

INGENIERIE OPTIQUE

Olivier J.F. Martin

Laboratoire de Nanophotonique et Métrologie
Septembre 2024



Table des matières

1	Optique Géométrique	1
1.1	Introduction	1
1.2	Postulats de l'optique géométrique	1
1.3	Réflexion par un miroir	3
1.4	Réfraction	4
1.5	Milieux inhomogènes et invisibilité	5
1.6	Miroirs	7
1.7	Lentilles	7
1.7.1	Lentilles minces	11
1.8	Tracé des rayons	13
1.9	Diaphragmes	15
1.10	L'oeil comme système optique	16
1.11	Systèmes optiques	18
1.12	Méthode matricielle $ABCD$	21
1.13	Aberrations chromatiques	23
1.14	Micro-optique et Optique diffractive	25
2	Optique Ondulatoire	27
2.1	Introduction	27

2.2	Equation d'onde	28
2.3	Ondes monochromatiques harmoniques	29
2.4	Représentation complexe	30
2.5	Onde plane, sphérique et évanescence	31
2.6	Accumulation de la phase	36
2.7	Interférence	37
2.8	Interféromètres	39
2.9	Interférence d'ondes obliques	40
2.10	Cohérence temporelle	41
2.11	Battements et vitesse de groupe	43
3	Optique de Fourier et diffraction	46
3.1	Principe de Huygens	47
3.2	Propagation dans l'espace libre – Cas 2D	50
3.3	Réseau de diffraction, équation de Bragg	54
3.4	Transformée de Fourier	55
3.5	Propagation dans l'espace libre – Cas général	56
3.6	Transformée de Fourier avec une lentille	63
3.7	Filtrage avec un système $4f$	64
3.8	Intégrale de Rayleigh–Sommerfeld	68
3.9	Approximation de Fraunhofer	69
3.10	Diffraction et transformée de Fourier	70
3.11	Diffraction par une fente	71
3.12	Diffraction par un réseau de fentes	72
3.13	Diffraction par une ouverture circulaire	73

3.14	Holographie	75
4	Optique de Maxwell et polarisation de la lumière	80
4.1	Le champ électromagnétique	80
4.2	Equations de Maxwell	81
4.3	Equation d'onde	83
4.4	Onde plane monochromatique harmonique	84
4.5	Conditions d'interface	86
4.6	Coefficients de Fresnel	87
4.6.1	Polarisation TE ou s	89
4.6.2	Polarisation TM ou p	90
4.7	Réflectance et transmittance	91
4.8	Polarisation de la lumière	92
4.8.1	Polarisation linéaire	95
4.8.2	Polarisation circulaire	96
4.9	Vecteurs et matrices de Jones	97
4.10	Matériaux anisotropes	101
4.11	Propagation le long d'un axe principal	102
4.12	Cristal uniaxial et biréfringence	103
4.13	Activité optique et effet Faraday	108
5	Guides d'ondes	110
5.1	Guide d'onde miroir planaire	111
5.1.1	Vecteurs de propagation	112
5.1.2	Distribution du champ	114
5.1.3	Nombre de modes	117

5.1.4	Relation de dispersion et vitesse de groupe	118
5.1.5	Modes TM	119
5.2	Guide miroir 2D	119
5.3	Guides diélectriques planaires	121
5.4	Fibres optiques	123
5.4.1	Fibres à saut d'indice	125
5.4.2	Fibres à profile d'indice	126
5.4.3	Profile du champ électrique	127
5.4.4	Atténuation	128
6	Photons et transitions optiques	131
6.1	Le photon	131
6.2	Radiation thermique	135
6.3	Niveaux d'énergie	138
6.3.1	Niveaux électroniques, modèle de Bohr	138
6.3.2	Etats vibrationnels et rotationnels	142
6.3.3	Solides et structure de bande	144
6.4	Effets thermiques	146
6.5	Luminescence, fluorescence et phosphorescence	148
6.6	Transitions optiques et coefficients d'Einstein	152
6.7	Colorimétrie	156
7	Lasers	159
7.1	Cavité laser	160
7.2	Amplification laser	162
7.3	Système laser	170

7.4	Différents types de lasers	173
7.4.1	Laser He–Ne	173
7.4.2	Laser à semiconducteur	174
8	Photodétecteurs	180
8.1	Effet photoélectrique	180
8.2	Photoconducteurs	184
8.3	Photodiodes	186

Chapitre 1

Optique Géométrique

1.1 Introduction

L'optique géométrique propose la description la plus élémentaire des phénomènes optiques. La lumière s'y propage sous forme de rayons et suit quelques règles simples décrivant le passage d'un milieu à l'autre. Bien que l'optique géométrique représente une approximation, elle permet de comprendre nombre de phénomènes rencontrés dans la vie courante. En particulier, l'optique géométrique permet de décrire la formation d'images avec des systèmes optiques complexes comprenant des éléments tels que miroirs et lentilles.

L'optique géométrique a cependant aussi ses limites et ne permet pas de d'étudier les phénomènes importants qui découlent du caractère ondulatoire de la lumière (interférence) ou dépendent des propriétés associées au champ électromagnétique (polarisation).

1.2 Postulats de l'optique géométrique

L'optique géométrique repose sur les quatre principes fondamentaux suivants.

Rayon La lumière se propage sous forme de rayons. Ceux-ci sont émis par une source et peuvent être mesurés par un détecteur.

Indice de réfraction Tout milieu est caractérisé par un indice de réfraction n (en général, $n \geq 1$ mais il existe des matériaux artificiels où n peut prendre une valeur plus petite que 1, voire une valeur négative). Pour la plupart des matériaux, l'indice de réfraction est complexe : $\tilde{n} = n + j\kappa$, la partie imaginaire étant liée à l'absorption dans le matériau. L'indice de réfraction correspond au rapport $n = c_0/c$ où c_0 est la vitesse de la lumière dans le vide et c la vitesse de la lumière dans le milieu. Le temps mis par la lumière pour parcourir une distance d est $d/c = nd/c_0$; ce temps est proportionnel au produit nd que l'on appelle le chemin optique L .

Chemin optique le long d'une courbe quelconque Dans un matériau inhomogène, l'indice de réfraction $n(\mathbf{r})$ dépend de la position $\mathbf{r} = (x, y, z)$ dans le milieu. Ainsi le chemin optique entre les points A et B s'obtient-il par intégration :

$$L_{AB} = \int_A^B n(\mathbf{r}) ds, \quad (1.1)$$

où l'on a introduit l'élément de chemin infinitésimal ds . Le temps mis par la lumière pour aller de A à B dépend du chemin optique.

Principe de Fermat Un rayon optique allant du point A au point B suit un chemin tel que le temps mis par la lumière est un extremum, par rapport aux chemins voisins. Mathématiquement, ce principe s'écrit

$$\delta \int_A^B n(\mathbf{r}) ds = 0, \quad (1.2)$$

où le symbole δ représente la variation du chemin optique. Remarquons qu'en mathématique un extremum peut être un minimum, un maximum ou un point d'inflexion (la dérivée est nulle pour chacun de ces cas) ; en optique cependant, la lumière choisit en général le chemin le plus court et on peut énoncer que la lumière se propage selon le chemin le plus rapide.

Le principe de Fermat a quelques conséquences importantes. Ainsi la lumière se propage-t-elle en ligne droite dans un milieu homogène, puisque la ligne droite est le chemin le plus rapide entre les points A et B .

Le principe de Fermat indique aussi que (dans un milieu homogène $n(\mathbf{r}) = n$) le chemin optique est le même, que la lumière se déplace de A vers B ou en sens inverse de B vers A . Ceci se démontre en écrivant simplement,

$$L_{AB} = \int_A^B n ds = \int_B^A n(-ds) = \int_B^A n ds' = L_{BA}, \quad (1.3)$$

où $ds' = -ds$ est l'élément de chemin orienté de B vers A . Il faut cependant se garder d'utiliser ce principe de retour inverse dans des situations où la polarisation de la lumière joue un rôle important. Ainsi existe-t-il des composants appelés "isolateurs optiques" qui ne laissent passer la lumière que dans un sens. Ces composants jouent un rôle important pour beaucoup d'expériences pratiques ; ils permettent par exemple d'éviter un retour de la lumière dans un laser, ce qui pourrait le déstabiliser. D'une façon générale, Eq. 1.3 est apparentée au principe de réciprocité qui dit que si la lumière prend un certain chemin de A à B , alors elle prendra le chemin inverse, de B à A si on inverse la flèche du temps.

La vitesse de la lumière dans le vide c_0 qui permet de définir l'indice de réfraction joue un rôle essentiel en optique. Nous verrons qu'elle se déduit des équations de Maxwell ; pour l'instant contentons-nous de noter qu'elle est indépendante de la direction de propagation et vaut approximativement $c_0 \simeq 2.998 \cdot 10^8$ m/s. Si la vitesse de la lumière semble extrêmement grande, on notera qu'il est possible avec un laser de créer des pulses très courts, de l'ordre de quelques fs = 10^{-15} s ; pendant un tel intervalle de temps la lumière ne parcourt qu'environ

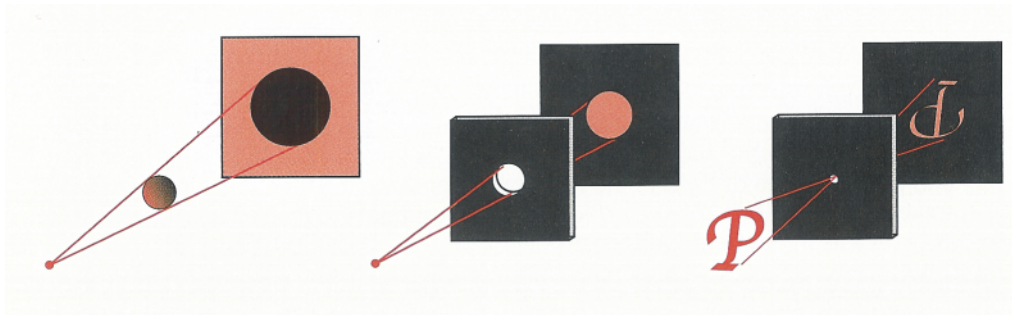


FIGURE 1.1 – Dans un milieu homogène, les rayons lumineux se propagent en ligne droite ; les ombres sont des projections parfaites des obstacles, d'après B.E.A Saleh et M.C. Teich, *Fundamentals of photonics*, 2nd Ed. (Wiley, Hoboken, 2007).

$0.3\,\mu\text{m}$. Des expériences utilisant de tels pulses laser permettent donc de sonder la matière à une très petite échelle.

La définition de l'indice de réfraction nous permet de déterminer la vitesse à laquelle se propage la lumière dans un matériau. Ainsi dans un morceau de verre d'indice $n = 1.5$, la lumière ne se déplace qu'à la vitesse $c = c_0/n \simeq 2.0 \cdot 10^8\text{ m/s}$ et dans du silicium aux fréquences correspondant à l'infrarouge proche utilisé pour les télécommunications (longueur d'onde $\lambda = 1.5\,\mu\text{m}$, $n = 3.5$) $c \simeq 0.86 \cdot 10^8\text{ m/s}$.

La figure 1.1 illustre la propagation de la lumière en ligne droite dans un milieu homogène.

Nous allons maintenant utiliser le principe de Fermat pour déterminer les lois de la réflexion et de la réfraction.

1.3 Réflexion par un miroir

Un miroir peut être fabriqué à l'aide d'un métal poli ou de couches minces diélectriques ou métalliques déposées sur un substrat. On peut aussi fabriquer des miroirs à l'aide d'une série de couches diélectriques successives d'indices de réfraction élevés et bas. Il faut garder à l'esprit qu'un miroir hautement réfléchissant pour la lumière visible peut se comporter comme un absorbant à d'autres longueurs d'ondes, par exemple dans l'infrarouge.

Le rayon réfléchi suit la loi suivante.

Loi de la réflexion Le rayon réfléchi se trouve dans le plan d'incidence ; l'angle de réflexion θ' est égal à l'angle d'incidence θ .

Pour démontrer l'égalité entre angles incidents et réfléchis, considérons la figure 1.2(a). Le rayon allant de A à C après réflexion sur le miroir en B effectue un chemin optique similaire au rayon $AB + BC'$ (on suppose le miroir infiniment mince). Ce rayon équivalent se propage dans un milieu homogène et donc la lumière prend le chemin le plus court (ligne droite) :

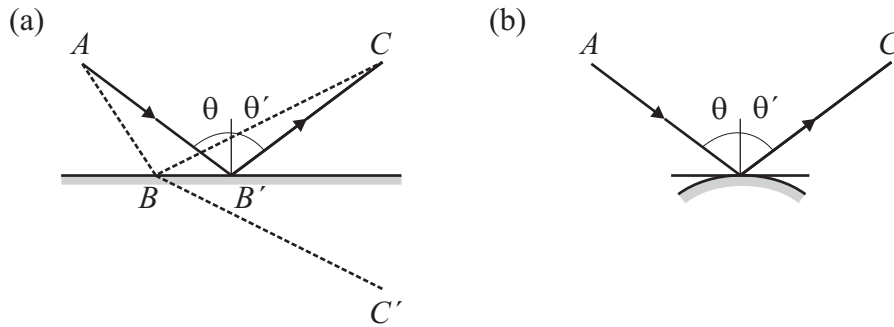


FIGURE 1.2 – Réflexion par un miroir. (a) Méthode de construction. (b) Miroir incurvé.

$AB' + B'C'$, Fig. 1.2(a). La réflexion en B' détermine l'égalité entre les angles θ et θ' .

Si la surface du miroir est courbe, on doit considérer la tangente à sa surface, comme indiqué Fig. 1.2(b).

1.4 Réfraction

Lorsque la lumière passe d'un milieu n_1 à un milieu n_2 , le rayon incident est transmis en suivant la loi suivante.

Loi de la réfraction (loi de Snell) Le rayon transmis (ou réfracté) se trouve dans le plan d'incidence ; l'angle de réfraction θ_2 dépend de l'angle d'incidence θ_1 et des indices des deux milieux selon la loi de Snell :

$$n_1 \sin \theta_1 = n_2 \sin \theta_2 . \quad (1.4)$$

La figure 1.3 illustre une méthode de construction simple pour obtenir l'angle de réfraction θ_2 . On commence par construire deux cercles concentriques centrés au point d'incidence B

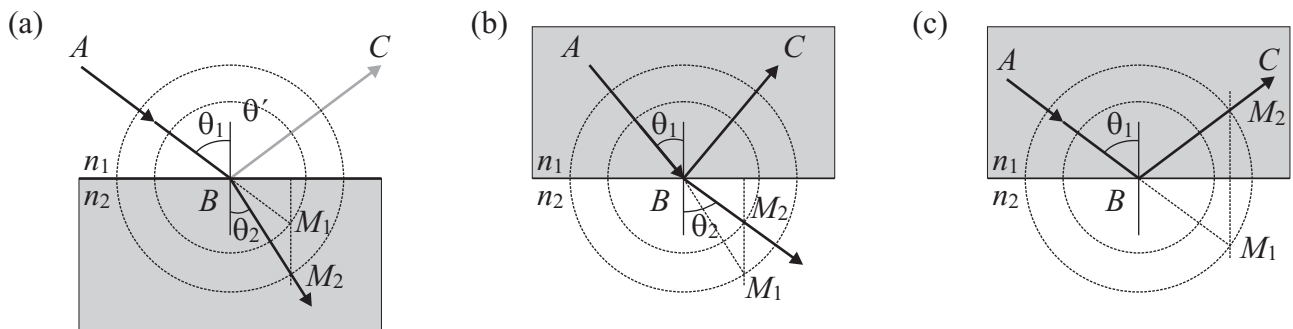


FIGURE 1.3 – Construction de la réfraction de la lumière à l'interface entre deux milieux n_1 et n_2 . (a) $n_1 < n_2$, (b) $n_1 > n_2$ et (c) $n_1 > n_2$ dans le cas de la réflexion total interne.

avec comme rayon respectif n_1 et n_2 . On prolonge le rayon incident (qui se propage donc dans le milieu n_1 en sorte qu'il intersecte au point M_1 le cercle de rayon n_1 dans le milieu n_2 . Pour obtenir le rayon réfracté dans le milieu n_2 , on projette ce point M_1 sur le deuxième cercle de rayon n_2 et on obtient le point M_2 par lequel passe le rayon réfracté, Fig. 1.3(a).

La procédure est la même lorsque $n_1 > n_2$, Fig. 1.3(b). Dans ce cas, on remarque qu'il existe un angle critique θ_c au-delà duquel le rayon réfracté n'existe plus. Ceci se traduit par le fait que lorsque l'on projette le point M_1 on obtient un point M_2 se trouvant dans le milieu incident, Fig. 1.3(c). Le rayon incident est alors entièrement réfléchi et il n'y a pas de rayon réfracté. On parle de réflexion interne totale. On remarque que l'angle de réfraction est toujours plus grand que l'angle incident lorsque $n_1 > n_2$. Le cas limite de réflexion interne totale s'obtient lorsque $\theta_2 = \pi/2$; cette condition donne la valeur de l'angle critique θ_c à l'aide d'Eq. (1.4) :

$$\theta_c = \arcsin(n_2/n_1). \quad (1.5)$$

Le rayon incident est entièrement réfléchi lorsque $\theta_1 > \theta_c$. Ce phénomène est à l'origine de l'optique guidée (guides d'ondes, fibres optiques), dans laquelle un rayon est guidé par une succession de réflexions internes totales. Les résultats de cette section nous indiquent donc qu'un guide d'onde nécessite un matériau d'indice plus élevé que son environnement.

Nous verrons que la loi de Snell peut se déduire des équations de Maxwell ; elle est l'expression d'une loi plus fondamentale : la conservation de la quantité de mouvement parallèle à l'interface pour les ondes incidentes, réfléchies et réfractées. Notons finalement que l'optique géométrique ne permet pas de déterminer quelle fraction de l'onde incidence est réfléchi et quelle fraction est réfractée ; seuls les angles des rayons peuvent être déterminés.

1.5 Milieux inhomogènes et invisibilité

Considérons maintenant un milieu inhomogène, comme par exemple une cuve optique remplie de glycérine ($n = 1.47$) et d'eau ($n = 1.33$), Fig. 1.4. Si on ne mélange pas les deux liquides,



FIGURE 1.4 – Lorsqu'il existe un gradient d'indice de réfraction (milieu inhomogène) la lumière ne se propage plus en ligne droite. Tel est le cas dans une cuve remplie de glycérine ($n = 1.47$) et d'eau ($n = 1.33$).

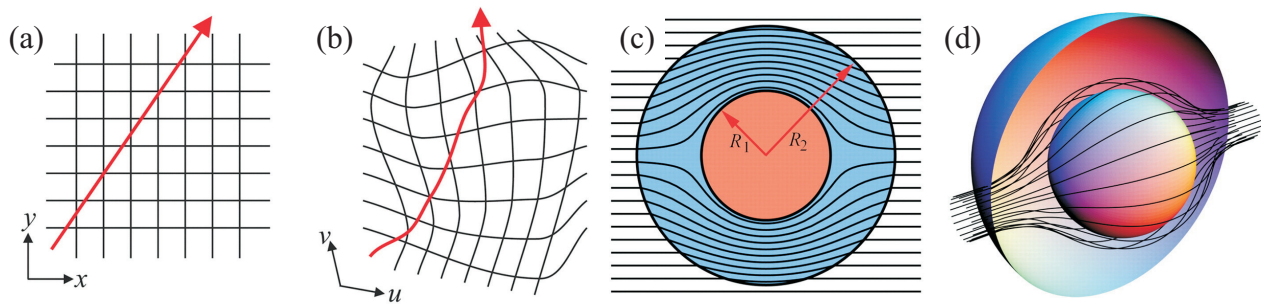


FIGURE 1.5 – (a) Dans un espace homogène la lumière se propage en ligne droite. (b) Tel n'est pas le cas dans un espace inhomogène. (c) En déformant l'espace, on peut forcer les rayons lumineux à éviter une région spécifique, dans laquelle on peut cacher un objet (d). D'après J.B. Pendry et al., *Science* **312**, 1780 (2006).

l'eau plus légère monte vers la surface. Il se crée donc un gradient d'indice de réfraction avec un indice élevé au fond de la cuve (glycérine) et un indice plus petit vers le haut de la cuve (eau). Pour tracer les rayons dans ce système, on peut le découper en tranches fines d'indice constant. Ainsi un rayon est-il réfracté de couche en couche lorsqu'il se propage dans ce milieu et décrit une trajectoire courbe, Fig. 1.4. On remarque que le rayon est toujours infléchi vers la direction du gradient d'indice. On peut formaliser ce comportement en introduisant l'équation eikonale qui décrit le trajet de la lumière.

Au tournant des années 2000, une nouvelle discipline de l'optique est apparue, basée sur la transformation de l'espace. Un milieu homogène d'indice n constant peut ainsi être représenté par un maillage régulier, Fig. 1.5(a). La lumière s'y propage en ligne droite. Le gradient d'indice dans la cuve Fig. 1.4 serait représenté par des lignes plus serrées dans la direction verticale. La figure 1.5(b) représente un matériau totalement inhomogène, avec des variations de l'indice dans toutes les directions. La lumière y prend un chemin compliqué pour aller des mêmes points que dans le cas de Fig. 1.5(a).

En déformant ainsi l'espace, on peut faire en sorte que les rayons lumineux évitent une région particulière, comme la partie centrale de Fig. 1.5(c). Il est alors possible de “cacher” un objet dans cette partie de l'espace, qui échappera totalement à un observateur extérieur, Fig. 1.5(d). Pour celui-ci, il n'est pas possible de savoir que les rayons ont pris un chemin compliqué autour de l'objet.

La réalisation pratique d'une telle cache est cependant compliquée et nécessite de contrôler l'indice de réfraction dans l'espace. De plus, l'indice de réfraction prend alors des valeurs complexes et l'introduction d'une partie imaginaire pour n donne lieu à des pertes par absorption de la lumière. Il est aussi difficile de réaliser de tels matériaux qui fonctionnent sur une large partie du spectre et puissent servir de cache optique pour toutes les longueurs d'onde. A ce jour, ce phénomène n'a été réalisé qu'à des fréquences microondes avec des matériaux artificiels qu'on nomme “métamatériaux”, parce qu'ils sont constitués de très petits éléments créant des indices de réfraction exotiques.

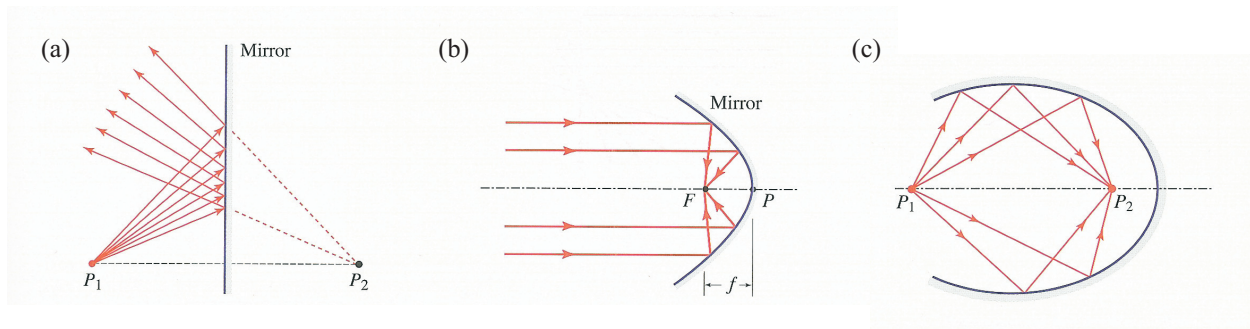


FIGURE 1.6 – Différents types de miroirs, d'après d'après B.E.A Saleh et M.C. Teich, *Fundamentals of photonics*, 2nd Ed. (Wiley, Hoboken, 2007).

1.6 Miroirs

Dans la suite de ce chapitre, nous allons nous concentrer sur des systèmes paraxiaux, c'est à dire que nous allons définir un axe optique et considérer les rayons qui se propagent avec des angles proches de cet axe optique.

Il existe une grande variété de miroirs. Un miroir plan réfléchit la lumière en sorte que les rayons issus de la source S_1 semblent provenir de son image S_2 , Fig. 1.6(a). Lorsque la surface du miroir est parabolique, tous les rayons parallèles à l'axe optique sont réfléchis en un seul point appelé foyer, Fig. 1.6(b). Pour un miroir elliptique, tous les rayons issus d'un foyer forment une image à l'autre foyer, Fig. 1.6(c). Le miroir elliptique donne lieu à un paradoxe important : si l'on place une source ponctuelle en un foyer, est-ce que l'on obtient une image ponctuelle à l'autre foyer ? Tel n'est pas le cas, car le champ évanescent associé à la source décroît rapidement et n'atteint pas l'image. Or ce champ évanescent est essentiel pour créer une image ponctuelle. On peut comprendre ceci en faisant une analogie avec la transformée de Fourier : le spectre d'une fonction de Dirac (source ponctuelle) est infini et pour recréer une source ponctuelle depuis son spectre, il faut utiliser toutes les composantes du spectre. En optique, une partie de ces composantes (le champ évanescent) ne se propagent pas et disparaît avant d'atteindre l'image. Celle-ci n'est donc pas ponctuelle et a une taille plus grande que la source.

1.7 Lentilles

Une lentille est un élément par lequel les rayons incidents sont réfractés en suite d'un changement d'indice du milieu. Une lentille modifie donc par transmission la géométrie d'un faisceau incident. Notons qu'il existe des lentilles pour toutes sortes d'ondes : optiques (de l'ultraviolet aux ondes radio), de pression (acoustiques) ; il existe aussi des lentilles pour les électrons – elles sont utilisées dans les microscopes électroniques – et en astrophysique on

rencontre des lentilles gravitationnelles créées par une masse importante dans l'univers dont la présence peut modifier le parcours de la lumière émise par une étoile.

Comme pour le miroir, la forme de la surface de la lentille joue un rôle déterminant dans la transformation des fronts d'onde, comme illustré Fig. 1.7. Bien que les surfaces hyperboloïdales et ellipsoïdales aient des propriétés très intéressantes, dans la pratique la plupart des lentilles sont réalisées à partir de surfaces sphériques. De telles surfaces sont en effet plus aisées à obtenir par polissage sur une forme.

La figure 1.8 montre une lentille sphérique convergente. Comme ses surfaces sont sphériques, la lentille est caractérisée par les rayons R_1 et R_2 . Il faut introduire un certain nombre de conventions afin de déterminer les caractéristiques d'une lentille. Dans la suite, on supposera toujours que la lumière est incidente depuis la gauche.

Pour mieux comprendre l'interaction de la lumière avec les lentilles, commençons par considérer le dioptré sphérique de la figure 1.9. Une source lumineuse ponctuelle S se trouvant sur l'axe du dioptré émet de la lumière. Considérons le rayon atteignant le dioptré au point A . Ce rayon est réfracté et croise l'axe optique au point P . On dit que les points S et P sont conjugués. On définit les distances $s_o = SV$ comme distance objet et $s_i = VP$ comme distance image. Le chemin optique pour ce rayon vaut donc $L_{SP} = n_1 l_o + n_2 l_i$. En utilisant le théorème du cosinus, on obtient les longueurs l_o et l_i :

$$l_o = [R^2 + (s_o + R)^2 - 2R(s_o + R) \cos \phi]^{1/2}, \quad (1.6)$$

$$l_i = [R^2 + (s_i - R)^2 + 2R(s_i - R) \cos \phi]^{1/2}. \quad (1.7)$$

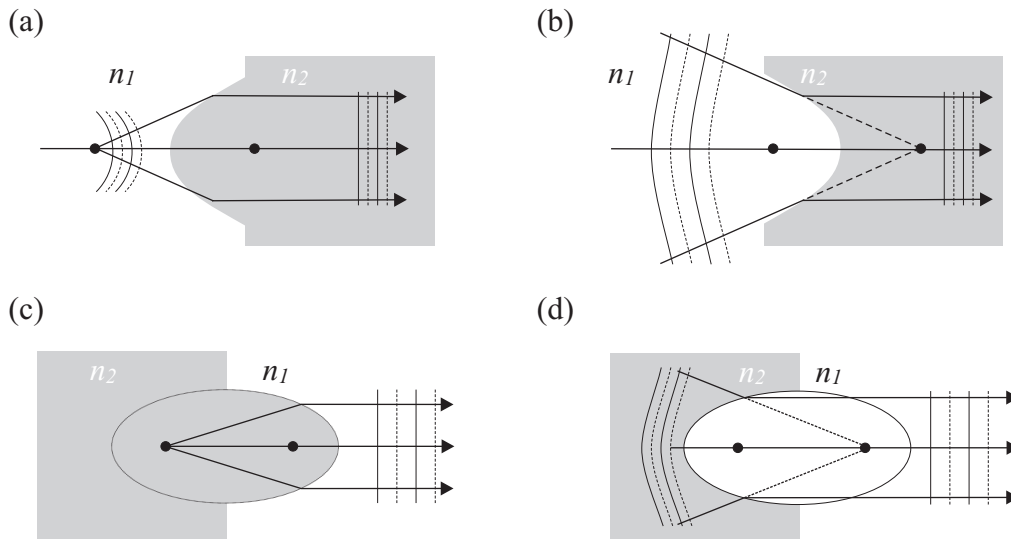


FIGURE 1.7 – Vue en coupe de différentes surfaces réfringentes : (a) et (b) : hyperboloïdales ; (c) et (d) ellipsoïdales. Remarquer comme le front d'onde est transformé par la réfraction à-travers la surface.

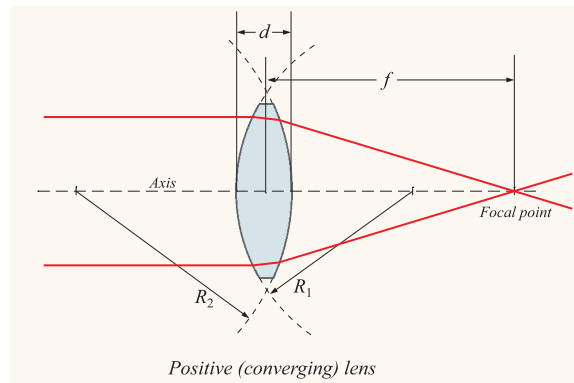


FIGURE 1.8 – Lentille sphérique convergente avec ses différent paramètres. D’après [www.wikipedia.org/wiki/Lens\(optics\)](http://www.wikipedia.org/wiki/Lens(optics)).

TABLE 1.1 – Convention de signe pour les dioptries sphériques (la lumière vient de la gauche)

Symbole	Convention
s_o, f_o	positif à gauche de V
s_i, f_i	positif à droite de V
R	positif si C est à droite de V

Ainsi le chemin optique s’écrit,

$$L_{SP} = n_1 [R^2 + (s_o + R)^2 - 2R(s_o + R) \cos \phi]^{1/2} + n_2 [R^2 + (s_i - R)^2 + 2R(s_i - R) \cos \phi]^{1/2}. \quad (1.8)$$

Tous les paramètres définis dans Fig. 1.9 sont positifs. Ceci découle des conventions qui déterminent le signe par rapport à la position des différent éléments, comme indiqué Table 1.1.

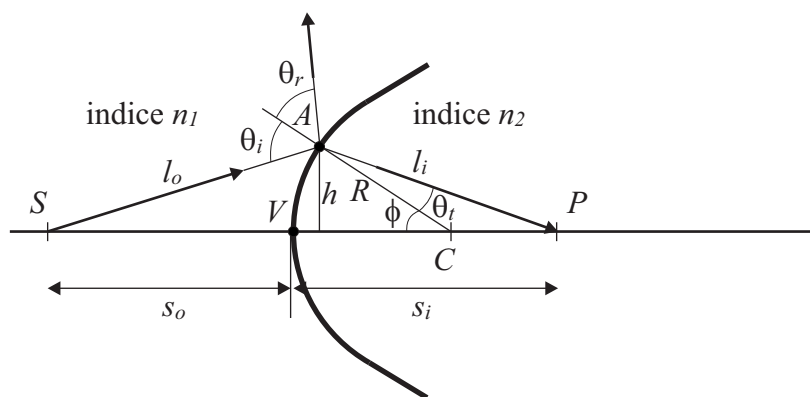


FIGURE 1.9 – Réfraction sur une surface sphérique.

L'angle ϕ détermine la position du point A et donc détermine le chemin optique. Le principe de Fermat requiert que la lumière choisisse un chemin qui corresponde à un extremum par rapport aux chemins voisins, i.e. $dL_{SP}/d\phi = 0$; il en découle,

$$\frac{n_1 R(s_o + R) \sin \phi}{2l_o} - \frac{n_2 R(s_i - R) \sin \phi}{2l_i} = 0, \quad (1.9)$$

d'où finalement

$$\frac{n_1}{l_o} + \frac{n_2}{l_i} = \frac{1}{R} \left(\frac{n_2 s_i}{l_i} - \frac{n_1 s_o}{l_o} \right). \quad (1.10)$$

La relation (1.10) est exacte, mais assez compliquée. Elle détermine les paramètres requis pour que la lumière aille du point S au point P par réfraction à travers la surface d'un dioptré sphérique de rayon R . Pour la suite, nous allons introduire une approximation essentielle, qui consiste à se limiter aux rayons proches de l'axe optique, i.e. A proche de V sur Fig. 1.9. Dans ce cas, l'angle ϕ est petit et on peut faire les développements en série suivants :

$$\begin{aligned} \cos \phi &\simeq 1 - \frac{\phi^2}{2!} + \frac{\phi^4}{4!} - \frac{\phi^6}{6!} + \dots \\ \sin \phi &\simeq \phi - \frac{\phi^3}{3!} + \frac{\phi^5}{5!} - \frac{\phi^7}{7!} + \dots \end{aligned} \quad (1.11)$$

En ne prenant que le premier terme, on peut approximer $\cos \phi \simeq 1$ ce qui implique $l_o \simeq s_o$ et $l_i \simeq s_i$, Fig. 1.9. L'équation (1.10) devient alors,

$$\frac{n_1}{s_o} + \frac{n_2}{s_i} = \frac{n_2 - n_1}{R}. \quad (1.12)$$

Cette formule est plus simple et permet de déterminer des situations importantes dans la pratique. Les rayons issus de S qui satisfont cette approximation (ϕ petit) sont appelés rayons paraxiaux puisqu'ils se propagent près de l'axe optique du système. La description détaillée de ces rayons fut faite en premier par Gauss en 1841. On parle donc souvent d'optique gaussienne ou d'optique de premier ordre puisque seuls les premiers termes sont retenus dans Eq. (1.11).

Si l'image de la source se trouvant à la distance s_o est à l'infini ($s_i = \infty$), on a

$$\frac{n_1}{s_o} + \frac{n_2}{\infty} = \frac{n_2 - n_1}{R}. \quad (1.13)$$

Cette distance objet s_o particulière s'appelle la distance focale objet f_o :

$$f_o = \frac{n_1}{n_2 - n_1} R. \quad (1.14)$$

Le point se trouvant à cette distance est le foyer objet F_o . On définit de même le foyer image F_i , point de l'axe optique où se forme l'image quand la source est à l'infini ($s_o = \infty$). On en déduit la distance focale image f_i :

$$f_i = \frac{n_2}{n_2 - n_1} R. \quad (1.15)$$

1.7.1 Lentilles minces

Considérons de nouveau le problème de la réfraction et varions la distance de l'image, Fig. 1.10 en utilisant Eq. (1.12) pour déterminer s_o et s_i . Si s_o est grand, s_i est petit, Fig. 1.10(a). Si s_o diminue, le point conjugué P s'éloigne, jusqu'au moment où $s_o = f_o$ et donc $s_i = \infty$, Fig. 1.10(b). Pour ce point, on a $n_1/s_o = (n_2 - n_1)/R$. Si la source se rapproche encore du dioptre (s_o diminue), s_i doit devenir négatif pour respecter Eq. (1.12). On parle alors d'une image virtuelle P' , Fig. 1.10(c).

Si l'on ajoute maintenant une seconde surface sphérique au système décrit par Fig. 1.10(c), on obtient une lentille d'indice n_l entourée d'un milieu homogène n_m , comme représentée Fig. 1.11. En utilisant Eq. (1.12) on a la relation suivante :

$$\frac{n_m}{s_{o1}} + \frac{n_l}{s_{i1}} = \frac{n_l - n_m}{R_1}, \quad (1.16)$$

qui donne la distance s_{i1} entre V_1 et le point P' où se croisent les rayons paraxiaux issus du point S . Ce point P' sert maintenant de nouveau point objet pour le second dioptre. Il est

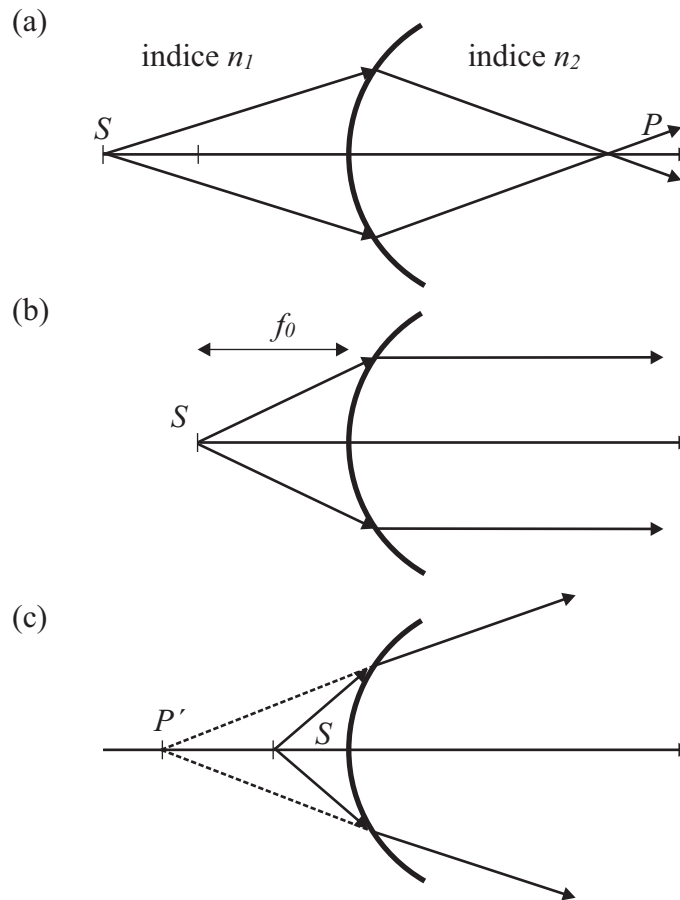


FIGURE 1.10 – Rayons réfractés sur un dioptre sphérique en fonction de la distance de la source.

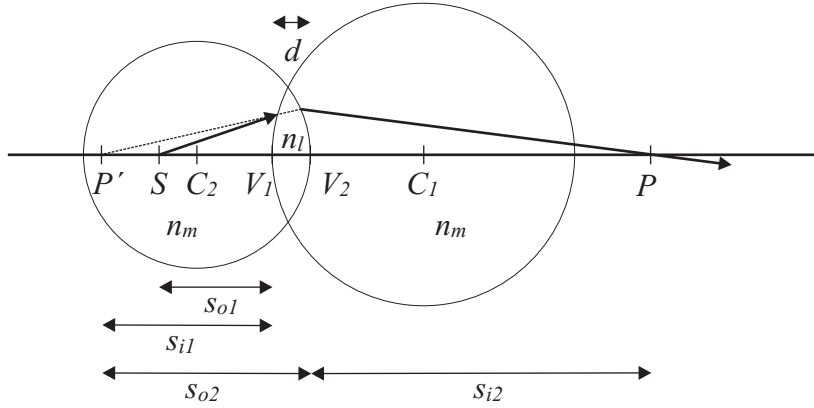


FIGURE 1.11 – Lentille sphérique.

situé à une distance s_{o2} du sommet correspondant V_2 et les rayons issus de P' évoluent dans un milieu d'indice n_l . En utilisant les conventions de signe, on a de plus $s_{o2} = -s_{i1} + d$ et Eq. (1.12) appliquée à la seconde surface donne,

$$\frac{n_l}{-s_{i1} + d} + \frac{n_m}{s_{i2}} = \frac{n_m - n_l}{R_2}, \quad (1.17)$$

En additionnant Eqs. (1.13) et (1.14) on obtient l'équation générale,

$$\frac{n_m}{s_{o1}} + \frac{n_m}{s_{i2}} = (n_l - n_m) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) + \frac{n_l d}{(s_{i1} - d)s_{i1}}. \quad (1.18)$$

Si la lentille est suffisamment mince ($d \rightarrow 0$), le dernier terme est nul. De plus, si le milieu environnant est de l'air ($n_m = 1$), Eq. (1.18) donne la formule des lentilles minces, ou des lunettes :

$$\frac{1}{s_o} + \frac{1}{s_i} = (n_l - 1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right), \quad (1.19)$$

où nous avons posé $s_{o1} = s_o$ et $s_{i2} = s_i$, puisque les points V_1 et V_2 se superposent lorsque $d \rightarrow 0$. La Fig. 1.12 indique les différents types de lentilles qui peuvent être ainsi réalisées en variant le signe de R_1 et de R_2 .

Pour un dioptré sphérique avec une seule interface, nous avons vu que pour s_o situé à l'infini, la distance image devient la distance focale image f_i et de même pour s_i situé à l'infini :

$$\begin{aligned} \lim_{s_o \rightarrow \infty} s_i &= f_i \\ \lim_{s_i \rightarrow \infty} s_o &= f_o. \end{aligned} \quad (1.20)$$

En introduisant successivement ces deux limites dans Eq. (1.19) on remarque que $f_i = f_o = f$ pour une lentille mince et on peut écrire :

$$\frac{1}{f} = (n_l - 1) \left(\frac{1}{R_1} - \frac{1}{R_2} \right) \quad (1.21)$$

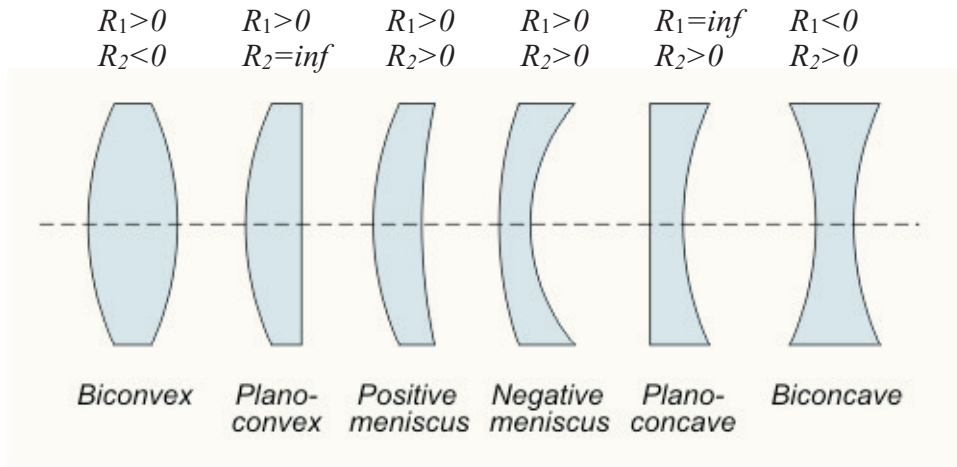


FIGURE 1.12 – Différents types de lentilles simples réalisées avec des surfaces sphériques concentriques ainsi que les valeurs des rayons R_1 et R_2 .

et

$$\frac{1}{s_o} + \frac{1}{s_i} = \frac{1}{f}. \quad (1.22)$$

Equation (1.22) est la relation de conjugaison des lentilles simples dans l'air.

Ces expressions sont utiles pour calculer la focale d'une lentille ainsi que pour déterminer la position d'une image en connaissant celle d'un objet. Ainsi, pour une lentille plan-convexe de rayon de courbure 50 mm et d'indice $n_l = 1.5$, si on considère que la lumière rencontre d'abord la face plane ($R_1 = \infty$, $R_2 = -50$), on a

$$\frac{1}{f} = (1.5 - 1) \left(\frac{1}{\infty} - \frac{1}{-50} \right) = 0.01 \text{ mm}^{-1}. \quad (1.23)$$

La focale est donc $f = 100$ mm. Avec cette lentille, la distance s_i de l'image d'un objet se trouvant à une distance $s_o = 600$ mm s'obtient en utilisant Eq. (1.22) :

$$s_i = \frac{s_o f}{s_o - f} = 120 \text{ mm}. \quad (1.24)$$

1.8 Tracé des rayons

Le tracé des rayons permet de déterminer les images pour des objets étendus. Plutôt que de considérer une source ponctuelle, on étudie ici des objets ayant une taille finie. On se limitera ici aux objets normaux à l'axe optique, dont l'image est aussi normale à l'axe optique.

La figure 1.13 indique la construction des images pour (a) une lentille convergente et (b) une lentille divergente. On notera que dans le premier cas l'image est renversée, alors qu'elle est

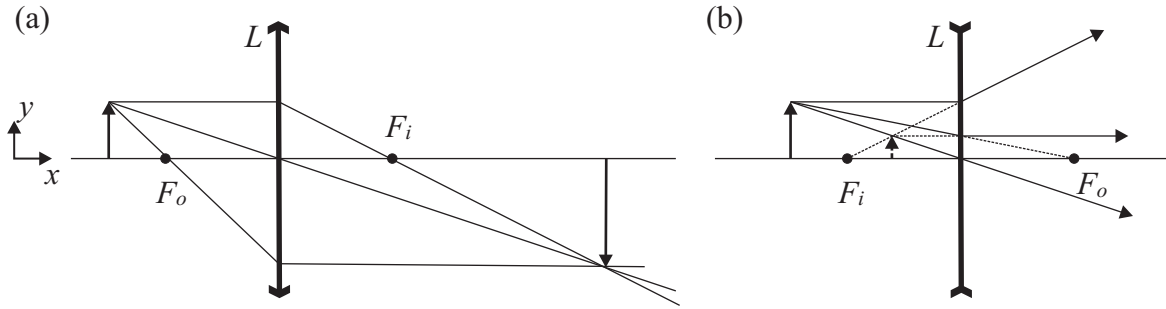


FIGURE 1.13 – Objet réel donnant (a) une image réelle avec une lentille convergente ; (b) une image virtuelle avec une lentille divergente.

droite (ou redressée) dans le second. Différents rapports peuvent être déterminés pour les positions des sources et objets. On en déduit la relation de conjugaison des positions :

$$x_o x_i = f^2. \quad (1.25)$$

Il faut prendre garde dans toutes ces formules d'utiliser les conventions de signe des Tables 1.1 et 1.2. On peut retenir que l'objet et l'image doivent être situés des côtés opposés par rapport à leurs foyers respectifs (par exemple l'objet à gauche de F_o et l'image à droite de F_i).

On définit aussi le grandissement transverse

$$M_T = \frac{y_i}{y_o} = -\frac{s_i}{s_o}; \quad (1.26)$$

dont la valeur est positive si l'image est droite et négative si l'image est renversée. On peut combiner les résultats précédents pour obtenir la relation de grandissement de Newton :

$$M_T = -\frac{x_i}{f} = -\frac{f}{x_o}. \quad (1.27)$$

En traçant les rayons on peut construire les images d'objets pour des systèmes arbitraires. La Table 1.3 indique les conventions de signe pour l'objet et son image. La Table 1.4 donne quant à elle le type d'image en fonction de la lentille utilisée et de la position de l'objet.

TABLE 1.2 – Conventions de signes supplémentaires pour le tracé des rayons (la lumière vient de la gauche)

Symbole	Convention
x_o	positif à gauche de F_o
x_i	positif à droite de F_i
y_o, y_i	positif au dessus de l'axe optique

TABLE 1.3 – Conventions de signes pour les paramètres des lentilles minces

Quantité	Signe +	Signe -
s_o	objet réel	objet virtuel
s_i	image réelle	image virtuelle
f	lentille convergente	lentille divergente
y_o	objet droit	objet renversé
y_i	image droite	image renversés
M_T	image droite	image renversée

TABLE 1.4 – Images d'objets réels formées par des lentilles minces

Lentille convexe				
objet position	image type	position	orientation	taille
$\infty > s_o > 2f$	réelle	$f < s_i < 2f$	renversée	réduite
$s_o = 2f$	réelle	$s_i = 2f$	renversée	identique
$f < s_o < 2f$	réelle	$\infty > s_i > 2f$	renversée	agrandie
$s_o = f$		$\pm\infty$		
$s_o < f$	virtuelle	$ s_i > s_o$	droite	agrandie
Lentille concave				
objet position	image type	position	orientation	taille
arbitraire	virtuelle	$ s_i < f , s_o > s_i $	droite	réduite

Notons finalement que si les rayons incidents sont parallèles (faisceau collimaté), l'image se fait dans le plan focal de la lentille.

1.9 Diaphragmes

Quel que soit le système optique, son extension latérale est finie : les lentilles ont un diamètre fini et sont placées dans des supports qui ne laissent pas passer la lumière. Ceci limite donc les rayons qui peuvent être transmis à travers le système. Ces limitations sont importantes pour l'utilisation pratique du système. Il faut aussi noter qu'en limitant les rayons incidents à de faibles angles, on satisfait mieux la condition de paraxialité sur laquelle repose les développements de l'optique géométrique.

On appelle diaphragme d'ouverture (en anglais *aperture stop*, AS) tout objet limitant la quantité de lumière contribuant à l'image finale. C'est par exemple le cas du diaphragme à iris utilisé dans les appareils photographiques réflex. L'élément qui limite la taille de l'image

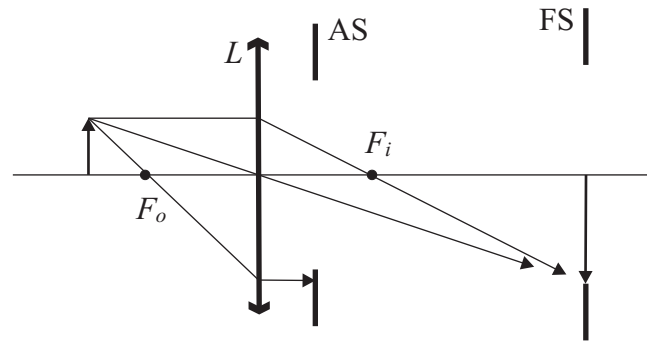


FIGURE 1.14 – Diaphragme d'ouverture (*aperture stop*, AS) et diaphragme de champ (*field stop*, FS).

dans le plan d'observation est appelé diaphragme de champ (en anglais *field stop*, FS). Ainsi le diaphragme d'ouverture contrôle le nombre de rayons qui atteignent le point image d'un point objet et le diaphragme de champ détermine quels points image sont vus ou non.

1.10 L'oeil comme système optique

L'oeil humain peut être considéré comme une association de deux lentilles convergentes qui produit une image réelle sur la rétine, Fig. 1.15. La lumière pénètre par la cornée qui constitue l'élément le plus convergent du système optique oculaire. Elle passe ensuite par la pupille qui fait office de diaphragme et contrôle la quantité de lumière. Finalement le cristallin est la lentille qui fait la mise au point grâce à son élasticité qui permet de modifier sa courbure

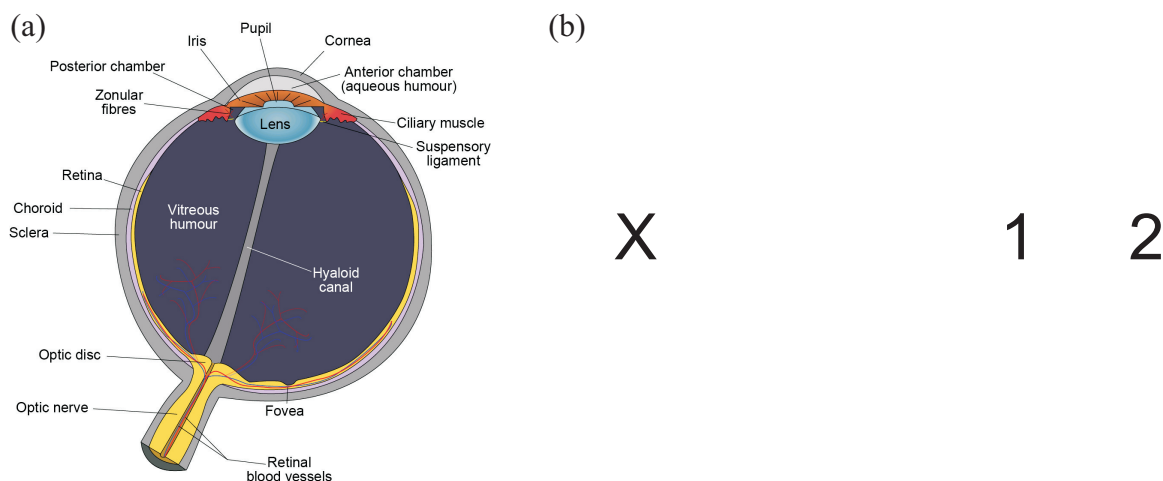


FIGURE 1.15 – (a) Structure de l'oeil (d'après www.wikipedia.org). (b) Test d'existence du point aveugle.

