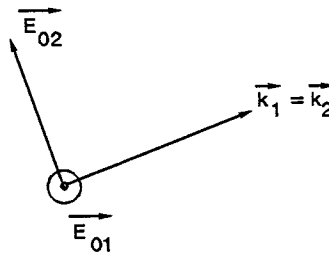
d'où aussi $\vec{E}_1 \cdot \vec{k}_2 = \vec{E}_1 \cdot \vec{k}_1 = 0$ et $\vec{E}_2 \cdot \vec{k}_1 = \vec{E}_2 \cdot \vec{k}_2 = 0$.

L'expression (A10.1) est nulle, ce qui montre que dans ce cas, on peut même écrire :

$$\vec{\Phi} = \vec{\Phi}_1 + \vec{\Phi}_2.$$



cas 1



cas 2

Annexe 11Quelques lemmes pour les réseaux anisotropes1. Décomposition d'un champ pseudo-périodique en ondes planes en milieu anisotrope.

Soit $\vec{E}(x,y)$ un champ pseudo-périodique, de période d , de coefficient de pseudo-périodicité $\exp(i\alpha_0 d)$, vérifiant l'équation $\text{rot rot } \vec{E} - k_0^2 [\varepsilon] \vec{E} = 0$. En posant $\alpha_n = \alpha_0 + n 2\pi/d$, $\vec{E}$ peut s'écrire :

$$\vec{E}(x,y) = \sum_{n \in \mathbb{Z}} \vec{E}_n(y) \exp(i\alpha_n x). \quad (\text{A11.1})$$

Nous en déduisons en s'inspirant du § 3.2.3 que $\vec{E}$ peut s'écrire sous la forme d'une combinaison d'ondes planes :

$$\vec{E}(x,y) = \sum_{\substack{n \in \mathbb{Z} \\ j=1,2,3,4}} C_{nj} \vec{E}_{nj} \exp(i\alpha_n x + i\beta_{nj} y) \quad (\text{A11.2})$$

où, pour n fixé, les β_{nj} sont les solutions de l'équation obtenue en remplaçant dans (3.28) α par α_n , et les $\vec{E}_{nj}$ sont obtenus en remplaçant dans (3.26) α par α_n . Les $\vec{E}_{nj}$ peuvent être choisis normés et les C_{nj} sont des constantes complexes.

2. Influence d'un changement de pseudo-périodicité et de la transposition de $[\varepsilon]$.

Soit $\vec{U}(x,y)$ un champ pseudo-périodique, de période d , de coefficient de pseudo-périodicité $\exp(-i\alpha_0 d)$, vérifiant l'équation $\text{rot rot } \vec{U} - k_0^2 {}^t[\varepsilon] \vec{U} = 0$. Soit $\tilde{\alpha}_n = -\alpha_0 + n 2\pi/d = -\alpha_{-n}$. $\vec{U}$ peut s'écrire :

$$\vec{U}(x,y) = \sum_{n \in \mathbb{Z}} \vec{U}_{-n}(y) \exp(i\tilde{\alpha}_n x) = \sum_{n \in \mathbb{Z}} \vec{U}_n(y) \exp(-i\alpha_n x).$$

En s'inspirant toujours du § 3.2.3, $\vec{U}$ peut se mettre sous la forme :

$$\vec{U}(x,y) = \sum_{\substack{n \in \mathbb{Z} \\ j=1,2,3,4}} D_{-nj} \vec{U}_{-nj} \exp(i\tilde{\alpha}_n x + i\tilde{\beta}_{nj} y) \quad (\text{A11.3})$$

où, pour n fixé, les $\tilde{\beta}_{nj}$ sont les solutions de l'équation obtenue en remplaçant dans (3.28) α par $-\alpha_{-n}$ et $[\varepsilon]$ par ${}^t[\varepsilon]$, et les $\vec{U}_{-nj}$ sont obtenus par (3.26) en y faisant les mêmes changements. L'observation attentive de (3.28) montre que les $\tilde{\beta}_{nj}$ sont liés aux β_{nj} par $\tilde{\beta}_{nj} = -\beta_{-nj}$. Nous pouvons donc réécrire (A11.3) sous la forme :

$$\vec{U}(x,y) = \sum_{\substack{n \in \mathbb{Z} \\ j=1,2,3,4}} D_{nj} \vec{U}_{nj} \exp(-i\alpha_n x - i\beta_{nj} y). \quad (\text{A11.4})$$