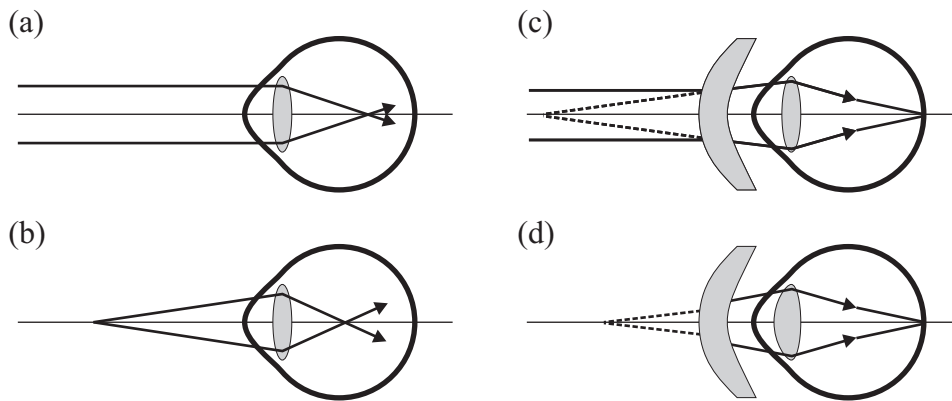


FIGURE 1.16 – La myopie et sa correction; (a), (b) oeil myope et (c), (d) correction avec une lentille. (a), (c) Objet à l'infini; (b) objet éloigné; (d) objet rapproché (nécessitant l'accommodation du cristallin).

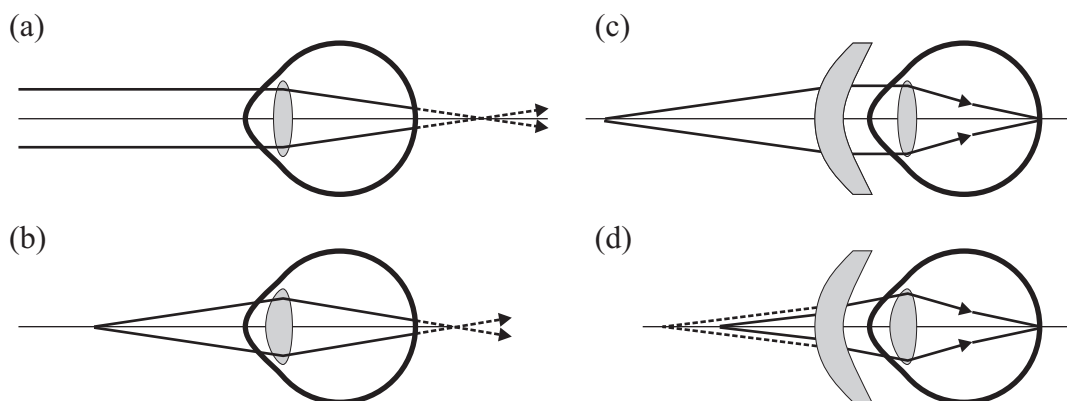


FIGURE 1.17 – L'hypermétropie et sa correction; (a), (b) oeil hypermétrope et (c), (d) correction avec une lentille. (a) Objet à l'infini; (b) objet proche (nécessitant l'accommodation du cristallin); (c) objet éloigné; (d) objet rapproché (nécessitant l'accommodation du cristallin).

et donc sa longueur focale. L'image se forme sur la surface de la rétine, sauf au point aveugle d'où part le nerf optique. Figure 1.15(b) permet de mettre en évidence ce point aveugle : observez le dessin à une distance d'environ 25 cm en fermant un oeil. Regardez attentivement la croix : le 2 disparaît. Rapprochez-vous du dessin : le 2 réapparaît tandis que le 1 disparaît à son tour.

La myopie est une anomalie de la vision où les rayons parallèles forment leur foyer en avant de la rétine, Fig. 1.16. Les objets éloignés apparaissent flous. Par contre un objet très proche peut être vu net. La myopie se corrige avec une lentille divergente en sorte qu'un objet à l'infini soit net. Pour un objet proche il est nécessaire d'accommoder, i.e. que les muscles de l'oeil contractent la cristallin en sorte d'augmenter sa convergence, Fig. 1.16(d).

A l'inverse de la myopie, l'hypermétropie est le défaut de l'oeil dans lequel les rayons parallèles provenant d'un objet à l'infini forment une image en arrière de la rétine (l'oeil est trop court et l'on doit constamment accommoder en contractant les muscles du cristallin). Ce défaut est corrigé à l'aide d'une lentille convergente placée devant l'oeil, Fig. 1.17

1.11 Systèmes optiques

Il existe une très grande variété de systèmes optiques, composés d'un nombre plus ou moins important de lentilles, nous nous contenterons ici de donner quelques exemples.

Une **loupe** est vraisemblablement le système optique le plus simple. Il permet d'agrandir un objet afin de l'examiner en détail. Cet effet a cependant une limite et si l'on approche l'objet au delà du *punctum proximum* il devient flou. Ce point représente la distance la plus petite où l'oeil voit net en accommodant au maximum. Cette distance varie avec l'âge, de 7 cm pour un adolescent à 12 cm pour un jeune adulte, 28 à 40 cm pour un quadragénaire et environ 1 m à 60 ans.

Un **oculaire** est un instrument visuel dont le rôle est de grossir l'image intermédiaire qu'un premier système optique (l'objectif) donne d'un objet réel. Ainsi l'oeil "regarde" dans l'oculaire qui "regarde" lui-même dans l'objectif. Ensemble l'objectif et l'oculaire forment un instrument optique tel un microscope ou un télescope. L'objet observé doit se trouver sur le foyer objet de l'oculaire et ce dernier doit conjuguer la pupille de sortie à une position convenable qui corresponde à l'endroit où l'observateur doit positionner son oeil. Idéalement, on souhaite que cette distance soit la plus grande possible, afin de donner à l'opérateur une certaine flexibilité.

L'oculaire le plus simple est celui de Huygens, Fig. 1.18(a). Le foyer objet est situé entre les deux lentilles. Cet objectif ne donne qu'une distance oculaire courte. L'oculaire de Ramsden n'a pas cet inconvénient et possède une pupille de sortie à plus grande distance δ .

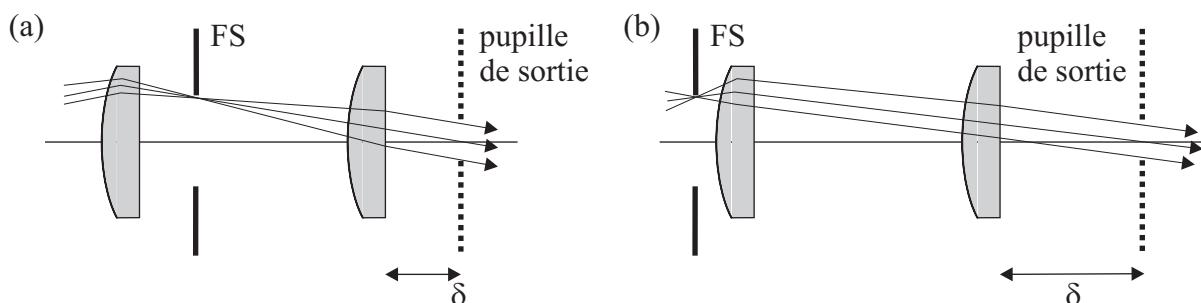


FIGURE 1.18 – Oculaire de (a) Huygens et (b) Ramsden. δ représente la distance oculaire, i.e. la distance à laquelle peut se placer l'oeil pour regarder l'image sans accommodation (*punctum remotum*).

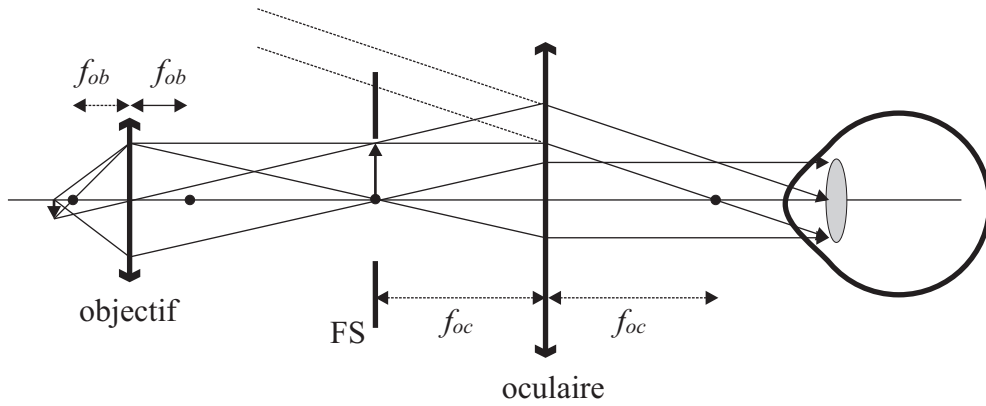


FIGURE 1.19 – Schéma de principe d'un microscope.

Le **microscope** va plus loin que la simple loupe et permet d'observer les détails d'objets rapprochés grâce à son fort grossissement. Dans sa plus simple expression il est constitué d'un objectif et d'un oculaire, Fig. 1.19. L'objectif forme une image réelle, inversée et agrandie de l'objet. Celle-ci est alors reprise par l'oculaire qui en forme une image agrandie à l'infini qui peut être observée par l'œil. L'image formée par l'objectif doit être dans le plan focal objet de l'oculaire ainsi que dans le plan du diaphragme de champ de l'oculaire. Les performances d'un microscope sont mesurées par son grossissement qui est le produit du grossissement transversal de l'objectif M_{To} par le grossissement de l'oculaire M_{Ae} :

$$MP = M_{To}M_{Ae} = -\frac{160}{f_o} \frac{254}{f_e}, \quad (1.28)$$

où l'on a utilisé le fait que la longueur du tube d'un microscope est standardisée à $L = 160$ mm et l'image doit se former pour un observateur normal à la distance du punctum remotum : 254 mm. En valeur absolue, si l'objectif et l'oculaire indiquent un grossissement respectivement de $5\times$ et $10\times$, le microscope a un grossissement de $50\times$.

Un paramètre essentiel pour les objectifs de microscope est l'ouverture numérique (*numerical aperture*, NA) qui dépend de l'indice de réfraction du milieu objet n_{obj} en contact avec la lentille et du demi-angle θ_{max} du plus gros cône de lumière que peut recevoir cette lentille (cet angle correspond donc à l'angle que fait un rayon marginal). L'ouverture numérique est généralement indiquée sur l'objectif, juste après le grossissement. Pour un objectif travaillant dans l'air $NA \leq 1$; par contre lorsque l'objet est immergé dans de l'eau NA peut atteindre des valeurs de 1.4 pour les meilleurs objectifs. Ceux-ci sont généralement composés d'un grand nombre de lentilles, comme illustré Fig. 1.20, qui permettent de diminuer les différentes aberrations de l'objectif.

La **lunette astronomique** permet l'observation d'objets très éloignés. A nouveau un oculaire refait une image créée par l'objectif. En général le foyer image de l'objectif coïncide avec le foyer objet de l'oculaire. De cette manière les rayons en ressortent parallèles entre eux et un œil normal peut ainsi observer confortablement l'image finale à l'infini, Fig. 1.21.

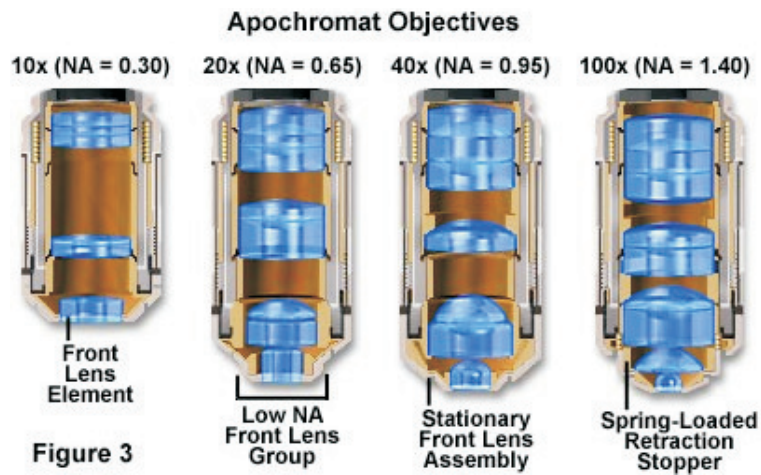


FIGURE 1.20 – Quelques exemples d'objectifs de microscope, d'après www.olympus.com.

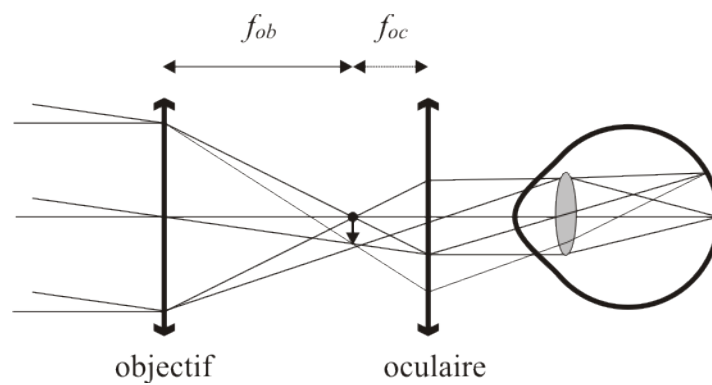


FIGURE 1.21 – Lunette astronomique afocale.

1.12 Méthode matricielle $ABCD$

Dans le cadre de l'optique géométrique, nous pouvons considérer qu'un système optique quelconque va transformer un rayon incident en un rayon transmis, Fig. 1.22. Ces deux rayons sont entièrement caractérisés par leurs hauteurs y_1 et y_2 ainsi que les angles θ_1 et θ_2 qu'ils forment avec l'axe optique du système. Notre objectif est de relier le rayon d'entrée (input) avec le rayon de sortie (output) à l'aide d'une relation matricielle.

En nous limitant à l'optique paraxiale, c'est à dire aux rayons se propageant le long de l'axe optique du système et ne déviant pas trop, nous pouvons faire l'hypothèse que les angles sont petits ($\sin \theta \simeq \theta$) ; la relation entre (y_2, θ_2) et (y_1, θ_1) est alors linéaire et peut être écrite sous la forme

$$\begin{aligned} y_2 &= Ay_1 + B\theta_1 \\ \theta_2 &= Cy_1 + D\theta_1, \end{aligned} \quad (1.29)$$

où A, B, C et D sont des nombres réels. Equation (1.29) peut s'écrire sous forme matricielle :

$$\begin{pmatrix} y_2 \\ \theta_2 \end{pmatrix} = \begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} y_1 \\ \theta_1 \end{pmatrix} = \mathbf{M} \begin{pmatrix} y_1 \\ \theta_1 \end{pmatrix}. \quad (1.30)$$

La matrice $\mathbf{M}$ caractérise entièrement le système optique et permet de déterminer (y_2, θ_2) connaissant (y_1, θ_1) ; on l'appelle matrice de transfert du système.

La Figure 1.23 décrit différents éléments optiques simples pour lesquels nous allons exprimer la matrice de transfert.

Propagation dans l'espace libre. Comme les rayons se propagent en ligne droite, un rayon parcourant une distance d est transformé de la façon suivante : $y_2 = y_1 + \theta_1 d$ et $\theta_2 = \theta_1$, Fig. 1.23(a). L'angle de propagation ne change pas et la matrice de transfert correspondante est donc,

$$\mathbf{M} = \begin{pmatrix} 1 & d \\ 0 & 1 \end{pmatrix}. \quad (1.31)$$

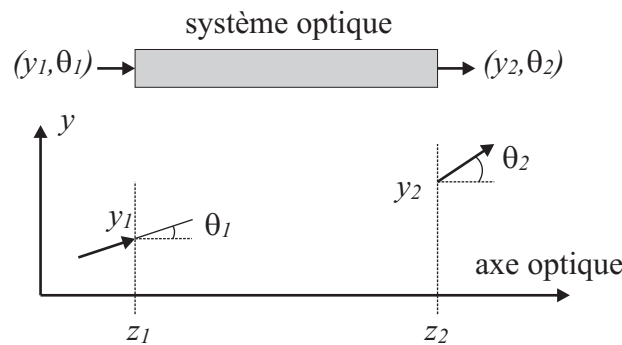


FIGURE 1.22 – Description des rayons d'entrée et de sortie d'un système optique

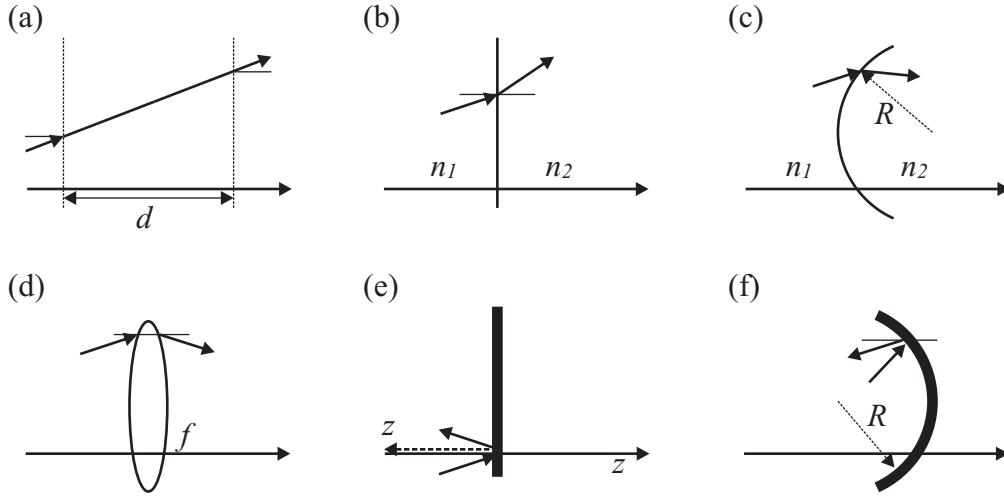


FIGURE 1.23 – Différents éléments optiques et les transformations qu'ils imposent à un rayon incident : (a) propagation dans l'espace libre ; (b) réfraction à un interface plan ; (c) réfraction à un interface sphérique ; (d) lentille mince ; (e) réflexion par un miroir plan et (f) réflexion par un miroir sphérique.

Réfraction à un interface. Les rayons incidents et réfractés satisfont la loi de Snell : $n_1 \sin \theta_1 = n_2 \sin \theta_2$ qui devient dans l'approximation paraxiale de l'optique géométrique $n_1 \theta_1 = n_2 \theta_2$, Fig. 1.23(b). On a donc la matrice de transfert,

$$\mathbf{M} = \begin{pmatrix} 1 & 0 \\ 0 & \frac{n_1}{n_2} \end{pmatrix}. \quad (1.32)$$

Réfraction à un interface sphérique. En utilisant les équations pour un dioptré sphérique on obtient la matrice suivante, Fig. 1.23(c) :

$$\mathbf{M} = \begin{pmatrix} 1 & 0 \\ -\frac{n_2 - n_1}{n_2 R} & \frac{n_1}{n_2} \end{pmatrix}; \quad (1.33)$$

remarquer que $R > 0$ pour un interface convexe et $R < 0$ pour un interface concave.

Transmission à travers une lentille fine. La matrice prend la forme suivante

$$\mathbf{M} = \begin{pmatrix} 1 & 0 \\ -\frac{1}{f} & 1 \end{pmatrix}, \quad (1.34)$$

où f représente la focale de la lentille ($f > 0$ pour une lentille convexe, $f < 0$ pour une lentille concave), Fig. 1.23(d).

Réflexion par un miroir plan. A nouveau la position verticale n'est pas changée entre le rayon d'entrée et le rayon de sortie : $y_2 = y_1$. Pour définir l'angle θ_2 on adopte la convention



FIGURE 1.24 – Transmission en série à travers plusieurs éléments optiques en utilisant le produit de leurs matrices de transfert.

suivante : l'axe z pointe toujours dans la direction de propagation du rayon. On assiste donc à un retournement de l'axe z pour le rayon réfléchi et la matrice $\mathbf{M}$ est la matrice unité, Fig 1.23(e),

$$\mathbf{M} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}. \quad (1.35)$$

Réflexion par un miroir sphérique. En utilisant les conventions précédentes, en particulier $R > 0$ pour un miroir convexe et $R < 0$ pour un miroir concave, Fig. 1.23(f), on obtient la matrice de transfert,

$$\mathbf{M} = \begin{pmatrix} 1 & 0 \\ \frac{2}{R} & 1 \end{pmatrix}. \quad (1.36)$$

Les matrices de transfert prennent tout leur intérêt lorsque l'on combine plusieurs éléments optiques en série pour former un système optique complexe, comme illustré Fig. 1.24. Ainsi, la transmission à travers N éléments optiques s'obtient-elle simplement en multipliant les matrices correspondantes :

$$\mathbf{M} = \mathbf{M}_N \mathbf{M}_{N-1} \cdots \mathbf{M}_2 \mathbf{M}_1. \quad (1.37)$$

On prendra cependant garde à l'ordre de multiplication : la première matrice rencontrée par le système est placée à droite en sorte qu'elle opère sur le vecteur du rayon incident. En général les matrices ne commutent pas et l'ordre est essentiel !

Notons que l'on peut aussi étendre cette approche à des systèmes périodiques dans lesquels un ou plusieurs éléments se répètent. Une version plus sophistiquée de ce formalisme, qui permet de calculer à la fois les ondes transmises et réfléchies donne lieu aux matrices de diffusion.

1.13 Aberrations chromatiques

Jusqu'à présent, dans le tracé des rayons, nous avons supposé qu'une lentille se comportait toujours de la même façon, quelle que soit la longueur d'onde utilisée. Ce n'est malheureusement pas le cas. On remarque par exemple dans la formule des lentilles minces Eq. (1.21) que la distance focale f dépend de l'indice de réfraction n_l de la lentille. Or, cet indice de réfraction, qui représente la réponse du matériau, varie avec la longueur d'onde. D'une façon générale, l'indice de réfraction diminue lorsque la longueur d'onde augmente. On peut le comprendre de façon intuitive en se rendant compte que plus la longueur d'onde est grande,

TABLE 1.5 – Indice de réfraction du verre BK7 selon Eq. (1.38) pour trois couleurs fondamentales

Couleur	$\lambda[\mu\text{m}]$	$n(\lambda)$
bleu	0.475	1.5232
vert	0.510	1.5208
rouge	0.650	1.5145

moins l'onde optique est en mesure d'observer les détails de la structure du matériau. À l'inverse, pour les courtes longueurs d'onde les variations du champ optique dans le matériau s'approchent des dimensions des atomes qui le constituent. Ainsi, dans le cas extrême des très courtes longueurs d'onde (rayons X), l'onde incidente interagit fortement avec la structure atomique de la matière, donnant lieu à la diffraction des rayons X qui permet de mettre en évidence la structure interne de la matière. La variation de l'indice de réfraction avec la longueur d'onde s'appelle la dispersion. En général, l'indice de réfraction augmente lorsque la longueur d'onde diminue (pour de petites longueurs d'onde l'interaction de la lumière avec la matière est plus forte).

Pour le verre BK7 qui est beaucoup utilisé en optique, la variation de l'indice de réfraction avec la longueur d'onde (exprimée en micromètres) peut s'exprimer de la manière suivante :

$$n(\lambda)^2 \approx 1 + \frac{1.03961212\lambda^2}{\lambda^2 - 0.00600069867} + \frac{0.231792344\lambda^2}{\lambda^2 - 0.0200179144} + \frac{1.01046945\lambda^2}{\lambda^2 - 103.560653}. \quad (1.38)$$

Cette équation est valable pour des longueurs d'onde entre $0.3 \mu\text{m}$ et $2.5 \mu\text{m}$.

La variation de l'indice de réfraction avec la longueur d'onde est cependant très petite, comme indiqué par la Table 1.5. Elle est cependant suffisante pour donner lieu à des effets tels que la décomposition de la lumière blanche dans un prisme, Fig. 1.25(a), ou le déplacement du point focal pour une lentille en fonction de la longueur d'onde, Fig. 1.25(b).

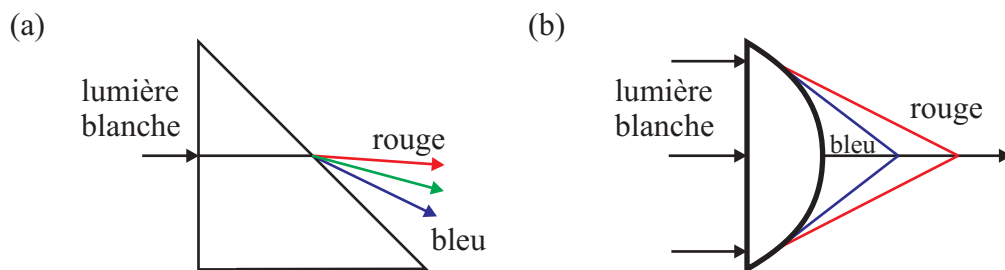


FIGURE 1.25 – Effet de la dispersion des matériaux : (a) Décomposition de la lumière blanche dans un prisme et (b) aberration chromatique d'une lentille.

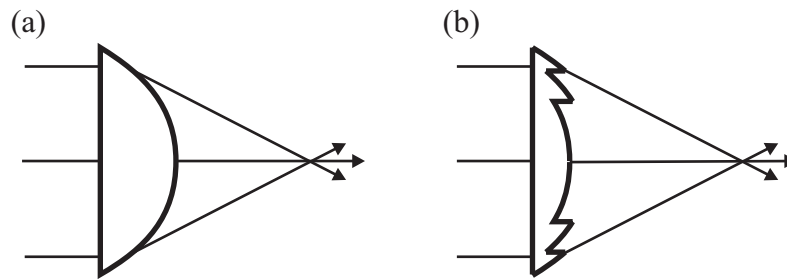


FIGURE 1.26 – (a) Lentille conventionnelle ou réfractive et (b) lentille diffractive ou micro-lentille.

1.14 Micro-optique et Optique diffractive

Pour terminer ce chapitre sur l'optique géométrique, arrêtons-nous sur une forme très intéressante des composants optiques permettant leur intégration dans des micro-systèmes. La figure 1.26(a) indique le tracé des rayons parallèles pour une lentille conventionnelle (on parle de lentille réfractive). En suivant les rayons, on s'aperçoit qu'ils ne subissent un changement de direction qu'au moment où ils traversent la face de sortie de la lentille et sont réfractés lors du passage du verre à l'air ; tout le chemin à l'intérieur de la lentille ne semble pas affecter la marche de ces rayons. Une lentille diffractive utilise ce principe pour diminuer l'épaisseur de la lentille en enlevant cette épaisseur intérieure "inutile". On obtient alors une lentille assez différente, dont la surface reproduit cependant la surface de la lentille originale morceau par morceau, Fig. 1.26(b). L'effet du profil de cette lentille sur les rayons s'étudie particulièrement bien dans le cadre de la diffraction qui permet de déterminer l'effet du profil d'indice de réfraction sur les rayons incidents. La diffraction sera étudiée dans un prochain chapitre. D'un point de vue technologique, il existe différentes façons de réaliser des lentilles diffractives. Les polymères se prêtent particulièrement bien à leur reproduction à large échelle pour réaliser des composants à faible coût.

Le fonctionnement d'une lentille diffractive se comprend aussi aisément en remarquant que l'effet d'une lentille conventionnelle, par exemple une lentille convergente Fig. 1.27(a), se déduit simplement avec la loi de Snell en remarquant que pour la deuxième surface de la lentille (du verre vers l'air) la perpendiculaire $\mathbf{p}$ utilisée dans la loi de Snell change d'orientation selon la position le long de la surface : elle est très inclinée pointant vers l'extérieur à la périphérie ($\mathbf{p}_1$), alors qu'elle est parallèle à l'axe optique au centre ($\mathbf{p}_3$). En chaque point l'angle de réfraction θ_2 sous lequel le rayon émerge de la lentille est donné par $\sin \theta_2 = n_1/n_2 \sin \theta_1$; comme le rayon passe du verre ($n_1 = 1.5$) à l'air ($n_2 = 1$), l'angle θ_2 est plus grand que l'angle θ_1 . De plus, comme la normale est plus inclinée par rapport à l'axe optique vers la périphérie de la lentille, l'angle θ_1 y est plus important ce qui rend l'angle θ_2 plus grand et produit l'effet de convergence des rayons observé. Un raisonnement similaire s'applique pour une lentille divergente.

Du point de vue de l'optique de Fourier que nous étudierons dans un prochain chapitre, on

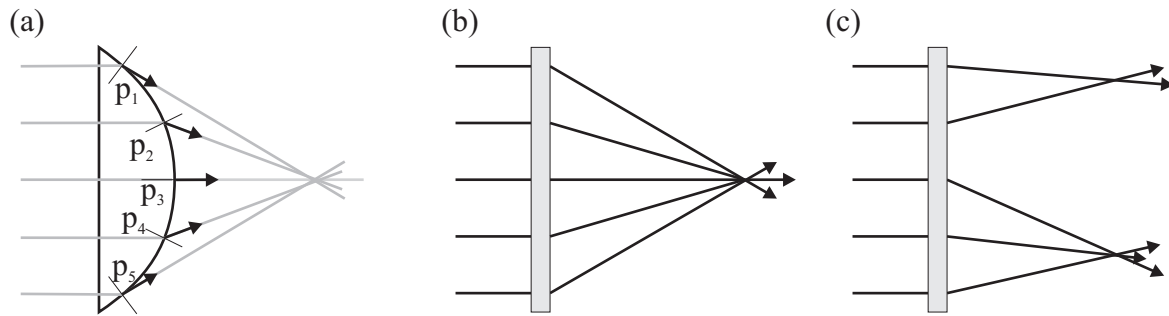


FIGURE 1.27 – (a) Lentille conventionnelle convergente et tracé des rayons et de leur réfraction sur la deuxième interface ; (b) élément optique diffractif convergent et (c) élément optique diffractif focalisant la lumière incidente en deux points.

considère pour chacune des ondes planes formant le faisceau incident sur la lentille que l'effet de la réfraction est de modifier sa direction de propagation. Ainsi, chaque onde incidente acquiert une certaine quantité de mouvement latérale via son interaction avec la lentille (noter le changement d'orientation des flèches, Fig. 1.27(a) alors que la longueur des flèches ne change pas). En optique de Fourier on dit que chaque onde acquiert une certaine quantité de mouvement latérale, dans le plan de la lentille. L'optique de Fourier permet de créer des composants optiques qui permettent de manipuler cette quantité de mouvement pour manipuler la propagation de l'onde émergente. Ainsi peut-on créer un composant optique plan qui a le même effet qu'une lentille convergente, Fig. 1.27(b), voire des composants plus sophistiqués qui permettent par exemple de focaliser les rayons incidents en deux points spécifiques, Fig. 1.27(c).

Chapitre 2

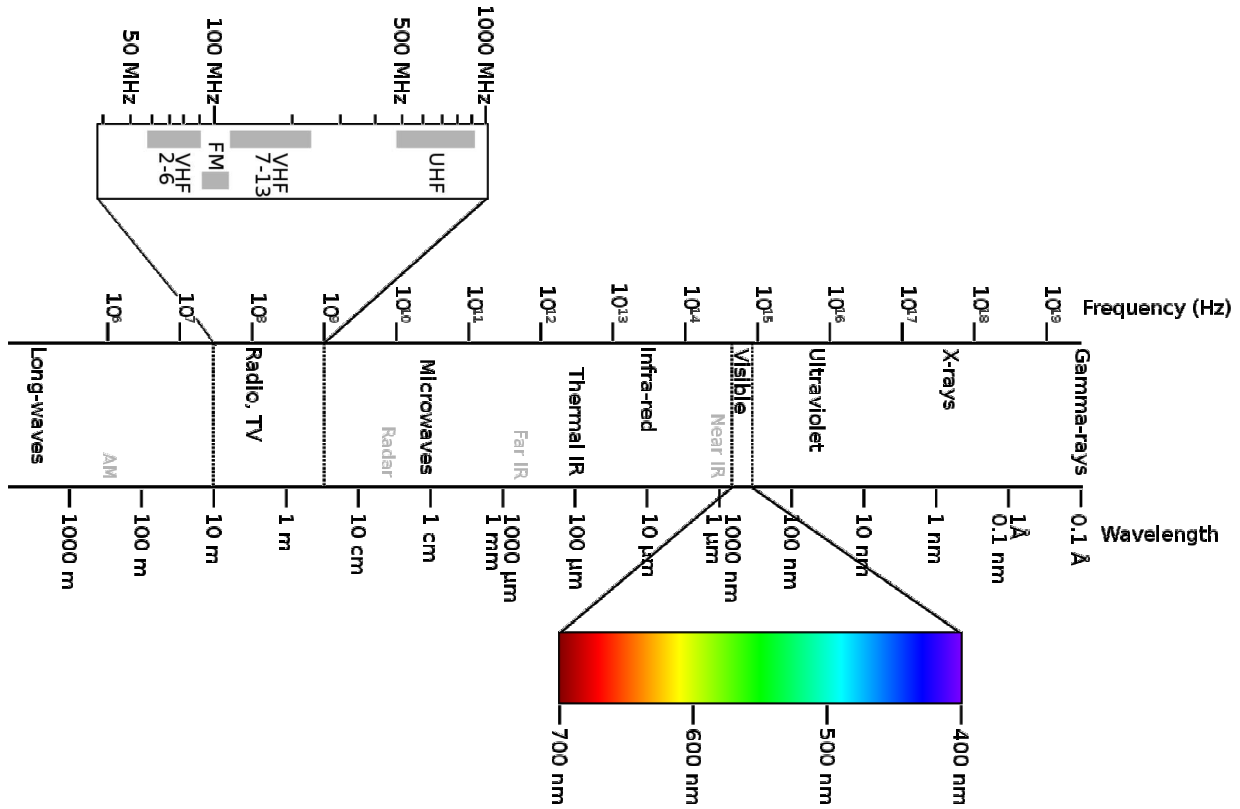
Optique Ondulatoire

2.1 Introduction

Comme nous l'avons vu, l'optique géométrique permet de décrire bon nombre de phénomènes optiques ainsi que de comprendre des composants tels les lentilles et les miroirs. Il existe cependant des phénomènes qui ne peuvent pas être décrits en termes de rayons, comme par exemple l'interférence de deux ondes. Pour étudier ces phénomènes nous introduisons dans ce chapitre l'optique ondulatoire, dans laquelle la propagation de la lumière se fait sous forme d'onde caractérisée par une longueur d'onde et une fréquence. Notons que l'optique géométrique n'est pas dissociée de l'optique ondulatoire, mais en représente un cas limite lorsque la longueur d'onde devient infiniment petite. Le domaine de validité de l'optique géométrique ne se limite cependant pas à ce cas limite et on peut généralement utiliser l'optique géométrique dans les situations où les objets considérés sont bien plus grands que la longueur d'onde.

Une des caractéristiques principales d'une onde est sa longueur d'onde λ_0 ou sa fréquence. Nous nous intéressons en particulier aux ondes visibles qui ont une longueur d'onde entre $\lambda_0 = 390 \text{ nm}$ (violet) et $\lambda_0 = 760 \text{ nm}$ (rouge foncé), Fig. 2.1. Le formalisme développé dans ce chapitre s'applique cependant à l'entier du spectre électromagnétique, de l'ultraviolet (UV) à l'infrarouge (IR), et bien au delà dans les micro-ondes. Dans le domaine des fréquences, le spectre de la lumière s'étend entre 10^{13} Hz (infrarouge) et 10^{16} Hz (ultraviolet), et s'exprime généralement en THz.

Relevons enfin que ce chapitre propose une approche scalaire des ondes ; i.e. on ne considère qu'un seul composant du champ électromagnétique. Celui-ci sera étudié plus en détail dans un prochain chapitre.

FIGURE 2.1 – Le spectre électromagnétique et sa partie visible, d'après www.wikipedia.org.

2.2 Equation d'onde

Une onde optique est décrite à un instant t en un point $\mathbf{r} = (x, y, z)$ de l'espace par une fonction réelle $u(\mathbf{r}, t)$ qui représente l'amplitude de l'onde. Cette fonction satisfait l'équation d'onde,

$$\nabla^2 u - \frac{1}{c^2} \frac{\partial^2 u}{\partial t^2} = 0, \quad (2.1)$$

où l'on a introduit la vitesse c de la lumière dans le milieu :

$$c = \frac{c_0}{n}; \quad (2.2)$$

$c_0 \simeq 3.0 \cdot 10^8$ m/s étant la vitesse de la lumière dans le vide et n l'indice de réfraction du milieu. Il est important de noter qu'Eq. (2.1) est valide quel que soit le système de coordonnées. Ainsi convient-il simplement d'utiliser le bon opérateur de Laplace ∇^2 pour obtenir l'équation d'onde dans un système de coordonnées quelconque. En coordonnées cartésiennes $u(x, y, z)$ on a :

$$\nabla^2 = \frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} + \frac{\partial^2 u}{\partial z^2}. \quad (2.3)$$

En coordonnées sphériques $u(r, \theta, \phi)$,

$$\nabla^2 = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial u}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial u}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \frac{\partial^2}{\partial \phi^2}; \quad (2.4)$$

et en coordonnées cylindriques $u(\rho, \phi, z)$,

$$\nabla^2 = \frac{1}{\rho} \frac{\partial}{\partial \rho} \left(\rho \frac{\partial u}{\partial \rho} \right) + \frac{1}{\rho^2} \frac{\partial^2 u}{\partial \phi^2} + \frac{\partial^2 u}{\partial z^2}. \quad (2.5)$$

Bien que ces formules puissent paraître un peu rébarbatives, les coordonnées sphériques sont particulièrement utiles pour étudier les ondes générées par une source ponctuelle comme une molécule fluorescente ; quant aux coordonnées cylindriques elles permettent d'étudier les guides d'ondes qui possèdent justement cette symétrie cylindrique (par exemple une fibre optique).

Une des propriétés les plus utiles de l'équation d'onde (2.1) est sa linéarité, donnant lieu au **principe de superposition** : si $u_1(\mathbf{r}, t)$ et $u_2(\mathbf{r}, t)$ sont des solutions de (2.1), alors $u(\mathbf{r}, t) = u_1(\mathbf{r}, t) + u_2(\mathbf{r}, t)$ est aussi une solution de l'équation d'onde. Ce principe se généralise évidemment à un nombre arbitraire d'ondes et permet par exemple d'utiliser l'analyse de Fourier pour étudier une onde quelconque.

2.3 Ondes monochromatiques harmoniques

Une onde monochromatique harmonique a la forme suivante,

$$u(\mathbf{r}, t) = a(\mathbf{r}) \cos(\omega t + \phi(\mathbf{r})), \quad (2.6)$$

où l'on a introduit l'amplitude $a(\mathbf{r})$, la phase $\phi(\mathbf{r})$, la pulsation ou fréquence angulaire $\omega = 2\pi\nu$ (rad/s). Celle-ci permet de définir la fréquence ν (Hz ou s^{-1}) et la période $T = 1/\nu = 2\pi/\omega$ (s). Ces différents paramètres sont indiqués sur Fig. 2.2.

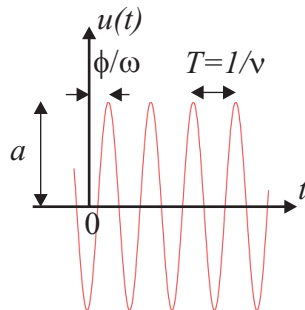


FIGURE 2.2 – Représentation d'une onde monochromatique harmonique et de ses différents paramètres.

2.4 Représentation complexe

Afin de bénéficier de l'outillage associé à l'analyse complexe, on introduit une fonction d'onde complexe $U(\mathbf{r}, t)$,

$$U(\mathbf{r}, t) = a(\mathbf{r})e^{j\phi(\mathbf{r})}e^{j\omega t}, \quad (2.7)$$

dont la fonction d'onde physique est la partie réelle :

$$u(\mathbf{r}, t) = \text{Re} \{U(\mathbf{r}, t)\} = \frac{1}{2} (U(\mathbf{r}, t) + U^*(\mathbf{r}, t)) . \quad (2.8)$$

On remarquera que la combinaison $U + U^*$ permet d'obtenir la partie réelle sans erreur, même lorsque l'on a des expressions compliquées. Notons finalement qu'il s'agit ici d'une convention, on aurait aussi pu considérer que la grandeur physique est la partie imaginaire de l'onde complexe.

L'amplitude complexe $U(\mathbf{r}) = a(\mathbf{r}) \exp(j\phi(\mathbf{r}))$ joue un rôle important en optique puisque la fonction d'onde complexe s'écrit

$$U(\mathbf{r}, t) = U(\mathbf{r})e^{j\omega t} = a(\mathbf{r})e^{j\phi(\mathbf{r})}e^{j\omega t} . \quad (2.9)$$

Ainsi, à un endroit donné $\mathbf{r}$ l'amplitude complexe $U(\mathbf{r})$ est un nombre complexe dont l'amplitude $|U(\mathbf{r})| = a(\mathbf{r})$ est l'amplitude de l'onde et dont l'argument $\arg\{U(\mathbf{r})\} = \phi(\mathbf{r})$ est la phase de l'onde, Fig. 2.3(a). À cet endroit $\mathbf{r}$ l'amplitude complexe de l'onde décrit donc un cercle avec la pulsation $j\omega t$; la phase au temps $t = 0$ vaut $\phi(\mathbf{r})$, Fig. 2.3(b). Il est important de garder à l'esprit que cette phase change avec le lieu $\mathbf{r}$.

En introduisant $U(\mathbf{r}, t)$ dans Eq. (2.1), on obtient une équation pour l'amplitude complexe :

$$\nabla^2 U + k^2 U = 0, \quad (2.10)$$

où l'on a introduit le nombre d'onde k :

$$k = \frac{\omega}{c} = \frac{\omega n}{c_0} = \frac{2\pi\nu}{c} = \frac{2\pi\nu n}{c_0} . \quad (2.11)$$

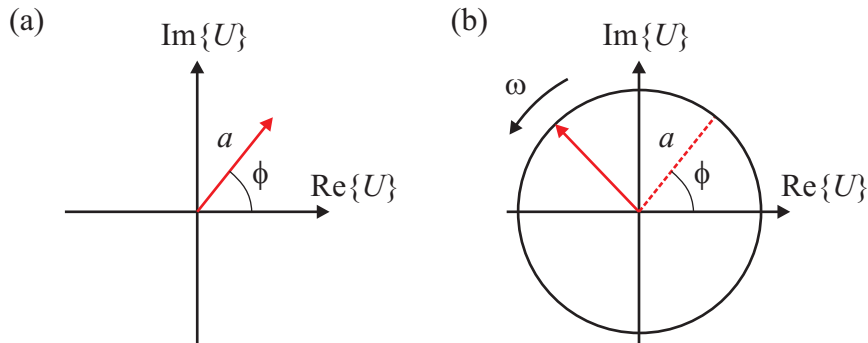


FIGURE 2.3 – (a) Représentation de l'amplitude complexe d'une onde en un point $\mathbf{r}$ donné. (b) Phaseur décrivant l'évolution de cette onde dans le temps.

Equation (2.10) s'appelle l'équation d'Helmholtz.

L'intensité lumineuse $I(\mathbf{r}, t)$ correspond à la puissance optique par unité de surface (W/cm^2) et joue un rôle important pour une onde polychromatique. Pour l'instant contentons-nous de la définition,

$$I(\mathbf{r}, t) = 2 \langle u^2(\mathbf{r}, t) \rangle . \quad (2.12)$$

Le symbole $\langle \rangle$ signifie le moyennage dans le temps. Pour celui-ci on considère un intervalle de temps bien supérieur à la période de l'onde, mais inférieur aux variations (modulations) du signal. Le facteur 2 introduit dans Eq. (2.12) permet de normaliser la puissance optique comme on le verra dans un prochain chapitre. Pour une onde monochromatique harmonique on a la propriété suivante :

$$I(\mathbf{r}) = |U(\mathbf{r})|^2 ; \quad (2.13)$$

et l'intensité d'une onde monochromatique harmonique est le carré de la norme de son amplitude complexe. Ainsi l'intensité d'une onde monochromatique harmonique est-elle constante dans le temps.

La puissance optique (W) s'écoulant à travers une surface A perpendiculaire à la direction de propagation de l'onde s'obtient en intégrant l'intensité lumineuse :

$$P(t) = \int_A I(\mathbf{r}, t) dA . \quad (2.14)$$

L'énergie s'obtient en intégrant durant l'intervalle de temps correspondant.

Un élément important dans la caractérisation d'une onde est son front d'onde. Celui-ci correspond à la surface définie par l'onde lorsque l'on considère une valeur de la phase constante : $\phi(\mathbf{r}) = \text{const}$. Pour représenter l'onde on utilise généralement les fronts d'onde correspondants à une phase égale à un multiple de 2π : $\phi(\mathbf{r}) = 2\pi q$, où q est un entier.

2.5 Onde plane, sphérique et évanescente

Pour un milieu homogène, les deux solutions les plus simples de l'équation d'Helmholtz (2.10) sont l'onde plane et l'onde sphérique. Pour l'onde plane, on définit l'amplitude complexe

$$U(\mathbf{r}) = Ae^{-j\mathbf{k}\cdot\mathbf{r}} = Ae^{-j(k_x x + k_y y + k_z z)} , \quad (2.15)$$

où $\mathbf{k} = (k_x, k_y, k_z)$ représente le vecteur d'onde. En introduisant (2.15) dans Eq. (2.9) on obtient la relation importante $k_x^2 + k_y^2 + k_z^2 = k^2$ qui indique que pour satisfaire l'équation d'Helmholtz la norme du vecteur $\mathbf{k}$ doit être égale au nombre d'onde.

Pour l'onde plane, les fronts d'onde sont des plans parallèles séparés par une distance $\lambda = 2\pi/k$ où

$$\lambda = \frac{c}{\nu} = \frac{c_0}{\nu n} = \frac{2\pi}{k} = \frac{2\pi}{k_0 n} \quad (2.16)$$

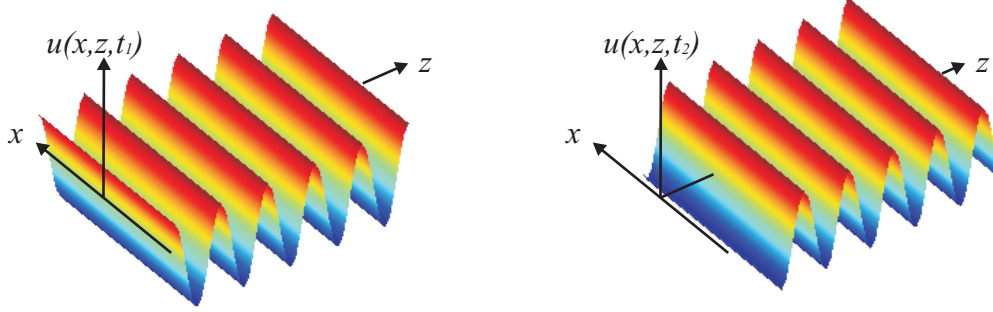


FIGURE 2.4 – Propagation d'une onde plane harmonique dans l'espace à deux instants t_1 et t_2 différents.

représente la longueur d'onde.

Si l'on revient à la fonction réelle, une onde plane se propageant dans la direction z ($\mathbf{k} = (k_x, k_y, k_z) = (0, 0, k)$) s'écrit

$$u(\mathbf{r}, t) = |A| \cos(\omega t - kz + \phi) = |A| \cos(\omega(t - z/c) + \phi). \quad (2.17)$$

Cette fonction est périodique dans le temps avec une période $T = 1/\nu$ et périodique dans l'espace avec une période $\lambda = 2\pi/k$, Fig. 2.4. Comme la phase de l'onde, $\arg\{U(\mathbf{r}, t)\} = \omega(t - z/c) + \phi$ varie dans le temps et l'espace en fonction de la variable $t - z/c$, on appelle c la vitesse de phase de l'onde.

Figure 2.4 indique que l'onde plane

$$U(\mathbf{r}, t) = Ae^{-j\mathbf{k}\cdot\mathbf{r} + j\omega t} \quad (2.18)$$

se propage dans la direction du vecteur $\mathbf{k}$ lorsque le temps t avance. Si au lieu d'avoir choisi le signe "-" pour le terme $j\mathbf{k}\cdot\mathbf{r}$, nous avions choisi le signe "+", nous aurions

$$U(\mathbf{r}, t) = Ae^{+j\mathbf{k}\cdot\mathbf{r} + j\omega t}. \quad (2.19)$$

Or Eq. (2.19) est aussi une solution de l'équation d'Helmholtz (2.10). Une telle onde ne se propage cependant pas dans la direction du vecteur $\mathbf{k}$, mais dans la direction du vecteur $-\mathbf{k}$ lorsque le temps avance. On parle parfois d'onde rétrograde pour Eq. (2.19) et d'onde progressive pour (2.18). Le signe de l'exposant complexe n'influence pas la solution de l'équation d'Helmholtz car celle-ci contient les deuxièmes dérivées dans le temps et l'espace et les signes s'annulent. Nous aurions donc aussi pu utiliser le signe "-" pour la partie temporelle de l'onde et nous aurions alors les deux solutions en onde plane suivantes :

$$U(\mathbf{r}, t) = Ae^{+j\mathbf{k}\cdot\mathbf{r} - j\omega t} \quad \text{onde progressive} \quad (2.20a)$$

$$U(\mathbf{r}, t) = Ae^{-j\mathbf{k}\cdot\mathbf{r} - j\omega t} \quad \text{onde rétrograde}. \quad (2.20b)$$

La convention pour le signe de l'exponentielle temporelle est donc arbitraire. Une fois celle-ci fixée, on a deux ondes planes possibles avec des signes différents pour la partie spatiale $j\mathbf{k} \cdot \mathbf{r}$. L'onde plane qui a le même signe que $j\omega t$ est rétrograde (elle se propage dans la direction $-\mathbf{k}$) et celle qui a un signe différent est progressive et se déplace dans la direction $+\mathbf{k}$. Pour le reste du cours, nous utiliserons toujours la convention $+j\omega t$ et Eq. (2.18) représentera une onde progressive. Le signe choisi pour l'exponentielle temporelle influence d'autres expressions décrivant la propagation des ondes, en particulier lorsque le milieu de propagation a des pertes ou du gain. Lorsque l'on consulte un livre il est important de commencer par s'assurer de la convention choisie pour pouvoir appliquer les formules sans soucis.

Lorsqu'une onde monochromatique harmonique se propage dans un milieu d'indice n sa fréquence ω reste inchangée, par contre les autres paramètres sont affectés de la manière suivante :

$$c = \frac{c_0}{n}, \quad \lambda = \frac{\lambda_0}{n}, \quad k = nk_0. \quad (2.21)$$

La figure 2.5 illustre ces différents phénomènes lorsqu'une onde plane est réfractée successivement par deux interfaces plans entre les milieux d'indices $n_1 < n_2$ puis $n_2 > n_1$. On remarque que l'onde qui émerge dans le milieu n_1 à droite a les mêmes caractéristiques que l'onde qui entre dans le milieu n_1 à gauche : même longueur d'onde (distance entre les maxima, $\lambda_1 = \lambda_0/n_1$) et même direction de propagation. Par contre, l'onde dans le milieu n_2 a une longueur d'onde plus petite : $\lambda_2 = \lambda_0/n_2$ et une direction de propagation différente. On peut comprendre ce changement de direction de propagation en remarquant que la phase est conservée le long de l'interface entre les milieux n_1 et n_2 . Il faut donc que les maxima de l'onde dans le milieu n_1 coïncident avec les maxima de l'onde dans le milieu n_2 ; comme ces derniers sont plus rapprochés, l'onde doit "se tordre" dans le milieu n_2 . Le même phénomène apparaît au deuxième interface entre n_2 et n_1 et l'onde reprend sa direction initiale.

Considérons maintenant une onde plane se propageant dans la direction z ($\mathbf{k} = (k_x, k_y, k_z) = (0, 0, k)$) et supposons que k est un nombre complexe avec une partie réelle et imaginaire

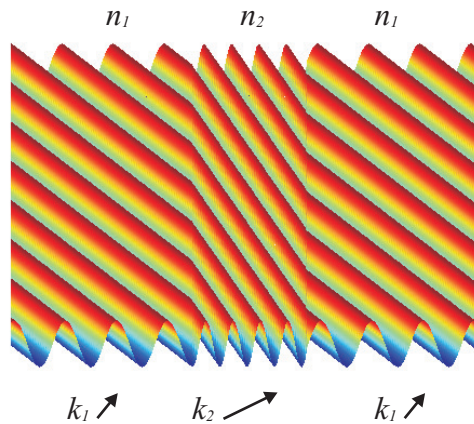


FIGURE 2.5 – Réfraction d'une onde plane à deux interfaces ($n_1 < n_2$).

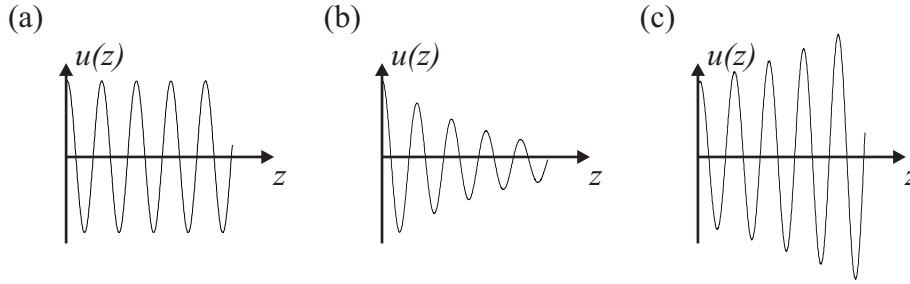


FIGURE 2.6 – Onde (a) propageante ($k'' = 0$), (b) atténuée ou évanescente ($k'' < 0$) et (c) amplifiée ($k'' > 0$).

$k = k' + jk''$:

$$U(\mathbf{r}, t) = Ae^{-jkz+j\omega t} = Ae^{k''z}e^{-jk'z+j\omega t}. \quad (2.22)$$

L'amplitude de l'onde varie maintenant avec la distance de propagation z . Si $k'' > 0$ l'amplitude augmente avec la distance de propagation : l'onde est amplifiée ; un tel phénomène peut se produire lorsque le milieu possède un certain gain optique, comme par exemple dans un laser. Si $k'' < 0$, alors l'amplitude de l'onde diminue au fur et à mesure de la propagation ; on parle d'onde atténuée ou évanescente. Ce phénomène peut se produire lorsque le matériau absorbe la lumière, comme par exemple un métal. Il se produit aussi à l'interface entre deux milieux d'indices différents : l'onde incidente peut alors être "totalement" réfléchi dans le milieu initial, avec une toute petite partie de l'onde transmise sous forme d'onde évanescence dans le deuxième milieu. Ces trois situations sont illustrées Fig. 2.6.

Une autre solution simple de l'équation d'Helmholtz (2.10) est l'onde sphérique dont l'amplitude complexe vaut

$$U(\mathbf{r}) = \frac{A_0}{r}e^{-jkr}, \quad (2.23)$$

où l'équation d'Helmholtz est exprimée en coordonnées sphériques et r représente la distance à l'origine ; $k = \omega/c$ est le nombre d'onde et A_0 est une constante. L'intensité d'une onde sphérique décroît avec le carré de la distance : $I(\mathbf{r}) = |A_0|^2/r^2$. Les fronts d'onde sont des surfaces sphériques concentriques, séparées par la distance λ .

A l'origine ($r = 0$) l'intensité de l'onde sphérique diverge, ce qui n'est pas possible physiquement (il faudrait une quantité d'énergie infinie à l'origine). Une telle onde n'est donc pas réaliste ; elle n'en demeure pas moins une approximation très utile, par exemple pour l'onde émise par une source ponctuelle ; elle forme aussi la base de la théorie de la diffraction. Notons que l'onde plane n'a pas non plus de sens physique réel : son extension latérale infinie lui confère aussi une énergie infinie. En superposant dans l'espace plusieurs ondes planes, on peut cependant limiter l'énergie dans une région finie de l'espace. Ce principe donne lieu à l'optique de Fourier que nous étudierons dans un prochain chapitre.

Si l'onde (2.23) est une onde qui se propage de l'origine vers l'infini (ou sortante), une onde

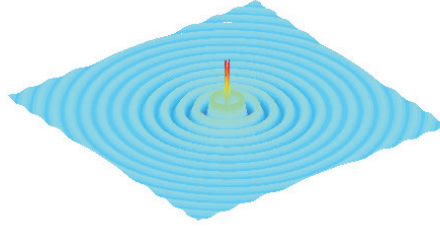


FIGURE 2.7 – Onde sphérique.

définie par l'amplitude complexe

$$U(\mathbf{r}) = \frac{A_0}{r} e^{+jkr}, \quad (2.24)$$

est une onde rentrante qui se propage de l'infini vers l'origine (en supposant évidemment que nous avons choisi $e^{+j\omega t}$ pour la partie temporelle. Notons finalement qu'une onde sphérique sortante ayant son origine au point $\mathbf{r}_0$ s'écrit

$$U(\mathbf{r}) = \frac{A_0}{|\mathbf{r} - \mathbf{r}_0|} e^{-jk|\mathbf{r} - \mathbf{r}_0|}. \quad (2.25)$$

Notons que les ondes planes et sphériques ne sont pas physiques : l'onde plane, avec son étendue infinie, a une énergie infinie ; alors que l'onde sphérique diverge à l'origine. Néanmoins, ces ondes sont très utiles conceptuellement.

Figure 2.8 fait le lien entre optique géométrique et optique ondulatoire. On note en particulier que les rayons optiques sont normaux aux fronts d'onde de l'optique ondulatoire.

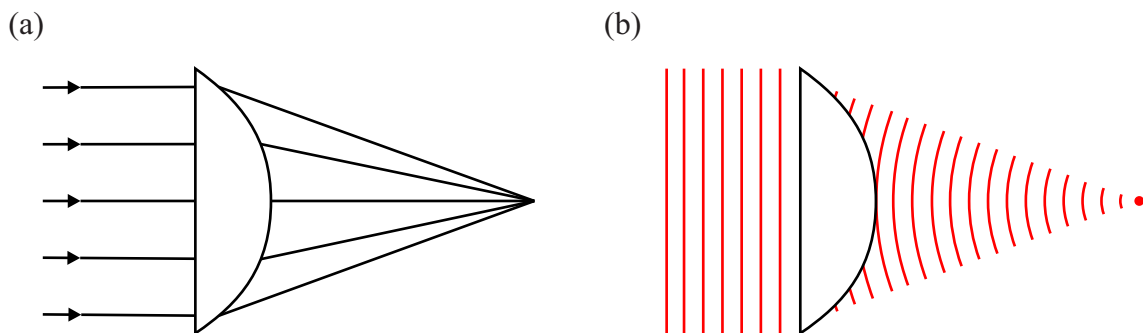


FIGURE 2.8 – Lien entre (a) optique géométrique et (b) optique ondulatoire.

2.6 Accumulation de la phase

Lorsqu'une onde se propage, elle acquiert une certaine quantité de phase au fur et à mesure de sa propagation. On dit que l'onde accumule de la phase. Par simplicité nous considérons ici une onde plane se propageant dans la direction z ; une onde sphérique accumule aussi de la phase d'une façon similaire. Pour déterminer la quantité de phase accumulée pendant une certaine distance, il suffit de compter le nombre de longueurs d'onde parcourues et de multiplier par 2π . Comme la longueur d'onde varie avec l'indice de réfraction, il faut tenir compte du milieu dans lequel l'onde se propage. Ainsi, dans un milieu d'indice élevé, la longueur d'onde est plus courte et, pour une distance parcourue donnée, l'onde accumule plus de phase que dans un milieu de faible indice.

Cet effet est illustré Fig. 2.9. Dans un milieu homogène d'indice n_1 , la phase accumulée dépend simplement de la distance parcourue z :

$$\phi(z) = k_1 z = \frac{2\pi}{\lambda} z = \frac{2\pi n_1}{\lambda_0} z. \quad (2.26)$$

On remarque l'intérêt d'avoir introduit le nombre d'onde k qui possède comme unité [rad/m]. Ainsi dans Fig. 2.9(a) la phase accumulée pendant la distance de propagation L_1 est simplement $\phi(L_1) = k_1 L_1$ et pour une distance de propagation plus longue L_2 , $\phi(L_2) = k_1 L_2$.

Si maintenant l'onde se propage dans différents milieux, comme illustré Fig. 2.9(b), elle accumule plus de phase dans le milieu d'indice élevé, où sa longueur d'onde est plus courte. On a ainsi maintenant pour la phase accumulée sur la distance L_2 :

$$\phi(L_2) = k_1 z_1 + k_2 (z_2 - z_1) + k_1 (L_2 - z_2) = \frac{2\pi n_1}{\lambda_0} z_1 + \frac{2\pi n_2}{\lambda_0} (z_2 - z_1) + \frac{2\pi n_1}{\lambda_0} (L_2 - z_2). \quad (2.27)$$

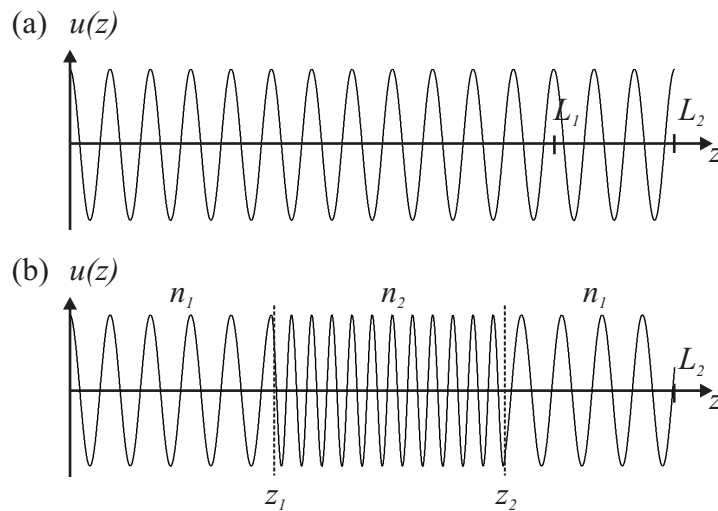


FIGURE 2.9 – Phase accumulée par une onde lorsqu'elle se propage (a) dans un milieu homogène et (b) dans différent milieux.

On a donc simplement découpé le trajet de l'onde en segments de matériaux homogènes n_1 , n_2 et n_1 et additionné les phases accumulées dans chaque segment. Il convient de remarquer que Fig. 2.9(b) ne montre que l'onde progressive dans le système. En réalité une partie de l'onde incidente sera réfléchiée aux interfaces n_1/n_2 et n_2/n_1 ; ces réflexions peuvent aussi donner lieu à un saut de phase dont il faut tenir compte dans la phase accumulée, comme nous le verrons dans le chapitre consacré à l'optique de Maxwell. De plus, l'amplitude de l'onde transmise est généralement plus petite que celle de l'onde incidente. La figure 2.9(b) est cependant parfaitement correcte si l'on ne considère que la phase de l'onde propageante.

2.7 Interférence

Dans cette section on s'intéresse à la superposition d'ondes monochromatiques de même fréquence. Lorsque deux ondes ou plus se trouvent en un point de l'espace, la fonction d'onde totale est la superposition de ces ondes. Ainsi, pour des ondes monochromatiques de même fréquence, la fonction d'onde complexe de l'onde totale est la somme des fonctions d'ondes complexes des ondes individuelles :

$$U(\mathbf{r}) = U_1(\mathbf{r}) + U_2(\mathbf{r}). \quad (2.28)$$

Si l'on a les ondes, $U_1(\mathbf{r}) = \sqrt{I_1} \exp(j\phi_1(\mathbf{r}))$ et $U_2(\mathbf{r}) = \sqrt{I_2} \exp(j\phi_2(\mathbf{r}))$, on obtient pour l'intensité $I(\mathbf{r})$ de la superposition (2.13) :

$$I = |U|^2 = |U_1 + U_2|^2 = |U_1|^2 + |U_2|^2 + U_1 U_2^* + U_1^* U_2, \quad (2.29)$$

où l'on a omis la dépendance spatiale pour alléger l'écriture. Equation (2.29) peut s'écrire

$$I = I_1 + I_2 + 2\sqrt{I_1 I_2} \cos \phi, \quad (2.30)$$

avec $\phi = \phi_2 - \phi_1$. L'équation d'interférence (2.30) indique que l'intensité de l'onde résultante ne dépend pas seulement des intensités des deux ondes mais aussi de la différence ϕ entre leurs amplitudes complexes. Ce phénomène se comprend bien en étudiant les phaseurs des ondes U_1 et U_2 , ainsi que le phaseur de l'onde résultante U , Fig. 2.10.

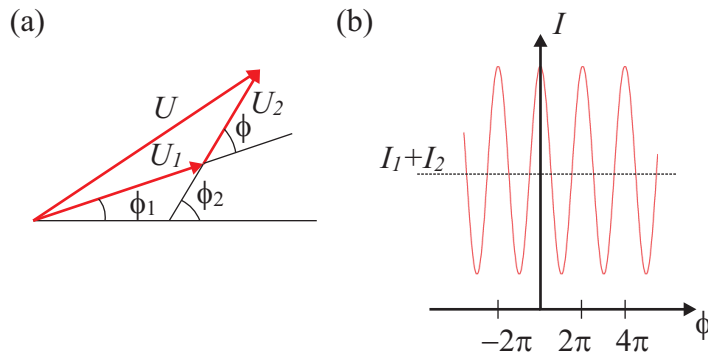


FIGURE 2.10 – Interférence de deux ondes.

Si on considère l'interférence d'ondes de même amplitude $I_1 = I_2 = I_0$, Eq. (2.30) devient

$$I = 2I_0(1 + \cos \phi) = 4I_0 \cos^2(\phi/2). \quad (2.31)$$

Ainsi pour $\phi = 0$ (les phaseurs des deux ondes sont alignés) l'intensité totale vaut quatre fois l'intensité de chacune des ondes superposées : $I = 4I_0$. Par contre, pour $\phi = \pi$ les deux ondes s'annulent et l'intensité totale vaut $I = 0$. Lorsque $\phi = \pi/2$ ou $\phi = 3\pi/2$ le terme d'interférence dans Eq. (2.30) s'annule et $I = 2I_0$. Différentes situations sont illustrées Fig. 2.11

La très forte dépendance de l'intensité I sur la différence de phase ϕ offre un moyen très précis de mesurer la phase d'une onde (par rapport à une onde de référence) en mesurant simplement une intensité.

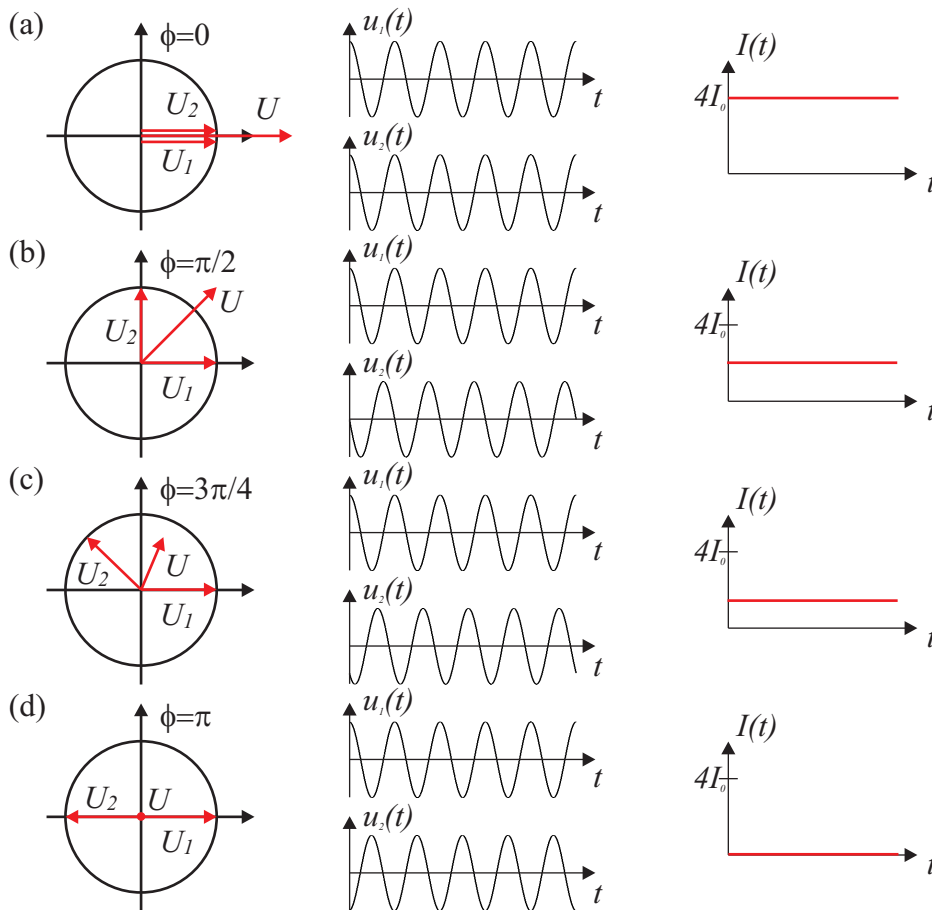


FIGURE 2.11 – Interférence de deux ondes pour différentes valeurs de la différence de phase ϕ . Pour chaque cas on représente les phaseurs des deux ondes U_1 et U_2 ainsi que le phaseur de l'interférence U ; l'amplitude des ondes et de l'interférence en fonction du temps.

2.8 Interféromètres

Un interféromètre est un instrument qui permet de mesurer avec grande précision des déplacements très petits. Son principe repose sur l'interférence de deux ondes co-planaires. Considérons deux ondes planes se propageant dans la direction z :

$$\begin{aligned} U_1 &= \sqrt{I_0} e^{-jkz}, \\ U_2 &= \sqrt{I_0} e^{-jk(z-d)}, \end{aligned} \quad (2.32)$$

telles qu'une onde est décalée d'une distance d par rapport à l'autre. Ceci se produit par exemple si une onde a parcouru une distance L_1 alors que l'autre a parcouru une distance $L_2 = L_1 + d$, comme illustré Fig. 2.9(a). L'équation d'interférence (2.30) s'écrit dans ce cas

$$I = 2I_0 \left[1 + \cos \left(2\pi \frac{d}{\lambda} \right) \right]. \quad (2.33)$$

La variation de l'intensité de l'onde résultante en fonction de d est représentée Fig. 2.12. On remarque des interférences constructives avec une intensité maximale $I = 4I_0$ lorsque d est un multiple de λ . À l'inverse, on a une interférence destructive avec une intensité nulle lorsque d est un multiple impair de $\lambda/2$. L'intensité moyenne est la somme des intensités : $I = 2I_0$.

Un interféromètre fonctionne exactement selon le principe illustré Fig. 2.12 : un faisceau incident est séparé en deux ondes dont l'une subit un détournement d . Ces deux ondes sont ensuite recombinaisonnées pour former une figure d'interférence. Comme l'intensité I est sensible à la phase $\phi = 2\pi d/\lambda = 2\pi dn/\lambda_0$, un interféromètre permet de mesurer un changement de distance d , un changement d'indice de réfraction n , ou un changement de longueur d'onde λ .

La figure 2.13 illustre différents types d'interféromètres. Notez que l'on peut aussi réaliser des interféromètres dans lesquels la lumière ne se propage pas dans l'espace libre, mais dans des fibres optiques ; voire d'intégrer complètement un interféromètre dans des semiconducteurs.

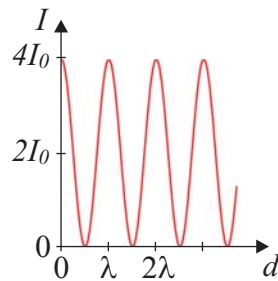


FIGURE 2.12 – Intensité de la superposition de deux ondes de même intensité I_0 en fonction de la différence de marche d .

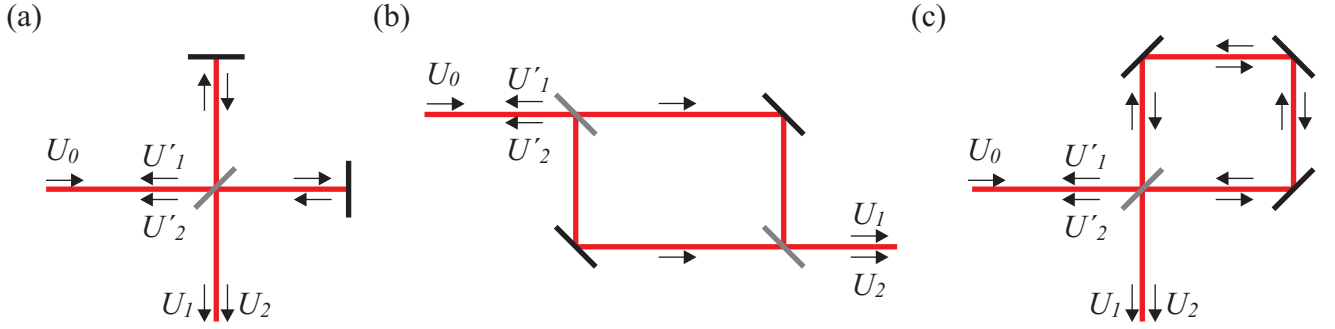


FIGURE 2.13 – Différent types d'interféromètres : (a) Michelson, (b) Mach-Zehnder, et (c) Sagnac.

2.9 Interférence d'ondes obliques

La figure 2.14 indique l'intensité résultant de l'interférence de deux ondes monochromatiques dont l'une se propage dans une direction faisant un angle θ avec la direction horizontale et l'autre un angle $-\theta$:

$$\begin{aligned} U_1 &= \sqrt{I_0} e^{-j(k \cos \theta x - k \sin \theta y)}, \\ U_2 &= \sqrt{I_0} e^{-j(k \cos \theta x + k \sin \theta y)}. \end{aligned} \quad (2.34)$$

En observant l'intensité de l'onde résultante, on remarque qu'elle est constante dans la direction de propagation x : l'onde résultante se propage dans cette direction. Par contre, dans

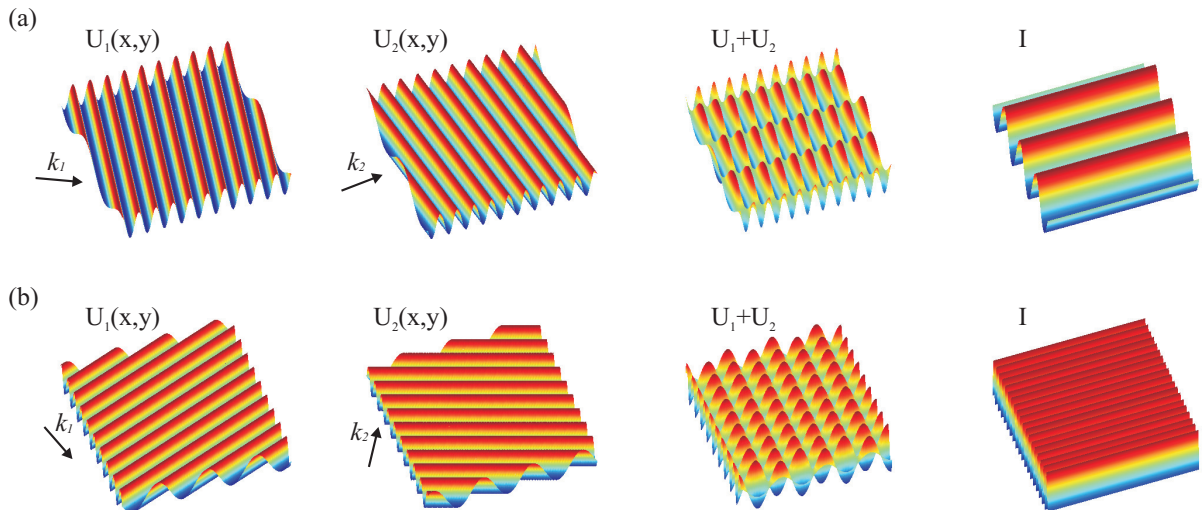


FIGURE 2.14 – Interférence de deux ondes se propageant avec un angle $\pm\theta$ par rapport à l'axe horizontal x . (a) $\theta = 10^\circ$ et (b) $\theta = 70^\circ$. Pour chaque valeur de θ on représente les ondes U_1 et U_2 , leur superposition $U_1 + U_2$ ainsi que l'intensité de la superposition.

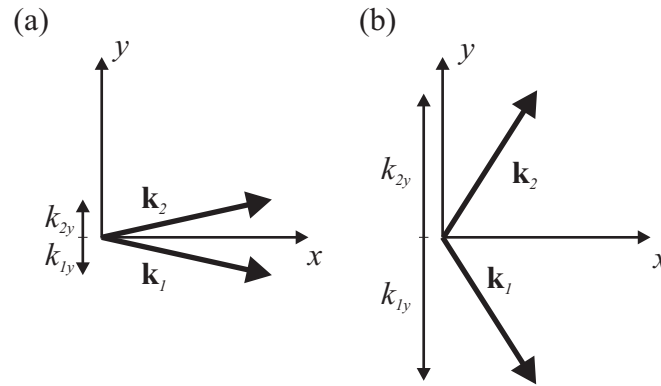


FIGURE 2.15 – Interférence de deux ondes : combinaison des vecteurs de propagation pour deux angles d'incidence différents (voir aussi la figure précédente).

la direction y l'intensité a un profil sinusoïdal dont la période dépend de l'angle d'incidence θ . Si l'angle d'incidence est grand, la période est petite, Fig. 2.14(b) et inversement. On peut comprendre ce comportement en considérant les vecteurs d'onde résultants, comme illustré Fig. 2.15. Les composantes dans la direction y des deux vecteurs d'onde $\mathbf{k}_1$ et $\mathbf{k}_2$ sont proportionnelles à $\sin \theta$; elles sont petites pour un petit angle d'incidence θ . En se souvenant que la longueur d'onde est inversement proportionnelle au vecteur d'onde, on comprend que dans ce cas la période dans la direction y est très grande, Fig. 2.14(a). Lorsque l'angle θ augmente, les composantes selon y des vecteurs de propagation augmentent, Fig. 2.15(b), correspondant alors à une petite longueur d'onde dans cette direction, Fig. 2.14(b).

Une autre caractéristique importante de l'onde résultante se déduit de Fig. 2.15 : alors que dans la direction de propagation les composantes x des vecteurs d'onde s'additionnent, ils s'annulent dans la direction y . L'onde ne se propage pas dans la direction y où l'on a affaire à une onde stationnaire (si on résonne en termes de photons, on peut imaginer que les quantités de mouvement dans la direction y de deux photons appartenant chacun à U_1 ou U_2 s'annulent ($\mathbf{k}_1$ est opposé à $\mathbf{k}_2$), alors que chaque photon se propage dans la direction x . On retrouvera ce genre d'ondes avec une composante stationnaire dans une direction et une composante propageante dans l'autre direction dans les guides d'ondes.

Notons finalement que la superposition des deux ondes permet de créer un profil sinusoïdal dans un photoresist, dont on peut ajuster la période en changeant l'angle d'incidence. Ce processus de lithographie par interférence est très utilisé pour fabriquer des nanostructures périodiques.

2.10 Cohérence temporelle

Dans notre discussion de l'interférence nous avons supposé que la source lumineuse était telle que l'on pouvait partager son faisceau en deux ondes, leur faire prendre des chemins

différents et les recombinaient pour observer la figure d'interférence. Dans une expérience avec une source réelle, en utilisant par exemple un interféromètre de Michelson – Fig. 2.13(a) – on s'aperçoit que lorsque la différence de chemin optique d augmente, l'amplitude des franges d'interférence diminue, comme illustré Fig. 2.16. On définit alors la visibilité des franges d'interférence :

$$\mathcal{V} = \frac{I_{\max} - I_{\min}}{I_{\max} + I_{\min}}, \quad (2.35)$$

où $I_{\max}$ et $I_{\min}$ se réfèrent à la figure d'interférence. On définit généralement la visibilité pour la frange centrale de l'interférogramme, mais elle peut aussi être définie en n'importe quel point de l'interférogramme, Fig. 2.16. Si la visibilité vaut $\mathcal{V} = 1$ il y a parfaite cohérence entre les deux faisceaux (dans ce cas $I_{\min} = 0$ et l'interférence est complète). Si $\mathcal{V} = 0$ les faisceaux sont incohérents (dans ce cas $I_{\min} = I_{\max}$ et il n'y a pas d'interférence).

On définit ainsi la longueur de cohérence L_c de la source comme la différence de chemin optique d sur laquelle la visibilité des franges sur un interférogramme de Michelson est supérieure ou égale à la moitié de la visibilité maximale à $d = 0$. La longueur de cohérence d'un laser He-Ne varie entre quelques dizaines de centimètres et plusieurs dizaines de mètres suivant que le laser est multi-mode ou single-mode, celle d'une diode laser peut atteindre une centaine de mètres, alors que celle d'une source thermique est quasi-nulle.

A partir de la longueur de cohérence, on définit le temps de cohérence :

$$\tau_c = \frac{L_c}{c}. \quad (2.36)$$

La longueur de cohérence est essentielle pour mettre en oeuvre une expérience d'interférence : il faut veiller à ce que la différence de chemin optique soit bien inférieure à L_c et choisir une source avec une longueur de cohérence appropriée pour un montage donné.

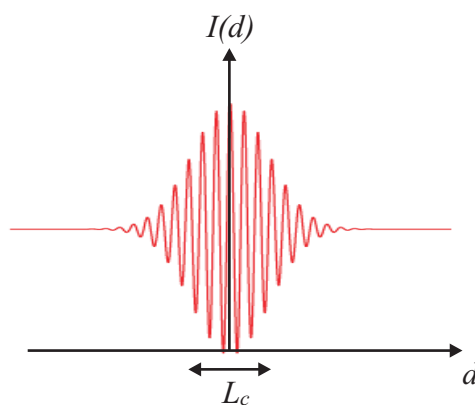


FIGURE 2.16 – Visibilité des franges d'interférence.

2.11 Battements et vitesse de groupe

Considérons maintenant le cas de la superposition de deux ondes planes harmoniques avec des pulsations ω_1 et ω_2 différentes et des vecteurs de propagation différents. Pour simplifier le calcul on considérera des ondes se propageant dans la même direction z , mais cette condition de co-linéarité n'est pas indispensable et l'on peut évidemment développer ce formalisme pour une superposition d'ondes quelconques. On supposera aussi que les amplitudes sont les mêmes ; on aura donc une modulation complète, Fig. 2.17(a). Les deux ondes sont

$$U_1(z, t) = U_0 e^{j\omega_1 t - jk_1 z} \quad (2.37a)$$

$$U_2(z, t) = U_0 e^{j\omega_2 t - jk_2 z} . \quad (2.37b)$$

Pour faire le calcul de la superposition, on introduit les grandeurs suivantes : $2k_0 = k_1 + k_2$, $2\omega_0 = \omega_1 + \omega_2$, $2\Delta k = k_2 - k_1$ et $2\Delta\omega = \omega_2 - \omega_1$.

Le champ total s'écrit

$$U(z, t) = U_1(z, t) + U_2(z, t) = U_0 \left(e^{j(\omega_0 - \Delta\omega)t - j(k_0 - \Delta k)z} + e^{j(\omega_0 + \Delta\omega)t - j(k_0 + \Delta k)z} \right) . \quad (2.38)$$

En factorisant le terme en k_0 ω_0 , on obtient

$$\begin{aligned} U(z, t) &= U_0 e^{j\omega_0 t - jk_0 z} \left(e^{j\Delta\omega t - j\Delta k z} + e^{-j\Delta\omega t - j\Delta k z} \right) \\ &= 2U_0 e^{j\omega_0 t - jk_0 z} \cos(\Delta\omega t - \Delta k z) . \end{aligned} \quad (2.39)$$

On remarque que l'onde résultante est le produit de deux ondes : l'onde porteuse d'amplitude $2U_0$ caractérisée par ω_0 et k_0 et l'onde de modulation ou de battement d'amplitude 1 caractérisée par $\Delta\omega$ et Δk , Fig. 2.17.

Dans l'exemple de Fig. 2.17(a), l'onde porteuse et l'enveloppe (la modulation) se propagent à la même vitesse. Ce n'est cependant pas le cas en général et cet effet nous permet d'introduire deux types de vitesses pour une onde : la vitesse de phase et la vitesse de groupe. Comme son nom l'indique, la vitesse de phase décrit la vitesse à laquelle l'onde accumule de la phase. La vitesse de groupe décrit la vitesse à laquelle l'onde transporte de l'information, par exemple la modulation sur Fig. 2.17(a). Lorsque l'on forme un pulse de lumière en superposant plusieurs ondes monochromatiques de fréquences différentes, l'enveloppe du pulse peut se propager à

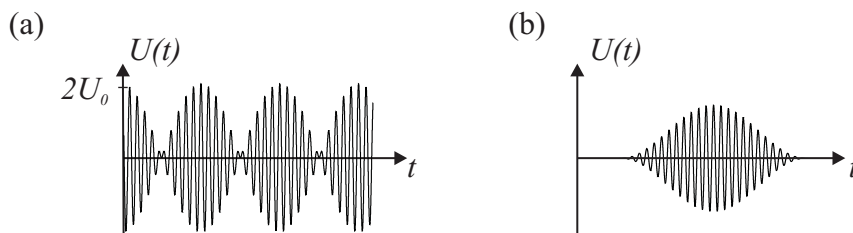


FIGURE 2.17 – (a) Superposition de deux ondes de fréquences différentes donnant lieu à un battement de l'onde résultante. (b) Pulse obtenu par la superpositions de nombreuses ondes.

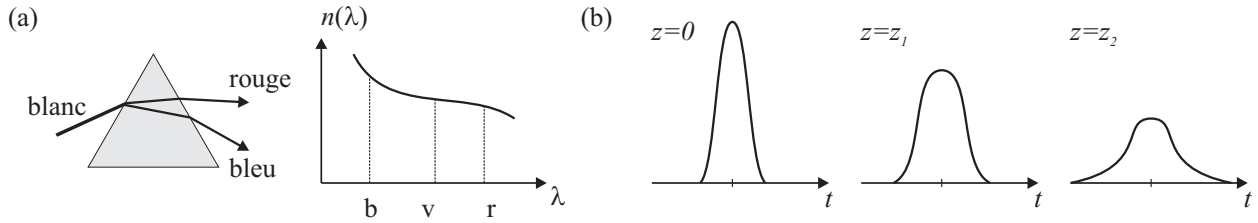


FIGURE 2.18 – (a) La dispersion du verre donne lieu à la séparation de la lumière blanche entre ses différentes composantes spectrales. (b) Un pulse de lumière se propageant dans un milieu dispersif s'étale durant la propagation.

une vitesse différente de l'onde porteuse, Fig. 2.17(b). Ce phénomène se produit lorsque le milieu dans lequel le pulse se propage est dispersif ; i.e. son indice de réfraction $n(\omega)$ dépend de la fréquence ω de l'onde. La plupart des milieux optiques sont dispersifs et même une très petite variation de l'indice de réfraction peut donner lieu au phénomène spectaculaire de la dispersion de la lumière blanche à travers un prisme, Fig. 2.18(a).

Pour un pulse de lumière composé de la superposition de nombreuses ondes de fréquences différentes, la dispersion du milieu de propagation donne lieu à un étalement du pulse durant sa propagation, Fig. 2.18(b). Cet étalement est très préjudiciable pour la bande passante du système, puisque plusieurs pulses se succédant à brefs intervalles vont se fondre les uns dans les autres au fur et à mesure de la propagation. Les systèmes de transmission de l'information font donc appel à des composants qui permettent de compenser cette dispersion et de remettre les pulses en forme après une certaine distance de propagation.

Introduisons la définitions de ces deux vitesses. La figure 2.19(a) illustre la relation qui existe entre ω et k pour les deux ondes de Fig. 2.17(a). La vitesse de phase est la vitesse usuelle que nous avons définie pour une onde :

$$c = \frac{\omega}{k} ; \quad (2.40)$$

dans le vide cette vitesse vaut c_0 .

La vitesse de groupe v est la vitesse à laquelle se déplace l'enveloppe. Elle est définie par

$$v = \frac{\partial \omega}{\partial k} . \quad (2.41)$$

Pour le battement (2.39) illustré Fig. 2.17(a) cette vitesse est la même que la vitesse de phase, mais en général tel n'est pas le cas. Comme c'est cette enveloppe qui transporte l'information associée avec l'onde (dans ce cas la modulation de son amplitude), la vitesse de groupe représente la vitesse à laquelle l'information est transportée.

Il est essentiel de bien distinguer ces deux vitesses ; en effet la vitesse de phase peut être dans certains cas supérieure à la vitesse de la lumière c , ce qui peut sembler violer le principe de causalité. En réalité tel n'est pas le cas, car l'information (ou l'énergie) se déplace avec la vitesse de groupe, qui est toujours inférieure à c .

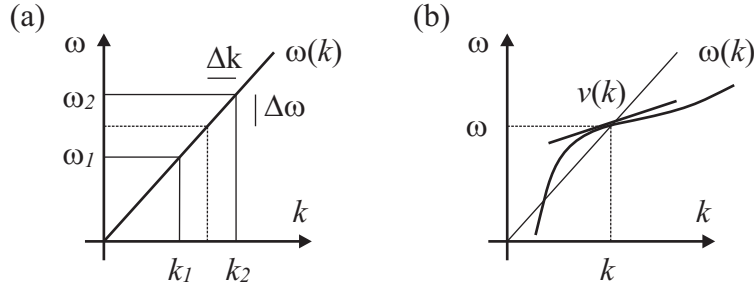


FIGURE 2.19 – Diagramme de dispersion (a) pour l'onde Eq. (2.39), (b) pour une onde se propageant dans un milieu dispersif.

Les vitesses de phase et de groupe diffèrent si la réponse du milieu est telle que la relation entre ω et k n'est pas linéaire, comme indiqué Fig. 2.19(b). Un tel milieu est dispersif et la vitesse de groupe est donnée par la pente de la fonction $\omega(k)$. Le diagramme Fig. 2.19 s'appelle un diagramme de dispersion. Il permet de comprendre la relation entre l'énergie (fréquence ou pulsation) et la quantité de mouvement (vecteur de propagation) d'une onde dans un milieu donné. Ce diagramme joue un rôle essentiel dans la description des matériaux optiques complexes tels les cristaux photoniques ou les métamatériaux.

Nous allons maintenant obtenir une expression générale pour la vitesse de groupe dans un milieu d'indice n en utilisant les définitions de l'indice de réfraction (2.2), de la vitesse de phase (2.41) et de la vitesse de groupe (2.41). En partant de

$$n = \frac{c_0}{c} = \frac{c_0 k}{\omega}, \quad (2.42)$$

on obtient

$$\frac{\partial n}{\partial \omega} = -\frac{c_0 k}{\omega^2} + \frac{c_0}{\omega} \frac{\partial k}{\partial \omega} = \frac{c_0}{\omega} \left(\frac{\partial k}{\partial \omega} - \frac{k}{\omega} \right) = \frac{c_0}{\omega} \left(\frac{1}{v} - \frac{1}{c} \right); \quad (2.43)$$

où on a remarqué que c_0 ne dépend pas de ω . On obtient ainsi pour v :

$$\frac{1}{v} = \frac{\omega}{c_0} \frac{\partial n}{\partial \omega} + \frac{1}{c}. \quad (2.44)$$

Soit finalement,

$$v = \frac{c}{1 + \frac{\omega}{n} \frac{\partial n}{\partial \omega}} = \frac{c_0}{n + \omega \frac{\partial n}{\partial \omega}}. \quad (2.45)$$

Equation (2.45) indique que si le milieu n'est pas dispersif, i.e. l'indice de réfraction ne dépend pas de la fréquence et le terme $\partial n / \partial \omega$ dans Eq. (2.45) est nul, alors la vitesse de groupe est égale à la vitesse de phase : $v = c_0 / n$. Dans un milieu dispersif par contre, l'indice de réfraction augmente en général avec la fréquence (la dérivée $\partial n / \partial \omega$ est positive, soit la dérivée $\partial n / \partial \lambda$ est négative, Fig. 2.17(a)) et la vitesse de groupe est plus petite que la vitesse de phase.

Chapitre 3

Optique de Fourier et diffraction

Dans ce chapitre nous allons étudier un phénomène essentiel de l'optique : la diffraction. Si nous considérons l'image projetée par une ouverture dans un écran opaque éclairé par une onde plane, l'optique géométrique nous indique que les zones d'ombre et de lumière sont clairement délimitées et s'obtiennent par le tracé des rayons, Fig. 3.1(a). Or la réalité est très différente, comme illustré Fig. 3.1(b) : si l'on observe bel et bien une zone claire au milieu de l'écran, celle-ci peut être entourée de plusieurs autres zones foncées et claires concentriques. De plus la transition entre ombre et lumière n'est pas abrupte. Ces différences sont la manifestation du phénomène de diffraction.

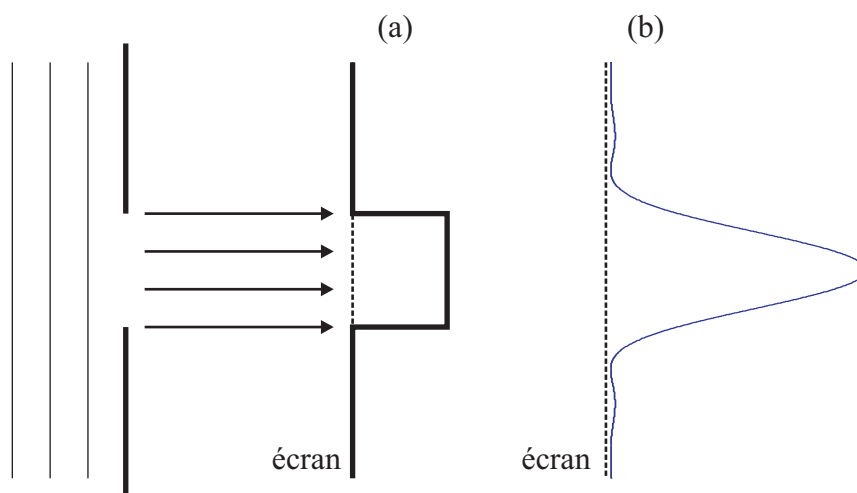


FIGURE 3.1 – Phénomène de la diffraction : d'après (a) l'optique géométrique et (b) l'optique ondulatoire.

3.1 Principe de Huygens

Le physicien hollandais Christiaan Huygens a énoncé en 1690 déjà le principe qui porte son nom : "Tout point d'un front d'onde se comporte comme une source secondaire d'ondelettes sphériques en sorte que l'enveloppe de ces ondelettes forme le front d'onde à un instant ultérieur." Ce principe est illustré Fig. 3.2 pour la diffraction par deux fentes A et B . Notons que la longueur d'onde de chaque ondelette est la même que la longueur d'onde de l'onde initiale.

Le phénomène de diffraction observé Fig. 3.1(b) peut alors s'expliquer en considérant que chaque point du front d'onde incident dans la fente est la source d'une ondelette sphérique secondaire. L'amplitude du champ optique en un point quelconque derrière l'écran est donc la superposition de toutes ces ondelettes obtenue en tenant compte de leur amplitude et de leur phase relative, Fig. 3.3.

Considérons la diffraction par une ouverture de dimension a dans un écran opaque illuminé par une onde plane incidente de la gauche, Fig. 3.3. Le principe de Huygens implique que chaque point de l'ouverture devient la source d'une ondelette sphérique de même fréquence. L'amplitude de l'onde transmise derrière l'écran est donc la superposition de toutes ces ondelettes. Sur l'axe de l'ouverture ($\theta \simeq 0$, par exemple au point P_1 Fig. 3.3) on retrouve essentiellement un front d'onde plan puisque les sources qui y contribuent dans l'ouverture sont symétriques par rapport à ce point. Tel n'est pas le cas sur les côtés, pour des angles θ plus larges, par exemple au point P_2 Fig. 3.3 : les ondes émises par le bord de l'ouverture ne sont pas compensées par des sources le long de l'écran pour former un front d'onde plan. Elles

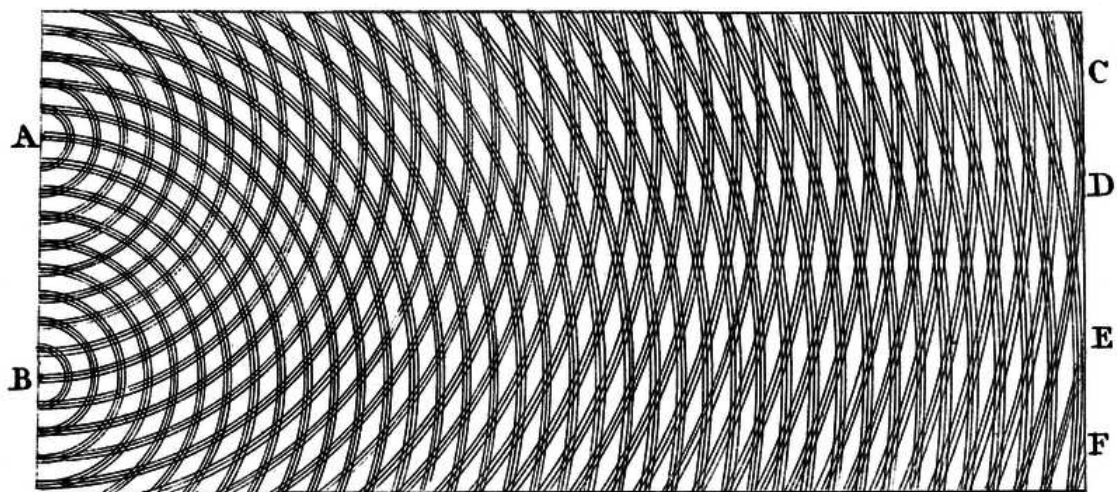


FIGURE 3.2 – Dessin original de Thomas Yung pour la diffraction par deux fentes selon le principe de Huygens : chaque fente A et B est la source d'ondelettes sphériques qui se propagent et interfèrent, donnant lieu à des lignes claires où les ondes interfèrent constructivement et de lignes foncées où elles interfèrent destructivement.

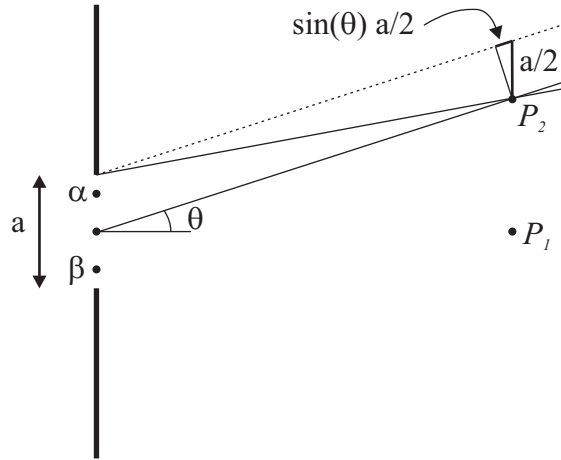


FIGURE 3.3 – Diffraction par une fente : approche simple dans le champ lointain en supposant les rayons parallèles.

forment un front d'onde courbe qui ouvre le faisceau et le fait diverger. De plus, la différence de chemin optique pour les ondelettes émises par deux points symétriques α et β peut donner lieu à une interférence destructive au point P_2 lorsque la différence de chemin optique est un multiple impair de $\lambda/2$. Si on suppose que l'on est très loin de la fente en comparaison de sa largeur, on peut considérer les trajets optiques comme parallèles et calculer la différence de chemin optique entre le rayon venant du centre de la fente et celui venant du bord supérieur de la fente (ces rayons sont donc espacés de $a/2$). Lorsque leur différence de chemin est égale à un multiple impair de $\lambda/2$ l'interférence est destructive. Cette condition s'écrit,

$$\sin \theta \frac{a}{2} = q \frac{\lambda}{2}, \quad q = 1, 3, 5, \dots \quad (3.1)$$

Ainsi la zone d'illumination principale transmise par la fente est-elle limitée par les angles $\pm\theta$ tels que (on pose simplement $q = 1$ dans l'équation ci-dessus) :

$$\sin \theta = \lambda/a. \quad (3.2)$$

Si l'ouverture est petite par rapport à la longueur d'onde la zone d'illumination centrale est large. On comprend dans ce cas que l'ouverture se comporte presque comme une source ponctuelle (on ne peut mettre qu'une source ponctuelle dans l'ouverture, Fig. 3.3) produisant une onde sphérique.

La discussion précédente met en évidence le rôle joué par les ondes secondaires lorsque l'on cherche à comprendre comment une image est formée après que la lumière a interagit avec des objets. Ainsi, dans le cas de Fig. 3.3, ce sont les ondelettes “émises” par les différents points de la fente qui donnent lieu à une figure d'interférence. Cette approche va nous conduire à la décomposition *dans l'espace* en série de Fourier d'un front d'onde quelconque. La décomposition d'un signal *temporel* $f(t)$ en une superposition d'harmoniques de fréquences différentes forme la base du traitement du signal, Fig. 3.4(a). On parle dans ce cas d'un signal à une dimension (par exemple le temps t), qui est évidemment la superposition d'harmoniques à une

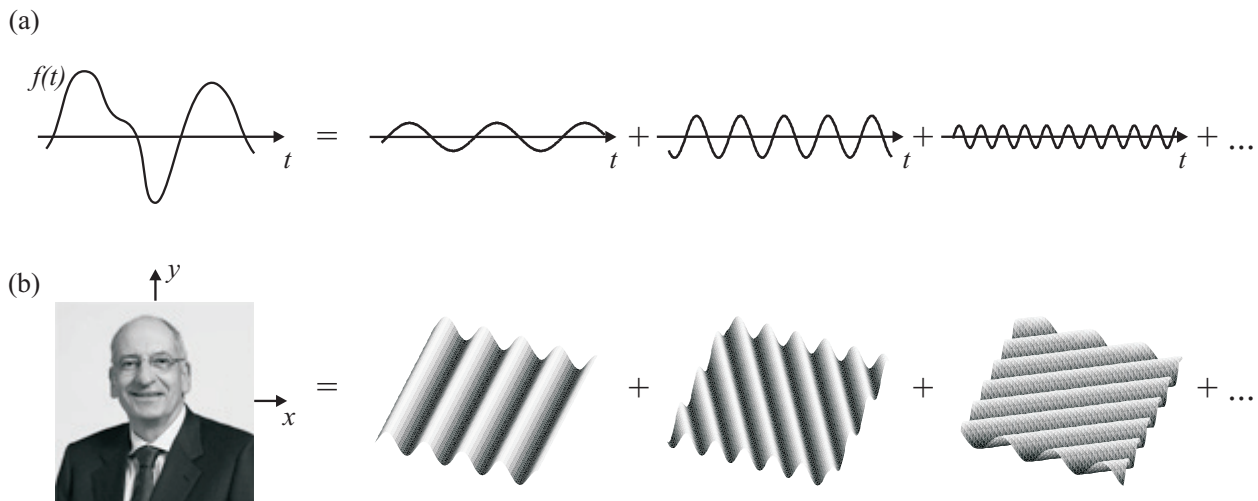


FIGURE 3.4 – (a) Un signal temporel peut se décomposer en une superposition d’harmoniques de fréquences différentes. L’amplitude de chaque harmonique détermine sa contribution au signal total. (b) De même, une image dans le plan (x, y) peut se décomposer en une superposition d’ondes planes de fréquences spatiales différentes se propageant dans des directions différentes.

dimension (toujours le temps t) ; chaque harmonique étant caractérisée par une fréquence donnée, Fig. 3.4(a). L’analyse de Fourier du signal original permet de déterminer l’amplitude de chaque harmonique, i.e. la contribution de chaque harmonique au signal original. Si le signal original est simplement sinusoïdal de fréquence ν , alors seulement cette harmonique de fréquence ν participe au signal et les amplitudes de toutes les autres harmoniques sont nulles. Si le signal original est plus compliqué, il faut superposer des harmoniques de fréquences diverses pour le reconstruire, Fig. 3.4(a).