3. Lemme 1

Les notations non précisées ici sont celles du § 6.3. Les champs $\vec{E}(x,y)$ et $\vec{U}(x,y)$ ont les mêmes pseudo-périodicités que précédemment. Nous supposons qu'ils vérifient les équations :

$$\text{rot rot } \vec{E} - k_0^2 [\varepsilon] \vec{E} = 0, \quad (\text{A11.5})$$

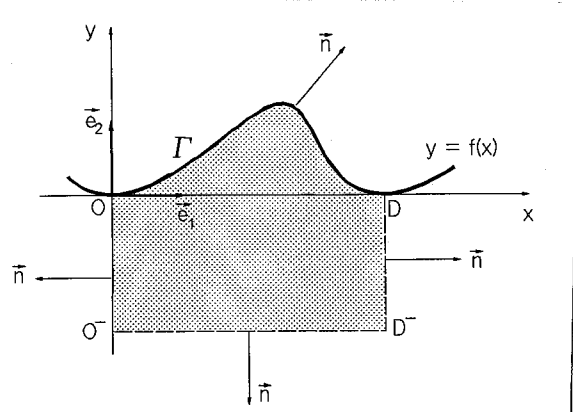
$$\text{rot rot } \vec{U} - k_0^2 {}^t[\varepsilon] \vec{U} = 0 \quad (\text{A11.6})$$

dans le domaine Ω^- ($y < f(x)$, figure A11.1). Alors :

$$\int_{\Gamma} \vec{n} \cdot (\vec{U} \wedge \text{rot } \vec{E} - \vec{E} \wedge \text{rot } \vec{U}) d\ell = \vec{e}_2 \cdot \int_{O^- D^-} (\vec{U} \wedge \text{rot } \vec{E} - \vec{E} \wedge \text{rot } \vec{U}) dx \quad (\text{A11.7})$$

Nous n'en reproduisons pas ici la démonstration qui suit la procédure utilisée au début du § 6.3.

Figure A11.1. Γ est l'arc OD de la courbe $y = f(x)$



4. Evaluation du second membre de (A11.7)

Les équations (A11.5) et (A11.6) étant vérifiées dans Ω^- , les développements (A11.2) et (A11.4) sont valables dans la zone $y < 0 = \min(f(x))$. Posons :

$$\vec{k}_{np} = \alpha_n \vec{e}_1 + \beta_{np} \vec{e}_2 . \quad (\text{A11.8})$$

Soit y_0 l'ordonnée du segment O^-D^- . On obtient après quelques calculs :

$$\vec{e}_2 \cdot \int_{O^-D^-} (\vec{U} \wedge \text{rot } \vec{E} - \vec{E} \wedge \text{rot } \vec{U}) dx = \sum_{\substack{n \in \mathbb{Z} \\ p=1,2,3,4 \\ q=1,2,3,4}} A_{npq} \exp[i(\beta_{np} - \beta_{nq})y_0] \quad (\text{A11.9})$$

avec :

$$\begin{aligned} A_{npq} &= i d C_{np} D_{nq} \vec{e}_2 \cdot \left\{ \vec{U}_{nq} \wedge (\vec{k}_{np} \wedge \vec{E}_{np}) + \vec{E}_{np} \wedge (\vec{k}_{nq} \wedge \vec{U}_{nq}) \right\} \\ &= i d C_{np} D_{nq} \{ (E_{npx} U_{nqx} + E_{npz} U_{nqz})(\beta_{np} + \beta_{nq}) - \alpha_n (E_{npx} U_{nqy} + E_{npy} U_{nqx}) \} \end{aligned} \quad (\text{A11.10})$$

où les indices x, y et z représentent les composantes des vecteurs sur $\vec{e}_1$, $\vec{e}_2$ et $\vec{e}_3$.