En optique, une image est généralement une distribution d’intensité lumineuse dans un plan (x, y) , comme par exemple les niveaux de gris dans la photographie Fig. 3.4(b). Ce signal à deux dimensions peut aussi se décomposer en fonctions harmoniques. Celles-ci sont alors définies dans le plan (x, y) et ont chacune une fréquence *spatiale* bien définie, correspondant à l’espacement entre les maxima, Fig. 3.4(b). Comme maintenant on travaille dans un espace à deux dimensions, en plus de la fréquence spatiale de l’harmonique, on peut aussi choisir sa direction de propagation dans le plan (x, y) . Ainsi les trois composantes spectrales illustrées Fig. 3.4(b) ont-elles non seulement des fréquences spatiales différentes, mais aussi des directions de propagation différentes. Chacune des ces trois harmoniques correspond à une onde plane se propageant dans le plan (x, y) . Comme nous le verrons c’est là le lien fondamental entre optique de Fourier et ondes planes.

3.2 Propagation dans l'espace libre – Cas 2D

L'optique de Fourier repose essentiellement sur les ondes planes pour lesquelles elle considère séparément la direction de propagation (nous choisirons la direction z) et la direction transverse. Dans le cas général, il y a deux directions transverses, x et y ; afin de mieux représenter ce qui se passe, nous allons cependant considérer dans un premier temps un système à 2 dimensions (2D), avec une direction transverse (la direction x) et une direction de propagation (la direction z). On omettra donc pour l'instant toute dépendance dans la coordonnée y . Une onde plane d'amplitude complexe s'écrit alors

$$U(x, z) = A \exp[-j(k_x x + k_z z)], \quad (3.3)$$

où on a omis la dépendance dans le temps $\exp(j\omega t)$ puisqu'on s'intéresse ici à la dépendance spatiale de ces ondes. Comme cette onde doit satisfaire l'équation d'Helmoltz (2.10), les composantes du vecteur d'onde 2D $\mathbf{k} = (k_x, k_z)$ satisfont Eq. (2.11) dont on prend le carré :

$$\sqrt{k_x^2 + k_z^2} = k = \frac{2\pi}{\lambda} = \frac{\omega}{c}. \quad (3.4)$$

L'amplitude du vecteur d'onde est donc déterminée par la longueur d'onde λ dans le milieu considéré (rappelons que $\lambda = \lambda_0/n$, où λ_0 est la longueur d'onde dans le vide et n l'indice de réfraction du milieu), ou – ce qui est équivalent – par la fréquence angulaire ω de l'onde. Ceci est illustré sur Fig. 3.5 par le cercle qui repère l'extrémité du vecteur d'onde $\mathbf{k}$. Une

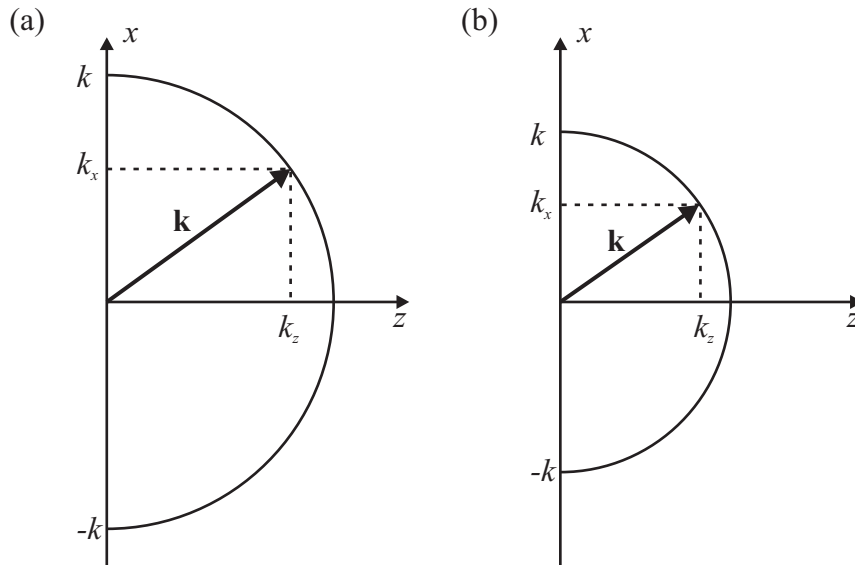


FIGURE 3.5 – L'extrémité du vecteur d'onde 2D se trouve sur un cercle dont le rayon est déterminé par la longueur d'onde de l'onde dans le milieu considéré. Deux différentes longueurs d'onde sont étudiées, (a) λ_1 et (b) λ_2 avec $\lambda_1 < \lambda_2$ (comme $k = 2\pi/\lambda$, le rayon du cercle est inversement proportionnel à λ).

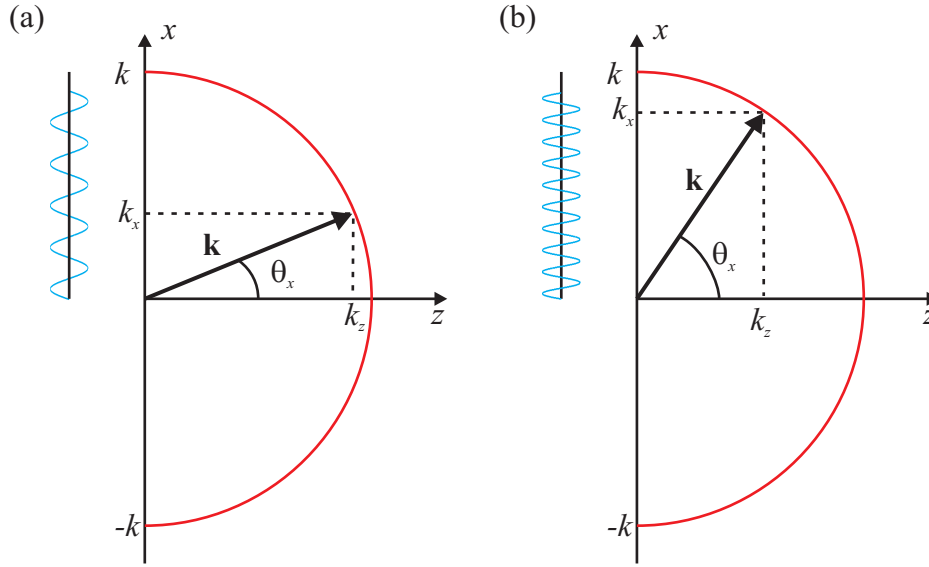


FIGURE 3.6 – Pour une longueur d'onde donnée, des ondes planes se propageant dans des directions différentes θ_x ont des composantes transverses k_x du vecteur d'onde différentes. (b) Lorsque k_x est grand, l'onde a une fréquence spatiale importante dans la direction transverse où elle oscille rapidement.

onde peut donc se propager dans n'importe quelle direction, pour autant que l'extrémité de son vecteur d'onde se trouve sur le cercle. Des ondes de longueur d'onde différente auront donc des cercles de rayons différents (plus la longueur d'onde est grande, plus le rayon est petit), comparer Figs. 3.5(a) et (b).

Pour une longueur d'onde donnée, toutes les ondes planes propageantes ont un vecteur $\mathbf{k}$ sur le cercle de rayon k . Suivant sa direction de propagation, le vecteur de propagation a cependant une composante k_x différente, comme illustré Fig. 3.6. Il est possible d'écrire cette composante en fonction de la fréquence spatiale $\nu_x = k_x/2\pi$ dans la direction x :

$$k_x = 2\pi\nu_x. \quad (3.5)$$

Il faut faire très attention à la définition de la fréquence *spatiale*, dont les unités sont des cycles/mètre (donc en relation avec l'espace), alors que la fréquence optique, $\nu = kc/2\pi$ a pour unité des cycles/seconde ou Hz (donc en relation avec le temps). Sur Fig. 3.6 on a indiqué la fréquence spatiale ν_x associée à chaque onde sous forme d'une modulation périodique dans la direction transverse x . On remarque que l'onde qui a la plus grande composante k_x a la plus grande fréquence spatiale ν_x , i.e. la plus rapide variation transverse. Ainsi, plus l'angle de propagation θ_x est important, plus la fréquence spatiale ν_x est importante, Fig. 3.6. Pour une onde plane se propageant dans la direction z ($k_x = 0$) la fréquence spatiale est nulle : $\nu_x = 0$. En d'autres termes, l'onde a dans ce cas une amplitude constante dans la direction transverse x .

Jusqu'à présent, nous avons considéré des ondes planes se propageant dans l'espace libre et fait des observations sur leur fréquence spatiale ν_x dans la direction transverse. En optique

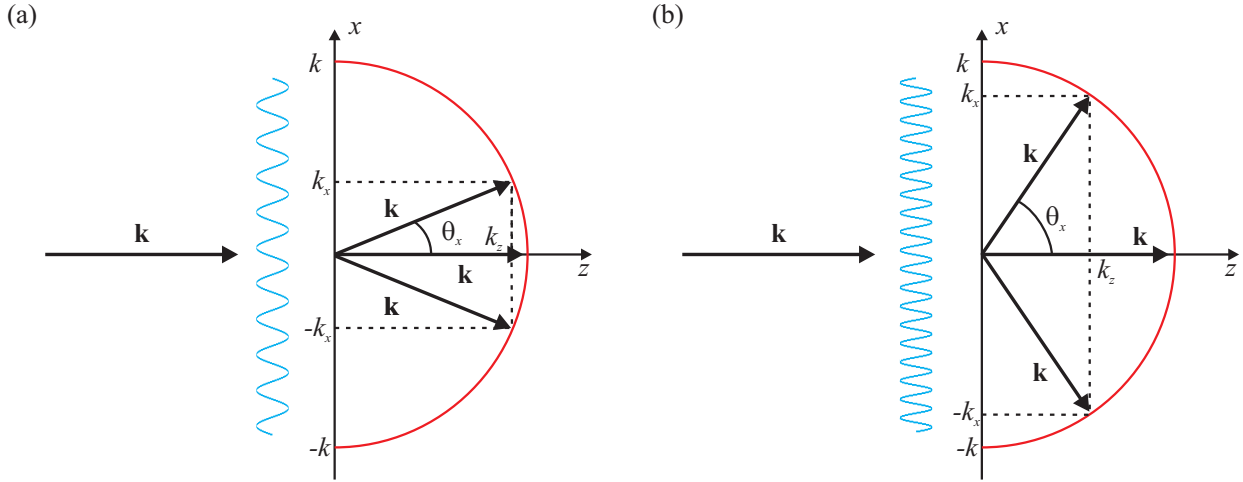


FIGURE 3.7 – Une onde plane incidente sur une surface transparente avec un profil périodique est déviée dans des directions transverses telles que la composante transverse k_x de l'onde transmise corresponde à la fréquence spatiale ν_x de la surface. Une partie de l'onde incidente n'est pas déviée et continue tout droit. Différentes fréquences spatiales pour la surface donnent lieu à différentes directions θ_x pour l'onde transmise.

de Fourier, on s'intéresse au processus inverse, dans lequel une onde plane est définie à partir de sa fréquence spatiale transverse. Ainsi, un objet ayant un profil donné dans la direction x sera décomposé en une série d'ondes planes se propageant dans des directions θ_x différentes ; chaque onde plane correspondant à une composante de Fourier de fréquence spatiale ν_x de la surface, Fig. 3.7. On peut dire que l'onde plane transmise emporte avec elle l'information sur la fréquence spatiale du profil de la surface et que cette information est codée dans la direction de propagation de l'onde. Sur Fig. 3.7 on remarque qu'une onde avec un vecteur de propagation $\mathbf{k}$ arrivant à incidence normale sur une surface sinusoïdale de fréquence spatiale ν_x repart dans une direction qui dépend de cette fréquence spatiale : plus elle est élevée, plus l'onde est fortement déviée. En utilisant Eq. (3.5), on remarque que dans Fig. 3.7(a) la composante transverse k_x du vecteur d'onde est petite, alors qu'elle est grande dans le cas de Fig. 3.7(b). En plus de l'onde sortante déviée vers le haut indiquée sur Fig. 3.7(a) ou (b), il existe aussi une onde symétrique, déviée vers le bas avec le même angle θ_x .

Jusqu'à présent, nous nous sommes limités à des ondes propagatees ; i.e. des ondes dont l'extrémité du vecteur de propagation se trouve sur le cercle de rayon k . Ces ondes ont donc une composante transverse du vecteur de propagation telle que $k_x \leq 2\pi/\lambda$, Eq. (3.4). En utilisant Eq. (3.5), cette condition s'écrit

$$\nu_x \leq 1/\lambda \text{ ou } \lambda \leq 1/\nu_x. \quad (3.6)$$

Cette simple équation a des implications extrêmement importantes pour la résolution spatiale des systèmes optiques. En effet, si on souhaite faire l'image d'une surface de profil sinusoïdal de très haute fréquence ν_x , il faut utiliser une longueur d'onde très courte. En d'autres termes, plus la longueur d'onde est petite, plus on sera à même d'observer les détails d'une surface.

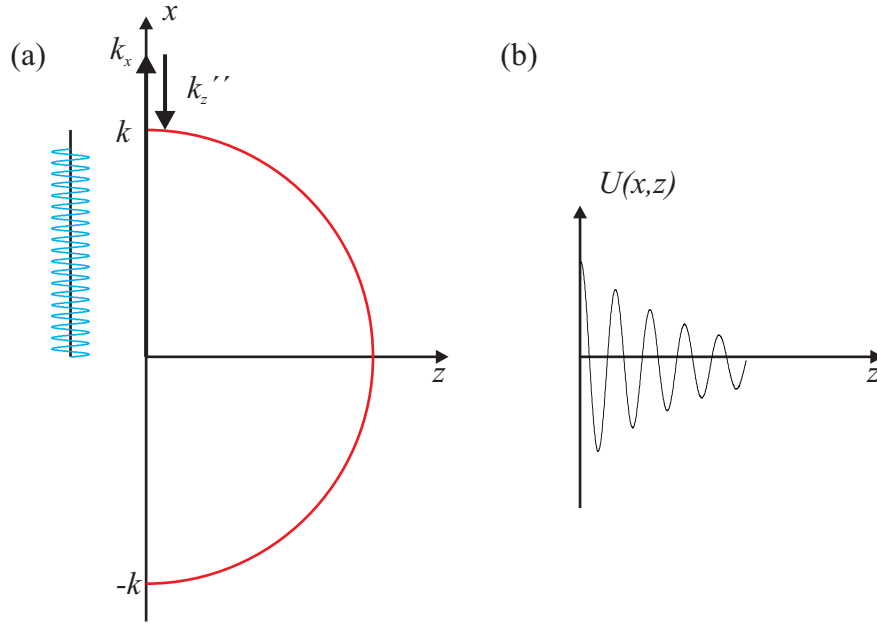


FIGURE 3.8 – (a) Une onde correspondant à une fréquence spatiale transverse ν_x très élevée a une composante k_z du vecteur de propagation purement imaginaire, afin de satisfaire la relation de dispersion. (b) Une telle onde ne se propage pas, mais décroît exponentiellement.

Si on considère maintenant que c'est le profil de la surface qui définit les directions de propagation des ondes, on a de la peine à comprendre que les fréquences spatiales décrivant une surface de profil arbitraire soient limitées par la longueur d'onde utilisée. Si par exemple une surface est très rugueuse, il faudra utiliser des fréquences spatiales élevées pour pouvoir la décrire et la limite imposée par Eq. (3.6) semble difficile à satisfaire, à moins que la longueur d'onde λ ne soit extrêmement petite. Si tel n'est pas le cas, nous allons voir que l'onde transmise deviendra évanescence. Considérons par exemple une onde plane ayant une composante transverse du vecteur de propagation $k_x > 2\pi/\lambda$. La composante k_z de cette onde est donnée par Eq. (3.4) :

$$k_z = \pm \sqrt{\frac{4\pi^2}{\lambda^2} - k_x^2} = \pm \sqrt{\frac{\omega^2}{c^2} - k_x^2} = -jk_z'', \quad (3.7)$$

où l'on a défini la partie imaginaire k_z'' de la composante k_z . En effet, dans ce cas, le nombre dans la racine carrée d'Eq. (3.7) est négatif et le résultat est purement imaginaire, Fig. 3.8 (on remarque que l'on a choisi le signe "<->" dans Eq. (3.7) pour garantir que l'onde n'est pas amplifiée – ce qui ne serait pas physique, comme discuté sur Fig. 2.6). Ainsi, l'onde ne se propage pas à l'infini dans la direction z mais décroît rapidement puisque son expression devient

$$U(x, z) = A \exp[-j(k_x x + k_z z)] = A \exp[-jk_x x] \exp[-k_z'' z]. \quad (3.8)$$

Il existe donc une limite pour la fréquence spatiale maximale qui peut être transmise par une onde de longueur d'onde λ : toutes les fréquences spatiales $\nu_x > 1/\lambda$ sont perdues lors

de la propagation. Comme ces fréquences spatiales correspondent à des détails très petits de l'objet, ceci explique qu'il existe une limite de résolution en optique. On remarque cependant que, comme ces composantes de haute fréquence s'évanouissent de façon exponentielle en s'éloignant de l'objet, on peut en retrouver une partie en mesurant la lumière à très petite distance de l'objet. Ce principe est mis en oeuvre en optique de champ proche, où une sonde locale est approchée de l'objet dont on souhaite imager les détails.

3.3 Réseau de diffraction, équation de Bragg

Sur Fig. 3.6, nous avons vu que, pour une longueur d'onde donnée, l'angle de déflexion θ_x dépend de la période de la surface utilisée. De même, si on considère une surface de période constante et qu'on y envoie deux ondes de longueur d'onde différentes, elles subiront des déflexions différentes, Fig. 3.9(a). Cette observation permet d'introduire un élément optique très utile : le réseau de diffraction.

Un réseau de diffraction est un composant optique dont la surface est modulée périodique-

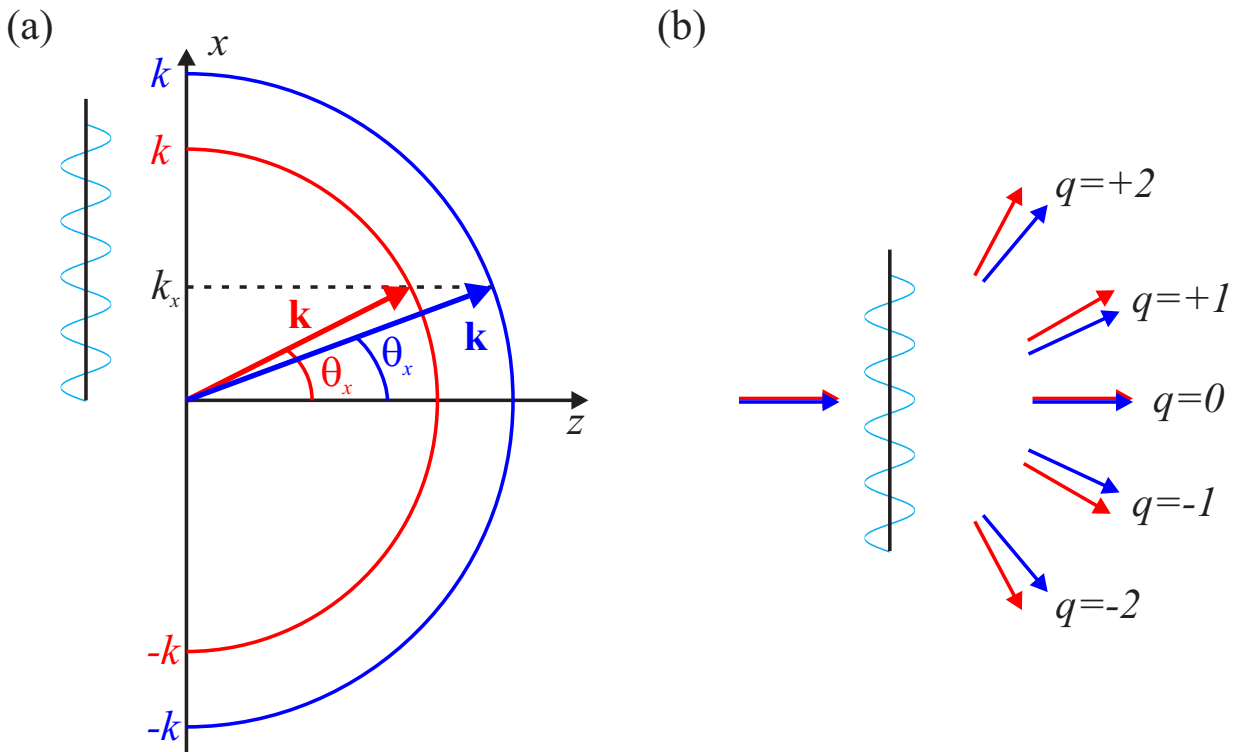


FIGURE 3.9 – Une onde plane incidente sur une surface transparente avec un profil périodique est déviée dans une direction telle que la composante transverse k_x de l'onde transmise corresponde à la fréquence spatiale ν_x de la surface. Différentes fréquences spatiales pour la surface donnent lieu à différentes directions θ_x pour l'onde transmise.

ment. Il peut s'agir d'une surface avec un profil sinusoïdal, carré ou triangulaire de période Λ ; parfois, on utilise aussi un réseau de fentes périodiques. Il peut s'utiliser en transmission ou en réflexion et sa fonction est de dévier l'onde incidente de longueur d'onde λ dans des directions θ_q définies par

$$\sin \theta_q = \sin \theta_i + q \frac{\lambda}{\Lambda}, \quad q = 0, \pm 1, \pm 2, \dots, \quad (3.9)$$

où θ_i est l'angle d'incidence et on a supposé que c'est le même matériau des deux côtés du réseau. Souvent, les réseaux sont utilisés avec de petits angles et on fait l'approximation $\sin \theta \simeq \theta$. Equation (3.9) devient alors

$$\theta_q = \theta_i + q \frac{\lambda}{\Lambda}, \quad q = 0, \pm 1, \pm 2, \dots. \quad (3.10)$$

Ce fonctionnement est illustré Fig. 3.9(b). Comme la longueur d'onde λ entre dans Eq. (3.9), on observe que les angles de déflexions sont différents pour chaque longueur d'onde. Ainsi, cet élément permet-il de séparer les différentes composantes spectrales contenues dans l'onde incidente. Il forme donc le cœur de tous les systèmes spectroscopiques.

Le fait que les angles de diffraction apparaissent par paire, $\pm q$, est une généralisation du phénomène que nous avons illustré Figs. 3.5–3.7. Dans ces figures, nous n'avons montré que la déflexion du faisceau incident vers le haut, mais dans la réalité, une partie de la lumière incidente est aussi défléchie vers le bas et une partie continue tout droit, comme illustré Fig. 3.9(b). Sur cette même figure, on remarque que l'angle entre les différentes composantes spectrales (par exemple le bleu et le rouge) augmente en fonction de l'ordre q .

Il existe en général plusieurs ordres de diffraction, correspondant à différentes valeurs de q dans Eq. (3.9). Leur nombre dépend du rapport λ/Λ : si ce rapport est petit, on a généralement beaucoup d'ordres de diffraction.

3.4 Transformée de Fourier

Avant d'aborder le cas général, il importe de rappeler brièvement la définition de la transformée de Fourier, pour le cas de deux variables d'espace x et y . La transformée de Fourier d'une fonction $f(x, y)$ s'écrit

$$F(\nu_x, \nu_y) = \int_{-\infty}^{+\infty} f(x, y) e^{j2\pi(\nu_x x + \nu_y y)} dx dy, \quad (3.11)$$

où ν_x et ν_y représentent les fréquences spatiales dans les directions x et y . La transformée de Fourier inverse s'écrit

$$f(x, y) = \int_{-\infty}^{+\infty} F(\nu_x, \nu_y) e^{-j2\pi(\nu_x x + \nu_y y)} d\nu_x d\nu_y. \quad (3.12)$$

Le signe choisi dans l'exponentielle de la transformée de Fourier est arbitraire ; la seule contrainte étant d'utiliser des signes différents pour la transformée de Fourier et son inverse. Le signe "-" dans l'exponentielle d'Eq. (3.12) indique que la fonction $f(x, y)$ s'obtient en intégrant sa transformée de Fourier $F(\nu_x, \nu_y)$ multipliée par une onde plane propageante $\exp[-j2\pi(\nu_x x + \nu_y y)]$, ce qui donne une représentation physique de la transformée de Fourier qui colle bien à l'optique.

3.5 Propagation dans l'espace libre – Cas général

Considérons maintenant le cas général à 3 dimensions et une onde plane d'amplitude complexe

$$U(x, y, z) = A \exp[-j(k_x x + k_y y + k_z z)] \quad (3.13)$$

dont les composantes du vecteur d'onde $\mathbf{k} = (k_x, k_y, k_z)$ satisfont,

$$\sqrt{k_x^2 + k_y^2 + k_z^2} = k = \frac{2\pi}{\lambda}. \quad (3.14)$$

Si on construit un parallélépipède en utilisant le système d'axes $Oxyz$ et le vecteur $\mathbf{k}$ comme diagonale, Fig. 3.10(a), on peut repérer la direction de propagation de l'onde en introduisant les angles $\theta_x = \sin^{-1}(k_x/k)$ et $\theta_y = \sin^{-1}(k_y/k)$ que fait le vecteur $\mathbf{k}$ avec les plans $y-z$ et $x-z$, Fig. 3.10(a).

Dans le plan $z = 0$, l'amplitude complexe $U(x, y, 0)$ est une fonction harmonique spatiale,

$$f(x, y) = A \exp[-j(k_x x + k_y y)] = A \exp[-j2\pi(\nu_x x + \nu_y y)], \quad (3.15)$$

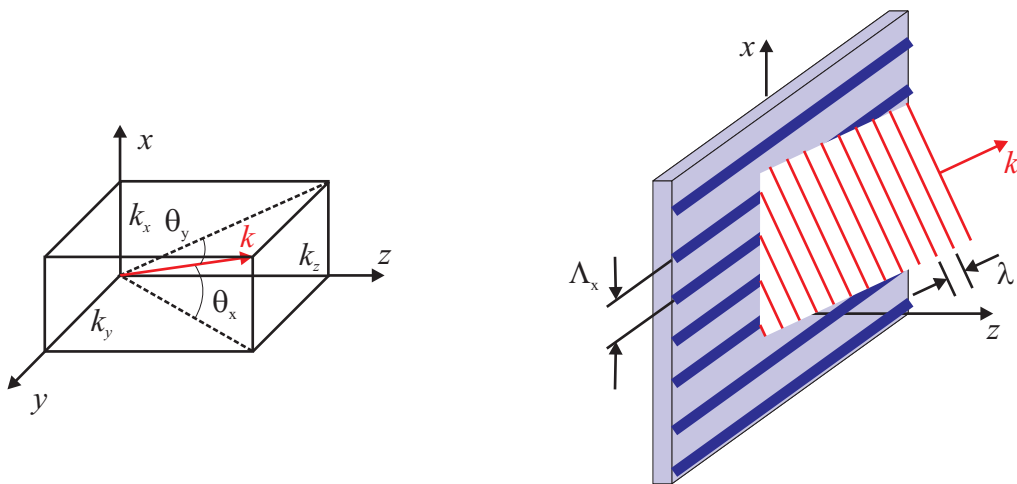


FIGURE 3.10 – (a) Une onde plane $U(x, y, z)$ se propageant dans l'espace dans la direction $\mathbf{k}$ définie par les angles θ_x et θ_y , (b) correspond à une fonction harmonique $f(x, y)$ de fréquences spatiales ν_x et ν_y dans le plan $z = 0$.

où l'on a introduit les fréquences spatiales dans les directions x et y : $\nu_x = k_x/2\pi$ et $\nu_y = k_y/2\pi$. A nouveau, on remarque que les fréquences spatiales sont en cycle/mètre. La direction dans laquelle se propage l'onde plane est donc en relation avec les fréquences spatiales de l'amplitude complexe dans le plan $z = 0$:

$$\theta_x = \sin^{-1} \left(\frac{k_x}{k} \right) = \sin^{-1} \left(\frac{2\pi\nu_x\lambda}{2\pi} \right) = \sin^{-1} \lambda\nu_x \quad \theta_y = \sin^{-1} \lambda\nu_y. \quad (3.16)$$

On remarquera que si l'onde se propage essentiellement dans la direction z ($k_x \ll k$ et $k_y \ll k$) on peut approximer les fréquences spatiales (3.16) par

$$\theta_x \simeq \lambda\nu_x \quad \theta_y \simeq \lambda\nu_y. \quad (3.17)$$

Cette approximation correspond à l'approximation paraxiale, dans laquelle l'onde se propage essentiellement dans la direction z avec des variations lentes dans les directions transverses.

Il est essentiel de remarquer qu'il existe une correspondance entre l'amplitude $U(x, y, z)$ de l'onde plane se propageant dans l'espace et la fonction harmonique $f(x, y)$ dans le plan $z = 0$, Fig. 3.10(a). Si on connaît une de ces deux fonctions, on peut déterminer l'autre, pour autant que l'on connaisse la longueur d'onde λ . Ainsi, si $f(x, y)$ est donné, on trouve $U(x, y, z) = f(x, y) \exp(-jk_z z)$, avec

$$k_z = \pm \sqrt{k^2 - k_x^2 - k_y^2}, \quad k = 2\pi/\lambda. \quad (3.18)$$

Deux remarques importantes s'imposent : le signe $+$ dans Eq. (3.18) correspond à une onde se propageant dans la direction z , le signe $-$ à une onde allant en arrière, dans la direction $-z$. Pour que k_z soit réel, il faut que $k_x^2 + k_y^2 < k^2$. Cette condition essentielle impose que

$$\lambda\nu_x < 1 \text{ et } \lambda\nu_y < 1, \quad (3.19)$$

qui est simplement la généralisation à deux dimensions d'Eq. (3.6). Si ces conditions ne sont pas remplies, k_z devient complexe, par exemple $k_z = k'_z - jk''_z$ où k'_z et k''_z sont les parties réelles et imaginaires. Dans ce cas, l'onde décroît au fur et à mesure de sa propagation,

$$U(x, y, z) = f(x, y) \exp(-jk'_z z) \exp(-k''_z z). \quad (3.20)$$

On parle alors d'onde évanescence. Equation (3.19) indique qu'une telle onde est associée avec les fréquences spatiales élevées. Comme nous le verrons, ces fréquences élevées correspondent aux détails fins de la structure. Ces détails sont donc perdus durant la propagation, puisque l'onde qui leur est associée décroît au fur et à mesure de sa propagation.

Nous allons maintenant considérer une onde plane d'amplitude unité se propageant dans la direction z et passant à-travers un écran mince dont la transmittance $f(x, y)$ est modulée dans les deux directions x et y : $f(x, y) = \exp[-j2\pi(\nu_x x + \nu_y y)]$, Fig. 3.11. Comme l'onde incidente est perpendiculaire au plan x - y et a une amplitude unité dans ce plan, l'onde transmise juste après l'écran vaut $U(x, y, 0) = f(x, y)$. Ainsi, l'onde plane incidente est transformée en une onde plane se propageant dans la direction définie par les angles $\theta_x =$

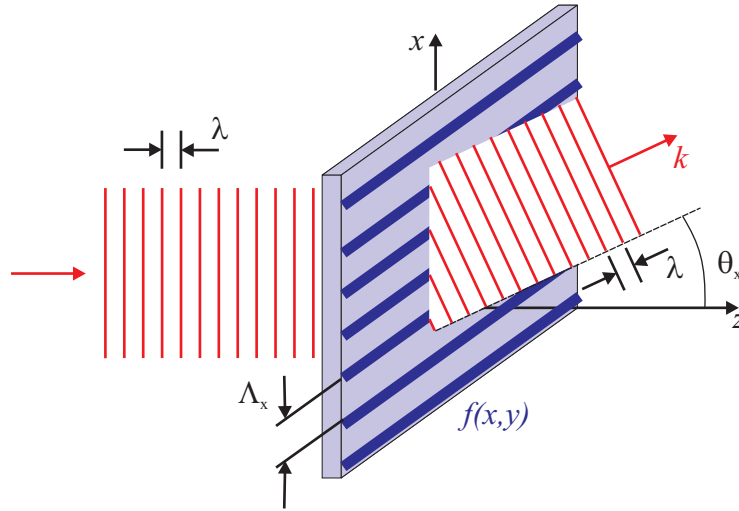


FIGURE 3.11 – Lorsqu’une onde plane arrive à incidence normale sur une surface dont la transmittivité est modulée périodiquement dans l’espace avec une période Λ_x , elle acquiert cette périodicité et ressort de la surface dans la direction correspondante (comparer avec Fig. 3.10).

$\sin^{-1} \lambda \nu_x$ et $\theta_y = \sin^{-1} \lambda \nu_y$. Si l’écran mince a une modulation $f(x, y) = \exp[j2\pi(\nu_x x + \nu_y y)]$, l’onde incidente est convertie en une onde se propageant dans la direction $-\theta_x, -\theta_y$.

On peut comprendre le changement de direction θ_x que subit une onde incidente de longueur d’onde λ lorsqu’elle passe à-travers une surface de période Λ_x par la nécessité de faire coïncider ces deux périodes spatiales juste à la sortie de la surface, comme illustré Fig. 3.11.

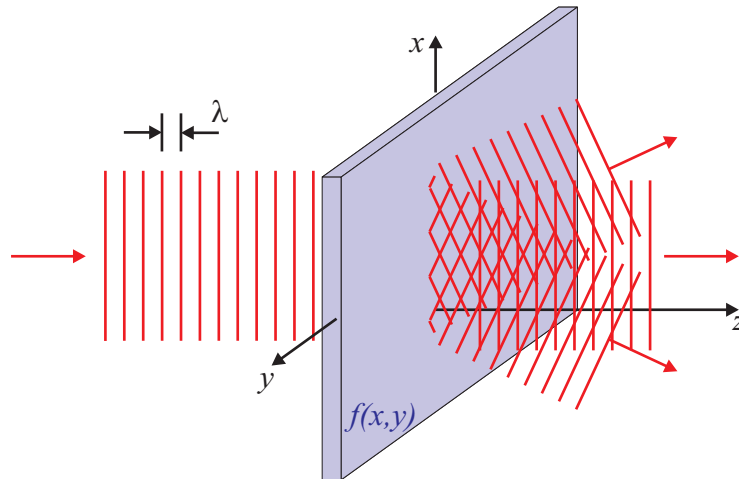


FIGURE 3.12 – Un élément optique mince avec une transmission $f(x, y)$ décompose une onde plane incidente en plusieurs ondes planes se propageant dans les directions θ_x, θ_y déterminées par la transformée de Fourier $F(\nu_x, \nu_y)$ de $f(x, y)$.

On peut associer des longueurs d'onde aux fréquences spatiales : $\Lambda_x = 1/\nu_x$ et $\Lambda_y = 1/\nu_y$ (remarquer l'unité de la fréquence spatiale $[\nu] = \text{m}^{-1}$ et la longueur d'onde Λ_x a donc comme unité $[\Lambda_x] = \text{m}$). Une autre façon de comprendre le changement de direction que fait l'onde en passant à travers la surface est de le considérer comme un phénomène d'interférence : dans la direction θ_x les ondes issues de deux points dans le plan x - y séparés par une distance Λ_x donnent lieu à une interférence constructive, car la différence de chemin optique est égale à λ .

Si maintenant la transmittance dans le plan $z = 0$ devient la somme de plusieurs fonctions harmoniques dans le plan x - y , alors l'onde transmise est la superposition des ondes planes associées à chacune de ces fonctions harmoniques, Fig. 3.12. Considérons quelques cas particuliers pour la fonction $f(x, y)$:

- Si $f(x, y) = \cos(2\pi\nu_x x) = 1/2[\exp(-j2\pi\nu_x x) + \exp(j2\pi\nu_x x)]$ alors l'onde incidente est décomposée en deux ondes transmises se propageant vers le haut et vers le bas, avec des angles $\pm \sin^{-1} \lambda\nu_x$.
- Si $f(x, y) = 1 + \cos(2\pi\nu_x x)$ alors l'onde incidente est décomposée en trois ondes transmises, deux dans les directions $\pm \sin^{-1} \lambda\nu_x$ et une qui va tout droit, Fig. 3.12.
- Si $f(x, y) = \mathcal{U}[\cos(2\pi\nu_x x)]$, où la fonction $\mathcal{U}(x)$ est la fonction escalier ($\mathcal{U}(x) = 1$ si $x > 0$ et $\mathcal{U}(x) = 0$ si $x < 0$). La fonction $f(x, y)$ représente donc une série de fentes parallèles de période Λ_x dans un écran opaque. On peut l'exprimer en terme de série de Fourier comme une somme d'harmoniques de fréquences $0, \pm\nu_x, \pm2\nu_x, \pm3\nu_x, \dots$. Ainsi cette série de fentes transmettra-t-elle des ondes planes dans les directions $0, \pm \sin^{-1} \lambda\nu_x, \pm2 \sin^{-1} \lambda\nu_x, \pm3 \sin^{-1} \lambda\nu_x, \dots$. A nouveau on peut comprendre ceci comme un phénomène d'interférence où les ondes transmises par chaque fente interfèrent de façon constructive dans les directions $0, \pm \sin^{-1} \lambda\nu_x, \pm2 \sin^{-1} \lambda\nu_x, \pm3 \sin^{-1} \lambda\nu_x, \dots$.

Considérons maintenant le cas le plus général, dans lequel la transmittance dans le plan est la somme (infinie) de fonctions harmoniques,

$$f(x, y) = \int \int_{-\infty}^{\infty} d\nu_x d\nu_y F(\nu_x, \nu_y) \exp[-j2\pi(\nu_x x + \nu_y y)], \quad (3.21)$$

où $F(\nu_x, \nu_y)$ représente l'amplitude de la composante de fréquence spatial ν_x, ν_y . L'onde transmise est alors la superposition des ondes planes,

$$U(x, y, z) = \int \int_{-\infty}^{\infty} d\nu_x d\nu_y F(\nu_x, \nu_y) \exp[-j2\pi(\nu_x x + \nu_y y)] \exp(-jk_z z), \quad (3.22)$$

où k_z est défini pour chaque composante spatiale (ν_x, ν_y) par la relation

$$k_z = \sqrt{k^2 - k_x^2 - k_y^2} = 2\pi\sqrt{\lambda^{-2} - \nu_x^2 - \nu_y^2}. \quad (3.23)$$

Il est important de remarquer que $F(\nu_x, \nu_y)$ est la transformée de Fourier de $f(x, y)$.

Rappelons brièvement quelques propriétés de la transformée de Fourier pour des fonctions de deux variables. Une fonction $f(x, y)$ peut être décomposée en une superposition de fonctions

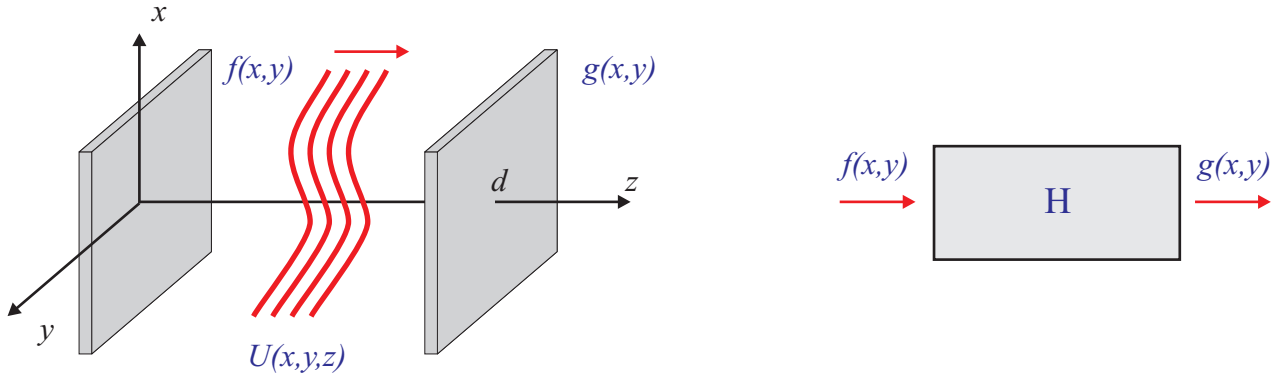


FIGURE 3.13 – La propagation d'une onde dans l'espace libre peut se traduire par une fonction de transfert H entre le champ $f(x, y)$ dans le plan d'entrée et le champ $g(x, y)$ dans le plan de sortie.

harmoniques :

$$f(x, y) = \int \int_{-\infty}^{\infty} d\nu_x d\nu_y F(\nu_x, \nu_y) \exp [-j2\pi(\nu_x x + \nu_y y)] . \quad (3.24)$$

où les coefficients $F(\nu_x, \nu_y)$ sont donnés par

$$F(\nu_x, \nu_y) = \int \int_{-\infty}^{\infty} dx dy f(x, y) \exp [j2\pi(\nu_x x + \nu_y y)] . \quad (3.25)$$

La fonction $F(\nu_x, \nu_y)$ est la transformée de Fourier de la fonction originale $f(x, y)$.

Le théorème de convolution indique que si la fonction $f(x, y)$ est la convolution des fonctions $f_1(x, y)$ et $f_2(x, y)$:

$$f(x, y) = \int \int_{-\infty}^{\infty} dx' dy' f_1(x', y') f_2(x - x', y - y') dx' dy' ; \quad (3.26)$$

alors la transformée de Fourier de $f(x, y)$ est donnée par

$$F(\nu_x, \nu_y) = F_1(\nu_x, \nu_y) F_2(\nu_x, \nu_y) , \quad (3.27)$$

où $F_1(\nu_x, \nu_y)$ et $F_2(\nu_x, \nu_y)$ sont les transformées de Fourier de $f_1(x, y)$ et $f_2(x, y)$.

Comme une fonction arbitraire peut être décomposée en une superposition d'ondes comme Eq. (3.21), la lumière transmise par un élément optique mince de transmittance arbitraire peut s'écrire comme une superposition d'ondes planes, Fig. 3.12. Il est important de noter la condition sur les fréquences spatiales composant l'élément optique mince : $\nu_x^2 + \nu_y^2 < \lambda^{-2}$. Ainsi, pour une longueur d'onde donnée, il existe des fréquences spatiales maximales qui peuvent être transmises. Les fréquences spatiales plus grandes que λ^{-1} ne sont pas transmises ! Ce concept est essentiel pour la définition de la résolution optique ; il est rattaché à la fonction de transfert de l'espace libre, illustrée Fig. 3.13.

Considérons la propagation d'une onde monochromatique de longueur d'onde λ et d'amplitude complexe $U(x, y, z)$ entre les plans $z = 0$ et $z = d$, Fig. 3.13. On désigne par $f(x, y) = U(x, y, 0)$ la fonction d'entrée et par $g(x, y) = U(x, y, d)$ la fonction de sortie. Notre objectif est de déterminer une fonction de transfert H , telle que

$$g(x, y) = H f(x, y). \quad (3.28)$$

Comme nous faisons ce développement dans le contexte des ondes planes qui doivent satisfaire l'équation de Helmholtz (2.10) qui est linéaire, la fonction de transfert H doit aussi être linéaire. Pour la déterminer, nous passons dans l'espace de Fourier et considérons une fonction input de forme harmonique : $f(x, y) = A \exp[-j2\pi(\nu_x x + \nu_y y)]$. Comme nous l'avons vu, ceci correspond à une onde plane de forme $U(x, y, z) = A \exp[-j(k_x x + k_y y + k_z z)]$ avec $k_x = 2\pi\nu_x$ et $k_y = 2\pi\nu_y$. Quant à k_z , il n'est pas indépendant et doit satisfaire la relation,

$$k_z = \sqrt{k^2 - k_x^2 - k_y^2} = 2\pi\sqrt{\lambda^{-2} - \nu_x^2 - \nu_y^2}. \quad (3.29)$$

En utilisant la même onde plane se propageant dans l'espace libre, on a pour la fonction de sortie : $g(x, y) = U(x, y, d) = A \exp[-j(k_x x + k_y y + k_z d)]$. En faisant le rapport $g(x, y)/f(x, y) = \exp[-jk_z d]$, on obtient la valeur de la fonction de transfert de l'espace libre :

$$H(\nu_x, \nu_y) = \exp \left[-j2\pi d \sqrt{\lambda^{-2} - \nu_x^2 - \nu_y^2} \right]. \quad (3.30)$$

La fonction de transfert de l'espace libre dépend donc des fréquences spatiales ν_x et ν_y . Pour les fréquences spatiales telles que $\nu_x^2 + \nu_y^2 \leq \lambda^{-2}$ (c'est à dire les fréquences spatiales se trouvant à l'intérieur d'un cercle de rayon λ^{-1} , l'amplitude de $H(\nu_x, \nu_y)$ est 1 et la phase dépend de ν_x et ν_y . Pour les fréquences spatiales $\nu_x^2 + \nu_y^2 > \lambda^{-2}$, l'argument de l'exponentielle dans Eq. (3.30) devient complexe et l'amplitude de $H(\nu_x, \nu_y)$ décroît. Ainsi ces fréquences spatiales ne sont-elles pas transmises complètement et diminuent durant la propagation dans l'espace. La fonction de transfert de l'espace libre est illustrée Fig. 3.14 : cette fonction est cylindrique de révolution et nous montrons sur la figure une coupe dans une direction particulière.

L'équation (3.30) peut être simplifiée si la fonction input ne contient pas de fréquences spatiales très élevées : $\nu_x^2 + \nu_y^2 \ll \lambda^{-2}$. En se souvenant de l'analogie avec les ondes planes, cette condition correspond à une onde se propageant selon des angles petits : $\theta_x \simeq \lambda\nu_x$, $\theta_y \simeq \lambda\nu_y$ (approximation paraxiale). En introduisant l'angle de propagation par rapport à l'axe optique : $\theta^2 = \theta_x^2 + \theta_y^2 \simeq \lambda^2(\nu_x^2 + \nu_y^2)$, la racine carrée dans Eq. (3.30) peut s'exprimer :

$$2\pi d \sqrt{\lambda^{-2} - \nu_x^2 - \nu_y^2} = 2\pi \frac{d}{\lambda} \sqrt{1 - \theta^2} \simeq 2\pi \frac{d}{\lambda} \left(1 - \frac{\theta^2}{2} + \frac{\theta^4}{8} - \dots \right). \quad (3.31)$$

En ne gardant que les deux premiers termes, on obtient l'approximation de Fresnel pour la fonction de transfert dans l'espace libre :

$$H(\nu_x, \nu_y) \simeq \exp(-jkd) \exp[j\pi\lambda d(\nu_x^2 + \nu_y^2)] = H_0 \exp[j\pi\lambda d(\nu_x^2 + \nu_y^2)]. \quad (3.32)$$

On voit que le terme $H_0 = \exp(-jkd)$ donne la propagation en ligne droite sur une distance d . On appelle Eq. (3.32) l'approximation de Fresnel ; elle est valide si le troisième terme dans

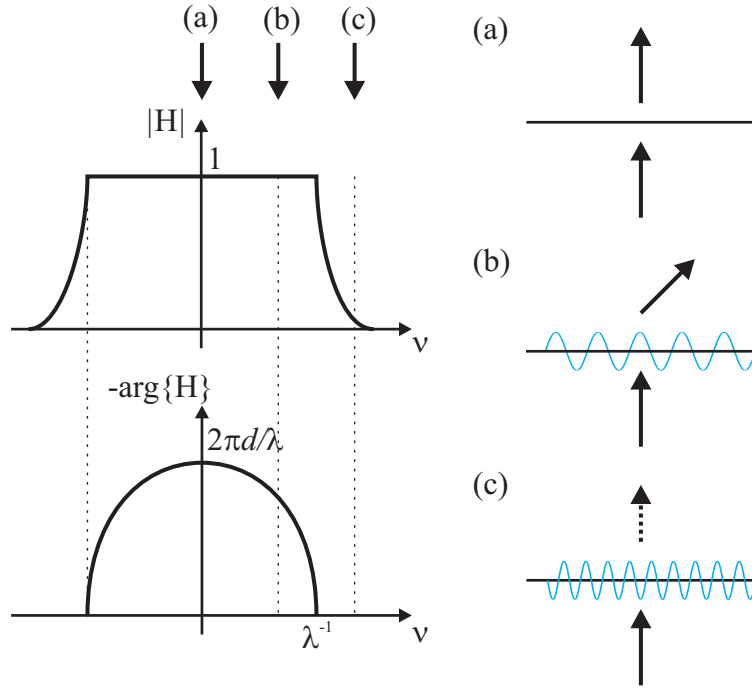


FIGURE 3.14 – Amplitude et phase de la fonction de transfert $H(\nu_x, \nu_y)$ de l'espace libre et illustration de son effet sur la propagation d'une onde plane de longueur d'onde λ à incidence normale sur une surface de profil harmonique ν (la fonction de transfert a une symétrie cylindrique et cette figure en montre une coupe).

Eq. (3.31) est beaucoup plus petit que π . Ceci est équivalent à la condition,

$$\frac{\theta^4 d}{4\lambda} \ll 1. \quad (3.33)$$

Pour une fonction d'entrée arbitraire $f(x, y)$, on calcule la fonction de sortie (après propagation dans l'espace libre sur une distance d) de la façon suivante :

- 1) On décompose $f(x, y)$ en ondes planes en calculant la transformée de Fourier :

$$F(\nu_x, \nu_y) = \int \int_{-\infty}^{\infty} dx dy f(x, y) \exp [j2\pi(\nu_x x + \nu_y y)] . \quad (3.34)$$

- 2) Chaque onde plane d'amplitude $F(\nu_x, \nu_y)$ est ensuite propagée dans l'espace libre sur une distance d . Le produit $H(\nu_x, \nu_y)F(\nu_x, \nu_y)$ donne l'amplitude complexe de cette onde plane dans le plan de sortie.
- 3) L'amplitude totale dans le plan de sortie est obtenue en superposant toutes ces ondes planes :

$$g(x, y) = \int \int_{-\infty}^{\infty} d\nu_x d\nu_y H(\nu_x, \nu_y) F(\nu_x, \nu_y) \exp [-j2\pi(\nu_x x + \nu_y y)] . \quad (3.35)$$

Si on introduit dans Eq. (3.35) l'approximation de Fresnel (3.32) pour la fonction de transfert H , on obtient la relation suivante :

$$g(x, y) = H_0 \int \int_{-\infty}^{\infty} d\nu_x d\nu_y F(\nu_x, \nu_y) \exp [j\pi\lambda d(\nu_x^2 + \nu_y^2)] \exp [-j2\pi(\nu_x x + \nu_y y)] . \quad (3.36)$$

3.6 Transformée de Fourier avec une lentille

Une lentille concentre une onde plane se propageant dans la direction (θ_x, θ_y) en un point dans le plan focal. Si l'onde se propage dans une direction proche de l'axe ($\sin \theta_x \simeq \theta_x$ et $\sin \theta_y \simeq \theta_y$), les coordonnées (x, y) de ce point dans le plan focal sont $(\theta_x f, \theta_y f)$, Fig. 3.15(a).