5. Lemme 2.

Soit un milieu anisotrope quelconque de permittivité $[\varepsilon]$. Pour α réel donné, considérons les deux ondes planes :

$$\vec{E}_1(x, y) = \vec{E}_{10} \exp(i\alpha x + i\beta_1 y) = \vec{E}_{10} \exp(i \vec{k}_1 \cdot \vec{r}) \quad \text{et}$$

$$\vec{E}_2(x, y) = \vec{E}_{20} \exp(-i\alpha x - i\beta_2 y) = \vec{E}_{20} \exp(i \vec{k}_2 \cdot \vec{r}),$$

respectivement solutions de :

$$\text{rot rot } \vec{E}_1 - k_0^2 [\varepsilon] \vec{E}_1 = 0, \quad (\text{A11.11})$$

$$\text{rot rot } \vec{E}_2 - k_0^2 {}^t[\varepsilon] \vec{E}_2 = 0 \quad (\text{A11.12})$$

et vérifiant une condition d'onde sortante dans le même demi-espace. En

notant $\vec{e}_2$ un vecteur directeur de l'axe des y, on a :

$$\vec{e}_2 \cdot \left\{ \vec{E}_1 \wedge (\vec{k}_2 \wedge \vec{E}_2) - \vec{E}_2 \wedge (\vec{k}_1 \wedge \vec{E}_1) \right\} = 0. \quad (\text{A11.13})$$

Nous n'avons pas réussi à montrer cette propriété dans le cas général. Pour des $[\varepsilon]$ quelconques, il est en effet difficile d'obtenir des propriétés des β_1 et β_2 , ainsi que des polarisation $\vec{E}_{10}$ et $\vec{E}_{20}$. La C.O.S. ne s'exprime pas facilement non plus (§ 3.3). Nous sommes cependant convaincus que la propriété (A11.13) est toujours vérifiée (notre conviction s'appuie aussi sur des essais numériques). Nous en donnons ci-dessous la démonstration dans le cas où $[\varepsilon]$ est diagonale, ce qui correspond à nos hypothèses du chapitre 6.

Nous supposons ici pour fixer les idées que $\vec{E}_1$ et $\vec{E}_2$ vérifient une C.O.S. quand $y \rightarrow -\infty$. Le lemme 1 montre que le premier membre de (A11.13) est indépendant de y. Il est donc nécessairement nul dans le cas de milieux absorbants, et aussi dans le cas de milieux sans pertes, lorsque l'une au moins des deux ondes planes est évanescence (lorsque β_1 ou β_2 possède une partie imaginaire). Il suffit donc d'étudier le cas de milieux sans pertes ($[\varepsilon]$ hermitienne, cf. annexe 6 ; $[\varepsilon]$ étant de plus diagonale, elle est réelle), lorsque β_1 et β_2 sont réels. Nous avons délibérément supposé que β_1 et β_2 ne sont pas nuls. Si cela était le cas, on serait placé sur une "anomalie de Rayleigh" (apparition ou disparition d'un ordre propagatif). Ces anomalies n'apparaissent que pour des valeurs particulières de l'angle d'incidence θ que nous éviterons.

On montre aisément (cf. A11.10) que :

$$\begin{aligned} \vec{e}_2 \cdot \left\{ \vec{E}_1 \wedge (\vec{k}_2 \wedge \vec{E}_2) - \vec{E}_2 \wedge (\vec{k}_1 \wedge \vec{E}_1) \right\} = \\ = - (\beta_1 + \beta_2) (E_{1x} E_{2x} + E_{1z} E_{2z}) + \alpha (E_{1x} E_{2y} + E_{2x} E_{1y}) \end{aligned} \quad (\text{A11.14})$$

Pour établir (A11.13), il suffit de montrer que le second membre de (A11.14) est nul. On peut pour cela mettre en évidence quelques propriétés de β_1 , β_2 , $\vec{E}_1$ et $\vec{E}_2$. Faisons-le par exemple pour β_1 et $\vec{E}_1$. L'équation (3.26) s'écrit ici :

$$(k_{xx}^2 - \beta_1^2) E_{1x} + \alpha \beta_1 E_{1y} = 0,$$

$$\alpha \beta_1 E_{1x} + (k_{yy}^2 - \alpha^2) E_{1y} = 0,$$

$$(k_{zz}^2 - \alpha^2 - \beta_1^2) E_{1z} = 0.$$

On a alors :

$$\cdot \text{ soit } \beta_1^2 = k_{zz}^2 - \alpha^2 \quad \text{et} \quad E_{1x} = E_{1y} = 0,$$

$$\cdot \text{ soit } \beta_1^2 = \frac{k_{xx}^2}{k_{yy}^2} (k_{yy}^2 - \alpha^2) \quad \text{et} \quad E_{1z} = 0, \quad \alpha \frac{k_{xx}^2}{k_{yy}^2} E_{1x} + \beta_1 E_{1y} = 0.$$

Les relations liant les composantes de $\vec{E}_2$ à α et β_2 sont identiques. De ces relations et de la détermination de β_1 et β_2 imposée par la C.O.S. ($\beta_1 < 0$ et $\beta_2 > 0$), on déduit facilement que le second membre de (A11.14) est nul.