Considérons maintenant un champ d'entrée arbitraire $f(x, y)$ formé de la superposition de plusieurs ondes planes se propageant dans les directions $\theta_x = \lambda\nu_x$, $\theta_y = \lambda\nu_y$, avec pour amplitude complexe $F(\nu_x, \nu_y)$. Après la lentille, chaque onde plane est concentrée dans le plan focal au point $x = \theta_x f = \lambda f \nu_x$, $y = \theta_y f = \lambda f \nu_y$. L'amplitude complexe dans le plan de sortie (plan focal) au point (x, y) est donc proportionnelle à la transformée de Fourier de $f(x, y)$ évaluée aux fréquences spatiales $\nu_x = x/\lambda f$, $\nu_y = y/\lambda f$, Fig. 3.15(b). Ainsi,

$$g(x, y) \propto F\left(\frac{x}{\lambda f}, \frac{y}{\lambda f}\right) . \quad (3.37)$$

Si on se place dans le cas particulier où la distance entre le plan d'entrée et la lentille est égale à la distance focale, alors le facteur de proportionnalité dans Eq. (3.37) prend la forme simple suivante :

$$g(x, y) = h_l F\left(\frac{x}{\lambda f}, \frac{y}{\lambda f}\right) , \quad (3.38)$$

où l'on a introduit la réponse impulsionnelle du système (i.e. le champ transmis par une source ponctuelle placée à l'origine) : $h_l = (j/\lambda f) \exp[-j2kf]$. On parle alors d'un système

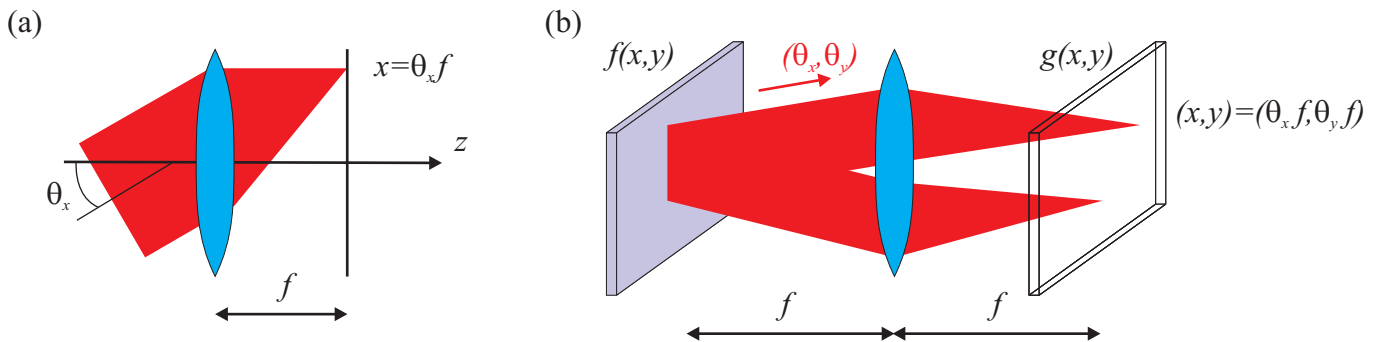


FIGURE 3.15 – (a) Une lentille concentre une onde plane se propageant dans la direction (θ_x, θ_y) en un point de coordonnées $(\theta_x f, \theta_y f)$ dans le plan focal. (b) Si le champ incident est composé de plusieurs ondes planes se propageant dans des directions différentes, chaque onde plane est focalisée en un point différent du plan focal.

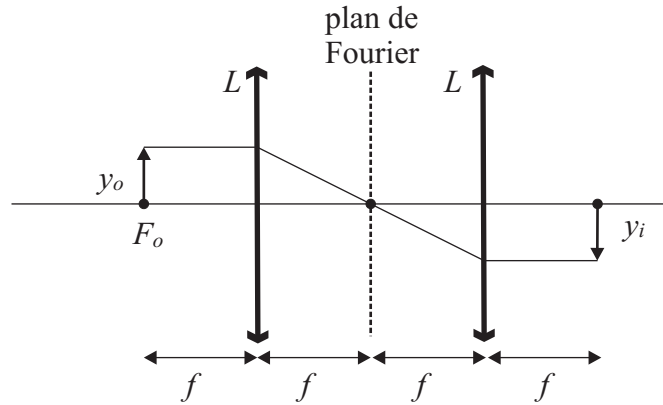


FIGURE 3.16 – Système $4f$: on a dans le plan intermédiaire, ou plan de Fourier, la transformée de Fourier de l'objet.

$2f$ puisque la distance entre entrée et sortie est deux fois la distance focale. Notons finalement que l'intensité dans le plan focal de sortie (en anglais *back focal plane*) est proportionnelle au carré de la transformée de Fourier de l'amplitude d'entrée :

$$I(x, y) = \frac{1}{(\lambda f)^2} \left| F \left(\frac{x}{\lambda f}, \frac{y}{\lambda f} \right) \right|^2. \quad (3.39)$$

3.7 Filtrage avec un système $4f$

Si l'arrangement $2f$ illustré Fig. 3.15(b) permet d'obtenir la transformée de Fourier du champ se trouvant dans le plan focal, il est clair que ce champ peut être à nouveau transformé avec une deuxième lentille, comme indiqué Fig. 3.16. On parle alors de système $4f$ puisque l'on a quatre fois la distance focale entre l'entrée et la sortie de l'arrangement.

Dans un système $4f$, le plan intermédiaire permet d'accéder à la transformée de Fourier du champ incident. En plaçant un filtre dans ce plan, on peut bloquer certaines composantes de Fourier du champ incident et procéder à son filtrage. Ceci est illustré sur Fig. 3.17 pour une photo utilisée souvent en traitement d'image. On remarque que l'image obtenue dans le plan de Fourier contient différentes composantes spectrales, un peu dans toutes les directions. Pour visualiser ce qui se passe dans le plan de Fourier, on peut mettre une autre lentille qui va focaliser sur ce plan de Fourier, ainsi peut-on mesurer expérimentalement le spectre d'une image.

Si on place un diaphragme dans le plan de Fourier, on ne laisse passer que les fréquences spatiales (ν_x, ν_y) qui sont petites et l'image perd en contraste, Fig. 3.17(c). Si on utilisait un diaphragme avec un très petit trou, on n'aurait qu'une image homogène, avec une teinte grisâtre correspondant à l'intensité moyenne de la photo. Plus on laisse passer des fréquences spatiales importantes, plus l'image gagne en contraste, Fig. 3.17(c).



FIGURE 3.17 – Filtrage avec un système $4f$: (a) image originale, (b) spectre dans le plan de Fourier et (c) effet de différents filtres passe-bas (la partie qui transmet la lumière est indiquée en blanc).

On peut aussi occulter le centre du plan de Fourier en utilisant un filtre passe haut qui coupe l'ordre zéro ($\nu_x = \nu_y = 0$) et les petites fréquences spatiales, Fig. 3.18. L'image résultante est très contrastée, comme si on avait mis en évidence les détours de la photo. Ce contraste dépend du seuil de fréquences spatiales au-delà duquel on laisse passer la lumière dans l'espace de Fourier (comparer les différentes images de Fig. 3.18).

Figure 3.19 illustre ce filtrage dans le plan de Fourier sur une grille infinie. On remarque que l'image dans le plan de Fourier est composé de points espacés régulièrement, Fig. 3.19(b). Chaque série de lignes dans l'espace direct (image) donne lieu à une série de points dans

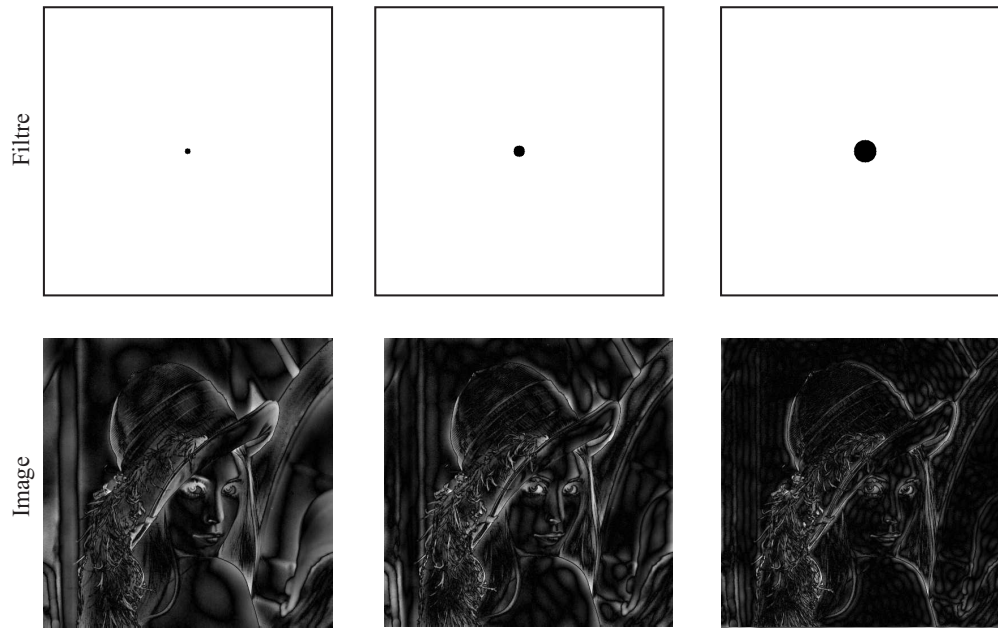


FIGURE 3.18 – Filtrage avec un système $4f$: effet de différents filtres passe-haut (la partie qui transmet la lumière est indiquée en blanc) sur l'image de Fig. 3.17(a).

l'espace de Fourier qui sont orientés perpendiculairement aux lignes. Ainsi, les lignes horizontales produisent une ligne de points verticaux dans le plan de Fourier et les lignes verticales produisent une ligne de points horizontaux dans le plan de Fourier. On remarque aussi que plus les lignes sont espacées dans l'espace direct et plus les points sont serrés dans l'espace de Fourier et inversement.

En mettant un filtre dans le plan de Fourier qui ne laisse passer qu'une série de points, on peut sélectionner telle, ou telle partie de la figure originale, Fig. 3.18(c).

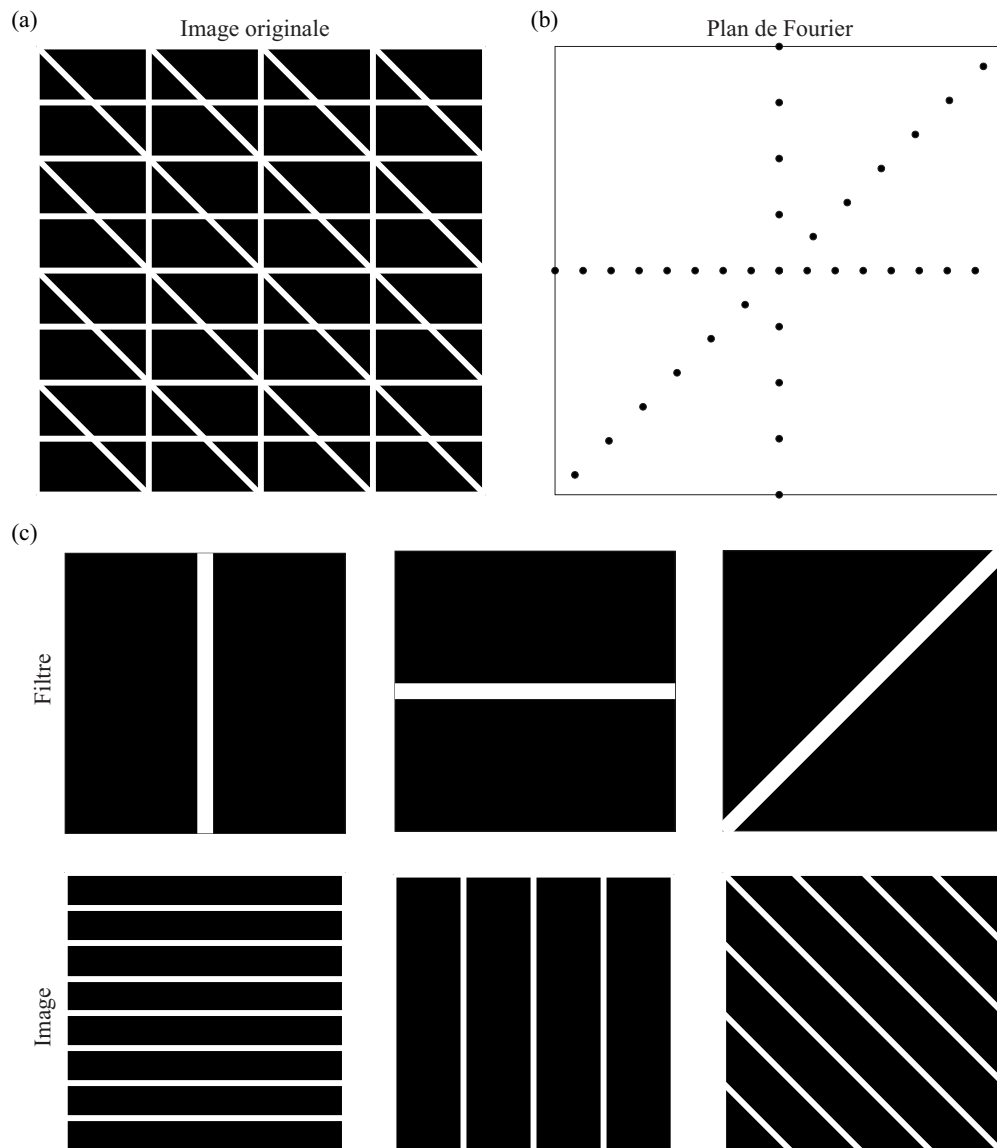


FIGURE 3.19 – Filtrage avec un système $4f$ d'une image formé d'une grille de lignes : (a) image originale et (b) image dans le plan de Fourier. (c) Effet de différents filtres (la partie du filtre qui transmet la lumière est indiquée en blanc).

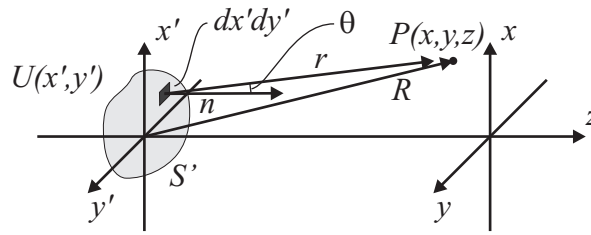


FIGURE 3.20 – Géométrie utilisée pour calculer la diffraction.

3.8 Intégrale de Rayleigh–Sommerfeld

Nous revenons maintenant au principe de Huygens pour étudier plus en détail la diffraction, dont une situation générale est illustrée Fig. 3.20 : une ouverture de forme arbitraire dans un écran opaque est illuminée par un champ d'amplitude $U(x', y')$ (on garde ici la variation du champ dans l'ouverture pour faire le lien avec la transformée de Fourier, mais souvent ce champ est constant dans l'ouverture : $U(x', y') = U_0$). On souhaite calculer la distribution spatiale du champ transmis à une certaine distance de cette ouverture. Le calcul est basé sur le principe de Huygens : chaque élément de surface $dx'dy'$ de l'ouverture est la source d'une onde sphérique infinitésimale. Le champ au point $P(x, y, z)$ s'obtient en combinant toutes ces sources infinitésimales, i.e. en intégrant toutes ces sources sur la surface de l'ouverture :

$$U(x, y, z) = \frac{1}{j\lambda} \int_{S'} U(x', y') \frac{e^{-jkr}}{r} \cos \theta dx'dy', \quad (3.40)$$

où pour simplifier nous avons supposé l'ouverture dans le plan $z = 0$ et omis la dépendance temporelle.

Le terme $\cos \theta$ dans Eq. (3.40) indique que l'on projette la surface S' dans la direction d'observation du point P . Ce terme assure que si l'on regarde la source (ouverture dans l'écran) sous un grand angle θ l'intensité lumineuse est plus petite que si l'on regarde la source en direct ($\theta = 0$). Une telle source, dont l'intensité radiée dépend du cosinus de l'angle d'observation, est dite lambertienne. Le terme $1/j\lambda$ dans Eq. (3.40) provient d'une formulation plus générale du problème, appelée équation de Fresnel–Kirchoff, et est associé aux conditions limites pour le champ électromagnétique dans l'écran opaque que l'on considère comme un conducteur parfait.

Dans le cas le plus général, il n'existe pas de solution simple à Eq. (3.40). Heureusement, il existe des situations particulières importantes pour lesquelles il est possible d'obtenir une telle solution. C'est par exemple le cas lorsque la distance d'observation R est grande par rapport aux dimensions latérales de l'ouverture : l'approximation de Fraunhofer.

3.9 Approximation de Fraunhofer

Considérons Fig. 3.20 et remarquons que dans Eq. (3.40) les termes en $1/r$ et en $\cos \theta$ varient lentement, comparés à la phase kr dans l'exponentielle. En coordonnées cartésiennes la distance r entre l'élément d'ouverture $dx'dy'$ et le point d'observation s'exprime

$$r = \sqrt{(x - x')^2 + (y - y')^2 + z^2} = z \sqrt{1 + \frac{(x - x')^2}{z^2} + \frac{(y - y')^2}{z^2}}. \quad (3.41)$$

Comme nous cherchons une approximation valable à larges distances de l'ouverture et sur l'axe du système (i.e. z grand par rapport à x et y), nous pouvons supposer que les termes divisés par z^2 dans Eq. (3.41) sont petits. On a donc dans cette équation une expression du genre

$$\sqrt{1 + \epsilon} \simeq 1 + \frac{\epsilon}{2}, \quad (3.42)$$

où l'on a utilisé le développement limité pour ϵ petit. Ainsi Eq. (3.41) devient

$$r = \sqrt{(x - x')^2 + (y - y')^2 + z^2} \simeq z \left[1 + \frac{1}{2z^2} ((x - x')^2 + (y - y')^2) \right]. \quad (3.43)$$

Comme de plus $R^2 = x^2 + y^2 + z^2$, on peut finalement écrire,

$$r = R \left[1 - \frac{xx' + yy'}{R^2} + \frac{x'^2 + y'^2}{2R^2} \right]. \quad (3.44)$$

L'approximation de Fraunhofer consiste à ne retenir que les deux premiers termes dans Eq. (3.44). On remarque que lorsque $R \rightarrow \infty$ les erreurs liées à cette approximation deviennent négligeables. Finalement, pour $R \gg x', y'$ on peut considérer le dénominateur $r \simeq R$ comme constant dans l'intégrale (3.40) et le sortir de l'intégrale. Dans les mêmes conditions le terme $\cos \theta$ peut être considéré comme constant et aussi sorti de l'intégrale. Ainsi l'approximation de Fraunhofer donne pour Eq. (3.40) :

$$U(x, y, R) \simeq \frac{e^{-jkR}}{j\lambda R} \cos \theta \int_{S'} U(x', y') e^{jk \frac{xx' + yy'}{R}} dx' dy', \quad (3.45)$$

Pour caractériser la limite de validité de l'approximation de Fraunhofer on introduit le nombre de Fresnel N ,

$$N = \frac{\rho^2}{\lambda z}, \quad (3.46)$$

où ρ est le "rayon maximal" de l'ouverture considérée (dans le cas d'une fente de largeur a , $\rho = a/2$; pour une ouverture circulaire dans le plan $x'-y'$, ρ correspond au rayon de l'ouverture) et z la distance à laquelle on calcule le champ. On distingue alors trois régimes :

- pour $N \ll 1$ (i.e. $\rho \ll \sqrt{\lambda z}$: régime de Fraunhofer où Eq. (3.45) s'applique ;
- pour $N \simeq 1$ (i.e. $\rho \simeq \sqrt{\lambda z}$: régime de Fresnel correspondant à des distances plus proches entre l'ouverture et l'écran d'observation. Equation (3.45) ne s'applique pas dans ce cas et il faudrait retenir le prochain terme dans Eq. (3.44) pour calculer le champ ;
- pour $N \rightarrow \infty$ (i.e. $\lambda \rightarrow 0$ ou $z \rightarrow 0$ ou $\rho \rightarrow \infty$) on est alors dans le domaine de l'optique géométrique et le phénomène de diffraction disparaît.

3.10 Diffraction et transformée de Fourier

Nous allons montrer que l'équation obtenue pour la diffraction de Fraunhofer peut s'exprimer comme la transformée de Fourier de l'amplitude dans le plan de l'écran. Pour ceci on commence par augmenter le support de définition de l'amplitude incidente $U(x', y')$ en supposant simplement que $U(x', y') = 0$ si l'on n'est pas dans l'ouverture de l'écran, Fig. 3.21. Ainsi peut-on étendre les domaines d'intégration sur x' et y' jusqu'à l'infini.

On introduit ensuite les fréquences spatiales p_x et p_y correspondant aux variables x et y :

$$p_x = \frac{x}{\lambda R} = \frac{\sin \theta_x}{\lambda} \simeq \frac{\theta_x}{\lambda} \quad (3.47a)$$

$$p_y = \frac{y}{\lambda R} = \frac{\sin \theta_y}{\lambda} \simeq \frac{\theta_y}{\lambda} . \quad (3.47b)$$

Les angles θ_x et θ_y représentent les angles d'inclinaison du vecteur R et sont définis Fig. 3.21 ; dans Eq. (3.47) on a utilisé le fait que ces angles sont petits pour approximer leur sinus par la valeur de l'argument. On peut de plus introduire l'approximation $\cos \theta \simeq 1$ dans Eq. (3.45) ; cette dernière équation devient alors

$$U(x, y, R) = \frac{e^{-jkR}}{j\lambda R} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} U(x', y') e^{j2\pi(p_x x' + p_y y')} dx' dy' . \quad (3.48)$$

Définissons la fonction

$$U(p_x, p_y) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} U(x', y') e^{j2\pi(p_x x' + p_y y')} dx' dy' = \mathcal{F} \{U(x', y')\} , \quad (3.49)$$

qui n'est autre que la transformée de Fourier de la distribution du champ $U(x', y')$ dans l'ouverture. On appelle cette fonction la densité d'énergie spectrale. Son carré $|U(p_x, p_y)|^2$ correspond à une distribution de puissance par unité d'angle solide $\theta_x \theta_y$.

Pour calculer le champ $U(x, y, R)$ derrière l'écran, il suffit donc de calculer la transformée de Fourier de la distribution $U(x', y')$ dans l'ouverture et de multiplier avec les coefficients selon Eq. (3.48) :

$$U(x, y, R) = \frac{e^{-jkR}}{j\lambda R} U(p_x, p_y) . \quad (3.50)$$

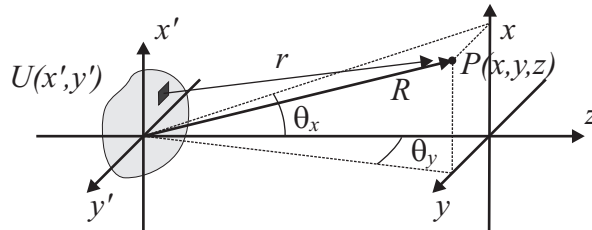


FIGURE 3.21 – Angles utilisés pour la définition des fréquences spatiales p_x et p_y .

La transformée de Fourier inverse permet de calculer la distribution du champ dans l'ouverture connaissant la densité d'énergie spectrale :

$$U(x', y') = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} U(p_x, p_y) e^{-j2\pi(p_x x' + p_y y')} dp_x dp_y = \mathcal{F}^{-1} \{U(p_x, p_y)\} . \quad (3.51)$$

La distribution d'intensité – dans l'approximation de Fraunhofer – s'obtient avec la formule suivante :

$$I(x, y, R) = U(x, y, R) U^*(x, y, R) = \frac{\cos \theta}{\lambda^2 R^2} U(p_x, p_y) U^*(p_x, p_y) . \quad (3.52)$$

On remarque que la décroissance en $1/R^2$ garantit la conservation de l'énergie. En effet, si on intègre l'intensité sur une sphère de rayon R , la surface d'intégration croît avec le carré de la distance. Pour une source donnée, on peut intégrer l'intensité transmise à-travers l'ouverture à différentes distances R et obtenir une valeur constante puisque la décroissance d'Eq. (3.52) compense la croissance de la surface.

3.11 Diffraction par une fente

Considérons la fente de largeur a représentée Fig. 3.22 ; il s'agit d'un problème à une dimension. Pour une illumination homogène, l'amplitude dans la fente s'écrit,

$$U(x) = \begin{cases} 1 & |x| \leq a/2 \\ 0 & |x| > a/2 \end{cases} = \text{rect} \left(\frac{x}{a} \right) , \quad (3.53)$$

où l'on a introduit la fonction rectangle $\text{rect}(x/a)$ dont la transformée de Fourier est la fonction $\text{sinc}(a\pi p_x)$:

$$U(p_x) = a \text{sinc}(a\pi p_x) = a \frac{\sin(a\pi p_x)}{a\pi p_x} . \quad (3.54)$$

Dans l'axe de la fente ($\theta = 0$) l'intensité vaut,

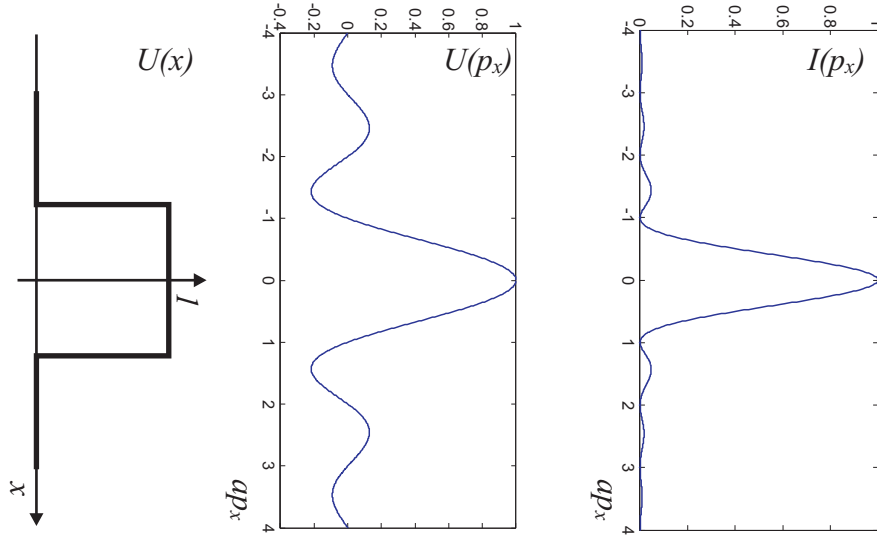
$$I(p_x) = \frac{a^2}{\lambda^2 R^2} \frac{\sin^2(a\pi p_x)}{(a\pi p_x)^2} . \quad (3.55)$$

Pour connaître la largeur du pic de diffraction, on peut calculer le premier zéro de cette fonction ; il s'obtient lorsque $p_x = 1/a$ en sorte que l'argument du sin dans Eq. (3.55) vaille π . En utilisant la définition de p_x (3.47a) on obtient l'angle auquel apparaît le premier minima :

$$\sin \theta_x = \frac{\lambda}{a} . \quad (3.56)$$

On retrouve le critère obtenu Eq. (3.2) en raisonnant sur l'interférence entre les ondes rayonnées de part et d'autre de la fente.

Les amplitudes et intensités diffractées par une fente sont représentées Fig. 3.22.

FIGURE 3.22 – Diffraction par une fente : $U(x)$, $U(p_x)$ et $I(p_x)$.

3.12 Diffraction par un réseau de fentes

Un réseau de transmission se construit en arrangeant N fentes de largeur a avec une période b , Fig. 3.23. L'amplitude $U(x)$ pour une illumination homogène s'écrit

$$U(x) = \sum_{q=-(N-1)/2}^{+(N-1)/2} \text{rect}\left(\frac{x + qb}{a}\right). \quad (3.57)$$

La transformée de Fourier de (3.57) s'écrit

$$U(p_x) = a \text{sinc}(a\pi p_x) \frac{\sin(bN\pi p_x)}{\sin(b\pi p_x)} e^{-j\pi b p_x/2}; \quad (3.58)$$

et l'intensité,

$$I(p_x) = \frac{a^2}{\lambda^2 R^2} \frac{\sin^2(a\pi p_x)}{(a\pi p_x)^2} \frac{\sin^2(bN\pi p_x)}{\sin^2(b\pi p_x)}. \quad (3.59)$$

Les deux premiers termes de cette équation sont semblables à ceux obtenus Eq. (3.55) pour une seule fente. Le dernier terme a la forme $\sin^2(N\xi)/\sin^2(\xi)$ et ressemble à un peigne avec des pics principaux qui apparaissent quand le dénominateur s'annule. Entre deux pics principaux il y a $N - 2$ pics secondaires puisque l'oscillation du numérateur est N fois plus rapide que celle du dénominateur. On peut montrer mathématiquement que la hauteur des pics principaux vaut N^2 . La figure de diffraction est donc constituée d'une série de résonances associées au troisième terme modulées par l'enveloppe de la diffraction d'une seule fente, Fig. 3.23.

La largeur des pics est donnée par la condition $bN\pi p_x = \pi$; en remplaçant p_x par sa définition (3.47a) on obtient

$$\sin \theta_x = \frac{\lambda}{bN}. \quad (3.60)$$

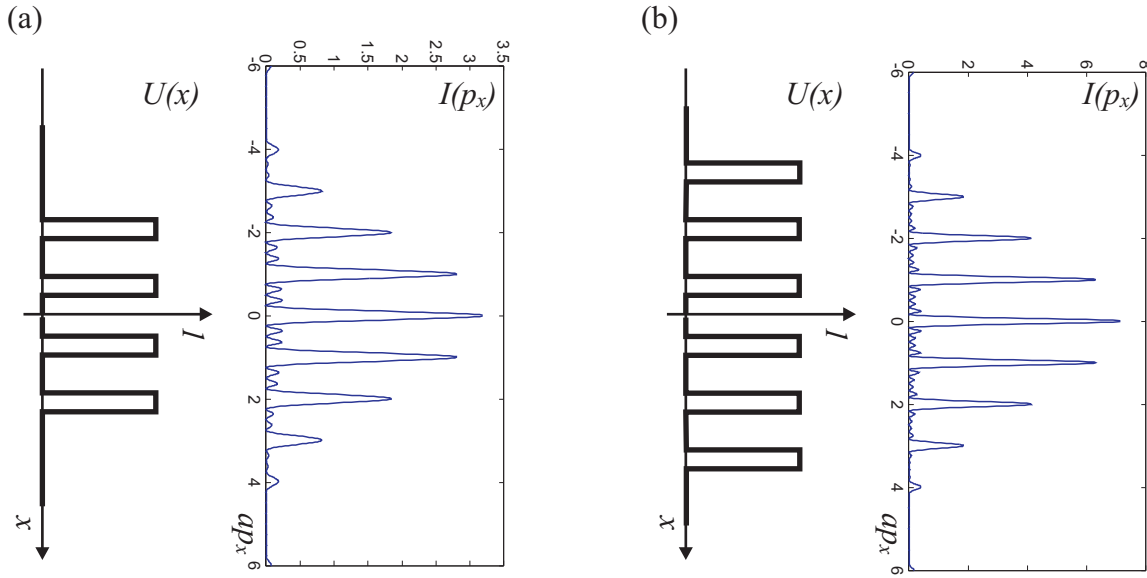


FIGURE 3.23 – Diffraction par un réseau composé de (a) $N = 4$ et (b) $N = 6$ fentes de largeur a et de période b : $U(x)$ et $I(p_x)$.

Remarquons que la diffraction par un système périodique – même très compliqué – donne toujours lieu au produit de deux termes, l'un issu de la périodicité du système et l'autre représentant la transformée de Fourier du motif répété périodiquement. On appelle souvent ce deuxième terme le facteur de forme.

3.13 Diffraction par une ouverture circulaire

Considérons maintenant une ouverture circulaire de diamètre D dans le plan $x'-y'$, Fig. 3.24. Dans le contexte de l'optique on appelle une telle ouverture une pupille ; elle peut par exemple correspondre au diamètre d'une lentille dans un système optique.

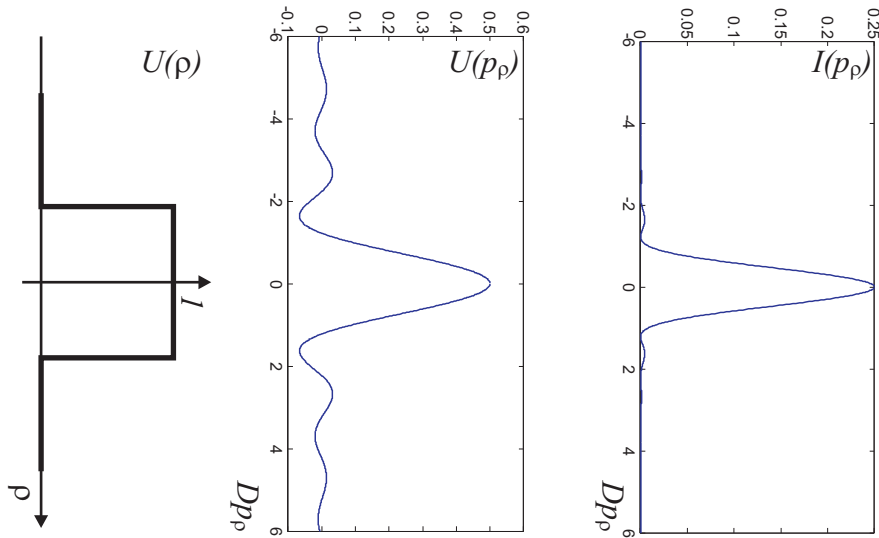
On suppose l'ouverture centrée en $(x', y') = (0, 0)$ et on introduit la distance ρ , $\rho^2 = x'^2 + y'^2$ ainsi que la fréquence spatiale $p_\rho = \sqrt{p_x^2 + p_y^2}$ pour tenir compte de la symétrie cylindrique du problème. La distribution du champ s'écrit,

$$U(\rho) = \begin{cases} 1 & \rho \leq D/2 \\ 0 & \rho > D/2. \end{cases} \quad (3.61)$$

La densité d'énergie spectrale correspondante vaut

$$U(p_\rho) = \frac{\pi D^2}{2} \frac{J_1(\pi D p_\rho)}{\pi D p_\rho}, \quad (3.62)$$

où l'on a introduit la fonction de Bessel d'ordre 1 de première espèce $J_1(x)$. Nous allons rencontrer cette fonction dans un prochain chapitre lorsque nous évoquerons les modes d'une

FIGURE 3.24 – Diffraction par une ouverture circulaire : $U(\rho)$, $U(p_\rho)$ et $I(p_\rho)$.

fibres optiques. Il est intéressant de remarquer que les fonctions $\sin x/x$ et $2J_1(x)/x$ sont assez proches l'une de l'autre, comme illustré Fig. 3.25.

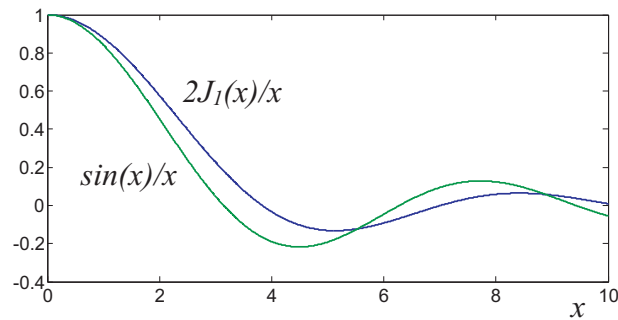
La distribution d'intensité vaut,

$$I(p_\rho) = \left(\frac{\pi D^2}{4\lambda R} \right)^2 \left(\frac{2J_1(\pi D p_\rho)}{\pi D p_\rho} \right)^2. \quad (3.63)$$

Le premier zéro de la fonction de Bessel $J_1(x)$ s'obtient pour $x = 3.83$. On peut donc calculer à l'aide de (3.47) et (3.63) la fréquence spatiale p_ρ correspondant à ce premier zéro : $\pi D p_\rho = 3.83$, d'où,

$$\sin \theta_\rho = \frac{3.83}{\pi} \frac{\lambda}{D} = 1.22 \frac{\lambda}{D}. \quad (3.64)$$

Les équations (3.62)–(3.64) sont extrêmement importantes pour les systèmes optiques. En effet, toute lentille de longueur focale f a une extension latérale finie et constitue donc une

FIGURE 3.25 – Comparaison des fonctions $2J_1(x)/x$ et $\sin x/x$.

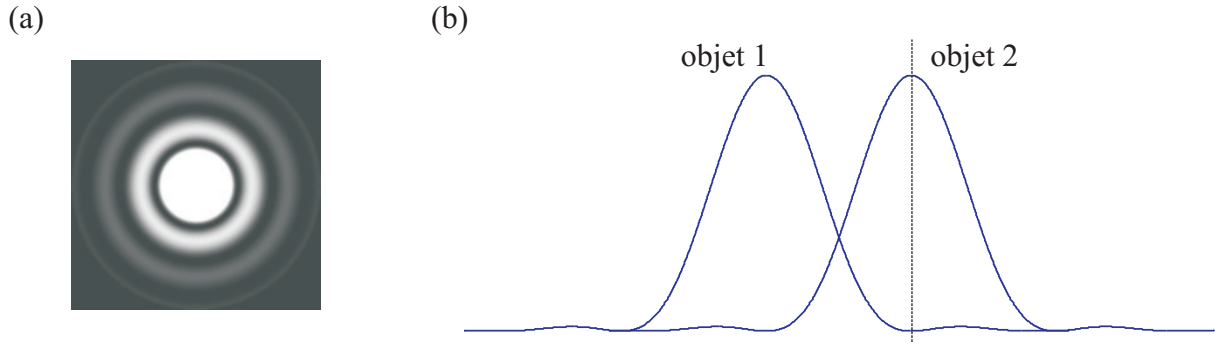


FIGURE 3.26 – (a) Disque d'Airy et (b) critère de résolution spatiale pour deux objets proches.

pupille qui se comporte selon ces équations. Ainsi une onde plane illuminant cette lentille crée une tache – dite d'Airy – dont le rayon est défini par

$$\rho = f\theta_\rho = 1.22 \frac{\lambda f}{D}, \quad (3.65)$$

comme illustré Fig. 3.26(a). Ces relations permettent aussi de donner un critère pour la résolution d'une lentille : Si deux objets sont très proches l'un de l'autre, leurs taches d'Airy se superposent et il ne sera pas possible de les résoudre. Généralement, on prend comme critère que le maxima du deuxième objet coïncide avec le minima du premier, Fig. 3.26(b).

3.14 Holographie

L'holographie consiste en l'enregistrement et la reproduction d'ondes optiques, un hologramme est un objet transparent dans lequel une onde optique est enregistrée. Il importe de noter que le caractère ondulatoire de l'onde est parfaitement reproduit dans un hologramme, c'est à dire l'amplitude et la phase. Par opposition, une photographie ordinaire ne permet de reproduire que l'intensité de l'onde.

Considérons une onde plane monochromatique dans un plan quelconque, par exemple le plan $z = 0$. L'amplitude complexe de cette onde vaut $U_o(x, y)$ dans le plan $z = 0$. Si nous sommes capables d'enregistrer dans un film une transmittance *complexe* $t(x, y)$, alors on peut reconstituer l'onde originale en envoyant sur ce film une onde plane d'amplitude unité. Dans ce cas, l'onde dans le plan $z = 0$ devient $U(x, y) = 1 t(x, y) = U_o(x, y)$ et ainsi, pour tous les points $z > 0$, l'onde reconstituée est égale à l'onde originale. Figure 3.11 illustre d'ailleurs ce principe, puisqu'une onde à incidence normale sur le milieu transparent donne lieu à une onde transmise dans une direction particulière k de l'autre côté du transparent. Ainsi, l'enregistrement dans le transparent permet-il de reconstruire l'onde originale se propageant dans la direction k .

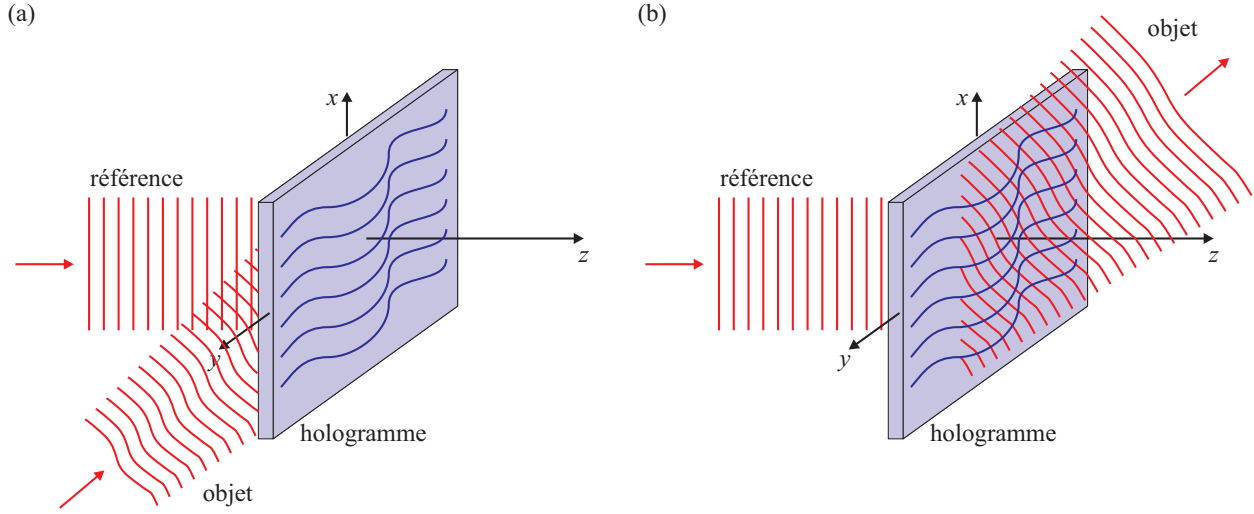


FIGURE 3.27 – (a) Enregistrement et (b) reconstruction d'une onde objet U_o au moyen d'une onde de référence U_r .

La question clef est d'arriver à enregistrer la transmission *complexe* dans le film transparent. Dans un hologramme, on réalise cet enregistrement en mélangeant l'onde incidente U_o (onde "objet") avec une onde de référence U_r et en enregistrant la figure d'interférence de ces deux ondes dans le plan $z = 0$. L'intensité de cette figure d'interférence est ainsi enregistrée et détermine la transparence de l'hologramme :

$$t \propto |U_o + U_r|^2 = |U_r|^2 + |U_o|^2 + U_r^* U_o + U_r U_o^* = I_r + I_o + U_r^* U_o + U_r U_o^*. \quad (3.66)$$

On remarque dans cette équation que la transmittance contient des informations sur l'intensité des deux ondes, ainsi que sur leur différence de phase (le terme $U_r^* U_o + U_r U_o^*$ peut en effet s'écrire en introduisant les intensités et arguments des ondes dans le plan $z = 0$: $2\sqrt{I_r I_o} \cos[\arg U_r - \arg U_o]$). Ce processus d'enregistrement est illustré Fig. 3.27(a).

Pour reconstruire l'onde objet U_o à partir de l'hologramme, il suffit de l'illuminer avec l'onde de référence U_r , comme indiqué Fig. 3.27(b). On obtient alors une onde de la forme suivante dans le plan $z = 0$:

$$U = U_r t \propto U_r I_r + U_r I_o + I_r U_o + U_r^2 U_o^*. \quad (3.67)$$

On remarque que le troisième terme représente l'onde originale objet enregistrée U_o multipliée par l'intensité de l'onde de référence. Si l'intensité de l'onde de référence est uniforme dans le plan $z = 0$ (c'est par exemple le cas pour une onde plane), alors le terme $I_r U_o$ représente bien l'onde originale. Malheureusement, cette onde originale est maintenant mélangée avec d'autres ondes : les deux premiers termes dans Eq. (3.67) correspondent à l'onde de référence modulée par la somme des intensités des ondes référence et objet ; le quatrième terme représente le complexe conjugué de l'onde objet, modulé par le carré de l'onde de référence. Tout l'art de l'holographie consiste donc à séparer ces différentes composantes pour extraire celle qui contient l'objet original. Différents montages optiques permettent d'y arriver, comme nous allons le voir ci dessous.

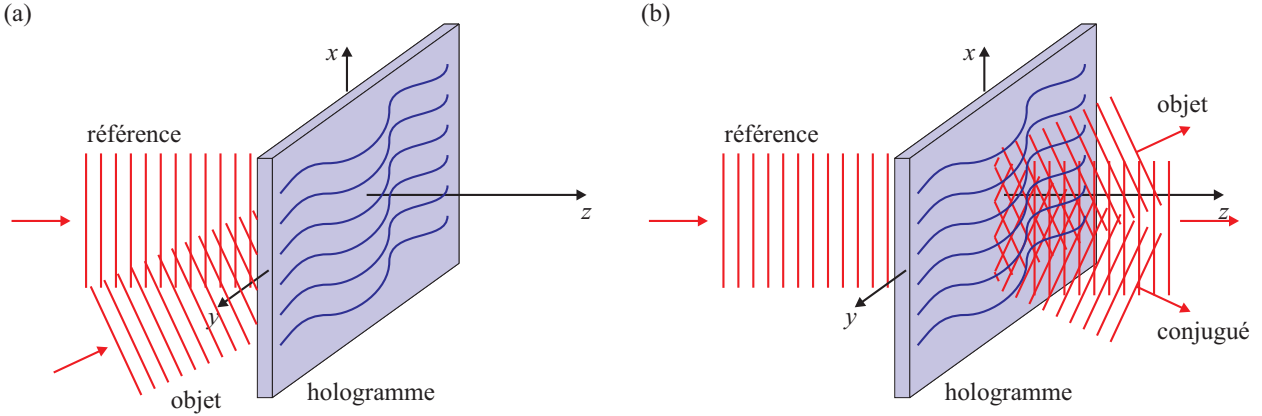


FIGURE 3.28 – (a) Onde objet consistant en une onde plane se propageant dans la direction θ par rapport à l'axe z . (b) Différentes composantes de l'onde reconstruite. Il apparaît en particulier une onde conjuguée se propageant dans la direction $-\theta$.

Pour simplifier, nous considérons que l'onde de référence est une onde plane se propageant dans la direction z et que nous écrivons $\sqrt{I_r} \exp(-jkz)$. Son intensité est donc constante dans le plan $z = 0$. On peut alors récrire Eq. (3.67) en la divisant par $U_r = \sqrt{I_r}$,

$$U(x, y) \propto I_r + I_o(x, y) + \sqrt{I_r} U_o(x, y) + \sqrt{I_r} U_o^*(x, y), \quad (3.68)$$

où l'on a mis en évidence les dépendances spatiales dans le plan $z = 0$.

Figure 3.28 décrit la situation où l'objet est une onde plane se propageant dans la direction θ par rapport à z , $U_o(x, y) = \sqrt{I_o} \exp(-jk \sin \theta x)$. L'onde reconstruite s'obtient à partir d'Eq. (3.68) :

$$U(x, y) \propto I_r + I_o + \sqrt{I_r I_o} \exp(-jk \sin \theta x) + \sqrt{I_r I_o} \exp(jk \sin \theta x). \quad (3.69)$$

Les deux premiers termes sont constants dans le plan $z = 0$ et correspondent à une onde se propageant dans la direction z . Le troisième terme correspond à l'onde objet originale se propageant dans la direction θ . Le dernier terme dans Eq. (3.69) correspond à une onde conjuguée de l'onde objet, se propageant dans la direction $-\theta$ par rapport à l'axe z , Fig. 3.28(b). Dans ce cas, l'onde objet reconstruite est parfaitement séparable des autres ondes et en regardant dans la direction adéquate, un observateur se trouvant en $z > 0$ ne peut pas distinguer l'onde reconstruite de l'onde objet originale. Cet observateur a donc l'impression de voir l'objet original, alors que celui-ci n'existe plus dans le montage de Fig. 3.28(b).

L'enregistrement et la reconstruction d'une onde sphérique issue d'une source ponctuelle sont illustrés Fig. 3.29. La source se trouve au point $\mathbf{r}_0 = (x, y, z) = (0, 0, -d)$ et l'onde objet prend dans le plan $z = 0$ la forme $U_o(x, y) = \exp(-jk |\mathbf{r} - \mathbf{r}_0|) / |\mathbf{r} - \mathbf{r}_0|$, avec $\mathbf{r} = (x, y, 0)$. En introduisant cette onde objet dans Eq. (3.68), on obtient à nouveau quatre termes pour l'onde reconstruite. Le premier, I_r , correspond à nouveau à une onde plane se propageant dans la direction z ; le second, $I_o = 1/|\mathbf{r} - \mathbf{r}_0|^2$ correspond à une onde quasi-plane, se propageant dans la direction z (la variation latérale du terme $1/|\mathbf{r} - \mathbf{r}_0|$ dans le plan (x, y)

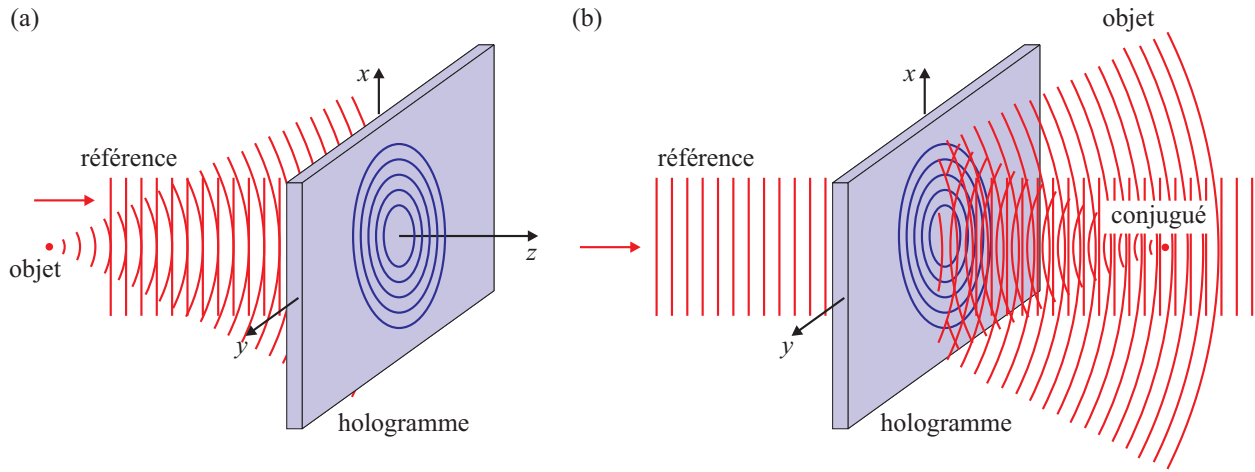


FIGURE 3.29 – (a) Onde objet sphérique issue d’une source ponctuelle et (b) sa reconstruction par holographie. L’onde conjuguée forme une image réelle de la source ponctuelle.

est lente); le troisième terme correspond à l’onde objet originale et représente une onde sphérique d’origine $(0, 0, -d)$; le dernier terme est proportionnel au complexe conjugué de l’onde objet : $U_o^* \propto \exp(jk|\mathbf{r} - \mathbf{r}_0|)/|\mathbf{r} - \mathbf{r}_0|$ et correspond à une onde convergeante vers le point $(0, 0, d)$. Cette dernière onde, conjuguée de l’onde objet originale, donne lieu à une image réelle de la source, Fig. 3.29(b).