6. Lemme 3

$\vec{E}(x,y)$ et $\vec{U}(x,y)$ vérifient les mêmes conditions que dans le lemme 1. Nous supposons de plus qu'ils vérifient une condition d'onde sortante quand $y \rightarrow -\infty$. Les deux membres de l'égalité (A11.7) sont alors nuls.

D'après le lemme 2, tous les A_{npq} (équation (A11.10)) sont nuls. L'égalité (A11.9) montre que les deux membres de (A11.7) sont nuls.

Annexe 12 . Expression des $[T_i]$ et de S_2 .

Le contenu de cette annexe a été déposé simultanément à la référence [76] aux Comptes Rendus de l'Académie des Sciences, et est conservé dans les archives de cette Académie .

Expression des matrices $[T_1]$, $[T_2]$, $[T_3]$ et $[T_4]$ intervenant dans la généralisation des formules de Kirchhoff - Helmholtz (§ 6.3).

Nous supposons ici que le milieu anisotrope considéré a pour permittivité relative $[\epsilon_2] = \begin{bmatrix} \epsilon_x & 0 & 0 \\ 0 & \epsilon_y & 0 \\ 0 & 0 & \epsilon_z \end{bmatrix}$. On pose

$$\beta_m^2 = \frac{\epsilon_x}{\epsilon_y} (k_0^2 \epsilon_y - d_m^2) , \quad \beta'_m{}^2 = k_0^2 \epsilon_z - d_m^2 , \text{ les}$$

racines (éventuellement complexes) étant déterminées par :

$$y \in \mathbb{C} , \quad \text{Im}(\sqrt{y}) > 0 \quad \text{ou bien} \quad \sqrt{y} > 0 .$$

Dans la suite, $\text{sgn}(y)$ représente la fonction égale à $+1$ si $y > 0$ et à -1 si $y < 0$, et $\sum_n$ désigne $\sum_{n \in \mathbb{Z}}$.

En se reportant aux équations (6.8) et (6.9) , et en se rappelant que les points M' et P ont pour coordonnées respectives $(x', f(x'))$ et (x, y) , les éléments des matrices ne dépendent que des variables $X = x' - x$ et $Y = f(x') - y$. Nous adoptons la notation suivante :

$$[T_l(M', P)] = [T_l(x, Y)] , \quad l = 1, 2, 3, 4 ,$$

et nous posons

$$\Psi_m(x, Y) = e^{i \beta_m |Y|} e^{-i d_m X} ,$$

$$\Psi'_m(x, Y) = e^{i \beta'_m |Y|} e^{-i d_m X} .$$

On a alors :

$$[T_1]_{11}(x, Y) = i \omega \mu_0 \sum_n \frac{i \beta_m}{2d k_0^2 \epsilon_x} \Psi_m(x, Y)$$

$$[T_1]_{12}(x, Y) = i \omega \mu_0 \sum_n \frac{i d_m}{2d k_0^2 \epsilon_y} \text{sgn}(Y) \Psi_m(x, Y)$$

$$[T_1]_{21}(x, Y) = [T_1]_{12}(x, Y)$$

$$[T_1]_{22}(X, Y) = i \omega \mu_0 \sum_n \frac{i \epsilon_x \alpha_n^2}{2d k_0^2 \epsilon_y^2 \beta_n} \Psi_n(X, Y)$$

$$[T_1]_{33}(X, Y) = i \omega \mu_0 \sum_n \frac{i}{2d \beta'_n} \Psi'_n(X, Y)$$

$$[T_1]_{13} = [T_1]_{31} = [T_1]_{23} = [T_1]_{32} = 0$$

$$[T_2]_{11} = [T_2]_{12} = [T_2]_{21} = [T_2]_{22} = [T_2]_{33} = 0$$

$$[T_2]_{13}(X, Y) = \sum_n \frac{1}{2d} \operatorname{sgn}(Y) \Psi_n(X, Y)$$

$$[T_2]_{31}(X, Y) = \sum_n \frac{-1}{2d} \operatorname{sgn}(Y) \Psi'_n(X, Y)$$

$$[T_2]_{23}(X, Y) = \sum_n \frac{1}{2d} \frac{\epsilon_x \alpha_n}{\epsilon_y \beta_n} \Psi_n(X, Y)$$

$$[T_2]_{32}(X, Y) = \sum_n \frac{-\alpha_n}{2d \beta'_n} \Psi'_n(X, Y)$$

$$[T_3]_{11} = [T_3]_{12} = [T_3]_{21} = [T_3]_{22} = [T_3]_{33} = 0$$

$$[T_3]_{13}(X, Y) = \sum_n \frac{1}{2d} \operatorname{sgn}(Y) \Psi'_n(X, Y)$$

$$[T_3]_{31}(X, Y) = \sum_n \frac{-1}{2d} \operatorname{sgn}(Y) \Psi_n(X, Y)$$

$$[T_3]_{23}(X, Y) = \sum_n \frac{\alpha_n}{2d \beta'_n} \Psi'_n(X, Y)$$

$$[T_3]_{32}(X, Y) = \sum_n \frac{-1}{2d} \frac{\epsilon_x \alpha_n}{\epsilon_y \beta_n} \Psi_n(X, Y)$$

$$[T_4]_{11}(X, Y) = \frac{1}{i \omega \mu_0} \sum_n \frac{i \beta'_n}{2d} \Psi'_n(X, Y)$$

$$[T_4]_{12}(X, Y) = \frac{1}{i \omega \mu_0} \sum_n \frac{i \alpha_n}{2d} \operatorname{sgn}(Y) \Psi'_n(X, Y)$$

$$[T_4]_{21}(X, Y) = [T_4]_{12}(X, Y)$$

$$[T_4]_{22}(X, Y) = \frac{1}{i \omega \mu_0} \sum_n \frac{i \alpha_n^2}{2d \beta'_n} \Psi'_n(X, Y)$$

$$[T_4]_{33}(X, Y) = \frac{1}{i \omega \mu_0} \sum_n \frac{i k_0^2 \epsilon_x}{2d \beta_n} \Psi_n(X, Y)$$

$$[T_4]_{13} = [T_4]_{31} = [T_4]_{23} = [T_4]_{32} = 0$$

Expression de l'opérateur S_2 attaché à la permittivité $[E_2]$.

Appliqué à la colonne $F(x) = \begin{pmatrix} F_1(x) \\ F_2(x) \\ F_3(x) \\ F_4(x) \end{pmatrix}$, l'opérateur S_2 donne la colonne $G(x) = \begin{pmatrix} G_1(x) \\ G_2(x) \\ G_3(x) \\ G_4(x) \end{pmatrix}$ qui s'exprime de la façon suivante (nous avons posé $\rho(x) = \sqrt{1 + f'(x)^2}$) :

$$G_1(x) = \int_0^d \left\{ N_{11}(x', x) F_1(x') + N_{14}(x', x) F_4(x') \right\} \rho(x') dx'$$

$$G_2(x) = \int_0^d \left\{ N_{22}(x', x) F_2(x') + N_{23}(x', x) F_3(x') \right\} \rho(x') dx'$$

$$G_3(x) = \int_0^d \left\{ N_{32}(x', x) F_2(x') + N_{33}(x', x) F_3(x') \right\} \rho(x') dx'$$

$$G_4(x) = \int_0^d \left\{ N_{41}(x', x) F_1(x') + N_{44}(x', x) F_4(x') \right\} \rho(x') dx'$$

En posant $\Psi_n(x, x') = e^{i \alpha_n (x - x') + i \beta_n |f(x) - f(x')|}$
et $\Psi'_n(x, x') = e^{i \alpha_n (x - x') + i \beta'_n |f(x) - f(x')|}$

les noyaux $N_{ij}(x', x)$ ont pour expression :

$$N_{11}(x', x) = \frac{-1}{d \rho(x)} \sum_n \left[\frac{\alpha_n}{\beta'_n} f'(x') + \operatorname{sgn}(f(x') - f(x)) \right] \Psi'_n(x, x')$$

$$N_{14}(x', x) = \frac{-\omega \mu_0}{d \rho(x)} \sum_n \frac{1}{\beta'_n} \Psi'_n(x, x')$$

$$N_{22}(x', x) = \frac{-1}{d \rho(x)} \sum_n \left[\frac{\epsilon_x \alpha_n}{\epsilon_y \beta_n} f'(x) + \operatorname{sgn}(f(x') - f(x)) \right] \Psi_n(x, x')$$