Une façon de séparer les quatre ondes reconstruites est de s’assurer qu’elles se propagent dans des directions bien différentes. Cette condition est remplie lorsque lors de l’enregistrement les ondes objet et référence se propagent dans des directions différentes, Fig. 3.30. Considérons

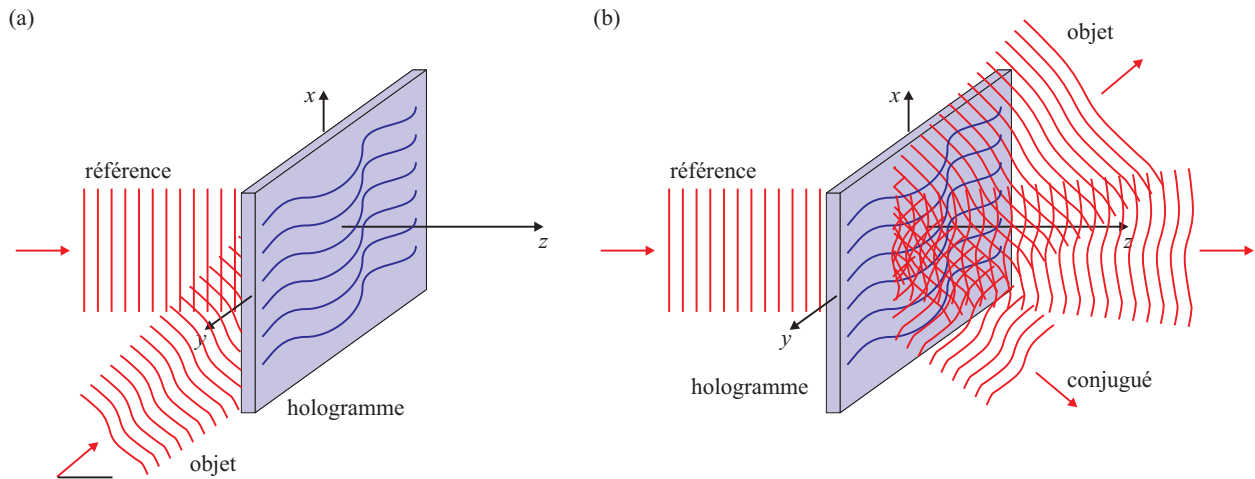


FIGURE 3.30 – (a) Enregistrement et (b) reconstruction d’un hologramme en utilisant des directions de propagation différentes pour l’objet et la référence. Les différentes ondes reconstruites sont bien séparées dans l’espace et peuvent être observées individuellement.

par exemple une onde objet comprenant une modulation spatiale $f(x, y)$ et se propageant dans la direction θ par rapport à l'axe z : $U_o(x, y) = f(x, y) \exp(-jk \sin \theta x)$ (on suppose que $f(x, y)$ varie lentement). L'onde reconstruite donnée par Eq. (3.68) devient

$$U(x, y) \propto I_r + |f(x, y)|^2 + \sqrt{I_r} f(x, y) \exp(-jk \sin \theta x) + \sqrt{I_r} f^*(x, y) \exp(jk \sin \theta x). \quad (3.70)$$

Le premier terme correspond à nouveau à une onde plane se propageant dans la direction z . Le deuxième terme à une onde quasi-plane, dans un cône centré autour de l'axe z . Le troisième terme correspond à l'onde objet et se propage dans la direction θ . Le dernier terme correspond au conjugué de l'objet et se propage dans la direction $-\theta$. Ainsi, suivant la direction dans laquelle on observe la reconstruction de l'hologramme, on peut “voir” telle ou telle composante. En choisissant la direction d'observation θ , on voit une onde qui correspond exactement à celle de l'objet original. Il n'est alors pas possible pour l'observateur de déterminer si l'objet est effectivement présent, ou remplacé par l'hologramme.

Au-delà de la reconstruction d'objets en trois dimensions, il a existé un intérêt marqué pour le stockage de l'information de façon holographique dans des structures tri-dimensionnelles. Cet intérêt a cependant été estompé par la difficulté de trouver des matériaux holographiques reconfigurables, permettant de stocker de l'information de façon dynamique en l'écrivant et en l'effaçant.

Chapitre 4

Optique de Maxwell et polarisation de la lumière

4.1 Le champ électromagnétique

Pour une description complète de la lumière et de sa propagation, il faut introduire le champ électromagnétique. Celui-ci comprend un ensemble de quatre champs, comme indiqué Tab. 4.1. Il est important de remarquer que nous utilisons le système d'unités MKSA (de Mètre, Kilo, Seconde et Ampère) connu aussi sous le nom de système international (SI). En effet, les équations de Maxwell prennent une forme différente dans d'autres systèmes d'unités.

En optique, le champ électrique $\mathbf{E}$ joue un rôle prépondérant car c'est lui qui interagit généralement avec la matière et permet par exemple de détecter une onde optique par la recombinaison de porteurs de charge dans un semiconducteur. Le champ $\mathbf{D}$ lui est étroitement associé et représente en quelque sorte la réponse de la matière à un champ électrique externe. Quant aux champs magnétique et d'induction magnétique $\mathbf{H}$ et $\mathbf{B}$ ils apparaissent dès qu'il existe une variation spatio-temporelle du champ électrique. Nous sommes donc contraints en optique d'étudier l'ensemble de ces quatre champs.

TABLE 4.1 – Les différents champs intervenant dans les équations de Maxwell

Symbole	Champ	Unité SI
$\mathbf{E}(\mathbf{r}, t)$	champ électrique	$\mathbf{V/m}$
$\mathbf{H}(\mathbf{r}, t)$	champ magnétique	$\mathbf{A/m}$
$\mathbf{D}(\mathbf{r}, t)$	champ de déplacement électrique	$\mathbf{As/m}^2$
$\mathbf{B}(\mathbf{r}, t)$	champ d'induction magnétique	$\mathbf{Vs/m}^2$

Notons que les champs indiqués Tab. 4.1 sont des champs vectoriels qui comportent trois composantes, dont la valeur dépend évidemment du système de coordonnées choisi :

$$\mathbf{E}(\mathbf{r}, t) = \begin{pmatrix} E_x(x, y, z, t) \\ E_y(x, y, z, t) \\ E_z(x, y, z, t) \end{pmatrix} \quad \mathbf{E}(\mathbf{r}, t) = \begin{pmatrix} E_r(r, \theta, \phi, t) \\ E_\theta(r, \theta, \phi, t) \\ E_\phi(r, \theta, \phi, t) \end{pmatrix} \quad \mathbf{E}(\mathbf{r}, t) = \begin{pmatrix} E_\rho(\rho, \phi, z, t) \\ E_\phi(\rho, \phi, z, t) \\ E_z(\rho, \phi, z, t) \end{pmatrix}. \quad (4.1)$$

Tous les opérateurs vectoriels que nous rencontrerons doivent aussi s'exprimer dans le système de coordonnées correspondant ; ainsi si des coordonnées cylindriques sont utilisées pour le champ, il faut utiliser les mêmes coordonnées pour ∇ , $\nabla \times$, $\nabla \cdot$ et ∇^2 comme ce fut par exemple le cas pour Eq. (2.5).

4.2 Equations de Maxwell

Les équations de Maxwell relient les variations spatiales et temporelles des quatre champs décrits Tab. 4.1. Dans le système MKSA elles ont la forme,

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\frac{\partial \mathbf{B}(\mathbf{r}, t)}{\partial t}, \quad (4.2a)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \frac{\partial \mathbf{D}(\mathbf{r}, t)}{\partial t} + \mathbf{J}(\mathbf{r}, t), \quad (4.2b)$$

$$\nabla \cdot \mathbf{D}(\mathbf{r}, t) = \rho(\mathbf{r}, t), \quad (4.2c)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}, t) = 0. \quad (4.2d)$$

Il est remarquable de noter que ces équations introduites au 19ème siècle n'ont jamais été démenties par la physique moderne et continuent à trouver un très vaste terrain d'applications.

Il existe deux termes source dans Eq. (4.2) : la densité de courant $\mathbf{J}$ (unités A/m²) et la densité de charge ρ (unités As/m³). Pour l'optique, on peut se limiter au cas sans source. En effet, le terme qui nous intéresse est le champ électromagnétique qui – comme nous le verrons – s'auto-propage de proche en proche et existe même en l'absence de sources comme des charges ou des courants. Notons que la densité de charge se rencontre surtout en électrostatique où elle est la source du champ électrique statique. Par contre des charges immobiles n'influencent pas le champ électromagnétique aux fréquences optiques. Quant à la densité de courant, elle peut jouer un rôle dans l'optique des métaux, en particulier comme source de dissipation du champ électromagnétique. Cependant, il s'agit de phénomènes particuliers et en général un courant ne peut pas créer directement un champ optique : les électrons ont un mouvement de fréquence bien inférieure à celle associée à la lumière. Dans ces conditions,

les équations de Maxwell s'écrivent sans source

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\frac{\partial \mathbf{B}(\mathbf{r}, t)}{\partial t}, \quad (4.3a)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \frac{\partial \mathbf{D}(\mathbf{r}, t)}{\partial t}, \quad (4.3b)$$

$$\nabla \cdot \mathbf{D}(\mathbf{r}, t) = 0, \quad (4.3c)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}, t) = 0. \quad (4.3d)$$

Les propriétés de la matière dans laquelle se développe le champ électromagnétique sont décrites par les relations constitutives :

$$\mathbf{D}(\mathbf{r}, t) = \epsilon_0 \mathbf{E}(\mathbf{r}, t) + \mathbf{P}(\mathbf{r}, t) = \epsilon_0 \epsilon_r(\mathbf{r}, t) \mathbf{E}(\mathbf{r}, t) = \epsilon_0 (1 + \chi(\mathbf{r}, t)) \mathbf{E}(\mathbf{r}, t), \quad (4.4)$$

et

$$\mathbf{B}(\mathbf{r}, t) = \mu_0 \mathbf{H}(\mathbf{r}, t) + \mu_0 \mathbf{M}(\mathbf{r}, t) = \mu_0 \mu_r(\mathbf{r}, t) \mathbf{H}(\mathbf{r}, t). \quad (4.5)$$

Deux nouveaux vecteurs ont été introduits : la densité de polarisation $\mathbf{P}(\mathbf{r}, t)$ (unités As/m²) et l'aimantation $\mathbf{M}(\mathbf{r}, t)$ (unités A/m). Cette dernière joue un rôle négligeable en optique puisqu'elle trouve son origine dans des courants électriques ; or ceux-ci ne peuvent pas se développer à des fréquences aussi élevées que celles de la lumière. La densité de polarisation joue par contre un rôle essentiel car elle est responsable de la réponse d'un matériau au champ extérieur. Ainsi, lorsqu'une onde de lumière arrive sur un morceau de matière, celui-ci se polarise et génère un champ $\mathbf{P}$ en sorte que le champ de déplacement est donné par $\mathbf{D}(\mathbf{r}, t) = \epsilon_0 \mathbf{E}(\mathbf{r}, t) + \mathbf{P}(\mathbf{r}, t)$.

Les paramètres ϵ_0 et μ_0 sont la permittivité et la perméabilité du vide. Les paramètres ϵ_r et μ_r représentent la permittivité et la perméabilité relative du milieu (par relatif on entend ramené aux valeurs de référence du vide). Pour la suite il est utile d'introduire la permittivité ϵ et la perméabilité μ qui combinent les valeurs relatives et celles du vide :

$$\epsilon = \epsilon_0 \epsilon_r, \quad (4.6a)$$

$$\mu = \mu_0 \mu_r. \quad (4.6b)$$

Dans Eq. (4.4) nous avons introduit une permittivité relative scalaire, qui indique que le champ $\mathbf{D}$ s'obtient simplement en multipliant le champ $\mathbf{E}$ avec un nombre ; ainsi $\mathbf{E}$ et $\mathbf{D}$ sont parallèles ; on parle alors de milieu isotrope. Tel n'est pas toujours le cas : dans un cristal la permittivité relative peut prendre la forme d'un tenseur (matrice 3×3) et les champs $\mathbf{E}$ et $\mathbf{D}$ ne sont alors pas nécessairement parallèles. De tels milieux anisotropes seront décrits ultérieurement.

Un autre paramètre a été introduit dans Eq. (4.4) : la susceptibilité χ qui est étroitement liée à la permittivité relative : $\epsilon_r = (1 + \chi)$. Si on récrit cette relation $\chi = \epsilon_r - 1$ on comprend que la susceptibilité exprime en quelque sorte comment la réponse d'un matériau

diffère de celle du vide, pour lequel $\epsilon_r = 1$. La susceptibilité joue un rôle important pour l'optique non-linéaire.

Dans le vide ($\epsilon_r = \mu_r = 1$), sans source, les équations de Maxwell prennent la forme

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\mu_0 \frac{\partial \mathbf{H}(\mathbf{r}, t)}{\partial t}, \quad (4.7a)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \epsilon_0 \frac{\partial \mathbf{E}(\mathbf{r}, t)}{\partial t}, \quad (4.7b)$$

$$\nabla \cdot \mathbf{E}(\mathbf{r}, t) = 0, \quad (4.7c)$$

$$\nabla \cdot \mathbf{H}(\mathbf{r}, t) = 0; \quad (4.7d)$$

et les relations constitutives deviennent

$$\mathbf{D}(\mathbf{r}, t) = \epsilon_0 \mathbf{E}(\mathbf{r}, t), \quad (4.8)$$

et

$$\mathbf{B}(\mathbf{r}, t) = \mu_0 \mathbf{H}(\mathbf{r}, t). \quad (4.9)$$

Notons que dans ce cas $\mathbf{P}(\mathbf{r}, t) = \mathbf{M}(\mathbf{r}, t) = \mathbf{0}$.

Les paramètres ϵ_0 et μ_0 ont des valeurs bien spécifiques dans le système MKSA :

$$\epsilon_0 = \frac{10^7}{4\pi c_0^2} \simeq 8.8541878 \cdot 10^{-12} \text{ F/m}, \quad (4.10)$$

$$\mu_0 = 4\pi \cdot 10^{-7} \text{ H/m}. \quad (4.11)$$

On remarque que l'on a $\epsilon_0 \mu_0 = 1/c_0^2$. Il faut cependant faire attention car cette relation n'est valable que dans le système MKSA. Dans d'autres systèmes d'unités utilisés en optique ces constantes prennent d'autres valeurs ce qui change la forme des équations de Maxwell, voire même la valeur de la vitesse de la lumière dans le vide c_0 . Ainsi, dans le système cgs (de Coulomb, Gauss, seconde) que l'on trouve dans beaucoup de livres d'optique, la vitesse de la lumière dans le vide a la valeur $c_0 = 1$.

4.3 Equation d'onde

Commençons par considérer le cas du vide. En appliquant le rotationnel à Eq. (4.7a) et en utilisant Eq. (4.7b), on obtient en omettant les dépendances spatio-temporelles pour alléger l'écriture

$$\nabla \times (\nabla \times \mathbf{E}) = -\mu_0 \frac{\partial (\nabla \times \mathbf{H})}{\partial t} = -\mu_0 \epsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2}. \quad (4.12)$$

En utilisant l'égalité vectorielle $\nabla \times (\nabla \times \mathbf{v}) = \nabla (\nabla \cdot \mathbf{v}) - \nabla^2$ et le fait que $\nabla \cdot \mathbf{E} = 0$ selon Eq. (4.7c), Eq. (4.12) devient

$$\nabla^2 \mathbf{E} - \frac{1}{c_0^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0, \quad (4.13)$$

où l'on a justement utilisé le fait que $\epsilon_0\mu_0 = 1/c_0^2$.

Il faut noter que l'opérateur $\nabla^2\mathbf{E}$ dans Eq. (4.13) représente le Laplacien vectoriel. Cela signifie que chaque composante du champ électrique satisfait une équation de la forme (4.13). Ainsi, par exemple pour la composante y du champ électrique :

$$\nabla^2 E_y - \frac{1}{c_0^2} \frac{\partial^2 E_y}{\partial t^2} = 0. \quad (4.14)$$

Nous avons donc déduit des équations de Maxwell l'équation d'onde (2.1) qui avait été introduite au chapitre 2.

Lorsque le milieu dans lequel se propage l'onde n'est pas le vide ($\epsilon_r \neq 1$ et/ou $\mu_r \neq 1$), des opérations similaires en partant des Eqs. (4.3) donnent l'équation d'onde suivante

$$\nabla^2 \mathbf{E} - \frac{\epsilon_r \mu_r}{c_0^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0. \quad (4.15)$$

On peut récrire cette équation sous la forme

$$\nabla^2 \mathbf{E} - \frac{1}{c^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0, \quad (4.16)$$

où l'on a introduit la vitesse de propagation c dans le milieu :

$$c = \frac{1}{\sqrt{\epsilon\mu}} = \frac{1}{\sqrt{\epsilon_0\epsilon_r\mu_0\mu_r}} = \frac{c_0}{\sqrt{\epsilon_r\mu_r}}. \quad (4.17)$$

En utilisant les notions introduites au début du chapitre 1, on peut définir l'indice de réfraction n du milieu :

$$n = \frac{c_0}{c} = \sqrt{\frac{\epsilon}{\epsilon_0} \frac{\mu}{\mu_0}}. \quad (4.18)$$

Comme la plupart des matériaux rencontrés en optique sont non-magnétiques ($\mu_r = 1$, soit $\mu = \mu_0$) on a pour l'indice de réfraction,

$$n = \sqrt{\frac{\epsilon}{\epsilon_0}} = \sqrt{\epsilon_r}. \quad (4.19)$$

4.4 Onde plane monochromatique harmonique

Si nous considérons maintenant des champs monochromatiques harmoniques avec une dépendance temporelle $\exp(j\omega t)$, les équations de Maxwell sans source prennent une forme particulièrement simple puisque chaque dérivée temporelle est remplacée par une multiplication par $j\omega t$:

$$\nabla \times \mathbf{E}(\mathbf{r}) = -j\omega \mathbf{B}(\mathbf{r}), \quad (4.20a)$$

$$\nabla \times \mathbf{H}(\mathbf{r}) = j\omega \mathbf{D}(\mathbf{r}), \quad (4.20b)$$

$$\nabla \cdot \mathbf{D}(\mathbf{r}) = 0, \quad (4.20c)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}) = 0. \quad (4.20d)$$

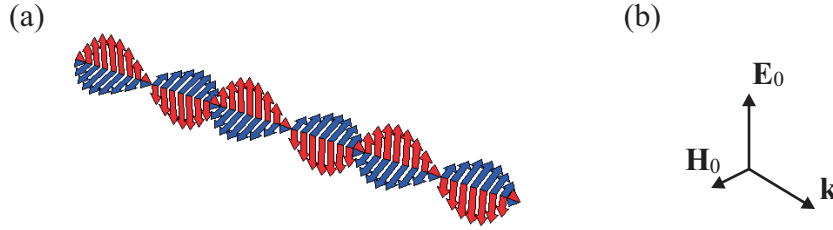


FIGURE 4.1 – Onde plane : (a) évolution spatiale des champs $\mathbf{E}$ et $\mathbf{H}$, (b) relation entre $\mathbf{E}_0$, $\mathbf{H}_0$ et $\mathbf{k}$.

En considérant de plus une onde plane monochromatique harmonique avec comme vecteurs champ électrique et magnétique

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0 e^{-j\mathbf{k} \cdot \mathbf{r} + j\omega t} \quad (4.21a)$$

$$\mathbf{H}(\mathbf{r}, t) = \mathbf{H}_0 e^{-j\mathbf{k} \cdot \mathbf{r} + j\omega t}, \quad (4.21b)$$

les dérivées spatiales $\nabla \times$ sont remplacées par des produits vectoriels $-j\mathbf{k} \times$ et Eqs. (4.20) deviennent

$$-j\mathbf{k} \times \mathbf{E}(\mathbf{r}) = -j\omega \mathbf{B}(\mathbf{r}), \quad (4.22a)$$

$$-j\mathbf{k} \times \mathbf{H}(\mathbf{r}) = j\omega \mathbf{D}(\mathbf{r}), \quad (4.22b)$$

$$-j\mathbf{k} \cdot \mathbf{D}(\mathbf{r}) = 0, \quad (4.22c)$$

$$-j\mathbf{k} \cdot \mathbf{B}(\mathbf{r}) = 0. \quad (4.22d)$$

Ces équations doivent être vérifiées en tout point de l'espace et en tout moment ; on obtient donc des relations entre $\mathbf{E}_0$, $\mathbf{H}_0$ et $\mathbf{k}$:

$$\mathbf{k} \times \mathbf{E}_0 = \omega \mu \mathbf{H}_0, \quad (4.23a)$$

$$\mathbf{k} \times \mathbf{H}_0 = -\omega \epsilon \mathbf{E}_0. \quad (4.23b)$$

Ces relations indiquent que $\mathbf{E}$, $\mathbf{H}$ et $\mathbf{k}$ forment un triplet droitier : en faisant $\mathbf{E} \times \mathbf{H}$ on obtient un vecteur parallèle à $\mathbf{k}$, Fig. 4.1(b).

Le rapport entre les amplitudes des champs électriques et magnétiques permet de définir l'impédance η du milieu :

$$\eta = \sqrt{\frac{\mu_0 \mu_r}{\epsilon_0 \epsilon_r}}. \quad (4.24)$$

On définit aussi l'impédance du vide,

$$\eta_0 = \sqrt{\frac{\mu_0}{\epsilon_0}} \simeq 376.6 \, \Omega. \quad (4.25)$$

Finalement, le flux instantané d'énergie associé à une onde plane est défini par le vecteur de Poynting,

$$\mathbf{S}(\mathbf{r}, t) = \mathbf{E}(\mathbf{r}, t) \times \mathbf{H}(\mathbf{r}, t). \quad (4.26)$$

En général on s'intéresse au flux d'énergie moyenné dans le temps (i.e. moyenné sur une période T). Pour une onde plane il s'écrit

$$\mathbf{S} = \frac{1}{2} \mathbf{E} \times \mathbf{H}^* . \quad (4.27)$$

L'introduction du complexe conjugué et du facteur $1/2$ correspondent à un moyennage dans le temps (on se souviendra que la moyenne sur une période d'une fonction harmonique comme $\cos^2 x$ vaut $1/2$). Les unités du vecteur de Poynting sont (W/m^2).

4.5 Conditions d'interface

Les équations de Maxwell en l'absence de source (4.3) imposent des conditions importantes aux composantes parallèles et perpendiculaires des champs lorsqu'une onde électromagnétique traverse l'interface entre les milieux 1 et 2 caractérisés par les permittivités ϵ_1, ϵ_2 et les perméabilités μ_1 et μ_2 . Ces conditions sont rassemblées dans Tab. 4.2 et illustrées Fig. 4.2.

TABLE 4.2 – Conditions d'interface pour les champs électromagnétiques

Champ	Composante continue
$\mathbf{E}$	parallèle
$\mathbf{H}$	parallèle
$\mathbf{D}$	perpendiculaire
$\mathbf{B}$	perpendiculaire

En combinant ces conditions d'interface avec les relations constitutives (4.4) et (4.5) on peut déterminer le champ dans le milieu 2 en le connaissant dans le milieu 1. Cette procédure est illustrée Fig. 4.2(b) pour le calcul du champ électrique : on considère deux milieux caractérisés par les permittivités ϵ_1 et ϵ_2 ; le champ incident $\mathbf{E}_1 = (E_{1x}, 0, E_{1z})$ dans le milieu

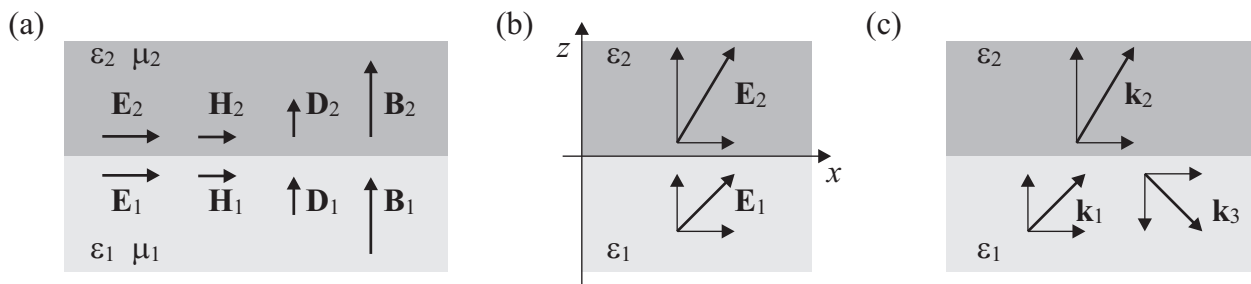


FIGURE 4.2 – (a) Continuité des champs électromagnétiques à une interface et (b) leur utilisation pour déterminer le champ dans le milieu 2 le connaissant dans le milieu 1 ; (c) continuité de la composante parallèle à l'interface du vecteur $\mathbf{k}$ pour une onde plane.

1 est connu, on cherche le champ dans le milieu 2 $\mathbf{E}_2 = (E_{2x}, E_{2y}, E_{2z})$. Par continuité, les composantes parallèles x et y sont les mêmes dans le milieu 2 : $\mathbf{E}_2 = (E_{1x}, 0, E_{2z})$. Pour déterminer la composante E_{2z} perpendiculaire à l'interface on utilise la continuité du vecteur $\mathbf{D}$: $\mathbf{D}_1 = (\epsilon_1 E_{1x}, 0, \epsilon_1 E_{1z})$ dans le milieu 1 et $\mathbf{D}_2 = (\epsilon_2 E_{2x}, \epsilon_2 E_{2y}, \epsilon_2 E_{2z})$ dans le milieu 2. Ainsi la continuité impose-t-elle $E_{2z} = \epsilon_1 E_{1z} / \epsilon_2$. Le champ électrique dans le milieu 2 vaut donc finalement $\mathbf{E}_2 = (E_{1x}, 0, \epsilon_1 E_{1z} / \epsilon_2)$.

Notons que les équations de Maxwell permettent aussi de dériver la loi de Snell. Celle-ci ne représente donc qu'une des conditions d'interface imposées par les équations de Maxwell, comme illustré Fig. 4.2(c). En l'occurrence il s'agit de la continuité de la composante parallèle du vecteur d'onde. En effet, dans le milieu 1 la norme du vecteur d'onde est donnée par $k_1 = n_1 k_0$ et dans le milieu 2 par $k_2 = n_2 k_0$. La loi de Snell (1.4) $n_1 \sin \theta_1 = n_2 \sin \theta_2$ peut se récrire en multipliant par k_0 : $n_1 k_0 \sin \theta_1 = n_2 k_0 \sin \theta_2$ qui n'est autre que la condition de continuité des composantes parallèles à l'interface des vecteurs $\mathbf{k}_1$ et $\mathbf{k}_2$ indiqués Fig. 4.2(c).

Ainsi, si l'on connaît le vecteur d'onde incident $\mathbf{k}_1$ on peut trouver par continuité la composante parallèle à l'interface du vecteur réfracté $\mathbf{k}_2$. La composante normale $\mathbf{k}_2$ s'obtient en s'assurant que la norme de $\mathbf{k}_2$ vaut $k_2 = n_2 k_0$.

D'un point de vue physique, la continuité de la composante parallèle du vecteur d'onde signifie pour les photons la conservation de la quantité de mouvement le long de l'interface.

4.6 Coefficients de Fresnel

Les conditions d'interface découlant des équations de Maxwell permettent de déterminer les amplitudes des ondes transmises et réfléchies lors de la réfraction d'une onde plane sur une surface. Comme ces conditions d'interface sont différentes pour les composantes du champ électromagnétique normales et parallèles à l'interface, on peut comprendre qu'il faudra distinguer deux types de polarisation incidente, comme indiqué Fig. 4.3. Ces polarisations sont définies par rapport au plan d'incidence π . Celui-ci est le plan qui contient le vecteur de propagation $\mathbf{k}$ et est perpendiculaire à l'interface. On définit alors la polarisation transverse électrique (TE), où le champ électrique est transverse par rapport au plan d'incidence. On

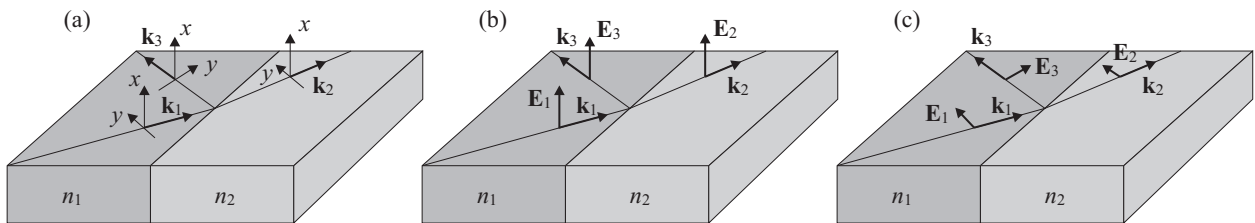


FIGURE 4.3 – (a) Système d'axe pour le calcul des ondes réfléchies et réfractées. Polarisation (b) TE (ou s), et (c) TM (ou p).

parle aussi de polarisation *s*, de l'allemand *senkrecht* pour perpendiculaire, puisque le champ électrique est perpendiculaire au plan d'incidence, Fig. 4.3(b). L'autre polarisation est la polarisation transverse magnétique (TM) ou *p* pour *parallel* en allemand. Dans ce cas, le champ magnétique est transverse par rapport au plan d'incidence et donc le champ électrique est parallèle au plan d'incidence, Fig. 4.3(c). Remarquer sur Fig. 4.3 que les ondes transmises et réfléchies gardent la même polarisation que l'onde incidente. Remarquer aussi sur cette même figure que les axes x - y pour l'onde incidente et l'onde réfléchie sont définis en sorte qu'ils coïncident lorsque $\theta_1 \rightarrow 0$. Cette convention a l'avantage de lever une ambiguïté sur les coefficients de réflexion à incidence normale. Il existe cependant des ouvrages où une autre convention est utilisée qui donne des résultats différents pour les coefficients de réflexion à incidence normale pour les polarisations TE et TM.

Les amplitudes des composantes du champ électrique transmis et réfléchi suivent les relations

$$\begin{aligned} E_{2x} &= t_x E_{1x} & E_{2y} &= t_y E_{1y} \\ E_{3x} &= r_x E_{1x} & E_{3y} &= r_y E_{1y} ; \end{aligned} \quad (4.28)$$

et les coefficients sont donnés par

$$r_x = \frac{\eta_2 \sec \theta_2 - \eta_1 \sec \theta_1}{\eta_2 \sec \theta_2 + \eta_1 \sec \theta_1} \quad t_x = 1 + r_x \quad \text{polarisation TE} \quad (4.29a)$$

$$r_y = \frac{\eta_2 \cos \theta_2 - \eta_1 \cos \theta_1}{\eta_2 \cos \theta_2 + \eta_1 \cos \theta_1} \quad t_y = (1 + r_y) \frac{\cos \theta_1}{\cos \theta_2} \quad \text{polarisation TM} . \quad (4.29b)$$

Rappelons que $\sec \theta = (\cos \theta)^{-1}$ (i.e. 1 divisé par la valeur du cosinus) et que l'impédance $\eta = \sqrt{\mu/\epsilon}$ peut être complexe, en particulier si ϵ est complexe comme c'est le cas pour un matériau avec des pertes. Pour un matériau non-magnétique ($\mu_r = 1$) sans perte, $\eta = \eta_0/n$ est réel avec $\eta_0 = \sqrt{\mu_0/\epsilon_0}$ et n l'indice de réfraction du milieu. Dans ce cas, Eqs. (4.29) sont connues sous le nom de **coefficients de Fresnel** et prennent la forme

$$r_x = \frac{n_1 \cos \theta_1 - n_2 \cos \theta_2}{n_1 \cos \theta_1 + n_2 \cos \theta_2} \quad t_x = 1 + r_x \quad \text{polarisation TE} \quad (4.30a)$$

$$r_y = \frac{n_1 \sec \theta_1 - n_2 \sec \theta_2}{n_1 \sec \theta_1 + n_2 \sec \theta_2} \quad t_y = (1 + r_y) \frac{\cos \theta_1}{\cos \theta_2} \quad \text{polarisation TM} . \quad (4.30b)$$

Connaissant n_1 , n_2 et θ_1 , les coefficients de Fresnel peuvent être déterminés en calculant $\sin \theta_2$ à l'aide de la loi de Snell (1.4). On en déduit ensuite

$$\cos \theta_2 = \sqrt{1 - \sin^2 \theta_2} = \sqrt{1 - (n_1/n_2)^2 \sin^2 \theta_1} . \quad (4.31)$$

Dans Eq. (4.31) l'argument de la racine carrée peut être négatif. Ainsi, en général, les coefficients de Fresnel sont complexes; i.e. ils possèdent une amplitude et une phase. D'un point de vue physique, cela signifie que l'onde réfléchie (ou transmise) subira un changement d'amplitude ainsi qu'un changement de phase. Ce déphasage peut jouer un rôle important pour des systèmes d'interférence multi-couches. Dans les sections suivantes, nous décrivons la valeur de coefficients de Fresnel en fonction de l'angle d'incidence pour les deux polarisations TE et TM. On distinguera aussi les cas de réflexion interne ($n_1 > n_2$) et externe ($n_1 < n_2$).

4.6.1 Polarisation TE ou s

Réflexion externe, ($n_1 < n_2$) l'onde se propage par exemple dans de l'air et arrive dans du verre, Fig. 4.4(a). Le coefficient de réflexion r_x est réel et négatif, correspondant à un saut de phase de $\phi_x = \pi$ pour l'onde réfléchie. L'amplitude de l'onde réfléchie augmente de $|r_x| = (n_2 - n_1)/(n_1 + n_2)$ à incidence normale, à $|r_x| = 1$ pour une incidence rasante ($\theta_1 \simeq 90^\circ$).

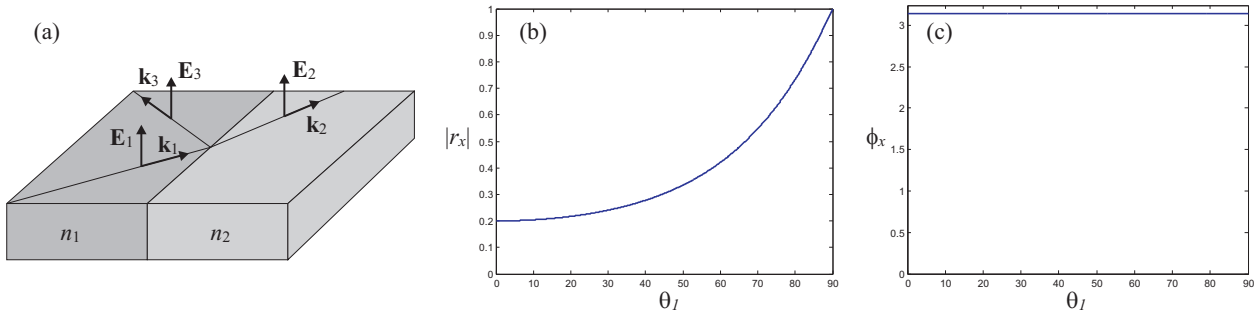


FIGURE 4.4 – Réflexion externe TE : (a) géométrie et (b) amplitude et (c) phase des coefficients de Fresnel.

Réflexion interne, ($n_1 > n_2$) l'onde se propage par exemple dans du verre et arrive dans l'air, Fig. 4.5(a). Pour des angles d'incidence plus petits que l'angle critique θ_c , le coefficient de réflexion est réel et positif. Il passe de $|r_x| = (n_1 - n_2)/(n_1 + n_2)$ à incidence normale, à $|r_x| = 1$ lorsque $\theta_1 = \theta_c$. Pour des angles supérieurs, l'amplitude de r_x reste égale à 1, correspondant à la réflexion interne totale. La phase de r_x est nulle pour $\theta_1 < \theta_c$ puis elle augmente progressivement pour $\theta_1 > \theta_c$. Cette phase joue un rôle important dans les guides d'onde.

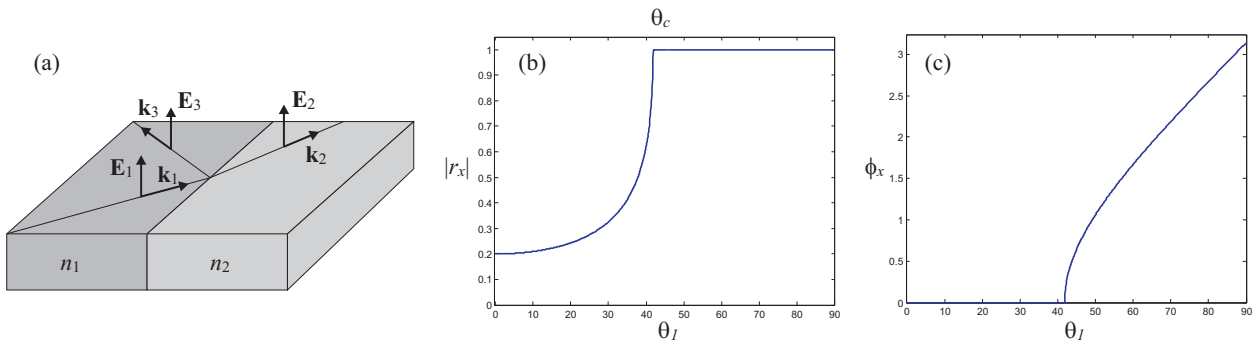


FIGURE 4.5 – Réflexion interne TE : (a) géométrie et (b) amplitude et (c) phase des coefficients de Fresnel.

4.6.2 Polarisation TM ou p

Réflexion externe, ($n_1 < n_2$) l'onde se propage par exemple dans de l'air et arrive dans du verre, Fig. 4.6(a). On s'intéresse maintenant au champ électrique parallèle au plan d'incidence. Le coefficient de réflexion est réel et négatif à incidence normale : $r_y = (n_1 - n_2)/(n_1 + n_2)$; il possède donc une phase de π . Sa valeur augmente — l'amplitude diminue, Fig.4.6(b) — jusqu'à l'angle défini par $n_1 \sec \theta_1 = n_2 \sec \theta_2$, Eq. (4.30b). Cet angle, connu sous le nom d'angle de Brewster, est défini par

$$\theta_B = \arctan(n_2/n_1). \quad (4.32)$$

Une onde polarisée TM incidente à l'angle $\theta_1 = \theta_B$ est donc entièrement transmise à l'interface, sa réflexion est nulle. Ce phénomène est utilisé pour réaliser des polariseurs.

Pour $\theta_1 > \theta_B$, r_y devient positif, ce qui se traduit par une phase nulle, Fig. 4.6(c). Comme pour la polarisation TE, l'amplitude du coefficient de Fresnel augmente jusqu'à l'unité pour une incidence rasante ($\theta_1 \simeq 90^\circ$).

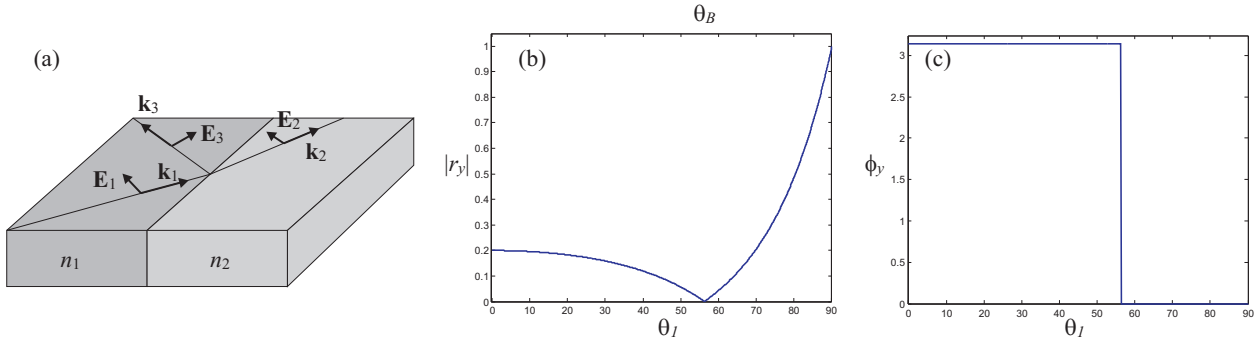


FIGURE 4.6 – Réflexion externe TM : (a) géométrie et (b) amplitude et (c) phase des coefficients de Fresnel.

Réflexion interne, ($n_1 > n_2$) l'onde se propage par exemple dans du verre et arrive dans l'air, Fig. 4.7(a). À incidence normale, le coefficient de Fresnel $r_y = (n_1 - n_2)/(n_1 + n_2)$ est réel et positif; sa phase est donc nulle, Fig. 4.7(b) et (c). L'amplitude diminue lorsque l'angle d'incidence augmente; elle s'annule pour $\theta_1 = \theta_B$. Le coefficient devient ensuite négatif ($\phi_y = \pi$) et son amplitude augmente jusqu'à atteindre 1 à l'angle critique θ_c . Pour $\theta_1 > \theta_c$, $|r_y| = 1$ et la phase décroît, Fig. 4.7(c).

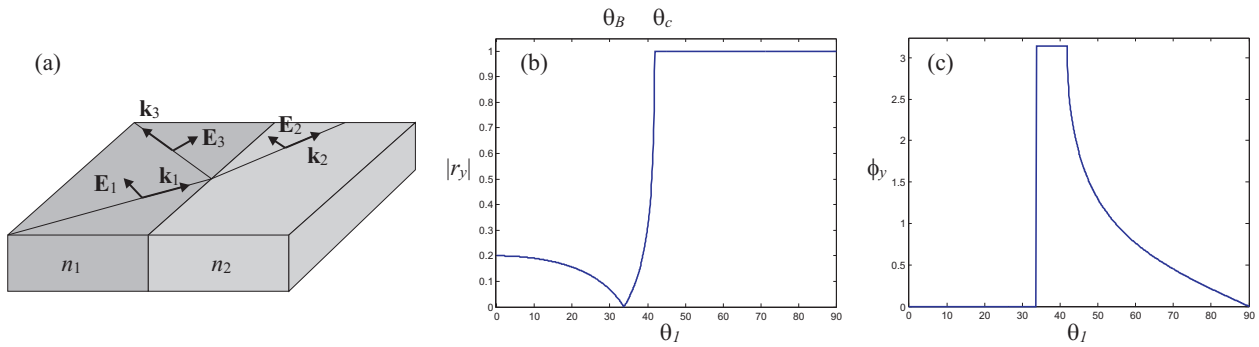


FIGURE 4.7 – Réflexion interne TM : (a) géométrie et (b) amplitude et (c) phase des coefficients de Fresnel.

4.7 Réflectance et transmittance

Alors que r et t représentent les rapports des amplitudes complexes, on définit la réflectance $\mathcal{R}$ et la transmittance $\mathcal{T}$ comme le rapport des puissances optiques transmises à travers l'interface selon une direction normale à l'interface. Pour obtenir ces valeurs, on part de la constatation que les ondes incidentes et réfléchies se propagent dans le même plan et dans le même milieu. Ainsi

$$\mathcal{R} = |r|^2. \quad (4.33)$$

La transmittance est obtenue en invoquant la conservation de la puissance optique. Dans le cas d'un milieu sans perte, on a alors,

$$\mathcal{T} = 1 - \mathcal{R}. \quad (4.34)$$

Notons qu'en général $\mathcal{T} \neq |t|^2$ puisque la puissance optique se propage selon des directions différentes dans les milieux n_1 et n_2 qui possèdent aussi des impédances différentes.

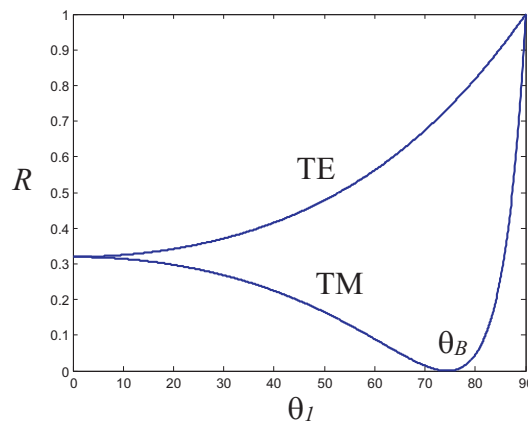


FIGURE 4.8 – Puissance réfléchie pour des ondes TE et TM entre de l'air ($n_1 = 1$) et un semiconducteur ($n_2 = 3.6$).

A incidence normal, la réflectance devient,

$$\mathcal{R} = \left(\frac{n_1 - n_2}{n_1 + n_2} \right)^2. \quad (4.35)$$

Sa valeur peut être beaucoup plus grande pour des incidences rasantes, lorsque $\theta_1 \rightarrow 90^\circ$, Fig. 4.8.

4.8 Polarisation de la lumière

En un point de l'espace, la polarisation de la lumière est déterminée par l'évolution temporelle du vecteur champ électrique. On se réfère spécifiquement à la trajectoire dans l'espace de l'extrémité du vecteur champ électrique en fonction du temps. Pour une onde monochromatique, dans le cas le plus général, cette trajectoire est elliptique et le plan qui la supporte change d'orientation de point en point. Si l'on a affaire à une onde plane monochromatique, les trajectoires elliptiques sont les mêmes dans l'espace, seule la position du vecteur électrique $\mathbf{E}$ en un instant donné est différente, comme indiqué Fig. 4.9.

La polarisation de la lumière joue un rôle essentiel dans l'interaction du champ optique avec la matière. En effet, nous avons déjà vu que les conditions limites imposées par les équations de Maxwell diffèrent suivant que le champ électrique est parallèle ou perpendiculaire à l'interface. Ainsi les coefficients de Fresnel ne sont pas les mêmes dans ces deux cas que nous avons décrits par les termes TM et TE.

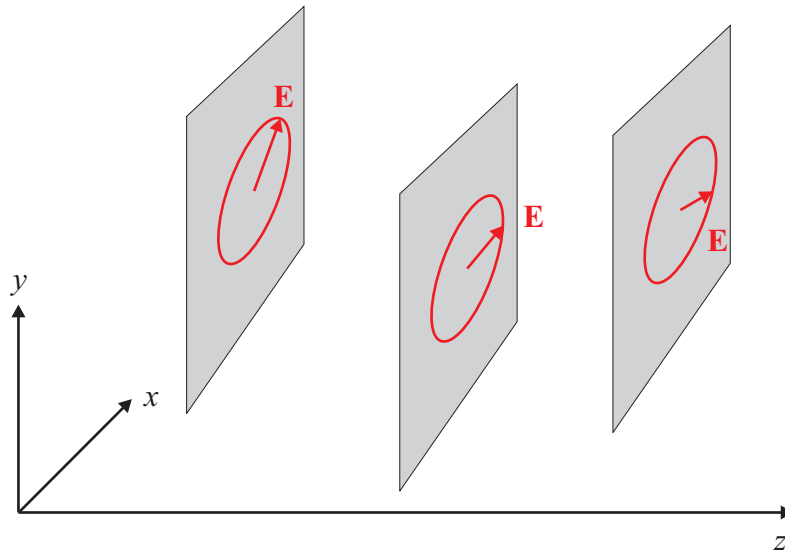


FIGURE 4.9 – Trajectoire du champ électrique à un instant donné en différents points de l'espace pour une onde plane se propageant dans la direction z .

Considérons une onde plane monochromatique se propageant dans la direction z et dont le champ électrique – qui se trouve dans le plan x - y – est donné par

$$\mathbf{E}(z, t) = \text{Re} \left\{ \mathbf{A} e^{j\omega(t - \frac{z}{c})} \right\}, \quad (4.36)$$

avec l'amplitude complexe

$$\mathbf{A} = A_x \mathbf{e}_x + A_y \mathbf{e}_y. \quad (4.37)$$

Pour décrire la polarisation de cette onde il faut étudier l'extrémité du vecteur électrique $\mathbf{E}$:

$$\mathbf{E}(z, t) = E_x \mathbf{e}_x + E_y \mathbf{e}_y. \quad (4.38)$$

En utilisant Eq. (4.36) on peut récrire ce vecteur,

$$E_x = a_x \cos \left[\omega \left(t - \frac{z}{c} \right) + \phi_x \right] \quad (4.39a)$$

$$E_y = a_y \cos \left[\omega \left(t - \frac{z}{c} \right) + \phi_y \right], \quad (4.39b)$$

où nous avons utilisé la représentation en amplitude et phase des composantes du vecteur complexe $\mathbf{A}$:

$$A_x = a_x e^{j\phi_x} \quad (4.40a)$$

$$A_y = a_y e^{j\phi_y}. \quad (4.40b)$$

On remarque que chaque composante du champ électrique donnée par Eq. (4.39) oscille avec une pulsation ω : une composante oscille le long de l'axe x et l'autre oscille le long de l'axe y . En un point z fixe (prenons $z = 0$ pour simplifier), à l'instant $t = 0$, ces deux composantes ont pour valeur $E_x = a_x \cos \phi_x$ et $E_y = a_y \cos \phi_y$. Dans les instants suivants $t > 0$, chaque composante oscille le long de son axe avec une amplitude propre : a_x pour la composante x et a_y pour la composante y . Ainsi la trajectoire décrite par l'extrémité du vecteur champ électrique est une ellipse. D'ailleurs Eq. (4.39) représente l'équation paramétrique d'une ellipse dont l'équation polaire est

$$\frac{E_x^2}{a_x^2} + \frac{E_y^2}{a_y^2} - 2 \cos \phi \frac{E_x E_y}{a_x a_y} = \sin^2 \phi, \quad (4.41)$$

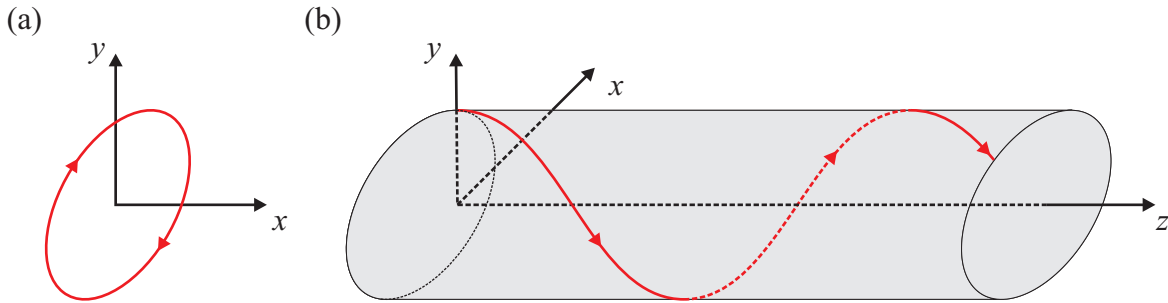


FIGURE 4.10 – (a) Trajectoire du champ électrique en fonction du temps. (b) Trajectoire du champ électrique en fonction de l'espace.

où l'on a introduit la différence de phase $\phi = \phi_y - \phi_x$. Comme les deux composantes oscillent à la même pulsation, on peut bien comprendre que c'est la différence de phase ϕ (i.e. quelle composante du champ "part en avance" par rapport à l'autre) qui va être déterminante pour cette trajectoire.

Si l'on se fixe à un endroit de l'espace (par exemple en $z = 0$), l'évolution temporelle de l'extrémité du vecteur champ électrique décrit dans le plan x - y une ellipse, comme indiqué Fig. 4.10(a). On remarquera que cette ellipse tourne dans le sens des aiguilles d'une montre. On dit que la lumière est polarisée circulairement à droite.

Si maintenant on regarde à un instant donné t fixe l'évolution spatiale de l'extrémité du vecteur champ électrique indiqué Fig. 4.10(b), on remarque qu'il suit aussi une trajectoire elliptique, mais que celle-ci tourne dans le sens opposé. Ceci provient de l'argument du cos dans Eq. (4.39) : $t - z/c$. Ainsi si z est fixe et t augmente, l'argument de cos augmente et la valeur de la fonction augmente ; par contre si t est fixe et z augmente, l'argument de cos diminue et la valeur de la fonction diminue.

Notons aussi que dans Fig. 4.10(a) nous observons la trajectoire du champ électrique depuis une position située en $z > 0$ (effectuer le produit vectoriel $\mathbf{e}_x \times \mathbf{e}_y = \mathbf{e}_z$ pour se convaincre de la direction de l'axe z). En quelque sorte nous regardons l'onde qui s'approche de nous. Si nous avons décidé de regarder la trajectoire depuis un point $z < 0$, i.e. d'observer l'onde qui s'éloigne de nous, nous verrions tourner le champ électrique dans l'autre sens (pour s'en convaincre on peut regarder Fig. 4.10(a) en transparence à travers la page).