$$N_{23}(x', x) = \frac{\omega \mu_0}{d \rho(x)} \sum_n \left[\frac{\alpha_n}{\epsilon_y k_0^2} (f'(x) + f'(x')) \operatorname{sgn}(f(x') - f(x)) + \frac{\epsilon_x \alpha_n^2}{\epsilon_y^2 k_0^2 \beta_n} f'(x) f'(x') + \frac{\beta_n}{\epsilon_x k_0^2} \right] \Psi_n(x, x')$$

$$N_{32}(x', x) = \frac{1}{\omega \mu_0 d \rho(x)} \sum_n \frac{\varepsilon_x k_0^2}{\beta_n} \Psi_n(x, x')$$

$$N_{33}(x', x) = \frac{-1}{d \rho(x)} \sum_n \left[\frac{\varepsilon_x \alpha_n}{\varepsilon_y \beta_n} f'(x') + \operatorname{sgn}(f(x') - f(x)) \right] \Psi_n(x, x')$$

$$N_{41}(x', x) = \frac{-1}{\omega \mu_0 d \rho(x)} \sum_n \left[\alpha_n (f'(x) + f'(x')) \operatorname{sgn}(f(x') - f(x)) + \frac{\alpha_n^2}{\beta_n'} f'(x) f'(x') + \beta_n' \right] \Psi_n'(x, x')$$

$$N_{44}(x', x) = \frac{-1}{d \rho(x)} \sum_n \left[\frac{\alpha_n}{\beta_n'} f'(x) + \operatorname{sgn}(f(x') - f(x)) \right] \Psi_n'(x, x')$$

Remarque : en dépit des apparences ;

- les noyaux N_{11} , N_{22} , N_{33} et N_{44} sont continus (Pavageau et Bousquet [52]),
- les noyaux N_{14} et N_{32} présentent une faible singularité intégrable (logarithmique, cf. Maystre dans [58] ou bien annexe 13),
- les noyaux N_{23} et N_{41} présentent une singularité de type partie finie de $\frac{1}{(x-x')^2}$

Annexe 13Sur le traitement numérique des noyaux des équations intégrales1. Position du problème

Les noyaux qui interviennent dans la résolution numérique du problème sont de la forme :

$$N_{32}^2(x, x') = \frac{1}{\omega \mu_0 d \rho(x)} \sum_{n \in \mathbb{Z}} \frac{\varepsilon_{xx} k_0^2}{\beta_{2,n}} \psi_n(x, x')$$

$$N_{33}^2(x, x') = \frac{-1}{d \rho(x)} \sum_{n \in \mathbb{Z}} \left[\frac{\varepsilon_{xx}}{\varepsilon_{yy}} \frac{\alpha_n}{\beta_{2,n}} f'(x') + \operatorname{sgn}(f(x') - f(x)) \right] \psi_n(x, x')$$

où $\alpha_n = \alpha_0 + n 2\pi/d = \alpha_0 + nK$,

$$\beta_{2,n} = \sqrt{\frac{\varepsilon_{xx}}{\varepsilon_{yy}} (k_0^2 \varepsilon_{yy} - \alpha_n^2)} \quad \text{avec } \beta_{2,n} > 0 \quad \text{ou} \quad \operatorname{Im}(\beta_{2,n}) > 0,$$

$$\psi_n(x, x') = e^{i\alpha_n(x-x') + i\beta_{2,n}|f(x)-f(x')|}.$$

Comme nous l'avons vu au § 6.5, nous devons calculer les coefficients de Fourier des fonctions périodiques $\frac{\rho(x) N_{3q}^2(x, x')}{\exp(i\alpha_0(x - x'))}$. Il est facile de voir

(il suffit de poser dans les expressions des noyaux $\sqrt{\frac{\varepsilon_{xx}}{\varepsilon_{yy}}} f(x) = g(x)$ pour s'en convaincre) que cela se ramène à l'étude de séries de la forme :

$$S(x, x') = \frac{1}{2id} \sum_{n \in \mathbb{Z}} \frac{1}{\beta_n} e^{inK(x-x') + i\beta_n|f(x)-f(x')|} \quad (\text{A13.1})$$

$$R(x, x') = \sum_{n \in \mathbb{Z}} \left[\frac{\alpha_n}{\beta_n} f'(x') + \operatorname{sgn}(f(x') - f(x)) \right] e^{inK(x-x') + i\beta_n|f(x)-f(x')|} \quad (\text{A13.2})$$

$$\text{où } \beta_n = \sqrt{k^2 - \alpha_n^2}.$$

Nous nous proposons d'en rechercher les coefficients de Fourier S_{pq} et R_{pq} :

$$S(x, x') = \sum_{p \in \mathbb{Z}} \sum_{q \in \mathbb{Z}} S_{pq} e^{ipKx} e^{-iqKx'} \quad (\text{A13.3})$$

$$R(x, x') = \sum_{p \in \mathbb{Z}} \sum_{q \in \mathbb{Z}} R_{pq} e^{ipKx} e^{-iqKx'}. \quad (\text{A13.4})$$

Nous procédons de la façon suivante :

- a) Extraction de la singularité de $S(x, x')$, qui sera mise sous la forme $S(x, x') = S_s(x, x') + S_r(x, x')$ où $S_r(x, x')$ est régulière et où la partie singulière $S_s(x, x')$ possède des coefficients de Fourier qui s'obtiennent analytiquement,
- b) Recherche de procédés destinés à rendre plus rapide le calcul numérique de $S_r(x, x')$ et $R(x, x')$ pour x et x' donnés,
- c) Obtention des coefficients de Fourier de $S_r(x, x')$ et $R(x, x')$ par transformation de Fourier discrète. Cette étape purement numérique ne sera pas décrite ici (voir par exemple [5]).

Pour simplifier l'écriture, nous supposons par la suite que $d = 2\pi$, d'où $K = 1$, $\alpha_n = \alpha_0 + n$ (on pourra toujours se ramener à cette situation par un changement d'échelle). Il est utile de savoir comment se comporte β_n lorsque $|n| \rightarrow \infty$. Pour les grandes valeurs de $|n|$, $\beta_n = i\sqrt{\alpha_n^2 - k^2}$, d'où l'on déduit :

$$\beta_n = i \left[|n| + \alpha_0 \operatorname{sgn} n - \frac{k^2}{2|n|} + \frac{\alpha_0 k^2}{2n|n|} + o\left(\frac{1}{n^3}\right) \right], \quad (\text{A13.5})$$

où $\operatorname{sgn} n$ représente le signe de n .

2. Extraction de la singularité de $S(x, x')$

Considérons la série :

$$S_s(x, x') = \frac{1}{2id} \sum_{n \in \mathbb{Z}^*} \frac{1}{i|n|} e^{inK(x-x')} = \sum_{n \in \mathbb{Z}^*} \frac{-1}{4\pi|n|} e^{in(x-x')}$$

dont les coefficients de Fourier sont :

$$S_{spq} = \frac{-1}{4\pi|p|} \delta_{pq} \delta_{p0}.$$

La série $S(x, x')$ possède la même singularité que $S_s(x, x')$ qui peut s'obtenir analytiquement :

$$S_s(x, x') = \frac{1}{2\pi} \text{Log} \left[2 \left| \sin \frac{x - x'}{2} \right| \right]$$

3. Etude de $S_r(x, x') = S(x, x') - S_s(x, x')$.

En posant $a = f(x) - f(x')$ et $b = x - x'$, on a :

$$S_r(x, x') = \frac{1}{4i\pi\beta_0} e^{i\beta_0|a|} + \frac{1}{4\pi} \sum_{n \in \mathbb{Z}^*} \left(\frac{e^{i\beta_n|a|}}{i\beta_n} + \frac{1}{|n|} \right) e^{inb}.$$

$$\text{Posons } S_1(x, x') = \sum_{n \in \mathbb{Z}^*} \left(\frac{e^{i\beta_n|a|}}{i\beta_n} + \frac{1}{|n|} \right) e^{inb}.$$

Pour sommer cette série, il est préférable de l'écrire sous la forme :

$$S_1(x, x') = S_2(x, x') + S_3(x, x')$$

où

$$S_2(x, x') = \sum_{n \in \mathbb{Z}^*} \left(\frac{e^{i\beta_n|a|}}{i\beta_n} + \frac{e^{-(|n| + \alpha_0 \text{sgn } n)|a|}}{|n|} \right) e^{inb} = \sum_{n \in \mathbb{Z}^*} u_n(a) e^{inb}$$

$$S_3(x, x') = \sum_{n \in \mathbb{Z}^*} \left(1 - e^{-(|n| + \alpha_0 \text{sgn } n)|a|} \right) \frac{e^{inb}}{|n|} \quad (\text{A13.6})$$