La détermination de la polarisation de la lumière et du sens de rotation du champ électrique peut donc donner lieu à beaucoup de confusion. Pour l'éviter nous adoptons la convention suivante pour déterminer la polarisation d'une onde plane :

- On se place en un point fixe sur l'axe de propagation et on regarde l'onde venir vers nous (par exemple l'axe $+z$, Fig. 4.10).
- On suit l'évolution temporelle de l'extrémité du champ électrique.
- Si ce champ tourne dans le sens des aiguilles d'une montre on dit que l'onde est polarisée à droite.
- Si ce champ tourne dans le sens contraire on dit que l'onde est polarisée à gauche.

Comme c'est le cas pour toute convention, il existe évidemment des ouvrages dans lesquels la convention inverse est choisie ! Il faut donc prendre soin de vérifier la convention choisie avant d'utiliser une formule.

Pour être complet, caractérisons encore l'ellipse définie par Eq. (4.41), comme indiqué Fig. 4.11. En introduisant le rapport des demi-axes $r = a_y/a_x$ et l'angle χ qui caractérise l'ellipsité,

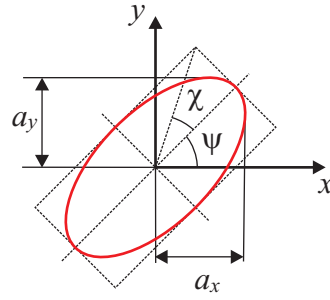


FIGURE 4.11 – Paramètres de l'ellipse de polarisation, selon Eq. (4.42).

on a

$$r = \frac{a_y}{a_x} \quad (4.42a)$$

$$\phi = \phi_y - \phi_x \quad (4.42b)$$

$$\tan 2\psi = \frac{2r}{1 - r^2} \cos \phi \quad (4.42c)$$

$$\tan 2\chi = \frac{2r}{1 + r^2} \sin \phi. \quad (4.42d)$$

Nous allons utiliser ces paramètres pour étudier deux cas particuliers importants : la polarisation linéaire et la polarisation circulaire.

4.8.1 Polarisation linéaire

Si une des composantes s'annule (par exemple $a_x = 0$) alors la lumière est polarisée linéairement dans l'autre direction (y dans ce cas). La lumière est aussi polarisée linéairement si $\phi = 0$ ou π et le lieu de l'extrémité du vecteur $\mathbf{E}$ est alors une ligne de pente a_y/a_x pour $\phi = 0$ et $-a_y/a_x$ pour $\phi = \pi$. L'ellipse qui supportait la trajectoire du champ est alors écrasée, $\chi = 0$. La figure 4.12 illustre cette situation.

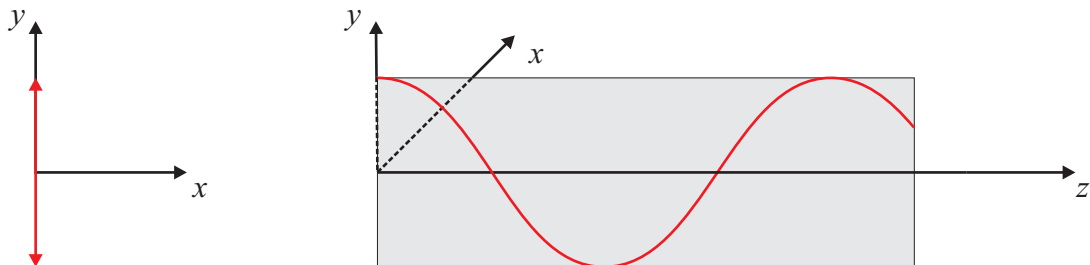


FIGURE 4.12 – Onde plane polarisée linéairement.

4.8.2 Polarisation circulaire

Si $\phi = \pm\pi/2$ et $a_x = a_y = a_0$, Eqs. (4.39) deviennent

$$E_x = a_0 \cos \left[\omega \left(t - \frac{z}{c} \right) + \phi_x \right] \quad (4.43a)$$

$$E_y = \mp a_0 \cos \left[\omega \left(t - \frac{z}{c} \right) + \phi_y \right], \quad (4.43b)$$

dont on déduit $E_x^2 + E_y^2 = a_0^2$ qui est l'équation géométrique d'un cercle. La surface elliptique qui supportait la trajectoire du champ devient un cylindre, comme illustré Fig. 4.13. Pour $\phi = \pi/2$ la lumière est polarisée circulairement à droite. Notons que dans ce cas on a par exemple à l'instant $t = 0$ une phase de $\phi_y = \pi/2$ pour la composante E_y qui vaut donc a_0 et une phase $\phi_x = 0$ pour la composante E_x qui vaut donc 0 (ceci correspond au point $(x, y) = (0, a_0)$ sur Fig. 4.13(a)). Ainsi, pour les instants $t > 0$ la composante E_y se trouve toujours "en avance" d'un quart de période par rapport à la composante E_x et le champ tourne dans le sens des aiguilles d'une montre.

Pour $\phi = -\pi/2$ la rotation s'effectue dans l'autre sens (dans ce cas c'est la composante E_x qui est "en avance" sur la composante E_y) et on parle de polarisation circulaire à gauche.

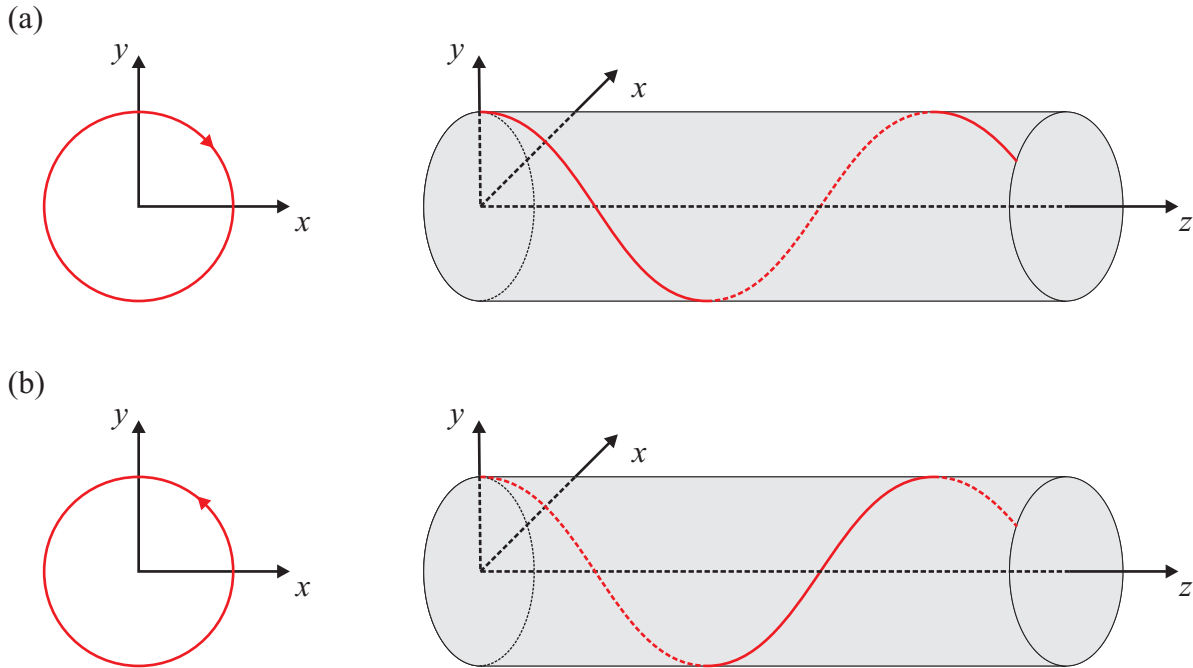


FIGURE 4.13 – Onde plane polarisée circulairement (a) à droite et (b) à gauche.

4.9 Vecteurs et matrices de Jones

La polarisation ne joue un rôle important que lorsque la lumière interagit avec des composants optiques. Il est alors essentiel de déterminer comment un composant change l'état de polarisation de la lumière. Cette situation est semblable à celle rencontrée en optique géométrique lorsque nous avons décrit un rayon (y, θ) avec son ordonnée et son angle puis introduit les matrices $ABCD$ pour caractériser les transformations que subissent un tel rayon au passage d'un élément d'optique. Dans cette section nous allons caractériser l'état de polarisation d'une onde par un vecteur de Jones et l'effet d'un élément optique par la matrice correspondante. Notons qu'il existe d'autres façons de caractériser la polarisation d'une onde, en particulier les paramètres de Stokes qui permettent d'indiquer un état de polarisation sur une sphère dite de Poincaré. Les paramètres de Stokes sont plus généraux que le vecteur de Jones et permettent aussi de caractériser le degré de polarisation d'une lumière qui n'est pas entièrement polarisée.

Les vecteurs de Jones ont l'avantage d'être simples et correspondent bien à la description que nous avons utilisée dans Sec. 4.8. Considérons une onde monochromatique se propageant dans la direction z avec l'enveloppe complexe suivante

$$A_x = a_x e^{j\phi_x} \quad (4.44a)$$

$$A_y = a_y e^{j\phi_y}, \quad (4.44b)$$

pour le champ électrique. Ces deux composantes complexes sont rassemblées dans le vecteur de Jones associé à cette onde :

$$\mathbf{J} = \begin{pmatrix} A_x \\ A_y \end{pmatrix}. \quad (4.45)$$

Noter que l'on peut reconstituer l'enveloppe du champ électrique à partir des valeurs complexes $(a_x, a_y, \phi_x$ et $\phi_y)$ contenues dans le vecteur de Jones. Les vecteurs de Jones correspondant à quelques cas particuliers illustrés dans Fig. 4.14, et pour lesquels on a choisi

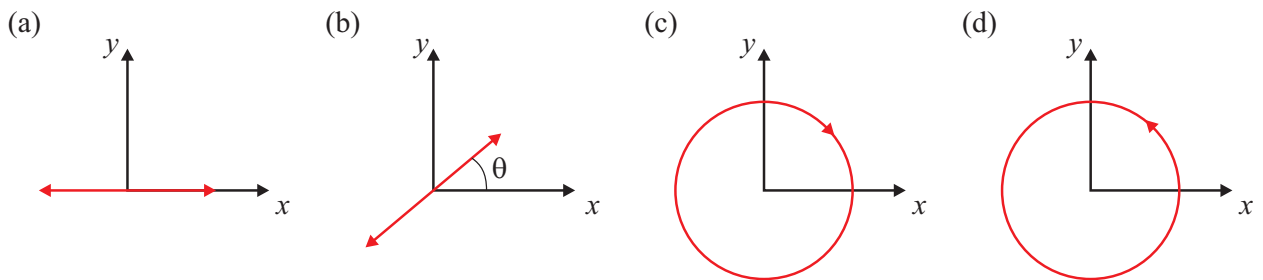


FIGURE 4.14 – Quelques états de polarisation particuliers correspondant aux vecteurs de Jones indiqués dans Eq. (4.46).

$|A_x|^2 + |A_y|^2 = 1$ et $\phi_x = 0$, sont les suivants :

$$\mathbf{J} = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \text{ linéaire direction } x \quad (4.46a)$$

$$\mathbf{J} = \begin{pmatrix} \cos \theta \\ \sin \theta \end{pmatrix} \text{ linéaire angle } \theta \quad (4.46b)$$

$$\mathbf{J} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ j \end{pmatrix} \text{ circulaire droite} \quad (4.46c)$$

$$\mathbf{J} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -j \end{pmatrix} \text{ circulaire gauche.} \quad (4.46d)$$

On dit que deux ondes ont des états de polarisation orthogonaux lorsque le produit intérieur (le produit scalaire en prenant le complexe conjugué du second élément) de leurs vecteurs de Jones vaut zéro :

$$(\mathbf{J}_1, \mathbf{J}_2) = A_{1x}A_{2x}^* + A_{1y}A_{2y}^* = 0. \quad (4.47)$$

Notons que deux vecteurs de Jones orthogonaux $\mathbf{J}_1$ et $\mathbf{J}_2$ peuvent servir de base pour développer le vecteur de Jones $\mathbf{J}$ de n'importe quelle onde : $\mathbf{J} = \alpha_1 \mathbf{J}_1 + \alpha_2 \mathbf{J}_2$. Les coefficients α_1 et α_2 s'obtiennent en faisant les produits intérieurs avec les vecteurs de base correspondants : $\alpha_1 = (\mathbf{J}, \mathbf{J}_1)$ et $\alpha_2 = (\mathbf{J}, \mathbf{J}_2)$ pour autant que $\mathbf{J}_1$ et $\mathbf{J}_2$ soient normalisés : $(\mathbf{J}_1, \mathbf{J}_1) = (\mathbf{J}_2, \mathbf{J}_2) = 1$.

Lorsqu'une onde plane dont la polarisation est décrite par le vecteur de Jones $\mathbf{J}_1$ (input) passe à travers un système optique linéaire sa polarisation change et devient $\mathbf{J}_2$ (output), Fig. 4.15. La relation linéaire entre $\mathbf{J}_1$ et $\mathbf{J}_2$ peut s'écrire

$$\mathbf{J}_2 = \mathbf{T} \cdot \mathbf{J}_1, \quad (4.48)$$

où l'on a introduit la matrice de Jones $\mathbf{T}$:

$$\mathbf{T} = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix}. \quad (4.49)$$

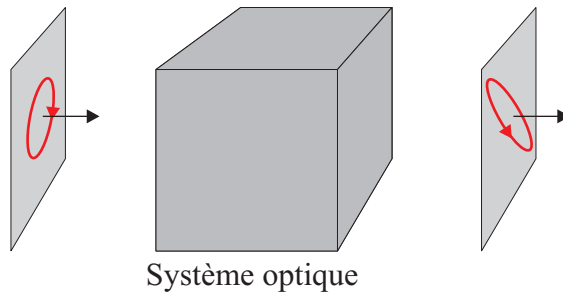


FIGURE 4.15 – Changement de la polarisation d'une onde plane à travers un système optique.

La structure de cette matrice détermine l'effet de l'élément optique sur la polarisation de l'onde incidente. Arrêtons-nous aux plus importants : les polariseurs, les retardeurs (ou retardateurs), et les rotateurs.

Un polariseur linéaire dans la direction x a comme matrice de Jones,

$$\mathbf{T} = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}. \quad (4.50)$$

On remarque qu'il supprime la composante y de l'onde incidente, ne laissant que la composante x .

Les retardeurs induisent – comme leur nom l'indique – un retard d'une composante du champ électrique par rapport à l'autre. La matrice de Jones correspondante s'exprime par

$$\mathbf{T} = \begin{pmatrix} 1 & 0 \\ 0 & e^{-j\Gamma} \end{pmatrix}. \quad (4.51)$$

Ainsi une onde plane incidente (A_{1x}, A_{1y}) est-elle transformée en $(A_{1x}, A_{1y}e^{-j\Gamma})$. La composante y subit donc un retard correspondant à une phase Γ . On dit que l'axe x du système est l'axe rapide (F de *Fast* en anglais) et l'axe y est l'axe lent (S de *Slow* en anglais).

Un déphasage quelconque Γ n'est pas nécessairement utile pour manipuler la polarisation de la lumière. On distingue cependant la lame quart d'onde ($\Gamma = \pi/2$, on se souviendra

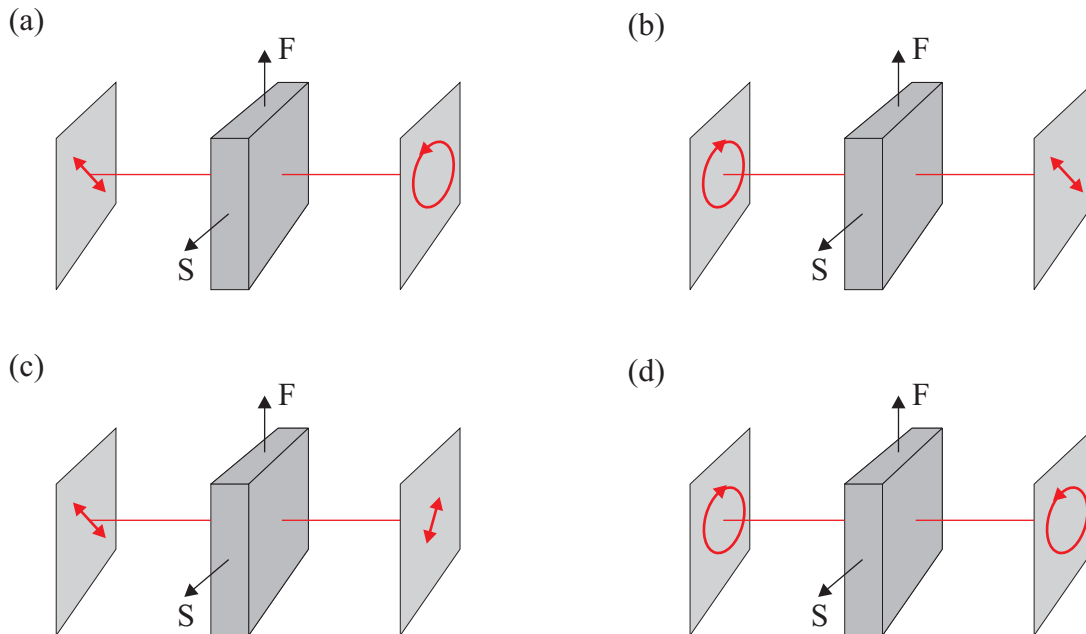


FIGURE 4.16 – Différent types de retardeurs : (a), (b) lame quart d'onde, (c), (d) lame demi-onde et leur effet sur des ondes polarisées linéairement (a), (c) et circulairement (b), (d).

qu'une longueur d'onde complète correspond à un déphasage de 2π) qui transforme une onde polarisée linéairement en une onde polarisée circulairement à gauche, Fig. 4.16(a) ; ou transforme une onde polarisée circulairement à droite en une onde polarisée linéairement, Fig. 4.16(b). On remarque que dans ce cas le déphasage d'une composante par rapport à l'autre correspond à un quart de la longueur d'onde.

La lame demi-onde ($\Gamma = \pi$) transforme une onde polarisée linéairement en une onde polarisée linéairement mais dont la polarisation est perpendiculaire à la première, Fig. 4.16(c) ; ou transforme une onde polarisée circulairement à droite en une onde polarisée circulairement à gauche, Fig. 4.16(d).

Alors que les lames quart- et demi-onde transforment l'état de polarisation, un rotateur de polarisation maintient la polarisation linéaire en la faisant simplement tourner d'un angle θ . La matrice de Jones correspondante s'écrit

$$\mathbf{T} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}. \quad (4.52)$$

Un tel rotateur transforme une onde d'état de polarisation $\mathbf{J}_1 = (\cos \theta_1, \sin \theta_1)^\top$ en une onde de polarisation $\mathbf{J}_2 = (\cos \theta_2, \sin \theta_2)^\top$, avec $\theta_2 = \theta_1 + \theta$.

Les matrices de Jones sont particulièrement utiles lorsque l'on combine plusieurs éléments optiques. Ainsi le vecteur de Jones $\mathbf{J}_{\text{out}}$ d'une onde incidente $\mathbf{J}_{\text{in}}$ après qu'elle a passé successivement à travers les éléments optiques décrits par les matrices de Jones $\mathbf{T}_1$, $\mathbf{T}_2$ et $\mathbf{T}_3$ s'obtient simplement par

$$\mathbf{J}_{\text{out}} = \mathbf{T}_3 \cdot \mathbf{T}_2 \cdot \mathbf{T}_1 \cdot \mathbf{J}_{\text{in}}. \quad (4.53)$$

On fait évidemment les multiplications de droite à gauche pour respecter l'ordre dans lequel l'onde rencontre les différents éléments optiques, comme ce fut le cas pour les matrices ABCD dans Chap. 1.

Pour terminer, relevons que les phénomènes de réflexion et réfraction étudiés Sec. 4.6 peuvent facilement s'exprimer dans ce formalisme. D'ailleurs, les axes x et y introduits dans Fig. 4.3 qui ont permis la définition des polarisations TE et TM ne sont qu'un exemple d'états de polarisation orthogonaux. Ainsi les vecteurs de Jones des ondes incidentes, transmises et réfléchies s'écrivent,

$$\mathbf{J}_1 = \begin{pmatrix} A_{1x} \\ A_{1y} \end{pmatrix}, \quad \mathbf{J}_2 = \begin{pmatrix} A_{2x} \\ A_{2y} \end{pmatrix}, \quad \text{et} \quad \mathbf{J}_3 = \begin{pmatrix} A_{3x} \\ A_{3y} \end{pmatrix}. \quad (4.54)$$

Quant aux matrices de Jones pour la transmission et la réflexion elles s'écrivent

$$\mathbf{T}_t = \begin{pmatrix} t_x & 0 \\ 0 & t_y \end{pmatrix} \quad \text{et} \quad \mathbf{T}_r = \begin{pmatrix} r_x & 0 \\ 0 & r_y \end{pmatrix}, \quad (4.55)$$

où les coefficients $t_{x,y}$ et $r_{x,y}$ sont donnés par Eq. (4.30).

4.10 Matériaux anisotropes

On dit qu'un matériau diélectrique est anisotrope si ses propriétés optiques dépendent de la direction, par exemple de la direction dans laquelle l'onde se propage. Pour cerner l'origine de cette anisotropie, il est important de déterminer l'échelle à laquelle on observe le matériau : pour l'optique c'est celle de la longueur d'onde. Ainsi, si l'on considère un gaz avec des atomes ou des molécules réparties de façon aléatoire, le matériau peut apparaître localement très anisotrope. A l'échelle de la longueur d'onde cependant ce matériau est isotrope : quelle que soit la direction dans laquelle une onde s'y propage, elle verra le même milieu, défini par une permittivité relative ϵ_r scalaire.

Si le matériau est formé de grains cristallins disjoints, les propriétés du matériau dépendent localement de l'orientation. A l'échelle de la lumière cependant, si les grains sont orientés de façon aléatoire, leur réponse se moyenne et le matériau apparaît aussi isotrope.

Lorsque les atomes ou molécules sont répartis de façon régulière en un cristal, les propriétés optiques dépendent de la direction de propagation et le matériau est anisotrope. Notons qu'un cristal n'a pas besoin d'être solide, mais on rencontre aussi des liquides dans lesquels un tel ordre existe : les cristaux liquides ou les assemblages supramoléculaires par exemple.

Comme nous l'avons vu, le déplacement électrique $\mathbf{D}$ décrit la réponse de la matière à un champ électrique $\mathbf{E}$ incident, selon Eq. (4.4) : $\mathbf{D} = \epsilon \mathbf{E}$. C'est donc la relation entre ces deux champs qui va caractériser l'anisotropie du milieu. Nous obtenons ainsi un résultat essentiel pour les milieux anisotropes : c'est la direction du champ électrique $\mathbf{E}$ dans le milieu et non la direction de propagation de l'onde qui déterminent la réponse du milieu. Ces deux directions sont évidemment liées, en particulier pour une onde plane, comme indiqué Fig. 4.1. Cependant, pour une direction de propagation fixe (par exemple la direction z pour l'onde plane polarisée linéairement sur Fig. 4.12), une onde plane polarisée linéairement peut avoir son plan de polarisation (direction du champ électrique) dans n'importe quelle direction dans le plan transverse à la direction de propagation (plan y - z sur Fig. 4.12).

Pour un matériau anisotrope la permittivité prend la forme d'un tenseur de rang deux (i.e. une matrice 3×3). Ainsi la relation $\mathbf{D} = \epsilon \mathbf{E}$ devient-elle

$$\mathbf{D} = \boldsymbol{\epsilon} \cdot \mathbf{E}, \quad (4.56)$$

ou exprimée en composantes

$$\begin{pmatrix} D_x \\ D_y \\ D_z \end{pmatrix} = \begin{pmatrix} \epsilon_{xx} & \epsilon_{xy} & \epsilon_{xz} \\ \epsilon_{yx} & \epsilon_{yy} & \epsilon_{yz} \\ \epsilon_{zx} & \epsilon_{zy} & \epsilon_{zz} \end{pmatrix} \cdot \begin{pmatrix} E_x \\ E_y \\ E_z \end{pmatrix}. \quad (4.57)$$

Les éléments du tenseur de permittivité $\boldsymbol{\epsilon}$ dépendent du système de coordonnées choisi, en particulier de sa relation avec les axes du cristal. Ainsi, si on tourne le cristal par rapport au système de coordonnées la valeur des éléments $\epsilon_{\alpha\beta}$ change. Il existe cependant un système

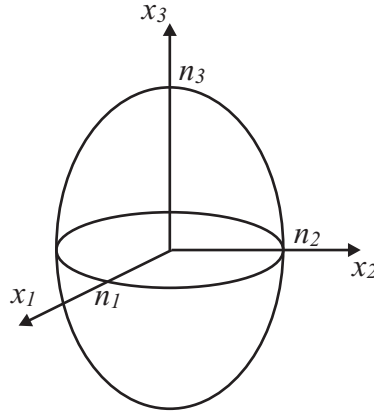


FIGURE 4.17 – Ellipsoïde d'indice pour un matériau anisotrope.

de coordonnées particulier dans lequel le tenseur est diagonal :

$$\begin{pmatrix} D_1 \\ D_2 \\ D_3 \end{pmatrix} = \begin{pmatrix} \epsilon_1 & 0 & 0 \\ 0 & \epsilon_2 & 0 \\ 0 & 0 & \epsilon_3 \end{pmatrix} \cdot \begin{pmatrix} E_1 \\ E_2 \\ E_3 \end{pmatrix}. \quad (4.58)$$

On dit que ce système correspond aux axes principaux du cristal. Nous allons nous limiter à ce système de coordonnées dans le reste de cette section.

A partir des permittivités ϵ_1 , ϵ_2 et ϵ_3 dans Eq. (4.58) on définit les indices de réfraction principaux du cristal, cf. Eq. (4.19) :

$$n_1 = \sqrt{\epsilon_1/\epsilon_0} \quad n_2 = \sqrt{\epsilon_2/\epsilon_0} \quad n_3 = \sqrt{\epsilon_3/\epsilon_0} \quad . \quad (4.59)$$

On peut obtenir une représentation graphique de ces indices en créant une surface définie par la relation

$$\frac{x_1^2}{n_1^2} + \frac{x_2^2}{n_2^2} + \frac{x_3^2}{n_3^2} = 1, \quad (4.60)$$

comme illustré Fig. 4.17. Cet ellipsoïde a pour demi-axes les indices de réfraction principaux n_1 , n_2 et n_3 .

Il existe trois types de cristaux suivant les valeurs de n_1 , n_2 et n_3 : les cristaux biaxiaux pour lesquels ces trois indices sont différents ; les cristaux uniaxiaux pour lesquels deux indices ont la même valeur ($n_1 = n_2 \neq n_3$) et les cristaux isotropes pour lesquels tous les indices sont les mêmes ($n_1 = n_2 = n_3$). Dans ce dernier cas l'ellipsoïde d'indices devient une sphère.

4.11 Propagation le long d'un axe principal

Pour un cristal biaxial décrit par l'ellipsoïde d'indice Fig. 4.17 il existe des directions privilégiées associées à la propagation d'une onde plane avec un vecteur de propagation spécifique,

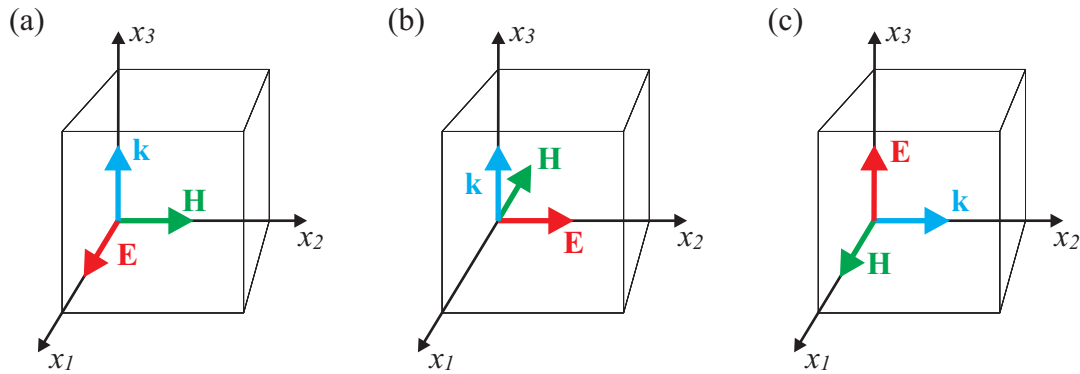


FIGURE 4.18 – Propagation d’ondes planes selon les axes principaux d’un cristal anisotrope. Les vitesses de phase et les nombres d’onde sont (a) $c = c_0/n_1$, $k = n_1k_0$; (b) $c = c_0/n_2$, $k = n_2k_0$ et (c) $c = c_0/n_3$, $k = n_3k_0$.

comme indiqué Fig. 4.18. Ainsi, si le champ électrique est parallèle à l’axe x_1 et l’onde se propage dans la direction x_3 , on a $D_1 = \epsilon_1 E_1$ et l’équation d’onde déduite des équations de Maxwell a pour solution une onde se propageant à la vitesse c_0/n_1 , Fig. 4.18(a). Pour une onde se propageant aussi dans la direction x_3 mais polarisée dans la direction x_2 , Fig. 4.18(b), la vitesse de l’onde est c_0/n_2 . On remarque donc que ce n’est pas la direction de propagation qui est déterminante, mais bien la direction de polarisation, i.e. la direction dans laquelle oscille le champ électrique. Finalement, pour une onde se propageant dans la direction x_2 avec son champ électrique parallèle à l’axe x_3 , la vitesse est c_0/n_3 . Le nombre d’onde est aussi différent pour chacun de ces cas : $k = n_1k_0$, $k = n_2k_0$ et $k = n_3k_0$.

On dit que les trois ondes décrites Fig. 4.18 représentent les modes principaux du cristal. Pour étudier la propagation d’une onde ayant une polarisation arbitraire, on doit la décomposer selon ces modes principaux et propager chaque mode indépendamment. Comme nous venons de le voir, chaque mode se propage alors avec une vitesse différente, donnant lieu à un déphasage entre les différents modes.

4.12 Cristal uniaxial et biréfringence

Les cristaux uniaxiaux sont très importants dans la pratique et il convient d’introduire un peu de nomenclature. On appelle indice ordinaire n_o les deux indices qui sont égaux : $n_o = n_1 = n_2$ et indice extraordinaire l’indice qui est différent des trois autres : $n_e = n_3$. L’axe x_3 qui correspond à l’indice extraordinaire n_e se nomme l’axe optique du cristal, Fig. 4.19(a). Si $n_e > n_o$ on dit que le cristal est uniaxial positif ; si $n_e < n_o$ il est uniaxial négatif.

Rappelons que lorsque le champ électrique se trouve dans le plan x_1-x_2 , l’onde associée se propage à la vitesse c_0/n_o , quelle que soit sa direction de propagation. On parle d’onde ordinaire, Fig. 4.19(b). Si par contre une onde a une composante du champ parallèle au

plan x_1 - x_2 et une composante parallèle à l'axe optique, on parle d'onde extraordinaire, Fig. 4.19(c).

En étudiant Figs. 4.19(a) et (c) on remarque qu'une onde extraordinaire se propageant parallèlement à l'axe optique ($\theta = 0$) voit un indice de réfraction n_o . A l'inverse, si elle se propage dans une direction perpendiculaire à l'axe optique avec le champ électrique parallèle à cet axe ($\theta = 90^\circ$), elle voit un indice n_e . Entre ces deux cas extrêmes, il faut décomposer le champ électrique en une composante "ordinaire" et une composante "extraordinaire". L'onde correspondante se propage avec un indice de réfraction $n(\theta)$ se trouvant sur l'ellipse indiquée Fig. 4.19(d) et dont l'équation s'obtient à partir de l'ellipsoïde d'indice (4.60) :

$$\frac{1}{n^2(\theta)} = \frac{\cos^2 \theta}{n_o^2} + \frac{\sin^2 \theta}{n_e^2}. \quad (4.61)$$

Pour un cristal uniaxial, Eq. (4.58) prend la forme

$$\begin{pmatrix} D_1 \\ D_2 \\ D_3 \end{pmatrix} = \begin{pmatrix} \epsilon_o & 0 & 0 \\ 0 & \epsilon_o & 0 \\ 0 & 0 & \epsilon_e \end{pmatrix} \cdot \begin{pmatrix} E_1 \\ E_2 \\ E_3 \end{pmatrix}. \quad (4.62)$$

Ainsi le vecteur $\mathbf{D}_o = (D_1, D_2, 0) = \epsilon_o(E_1, E_2, 0)$ associé à une onde ordinaire est-il parallèle au champ électrique, Fig. 4.19(b). Par contre, pour l'onde extraordinaire le vecteur déplacement électrique $\mathbf{D} = (D_1, D_2, D_3) = (\epsilon_o E_1, \epsilon_o E_2, \epsilon_e E_3)$ n'est généralement pas parallèle au champ électrique, Fig. 4.19(c).

D'une façon générale, une onde est caractérisée par son vecteur de propagation $\mathbf{k}$, les champs $\mathbf{E}$, $\mathbf{D}$, $\mathbf{H}$ et $\mathbf{B}$ ainsi que le vecteur de Poynting moyenné dans le temps $\mathbf{S} = \frac{1}{2} \mathbf{E} \times \mathbf{H}^*$. Ces vecteurs sont reliés par Eqs. (4.22) que nous pouvons récrire en omettant la dépendance

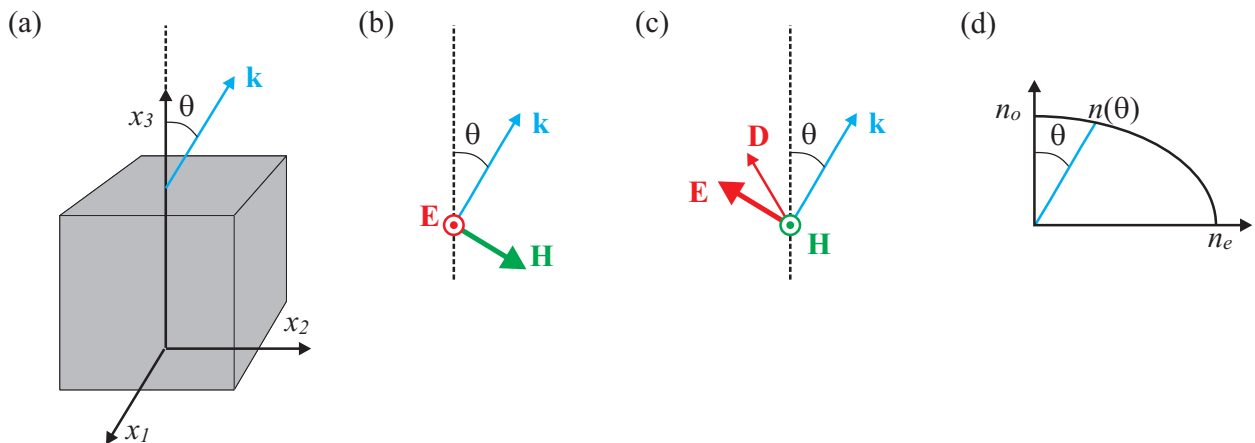


FIGURE 4.19 – (a) Cristal uniaxial et son axe optique, (b) onde ordinaire, (c) onde extraordinaire et (d) variation de l'indice de réfraction en fonction de l'angle de propagation θ pour l'onde extraordinaire.

spatiale pour alléger l'écriture,

$$\mathbf{k} \times \mathbf{E} = \omega \mu_0 \mathbf{H}, \quad (4.63a)$$

$$\mathbf{k} \times \mathbf{H} = -\omega \mathbf{D}. \quad (4.63b)$$

Equation (4.63b) indique que $\mathbf{D}$ est perpendiculaire à $\mathbf{k}$ et $\mathbf{H}$; Eq. (4.63a) indique que $\mathbf{H}$ est perpendiculaire à $\mathbf{k}$ et $\mathbf{E}$. Ces relations géométriques sont illustrées Fig. 4.20, qui met en évidence le fait que $\mathbf{S}$ est aussi perpendiculaire à $\mathbf{E}$ et $\mathbf{H}$. Ainsi $\mathbf{D}$, $\mathbf{E}$, $\mathbf{k}$ et $\mathbf{S}$ sont dans un plan; $\mathbf{H}$ et $\mathbf{B}$ sont perpendiculaires à ce plan. De plus $\mathbf{D}$ est perpendiculaire à $\mathbf{k}$, et $\mathbf{S}$ est perpendiculaire à $\mathbf{E}$, mais $\mathbf{D}$ n'est pas nécessairement parallèle à $\mathbf{E}$ et de même $\mathbf{S}$ n'est pas nécessairement parallèle à $\mathbf{k}$, Fig. 4.20. Ainsi dans un cristal anisotrope, les fronts d'onde (dans la direction de $\mathbf{k}$) ne sont pas nécessairement dans la même direction que le flux d'énergie qui lui est parallèle à $\mathbf{S}$.

En combinant Eqs. (4.63a) et (4.63b) et en utilisant Eq. (4.56), on obtient

$$\mathbf{k} \times (\mathbf{k} \times \mathbf{E}) + \omega^2 \mu_0 \epsilon \cdot \mathbf{E} = 0. \quad (4.64)$$

Ainsi le champ électrique doit satisfaire Eq. (4.64), qui permet d'écrire trois équations, une pour chaque composante E_1 , E_2 et E_3 du champ électrique le long des trois axes du cristal. Les solutions de ce système d'équations déterminent les valeurs des composantes k_1 , k_2 et k_3 du vecteur de propagation $\mathbf{k}$ selon les trois axes du cristal. Ces solutions d'Eq. (4.64) définissent une surface, que l'on appelle la surface k du cristal, qui donne une relation entre la pulsation ω de l'onde et les composantes du vecteur $\mathbf{k}$: $\omega = \omega(k_1, k_2, k_3)$.

Pour un cristal uniaxial, la surface k prend une forme simple et Eq. (4.64) devient

$$(k^2 - n_o^2 k_0^2) \left(\frac{k_1^2 + k_2^2}{n_e^2} + \frac{k_3^2}{n_o^2} - k_0^2 \right) = 0. \quad (4.65)$$

Equation (4.65) possède deux solutions pour le vecteur $\mathbf{k}$ qui correspondent à l'annulation de chacun des termes du produit : la première solution est une sphère définie par

$$k = n_o k_0, \quad (4.66)$$

la seconde solution est un ellipsoïde de révolution défini par

$$\frac{k_1^2 + k_2^2}{n_e^2} + \frac{k_3^2}{n_o^2} = k_0^2. \quad (4.67)$$

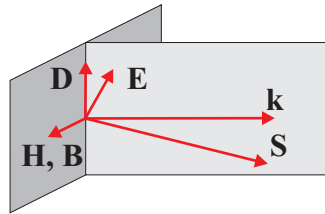


FIGURE 4.20 – Les vecteurs $\mathbf{D}$, $\mathbf{E}$, $\mathbf{k}$ et $\mathbf{S}$ sont dans un plan; $\mathbf{H}$ et $\mathbf{B}$ sont normaux à ce plan. $\mathbf{S}$ et $\mathbf{k}$ ne sont pas nécessairement parallèles dans un matériau anisotrope.

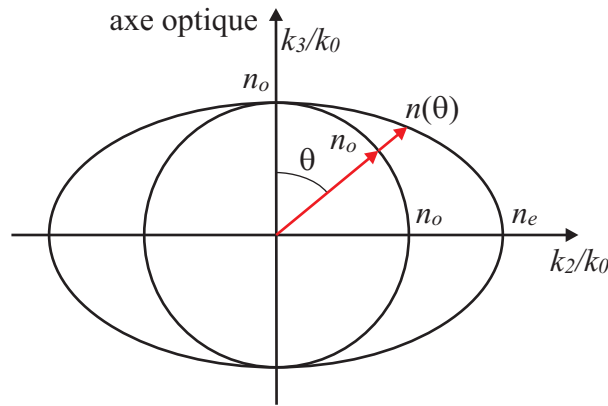


FIGURE 4.21 – Surface k pour un cristal uniaxial positif ($n_e > n_o$) dessiné dans le plan $y-z$. Pour une direction de propagation θ donnée, il existe deux ondes possibles, suivant la polarisation du champ électrique. L'onde ordinaire qui se propage avec un indice de réfraction n_o et l'onde extraordinaire qui se propage avec un indice de réfraction $n(\theta)$.

Comme un cristal uniaxial est symétrique autour de l'axe optique z , on peut se limiter sans perte de généralité au cas où le vecteur $\mathbf{k}$ se trouve dans le plan $y-z$. Sa direction est alors caractérisée par l'angle θ qu'il fait avec l'axe optique. On peut alors dessiner les surfaces k seulement dans le plan $y-z$ où elles se réduisent à un cercle et une ellipse, comme indiqué Fig. 4.21.

Les deux modes normaux pour un cristal uniaxial positif se propageant dans la direction θ par rapport à l'axe optique s'obtiennent par l'intersection avec les surfaces k , Fig. 4.21. Ces modes sont illustrés Fig. 4.22 : pour le mode ordinaire, les fronts d'onde et le vecteur de Poynting sont dans la même direction, Fig. 4.22(a), tel n'est pas le cas pour l'onde extraordinaire, Fig. 4.22(b). La direction des fronts d'onde pour cette dernière est déterminée par le vecteur

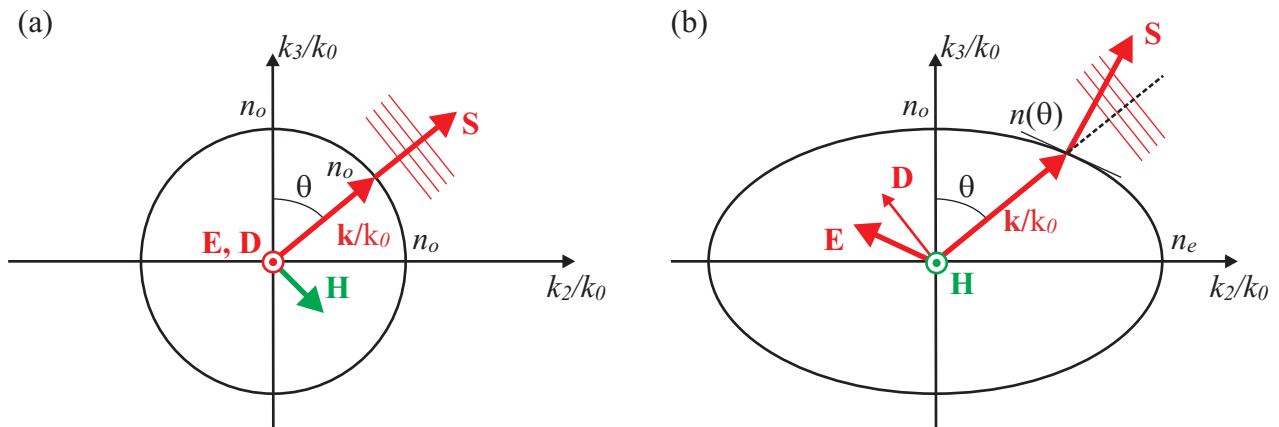


FIGURE 4.22 – (a) Onde ordinaire et (b) onde extraordinaire se propageant dans un cristal uniaxial. Les ondes se propagent avec un angle θ par rapport à l'axe optique.

$\mathbf{k}$, alors que la direction du vecteur de Poynting est déterminée par la tangente à l'ellipse.

Considérons maintenant une onde incidente sur un cristal uniaxial, Fig. 4.23. Comme nous l'avons vu, les fronts d'onde doivent correspondre à l'interface entre le milieu incident – que nous supposons être de l'air – et le cristal. Comme deux modes peuvent se propager dans ce dernier, l'onde incidente créera deux ondes dans le cristal, qui se propageront dans des directions différentes et avec des indices différents. Nous considérons que l'axe optique du cristal fait un angle θ_a avec la normale à l'interface entre les deux milieux. La construction des ondes transmises dans le cristal se fait en généralisant la méthode illustrée Fig. 1.3 : pour l'onde incidente, on trace le cercle de rayon k_0 (on suppose être dans l'air avec un indice $n = 1$) ; dans le cristal, on trace les deux surfaces k illustrées Fig. 4.21 en prenant soin d'orienter l'axe du cristal en conséquence. On utilise ensuite simplement la conservation de la composante du vecteur $\mathbf{k}$ parallèle à l'interface pour déterminer les directions des ondes ordinaires et extraordinaires. Ainsi l'onde ordinaire est réfractée dans le cristal à l'angle θ_o tel que,

$$\sin \theta_1 = n_o \sin \theta_o. \quad (4.68)$$

L'onde extraordinaire est réfractée à l'angle θ_e défini par,

$$\sin \theta_1 = n(\theta_a + \theta_e) \sin \theta_e, \quad (4.69)$$

où la valeur de l'indice $n(\theta_a + \theta_e)$ est donnée par Eq. (4.61).

Une fois de plus, il faut veiller au fait que c'est la polarisation de l'onde incidente qui détermine l'excitation des modes ordinaires et extraordinaires dans le cristal : l'onde ordinaire a une polarisation TE et l'onde extraordinaire a une polarisation TM. Lorsque ces deux ondes émergent du cristal, elles se propagent parallèlement, comme illustré Fig. 4.24. Elles sont cependant spatialement décalées et ont des polarisations différentes.

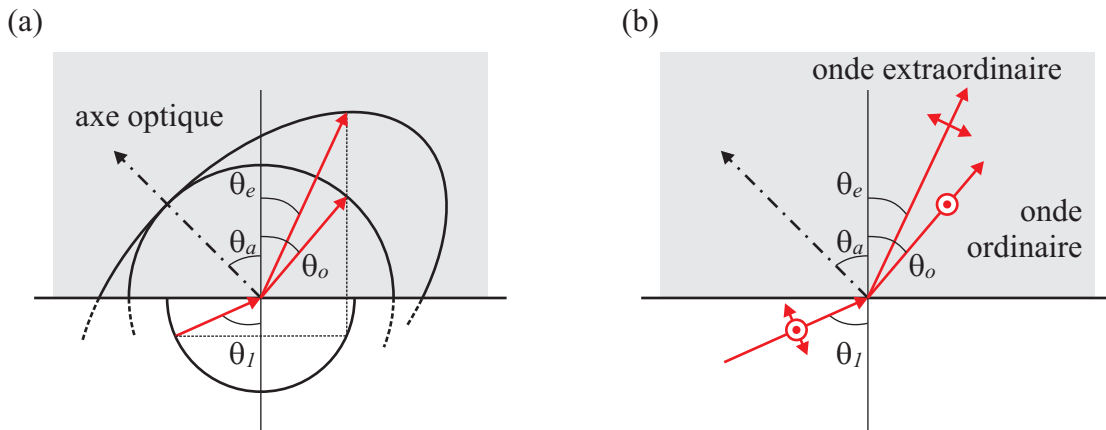


FIGURE 4.23 – (a) Construction de la réfraction à l'interface entre l'air et un cristal uniaxial. La conservation de la composante parallèle à l'interface du vecteur $\mathbf{k}$ permet de déterminer les directions des ondes excitées dans le cristal. (b) Les ondes ordinaires et extraordinaires ont des polarisations différentes.

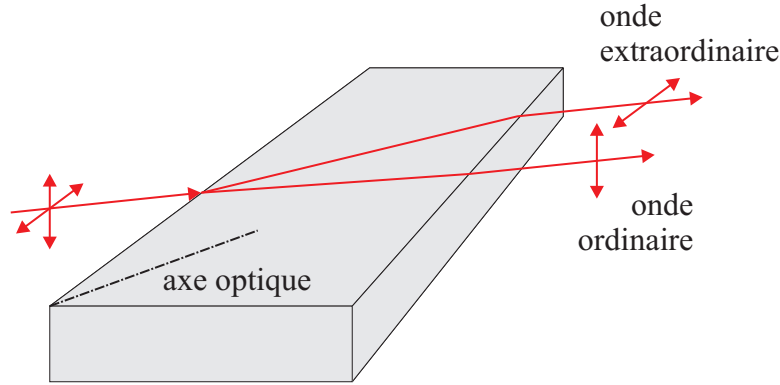


FIGURE 4.24 – Double réfraction à travers un cristal uniaxial. Les ondes émergentes sont décalées et ont des polarisations différentes.

4.13 Activité optique et effet Faraday

Certains matériaux ont la propriété de faire tourner le plan de polarisation, on dit qu'ils sont optiquement actifs. On introduit les indices de réfraction n_+ et n_- pour la propagation d'une onde circulairement polarisée à droite et à gauche. Ainsi la vitesse de phase d'une onde circulairement polarisée à droite est c_0/n_+ et celle d'une onde polarisée circulairement à gauche c_0/n_- . Une onde plane incidente verra son plan de polarisation tourner d'un angle

$$\phi/2 = \pi(n_- - n_+) \frac{d}{\lambda_0}, \quad (4.70)$$

après propagation à travers une épaisseur d de ce matériau, Fig. 4.25. On introduit aussi la puissance de rotation,

$$\rho = \frac{\pi}{\lambda_0}(n_- - n_+). \quad (4.71)$$

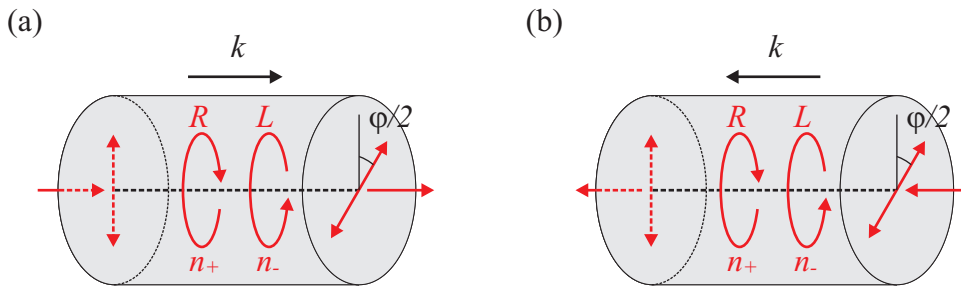


FIGURE 4.25 – (a) La rotation du plan de polarisation par un matériau optiquement actif résulte de la différence de vitesse de propagation pour les deux rotations à droite (R) et à gauche (L). Dans ce cas R est plus rapide et l'onde ressort avec une rotation $\phi/2$ à droite. (b) Si cette onde est réfléchi et revient vers le système, elle en ressort de l'autre côté avec la polarisation originale.

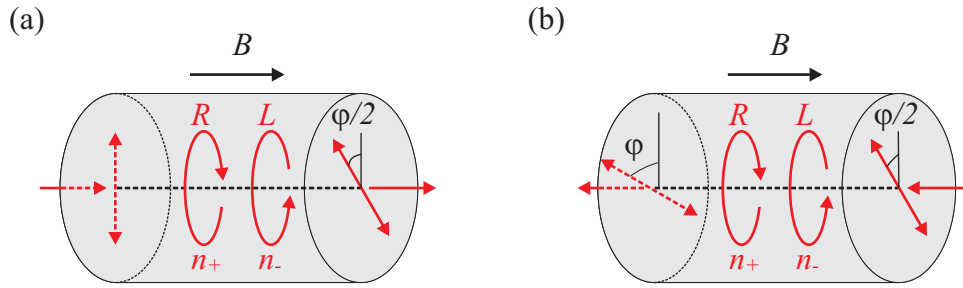


FIGURE 4.26 – (a) Polarisation exercée par l'effet Faraday. (b) Le sens de rotation est déterminé par la direction du champ magnétique (règle de la vis) et est invariant par rapport à la direction de propagation de l'onde.

Certains matériaux exhibent une activité optique lorsqu'ils sont soumis à un champ d'induction magnétique statique $\mathbf{B}$. On parle alors d'effet Faraday ou d'effet magnéto-optique et la puissance de rotation est alors proportionnelle au champ appliqué :

$$\rho = \mathcal{V}B, \quad (4.72)$$

où l'on a introduit la constante de Verdet $\mathcal{V}$.