La série S_2 converge bien plus rapidement que la série S_1 . On montre en effet que $u_n(a) = e^{-|na|} O\left(\frac{1}{n^2}\right)$. On en déduit la convergence uniforme de S_2 qui est par conséquent continue. Quant à S_3 , elle s'écrit sous forme analytique en utilisant le fait que (cf. [32], p.67) $\sum_{n \in \mathbb{N}^*} \frac{z^n}{n} = -\text{Log}(1 - z)$ dans tout le disque ouvert $|z| < 1$ en choisissant pour la détermination $\text{Log } 1 = 0$ et pour coupure le demi-axe $\text{Re}(z) \leq 0$ (qui est la détermination

utilisée par la fonction Fortran CLOG). On obtient ainsi, si a et b ne sont pas nuls, c'est à dire si $M' \neq M$ et $M' \neq M_0$ (figure A13.1), ou encore si $x' \neq x$ et $x' \neq x_0$:

$$S_3(x, x') = -\operatorname{Log} \left[\left(2 \sin \frac{b}{2} \right)^2 \right] + e^{-\alpha_0 |a|} \operatorname{Log}(1 - e^{-|a| + i b}) + e^{\alpha_0 |a|} \operatorname{Log}(1 - e^{-|a| - i b}). \quad (\text{A13.7})$$

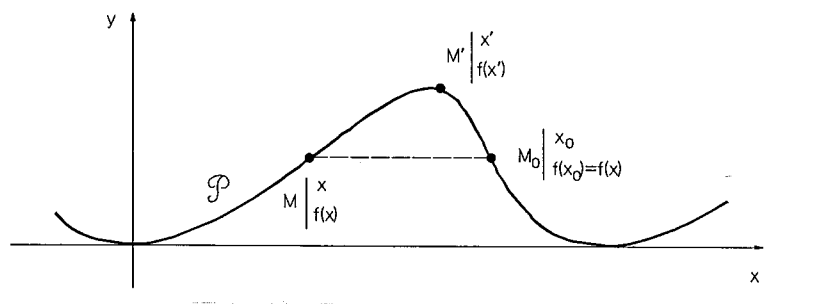


Figure A13.1. M étant fixé sur $\mathcal{P}$, il existe dans la même période du profil un point M_0 de même ordonnée.

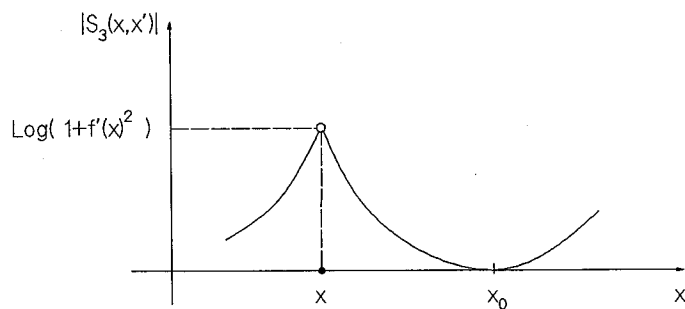
Remarque : $S_3(x, x)$ et $S_3(x, x_0)$ sont nuls (vérification directe en faisant $a = 0$ dans (A13.6)). On peut aussi calculer les limites de $S_3(x, x')$ lorsque $x' \rightarrow x$ et $x' \rightarrow x_0$ au moyen de l'expression (A13.7). On obtient après quelques développements limités fastidieux :

$$S_3(x, x') \xrightarrow{x' \rightarrow x} \operatorname{Log}(1 + f'(x)^2) \quad (\text{A13.8})$$

$$S_3(x, x') \xrightarrow{x' \rightarrow x_0} 0. \quad (\text{A13.9})$$

La distribution $S_3(x, x')$ n'est donc pas continue pour $x' = x$ (Fig. A13.2). C'est bien entendu la limite donnée par (A13.8) qu'il faudra retenir pour poursuivre l'étude numérique (rappelons que $S_3(x, x')$ est un "morceau" du noyau d'une équation intégrale).

Figure A13.2



4. Etude de $R(x, x')$

Le noyau $R(x, x')$ est continu. Son étude a déjà été réalisée par D. Maystre. Pour plus de précisions, se reporter à [37].