Remarquons que dans ce cas la direction de rotation du plan de polarisation est donnée par le sens du champ magnétique en utilisant la règle du tire-bouchon ou de la vis : pour enfoncer une vis dans le sens de $\mathbf{B}$ il faut la tourner dans un sens qui détermine la direction de rotation du plan de polarisation, Fig. 4.26(a). Contrairement à l'exemple de Fig. 4.25, la direction de rotation est maintenant la même si l'onde se propage dans l'autre sens, Fig. 4.26(b). Ce principe est utilisé pour créer des isolateurs optiques.

Chapitre 5

Guides d'ondes

Jusqu'à présent nous avons considéré la propagation dans l'espace libre. Dans ce chapitre nous nous concentrons sur des composants essentiels de l'optique moderne : les guides d'ondes. Comme leur nom l'indique, ceux-ci permettent de guider la lumière afin de transmettre l'information d'un point à un autre sur de grandes distances. Les guides d'ondes sont aussi les éléments clefs pour quantité de dispositifs permettant de traiter l'information sous forme optique.

Le principe fondamental d'un guide d'onde est illustré Fig. 5.1 : un rayon lumineux se propage de proche en proche, par réflexions successives sur des miroirs parfaits. Remarquons que la trajectoire en zig-zag de la lumière est à l'origine du modèle important que nous allons utiliser pour décrire la propagation dans les guides d'ondes. L'arrangement de Fig. 5.1 est cependant loin d'être parfait : il existe des angles d'incidence θ pour lesquels le rayon n'est pas réfléchi par le premier miroir et donc échappe au guide. On dit que les rayons qui se propagent dans le guide sont les modes du guide d'onde.

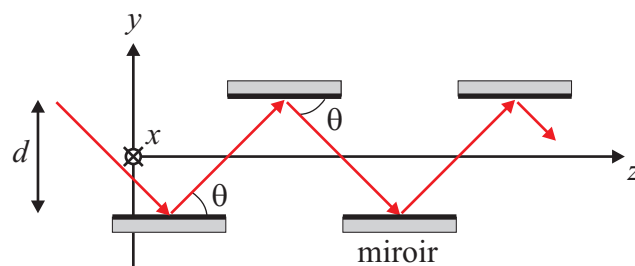


FIGURE 5.1 – Principe d'un "guide d'ondes" dans lequel un rayon lumineux se propage par réflexions successives sur des miroirs.

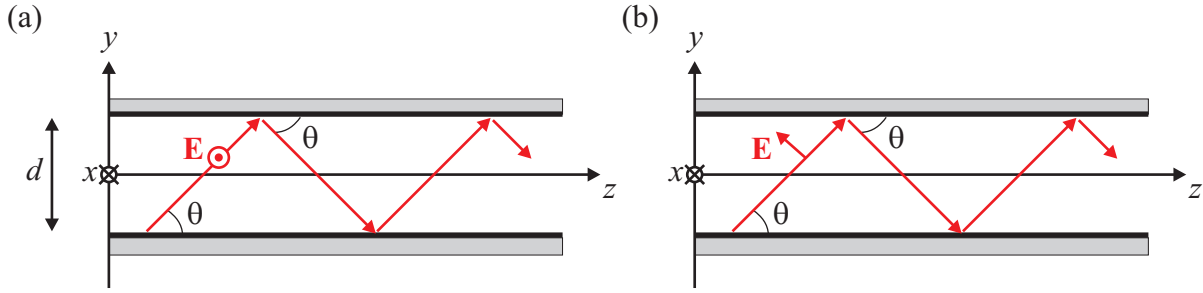


FIGURE 5.2 – Guide d'onde miroir planaire et les deux modes qui peuvent s'y propager : (a) transverse électrique (TE) et (b) transverse magnétique (TM).

5.1 Guide d'onde miroir planaire

Considérons maintenant un guide formé par deux miroirs parfaits infinis séparés par une distance d , Fig. 5.2. Etudions un rayon lumineux se propageant dans le plan y - z . Comme les miroirs sont parfaits, ce rayon est réfléchi sans perte et se propage donc dans la direction z jusqu'à l'infini. On dit que ce rayon est guidé par le guide d'onde. Dans la pratique il est extrêmement difficile de réaliser un tel guide miroir. Cependant, cette géométrie contient toute la physique des guides d'ondes et ses propriétés sont particulièrement simples. Nous allons donc l'utiliser pour étudier en détail ce qu'est un mode optique.

Nous avons vu que la réflexion de la lumière dépend de sa polarisation ; il faut donc définir la polarisation du mode considéré. Dans un premier temps nous allons nous concentrer sur les modes transverse électriques (TE) pour lesquels le champ électrique est normal au plan de propagation, Fig. 5.2(a). Il existe aussi des modes transverse magnétiques (TM), pour lesquels le champ magnétique est normal au plan de propagation (donc le champ électrique se trouve dans le plan de propagation), Fig. 5.2(b). Remarquons que la polarisation est conservée par un mode : si en début de propagation le mode a une polarisation spécifique, il la conserve durant la propagation. On peut donc décomposer tout mode en ses composantes TE et TM puis étudier séparément la propagation de chaque composante.

Considérons une onde plane monochromatique de longueur d'onde $\lambda = \lambda_0/n$, de nombre d'onde $k = nk_0$ et de vitesse de phase $c = c_0/n$, où n représente l'indice de réfraction du milieu remplissant le guide. Si l'onde se propage avec un angle θ , elle est réfléchiée avec un angle $-\theta$ lorsqu'elle atteint le miroir supérieur, Fig. 5.2(a). Pour un miroir parfait, l'onde subit aussi un déphasage de π lors de la réflexion. En effet, pour un miroir parfait le champ électrique est entièrement réfléchi ; ainsi le champ électrique dans le miroir est nul. Comme ce champ doit être continu en vertu des équations de Maxwell (Tab. 4.2), le champ électrique dans le guide juste à l'extérieur du miroir doit aussi être nul. Or ce champ est la superposition du champ électrique incident et du champ électrique réfléchi. Il est clair que ni le champ incident ni le champ réfléchi sont nuls. Pour que leur somme s'annule, il faut que le phaseur de l'un soit décalé de π par rapport au phaseur de l'autre. Notons que ce déphasage de π apparaît aussi dans les coefficients de Fresnel pour la réflexion externe d'un champ TE,

comme indiqué Fig. 4.4(c)).

5.1.1 Vecteurs de propagation

Notre objectif est de déterminer les modes se propageant dans le guide. Un mode est une forme particulière du champ qui maintient la même distribution transverse dans le guide, quelle que soit la position le long de la direction de propagation. On peut donc dire que le mode se propage de façon monobloc dans la direction de propagation. Ainsi cherchons-nous des formes du champ telles qu'il se propage dans une direction spécifique (disons z , Fig. 5.2) sans se modifier dans les directions transverses (disons x - y , Fig. 5.2). Pour obtenir de tels modes, nous allons imposer que le champ qui a été réfléchi deux fois par les surfaces du guide soit identique à un champ qui n'aurait pas été réfléchi et se serait propagé tout droit, comme illustré Fig. 5.3. Le champ qui satisfait cette condition est un mode du guide d'onde.

Observons les fronts d'onde sur Fig. 5.3 ; nous voyons qu'il faut imposer que la phase accumulée par l'onde initiale lorsqu'elle voyage de A à B soit égale à la phase accumulée par l'onde lorsqu'elle est réfléchiée en A , se propage de A à C et est finalement réfléchiée en C . Il ne s'agit évidemment pas d'une condition d'égalité stricte, mais il faut que la différence de phase $\Delta\phi$ entre ces deux chemins optiques soit un multiple entier de 2π :

$$\Delta\phi = kAC - 2\pi - kAB = k(AC - AB) - 2\pi = 2\pi q, \quad q = 0, 1, 2, \dots \quad (5.1)$$

Noter que dans Eq. (5.1) on a utilisé un saut de phase de $-\pi$ lors de chaque réflexion. Ce choix est arbitraire, on peut prendre en effet $\pm\pi$. Avec Fig. 5.3 on observe de plus que $AC - AB = 2d \sin \theta$, où d est l'espace entre les miroirs. On a donc $k2d \sin \theta = 2\pi(q + 1)$, soit

$$k2d \sin \theta = 2\pi m, \quad m = 1, 2, 3, \dots, \quad (5.2)$$

avec $m = q + 1$. Ainsi, si l'on souhaite satisfaire la condition que la phase accumulée est la même pour les chemins AB et AC , il faut que l'angle de propagation θ satisfasse Eq. (5.2),

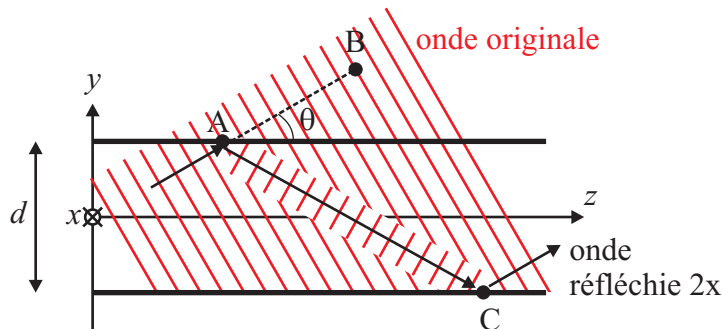


FIGURE 5.3 – Pour définir un mode, on impose que le champ réfléchi deux fois par le guide soit identique à un champ qui se serait propagé directement sans réflexion.

il doit donc prendre les valeurs discrètes θ_m définies par

$$\sin \theta_m = m \frac{\lambda}{2d}, \quad m = 1, 2, 3, \dots \quad (5.3)$$

Ainsi chaque entier m correspond à un angle de propagation spécifique θ_m . Le champ correspondant est le mode m du guide. Le mode $m = 1$ est le mode fondamental, il a le plus petit angle de propagation : $\theta_1 = \sin^{-1}(\lambda/2d)$. Les modes d'ordre supérieur se propagent en zig-zag avec des angles θ_m plus importants. Notons finalement que la longueur d'onde intervenant dans Eq. (5.3) est la longueur d'onde effective qui dépend de l'indice n dont est rempli le guide.

Un mode se propageant dans un guide peut donc être visualisé comme la superposition de deux ondes planes, une se propageant vers le haut, l'autre vers le bas, Fig. 5.4. Remarquer sur cette figure que la distribution du champ résultant est invariante dans la direction de propagation z .

La composante y de la constante de propagation est donnée par $k_y = nk_0 \sin \theta$. Comme seules des valeurs spécifiques de θ sont possible pour les modes, il existe un nombre limité de valeurs de k_y qui correspondent à un mode. En utilisant Eq. (5.3) on obtient

$$k_{ym} = m \frac{\pi}{d}, \quad m = 1, 2, 3, \dots \quad (5.4)$$

Ainsi, les valeurs possibles de k_y sont-elles discrètes. On dit que k_y est la composante transverse du vecteur de propagation. La composante longitudinale (dans la direction de propagation, appelée donc constante de propagation) est généralement représentée avec la lettre β . On a par projection : $\beta = k_z = k \cos \theta = nk_0 \cos \theta$. Ainsi la constante de propagation β est aussi quantifiée : $\beta_m = k \cos \theta_m$, d'où $\beta_m^2 = k^2(1 - \sin^2 \theta_m)$ et en utilisant à nouveau Eq. (5.3),

$$\beta_m^2 = k^2 - \frac{m^2 \pi^2}{d^2}, \quad m = 1, 2, 3, \dots \quad (5.5)$$

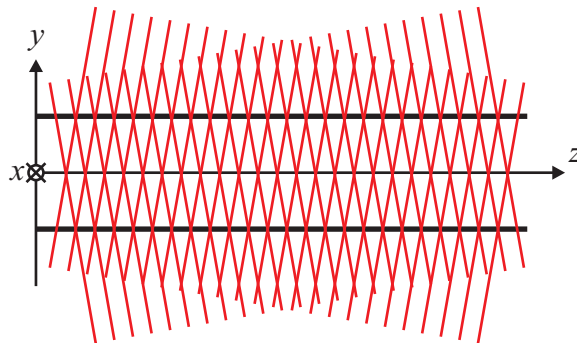


FIGURE 5.4 – Superposition de deux ondes se propageant selon $+\theta$ et $-\theta$. (a) Lorsque ces ondes forment un mode, la figure d'interférence est invariante dans la direction de propagation. (b) Tel n'est pas le cas pour des ondes qui ne forment pas un mode.

On peut facilement construire les différents vecteurs de propagation associés avec les modes d'un guide d'onde, Fig. 5.5. On trace un cercle de longueur nk_0 . Pour satisfaire l'équation d'onde l'extrémité du vecteur d'onde doit se trouver sur ce cercle : $\beta_m^2 + k_{ym}^2 = k^2 = nk_0^2$. On remarque ensuite qu'Eq. (5.4) indique que les composantes transverses sont régulièrement espacées avec une distance π/d . En traçant les lignes horizontales correspondantes, on obtient à l'intersection avec le cercle l'extrémité de chaque vecteur d'onde, que l'on peut donc décomposer en vecteur de propagation β_m et en vecteur transverse k_{ym} .

5.1.2 Distribution du champ

Le champ associé à un mode m se propageant dans le guide est la superposition des deux ondes planes, l'une se propageant vers le haut avec un vecteur de propagation $k_{\uparrow} = (k_y, k_z) = (k_{ym}, \beta_m)$ et l'autre se propageant vers le bas avec le vecteur de propagation $k_{\downarrow} = (k_y, k_z) = (-k_{ym}, \beta_m)$. Ces deux ondes ont pour amplitudes

$$E_{x\uparrow}(y, z) = A_m e^{-jk_{ym}y - j\beta_m z}, \quad (5.6a)$$

$$E_{x\downarrow}(y, z) = e^{j(m-1)\pi} A_m e^{+jk_{ym}y - j\beta_m z}, \quad (5.6b)$$

où l'on a omis la dépendance temporelle en $\exp(j\omega t)$. Le facteur de phase $\exp j(m-1)\pi$ indique qu'en $y = 0$ les deux ondes ont une différence de phase de $(m-1)\pi$. Suivant la valeur de m , les deux ondes vont donc s'additionner ou se soustraire. Ceci donne lieu à des modes symétriques ou anti-symétriques par rapport au plan $y = 0$. Finalement, les champs s'écrivent

$$E_x(y, z) = a_m u_m(y) e^{-j\beta_m z}, \quad (5.7)$$

avec

$$u_m(y) = \begin{cases} \sqrt{\frac{2}{d}} \cos\left(m\pi \frac{y}{d}\right), & m = 1, 3, 5, \dots \\ \sqrt{\frac{2}{d}} \sin\left(m\pi \frac{y}{d}\right), & m = 2, 4, 6, \dots \end{cases} \quad (5.8)$$

avec $a_m = \sqrt{2d}A_m$ pour m impaire et $a_m = j\sqrt{2d}A_m$ pour m paire.

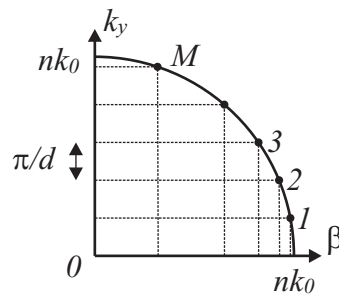


FIGURE 5.5 – Représentation du vecteur de propagation pour différents modes.

Il est très important de remarquer que dans Eq. (5.7) on a décomposé chaque mode en le produit de deux fonctions : l'une ne dépend que de la variable transverse y , alors que l'autre décrit la propagation selon l'axe z . Il s'agit d'une propriété essentielle des modes : l'enveloppe du champ électrique dans le guide ne dépend pas de la section $z = \text{const.}$ où l'on observe le mode. Par contre, l'amplitude va varier dans les limites de cette enveloppe, au fur et à mesure de la propagation z , comme pour une onde plane, dont l'amplitude varie de façon harmonique, comme illustré Fig. 4.1.

Souvent on normalise les fonctions transverses en sorte que leur intégrale sur la section du guide vaille 1 :

$$\int_{-d/2}^{d/2} u_m^2(y) dy = 1, \quad (5.9)$$

ainsi a_m dans Eq. (5.7) représente l'amplitude du mode m .

On a aussi une propriété d'orthogonalité des modes :

$$\int_{-d/2}^{d/2} u_l(y) u_m(y) dy = 0, \quad \text{si } l \neq m. \quad (5.10)$$

Cette propriété a des implications importantes d'un point de vue pratique : si on place de l'énergie dans un mode donné, cette énergie reste dans ce mode pendant toute la propagation. D'une façon générale, si on excite quelques modes spécifiques d'un guide d'ondes, l'énergie demeure dans ces seuls modes et il n'y a pas de transfert d'énergie d'un mode à l'autre.

Figure 5.6 indique l'amplitude des fonctions $u_m(y)$ pour les quelques premiers modes d'un guide miroir. On remarque que les modes correspondant à m impaire sont symétriques par rapport au plan $y = 0$ ($u_m(y) = u_m(-y)$), alors que les modes correspondant à m paire sont anti-symétriques par rapport à ce plan ($u_m(y) = -u_m(-y)$). Remarquer aussi que $u_m(\pm d/2) = 0$. Ainsi les conditions limites requises à l'interface du miroir sont automatiquement satisfaites puisque le champ s'y annule. On observe finalement que les modes d'ordre élevé (m grand) varient plus rapidement dans la direction transverse. Si l'on se souvient que la distribution transverse est le fruit du vecteur transverse k_{ym} , et que pour un mode d'ordre élevé k_{ym} est grand comme indiqué Fig. 5.5, alors il est clair que la "longueur d'onde"

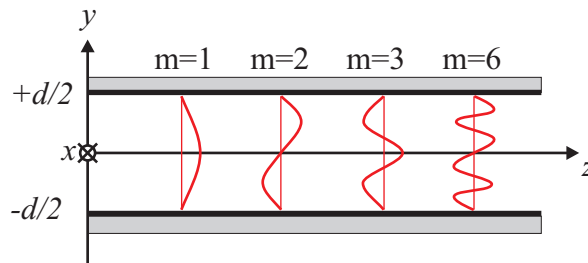


FIGURE 5.6 – Distribution du champ électrique pour quelques modes d'un guide miroir planaire.

associée à cette variation transverse est petite (à un grand vecteur k correspond une petite longueur d'onde $\lambda = 2\pi/k$) et donc le champ varie rapidement dans la direction transverse.

La figure 5.7 représente l'amplitude des trois premiers modes associés avec un guide d'onde

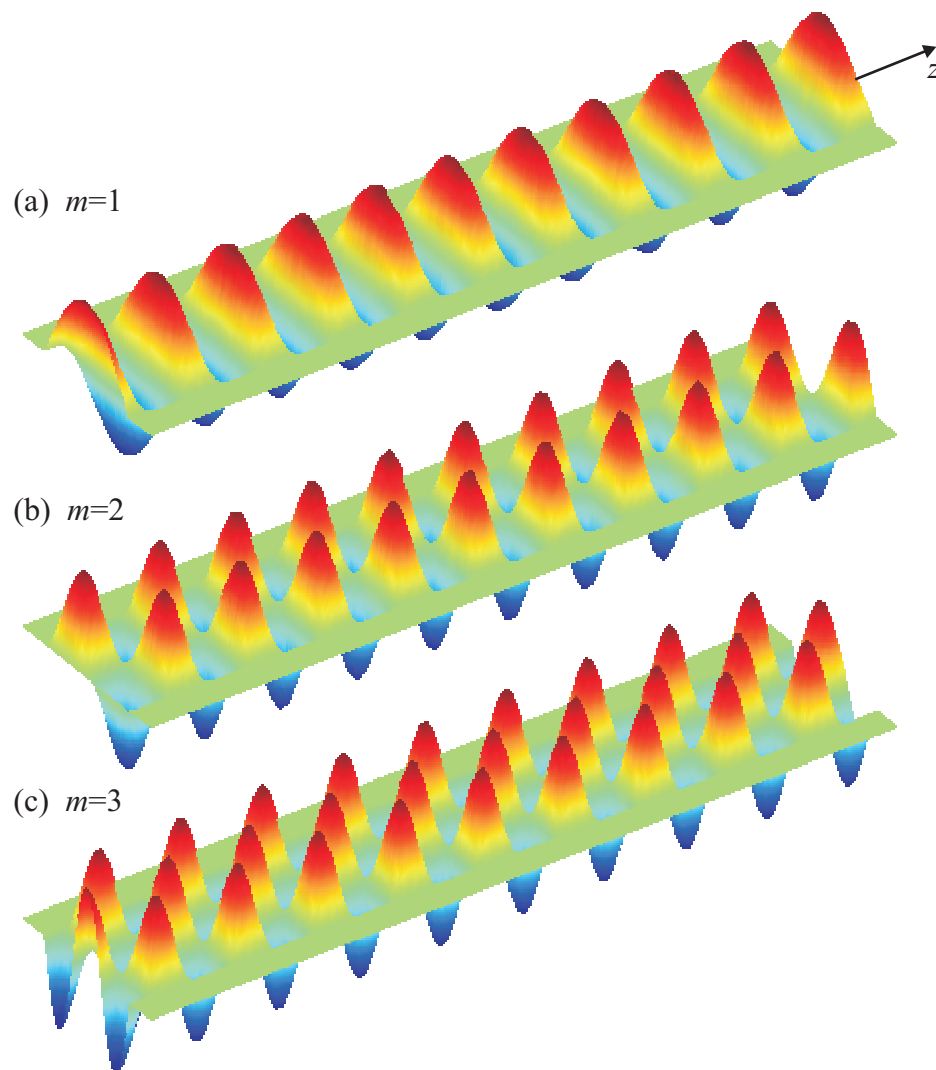


FIGURE 5.7 – Evolution du champ électrique pour trois modes d'un guide miroir plan.

miroir ($m = 1, 2$ et 3); la longueur de propagation est la même pour chaque guide. On remarque que lorsque l'indice du mode augmente, la variation spatiale dans la direction transverse est plus rapide. Par contre, la variation dans la direction de propagation z est alors plus lente : on compte 10 maxima le long de la direction de propagation sur Fig. 5.7(a), 9 maxima sur Fig. 5.7(b) et seulement 8 sur Fig. 5.7(c). On en déduit que la longueur d'onde effective est plus petite pour le premier mode et augmente avec l'indice m des modes ; en d'autres termes, le vecteur de propagation β diminue avec l'indice m des modes. Ceci est en accord avec Eq. (5.5), puisque k_{ym} augmente lorsque m augmente. Si nous reprenons l'image d'un mode comme étant une onde se propageant en zig-zag dans le guide, un mode d'indice m élevé se propage à un angle θ plus important qu'un mode d'indice m bas ; il fait donc plus d'aller-retours entre les deux miroirs, Fig. 5.2.

5.1.3 Nombre de modes

Les angles de propagation sont définis par Eq. (5.3). Si l'on requiert que $\sin \theta_m < 1$, il faut que m soit plus petit que $1/(\lambda/2d)$. Ainsi le nombre de modes vaut-il

$$M \doteq \frac{2d}{\lambda}, \quad (5.11)$$

où le signe $\doteq$ signifie que l'on prend la partie entière par défaut (par exemple 0 pour 0.9, 0 pour 1, ou 1 pour 1.1).

On observe que le nombre de modes dépend de la relation entre la largeur du guide d et la longueur d'onde effective. Pour qu'un mode puisse se propager, il faut que le guide soit suffisamment "large" pour qu'un peu plus qu'une demi-longueur d'onde y trouve place. Ainsi, si $d \leq \lambda/2$ alors $M = 0$ et le guide ne peut pas supporter de mode. Inversement, si la largeur d du guide est fixe, on obtient une condition sur la longueur d'onde : la longueur d'onde $\lambda_c = 2d$ s'appelle la longueur d'onde de coupure (en anglais *cutoff wavelength*) du guide. C'est la longueur d'onde la plus grande qui peut se propager sous forme de mode dans le guide. On peut évidemment définir la fréquence de coupure (*cutoff frequency*) :

$$\nu_c = \frac{c}{2d}. \quad (5.12)$$

Cette fréquence correspond à la plus petite fréquence qui peut être guidée par la structure. Figure 5.8(a) indique le nombre de modes en fonction de la pulsation ω . On remarque que les modes sont espacés régulièrement à $\omega_c, 2\omega_c, 3\omega_c, \dots$. On a aussi indiqué sur cette figure qu'il n'existe pas de mode pour $\omega < \omega_c$. Parfois on nomme cette région la région interdite.

Il est important de garder à l'esprit que ces valeurs de coupure dépendent de l'indice de réfraction n du matériau dont est rempli le guide. Ainsi, si on augmente cet indice de réfraction on diminue la fréquence de coupure $\nu_c = c_0/(2dn)$; i.e. des longueurs d'ondes dans le vide λ_0 plus importantes (i.e. des fréquences plus basses) peuvent alors se propager dans le guide.

5.1.4 Relation de dispersion et vitesse de groupe

Nous avons déjà vu que la relation entre nombre d'onde k et pulsation ω joue un rôle important pour les phénomènes de propagation. Pour une onde plane, on a $k = \omega/c$ i.e. une relation linéaire entre k et ω . Lorsque nous avons étudié la dispersion dans Fig. 2.19 nous avons mentionné qu'il existe des cas où cette relation peut être compliquée et donner lieu à une vitesse de groupe $v = d\omega/dk$ différente de la vitesse de phase $c = \omega/k$. Dans le cas d'un mode, la propagation se fait selon le vecteur de propagation β ainsi on définit dans ce cas les vitesses de phase et de groupe par rapport à β : $c = \omega/\beta$ et $v = d\omega/d\beta$.

Pour un milieu homogène on a simplement $\omega = c\beta$. Pour un guide miroir planaire, Eq. (1.5) donne

$$\beta_m^2 = \frac{\omega^2}{c^2} - \frac{m^2\pi^2}{d^2}. \quad (5.13)$$

Par tradition, on remplace dans Eq. (5.13) l'épaisseur du guide d par la pulsation de coupure en utilisant Eq. (5.12) : $\omega_c = 2\pi\nu_c = \pi c/d$. Ainsi on obtient des relations très générales que l'on peut utiliser pour des géométries différentes en les ramenant simplement à la fréquence de coupure. On obtient alors

$$\beta_m = \frac{\omega}{c} \sqrt{1 - m^2 \frac{\omega_c^2}{\omega^2}}. \quad (5.14)$$

Figure 5.8(b) représente la relation de dispersion pour les premiers modes du guide. On remarque que pour les petite valeurs de ω chaque courbe tend vers la valeur $m\omega_c$. Pour les grandes valeurs de ω la courbe de dispersion s'approche de la "ligne de lumière" $\omega = c\beta$. Ceci se comprend facilement : la ligne de lumière correspond à la propagation dans l'espace homogène infini ; or, lorsque ω devient grand, le mode a une longueur d'onde si petite qu'il ne sent (presque) pas le bord du guide et tend à se propager comme dans un espace homogène infini.

Pour obtenir la vitesse de groupe du mode m , on applique la définition de la vitesse de

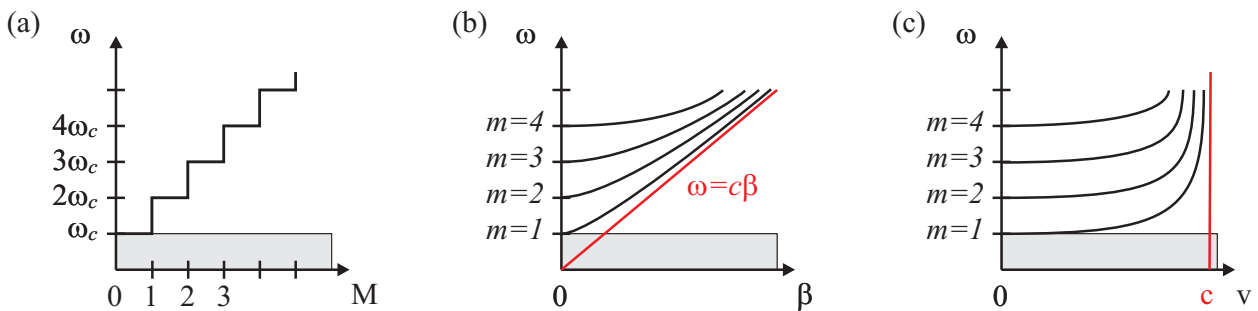


FIGURE 5.8 – (a) Nombre de mode M en fonction de la fréquence angulaire ω . (b) Relations de dispersion pour les premiers modes d'un guide miroir. (c) Vitesse de groupe pour ces modes.

groupe $v_m = d\omega/d\beta_m$ après avoir différencié par rapport à ω Eq. (5.13) des deux côtés :

$$2\beta_m \frac{d\beta_m}{d\omega} = \frac{2\omega}{c^2}. \quad (5.15)$$

En réarrangeant Eq. (5.15) pour faire apparaître $v_m = d\omega/d\beta_m$ on obtient

$$v_m = \frac{c^2\beta_m}{\omega} = c\sqrt{1 - m^2\frac{\omega_c^2}{\omega^2}}, \quad (5.16)$$

où l'on a utilisé la définition de β_m , Eq. (5.14). La vitesse de groupe correspondant aux premiers modes du guide est indiquée Fig. 5.8(c). On remarque que pour un mode donné, v tend vers c lorsque ω augmente. A une fréquence angulaire donnée ω , on remarque aussi que la vitesse de groupe diminue lorsque le numéro du mode augmente. Ainsi, si on transmet de l'information avec un guide, cette information est transmise à plus petite vitesse lorsque l'on utilise un mode élevé. On peut comprendre ce résultat en se souvenant qu'un mode élevé fait plus de zig-zags qu'un mode bas.

5.1.5 Modes TM

Tous les résultats énoncés précédemment sont aussi valables pour les modes polarisés TM. Ceux-ci ne diffèrent des modes TE que par l'orientation du champ électrique, comme indiqué Fig. 5.2. Les composantes du champ électrique sont alors,

$$E_z(y, z) = \begin{cases} a_m \sqrt{\frac{2}{d}} \cos\left(m\pi\frac{y}{d}\right) e^{-j\beta_m z}, & m = 1, 3, 5, \dots \\ a_m \sqrt{\frac{2}{d}} \sin\left(m\pi\frac{y}{d}\right) e^{-j\beta_m z}, & m = 2, 4, 6, \dots, \end{cases} \quad (5.17)$$

et

$$E_y(y, z) = \begin{cases} a_m \sqrt{\frac{2}{d}} \cot\theta_m \cos\left(m\pi\frac{y}{d}\right) e^{-j\beta_m z}, & m = 1, 3, 5, \dots \\ a_m \sqrt{\frac{2}{d}} \cot\theta_m \sin\left(m\pi\frac{y}{d}\right) e^{-j\beta_m z}, & m = 2, 4, 6, \dots, \end{cases} \quad (5.18)$$

avec $a_m = \sqrt{2d}A_m$ pour m impaire et $a_m = j\sqrt{2d}A_m$ pour m paire.

5.2 Guide miroir 2D

Le guide miroir planaire, Fig. 5.1 permet de confiner la lumière dans la direction verticale y alors qu'elle se propage dans la direction z . Par contre le mode s'étend à l'infini dans la direction x . La généralisation la plus simple de cette géométrie est le guide miroir carré de Fig. 5.9(a). Dans ce cas le mode est confiné transversalement dans les deux directions x et y alors qu'il se propage dans la direction z .

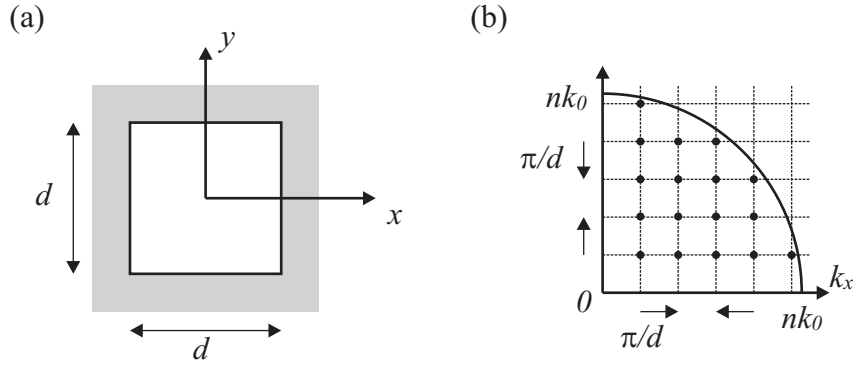


FIGURE 5.9 – Guide miroir carré : (a) section et (b) distribution des modes dans l'espace k_x - k_y .

Le vecteur de propagation a trois composantes $\mathbf{k} = (k_x, k_y, k_z) = (k_x, k_y, \beta)$, où l'on a introduit la constante de propagation β . La relation de dispersion pour un tel mode s'écrit,

$$k_x^2 + k_y^2 + \beta^2 = n^2 k_0^2 = \frac{n^2 \omega^2}{c_0^2}. \quad (5.19)$$

En généralisant Eq. (5.4) on obtient les conditions sur k_x et k_y :

$$2k_x d = 2\pi m_x, \quad m_x = 1, 2, 3, \dots \quad (5.20a)$$

$$2k_y d = 2\pi m_y, \quad m_y = 1, 2, 3, \dots \quad (5.20b)$$

Equation (5.19) donne une condition sur les valeurs possibles des composantes transverses k_x et k_y du vecteur de propagation $\mathbf{k}$: il faut que $k_x^2 + k_y^2 \leq n^2 k_0^2$ pour qu'un mode existe. Cette condition sert de base à la construction de Fig. 5.9(b). Le nombre de modes peut s'obtenir dans ce cas en comptant les couples (k_x, k_y) se trouvant dans le cercle de rayon nk_0 . Pour un guide comptant un nombre important de modes, on peut approximer le nombre de modes

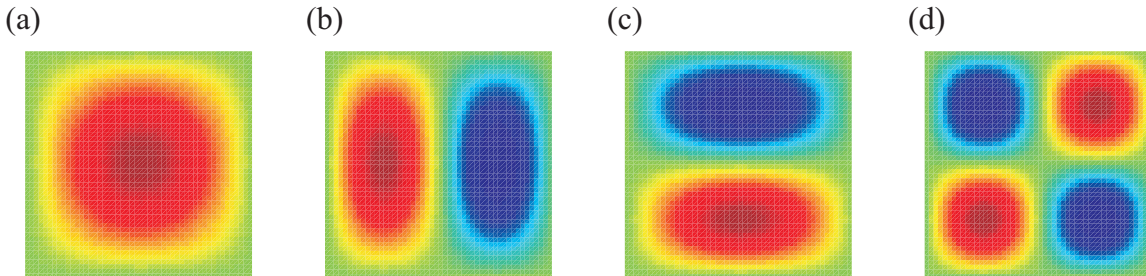


FIGURE 5.10 – Champ associé aux modes d'un guide miroir carré 2D. On peut associer à chaque mode un couple d'indices (l, m) correspondant à sa variation spatiale dans la direction x et y . Les modes représentés correspondent à (a) $(l, m) = (1, 1)$, (b) $(2, 1)$, (c) $(1, 2)$ et (d) $(2, 2)$.

par le rapport entre la surface du quart de disque $\pi(nk_o)^2/4$ et l'aire associée à un mode $(\pi/d)^2$:

$$M \simeq \frac{\pi}{4} \left(\frac{2d}{\lambda} \right)^2. \quad (5.21)$$

Les distributions du champ associé aux premiers modes d'un guide miroir carré sont illustrées Fig. 5.10. On peut classer ces modes à l'aide d'indices (l, m) qui indiquent l'ordre dans les directions x et y , de façon analogue aux modes du guide miroir 1D dont les modes sont représentés Fig. 5.6. Ainsi, chaque mode (l, m) du guide 2D peut se comprendre comme le produit d'un mode l avec un mode m du guide 1D correspondant au côté considéré.

5.3 Guides diélectriques planaires

Il est très difficile de réaliser des miroirs parfaits de grande dimension. Aussi, les guides planaires miroirs ne sont pas utilisés pour les longueurs d'ondes optiques (ils le sont cependant pour les microondes puisque dans ce cas il est plus facile de réaliser des surfaces polies à l'échelle de la longueur d'onde ; de plus, aux fréquences microondes beaucoup de métaux se comportent comme des réflecteurs parfaits ce qui n'est pas le cas aux fréquences optiques où l'absorption dans le métal devient importante).

Pour la lumière on utilise plutôt des guides diélectriques fabriqués avec plusieurs couches de matériaux transparents comme de la silice dans le visible ou des semiconducteurs dans l'infrarouge proche. Le matériau intérieur est toujours celui qui possède le plus grand indice de réfraction n_1 ; il est entouré de matériaux identiques d'indice n_2 dans le cas d'un guide symétrique, Fig. 5.11(a). Dans beaucoup de configurations pratiques on utilise des matériaux différents au dessus et au dessous de la couche intérieure, Fig. 5.11(b). C'est par exemple le cas lorsqu'on croûte la couche du guide n_1 sur un substrat d'indice n_2 et qu'on laisse de l'air au dessus du guide, indice n_3 sur Fig. 5.11(b). Quelle que soit la configuration (symétrique ou asymétrique), le champ se propage de proche en proche par réflexion total aux interfaces du milieu d'indice élevés, Fig. 5.11. Il existe donc des rayons incidents qui ne sont pas guidés par une telle structure, lorsque l'angle d'incidence est supérieur à l'angle d'acceptance θ_a , Fig. 5.11(a).

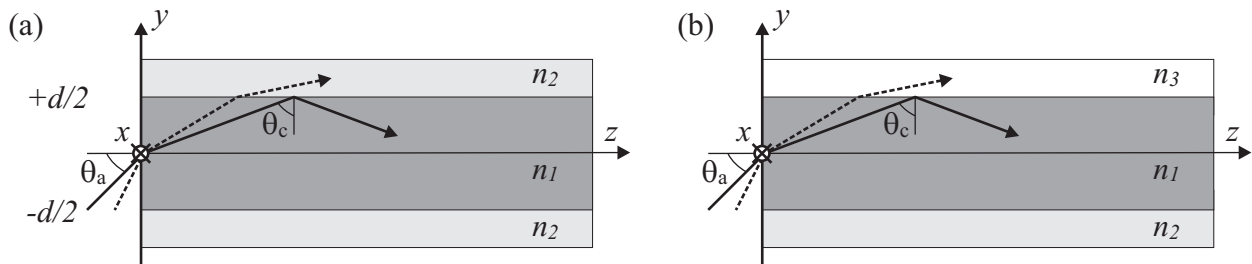


FIGURE 5.11 – Guide diélectrique planaire (a) symétrique et (b) asymétrique.

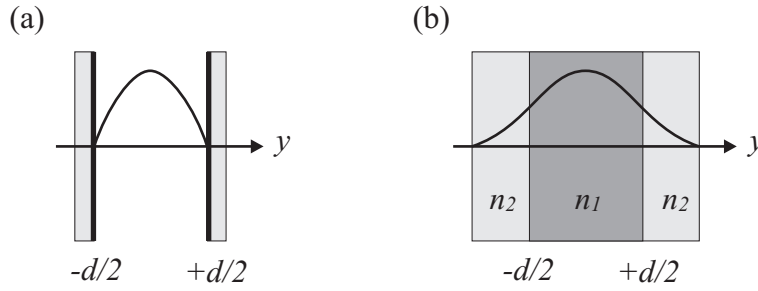


FIGURE 5.12 – Comparaison du profil du champ électrique pour (a) un guide planaire miroir et (b) un guide diélectrique planaire symétrique.

Le calcul des modes dans un guide diélectrique planaire est relativement compliqué et nous nous bornons à expliquer ici le principe. Celui-ci peut se comprendre en considérant Fig. 5.12. Pour un guide miroir, nous avons vu que les conditions au bord imposées au champ électrique sont que celui-ci s'annule sur la surface du miroir, Fig. 5.12(a). De plus il n'y a pas de champ dans le miroir. Pour le guide diélectrique les conditions au bord sont différentes. Tout d'abord le champ n'est pas nul dans les couches environnantes : simplement son amplitude décroît en sorte que le champ s'annule à l'infini (il ne serait en effet pas physique que de l'énergie s'accumule ainsi à l'infini). Dans le guide, le champ prend une forme en sinus ou cosinus comme pour le guide miroir, Fig. 5.12(b). Pour déterminer les modes du guide on impose donc la continuité du champ électrique entre les matériaux n_1 et n_2 comme requis par les équations de Maxwell, Tab. 4.2.

Dans le cas du guide miroir, satisfaire les conditions d'interface donne lieu aux relations (5.8) qui permettent de déterminer le profil du champ électrique (et le vecteur de propagation) de chaque mode. La situation est plus compliquée pour un guide diélectrique planaire : satisfaire les conditions d'interface donne lieu à une condition complexe que l'on peut exprimer en

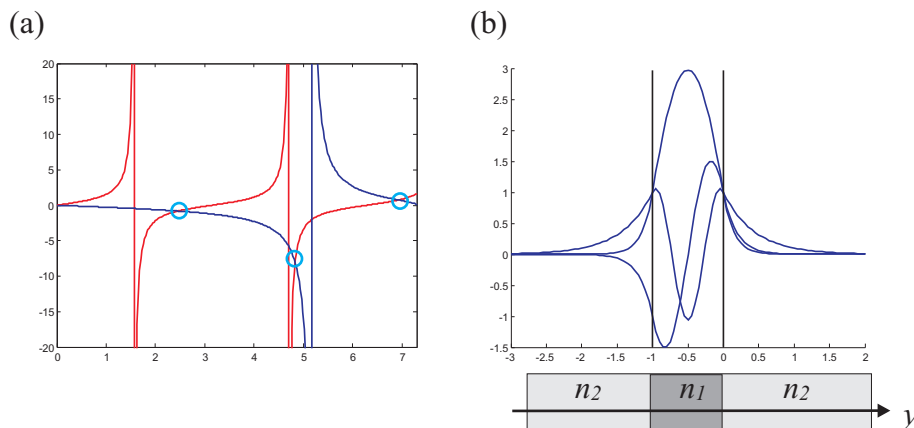


FIGURE 5.13 – Calcul des modes pour un guide diélectrique symétrique : (a) équation transcendante et (b) profil du champ électrique pour les 3 modes TE existant dans ce cas.

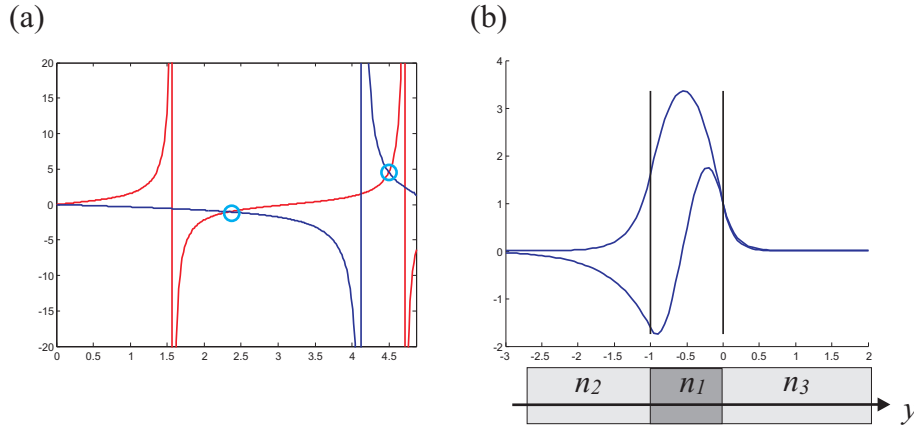


FIGURE 5.14 – Calcul des modes pour un guide diélectrique asymétrique : (a) équation transcendante et (b) profile du champ électrique pour les 2 modes TE existant dans ce cas.

fonction du vecteur de propagation β . Cette relation a une forme transcendante et s'écrit formellement,

$$\arctan \beta = \frac{1}{\sqrt{(\beta - \gamma)}}. \quad (5.22)$$

Equation (5.22) ne peut pas se résoudre analytiquement et l'on doit utiliser des méthodes numériques pour obtenir les valeurs de β la satisfaisant. Cette procédure est illustrée Fig. 5.13 pour un guide symétrique avec pour indices $n_1 = 1.6$, $n_2 = 1.1$ et une épaisseur $d = \lambda_0$ égale à la longueur d'onde considérée. Les intersections des deux courbes correspondant à Eq. (5.22) sont indiquées par des cercles, Fig. 5.22(a). Ces solutions donnent les constantes de propagation des modes et permettent de calculer les distributions des champs électriques correspondants, Fig. 5.22(b).

Pour un guide asymétrique, la distribution du champ n'est pas symétrique comme indiqué Fig. 5.14. Cet exemple correspond à un guide d'indice $n_1 = 1.6$ déposé sur un substrat d'indice $n_2 = 1.4$ avec comme couverture de l'air ($n_3 = 1$) ; l'épaisseur est à nouveau $d = \lambda_0$. On remarque dans ce cas que le champ pénètre davantage dans le substrat que dans l'air puisque la différence d'indice est plus faible entre le guide et le substrat qu'entre le guide et l'air.

Notons finalement que les guides diélectriques planaires supportent des modes TE et TM, comme les guides miroirs. On divise donc aussi les modes d'un guide diélectrique planaire en ces deux familles de polarisation.

5.4 Fibres optiques

Une fibre optique est un guide d'onde diélectrique cylindrique formé d'un coeur (en anglais *core*) d'indice n_1 entouré par une gaine (en anglais *cladding*) d'indice $n_2 < n_1$. Comme le

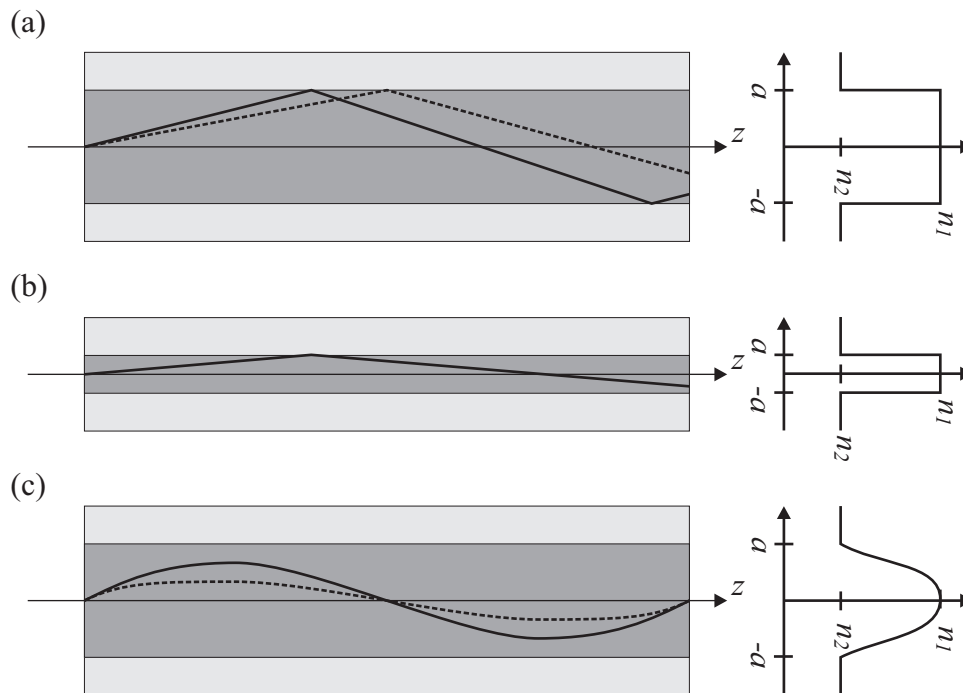


FIGURE 5.15 – Trois types de fibres optiques : fibre à saut d'indice (a) multimode et (b) monomode ; (c) fibre à profil d'indice.

coeur a un indice plus grand, le phénomène de réflexion interne totale peut se produire à l'interface entre le coeur et la gaine, Fig. 5.15(a). Ainsi la lumière est guidée de proche en proche dans le coeur de la fibre par réflexions successives. Le coeur d'une fibre optique est généralement fabriqué en silice (SiO_2), matériau d'une très grande pureté avec une très faible absorption en particulier dans l'infrarouge proche. La technologie moderne a atteint un niveau tel de précision et de pureté des matériaux qu'il est possible de fabriquer des fibres avec des pertes aussi petites que 0.15 dB par kilomètre (ceci correspond à une perte d'intensité de l'ordre de 3.4%).

La figure 5.15 illustre différentes géométries de fibres optiques ainsi que leur profil d'indice de réfraction : la fibre à saut d'indice multimode permet de guider plusieurs modes simultanément, Fig. 5.15(a). On remarque que chaque mode se propage en zig-zag avec un angle de propagation différent. Si on resserre le coeur de la fibre, on arrive à une situation où un seul mode peut se propager ; on parle de fibre monomode, Fig. 5.15(b).

Plutôt que d'avoir une transition abrupte d'indice de réfraction, on réalise parfois des fibres à gradient d'indice, dans lesquelles la lumière se propage sous forme de rayon courbe, Fig. 5.15(c). Nous verrons qu'une telle fibre permet de diminuer la dispersion entre les modes, un phénomène qui limite la bande passante d'une fibre.

Il est important de remarquer que les géométries présentées Fig. 5.15 sont des coupes dans une géométrie cylindrique.

5.4.1 Fibres à saut d'indice

Dans une fibre à saut d'indice, le profil d'indice de réfraction prend des valeurs constantes dans le coeur et dans la gaine et fait un saut abrupte entre ces deux régions, Fig. 5.15(a). Comme la propagation s'effectue par réflexions successives à l'interface coeur-gaine, il existe un angle d'acceptance maximum θ_a au delà duquel l'onde n'est pas guidée mais va se perdre dans la gaine, Fig. 5.16. Remarquer aussi sur cette figure la réfraction à l'interface air-fibre.

L'angle d'acceptance s'obtient en utilisant la loi de Snell deux fois. Une première fois pour calculer l'angle de réflexion interne total selon Eq. (1.5) et Fig. 5.16 : $\sin \theta_c = n_2/n_1$ et une deuxième fois pour tenir compte de la réfraction à l'interface air-fibre :

$$\sin \theta_a = n_1 \cos \theta_c = n_1 \sqrt{1 - \sin^2 \theta_c} = \sqrt{n_1^2 - n_2^2}. \quad (5.23)$$

Ainsi, une onde incidente sur la fibre avec un angle d'incidence $\theta > \theta_a$ n'est pas guidée par la fibre, mais irrémédiablement perdue dans la gaine, Fig. 5.16. Remarquons en passant que si $n_1 \rightarrow n_2$ alors le guide devient très faiblement guidant et seulement une onde incidente perpendiculairement à l'interface ($\theta \simeq 0$) peut être guidée.

On définit aussi l'ouverture numérique de la fibre (en anglais *numerical aperture*),

$$\text{NA} = \sin \theta_a = \sqrt{n_1^2 - n_2^2}. \quad (5.24)$$

Cette ouverture numérique est importante pour le couplage entre une fibre et un composant optique.

Une fibre est généralement caractérisée par la différence relative d'indice,

$$\Delta = \frac{n_1^2 - n_2^2}{2n_1^2} \simeq \frac{n_1 - n_2}{n_1} \ll 1, \quad (5.25)$$

où l'on a utilisé le fait que n_1 et n_2 sont très proches. Aux longueurs d'ondes utilisées pour les télécommunications, n_1 est compris entre 1.44 et 1.46 ; quant à Δ il prend généralement des valeurs entre 0.001 et 0.02.

On déduit des Eqs. (5.24) et (5.25) la formule suivante pour l'ouverture numérique d'une fibre :

$$\text{NA} \simeq n_1 \sqrt{2\Delta}. \quad (5.26)$$

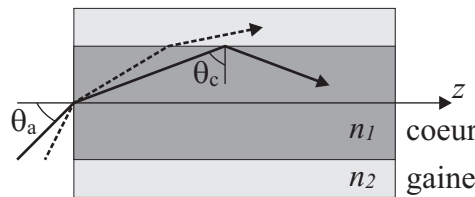


FIGURE 5.16 – Les modes se propageant dans une fibre à saut d'indice ont un angle d'incidence plus petit que θ_a .

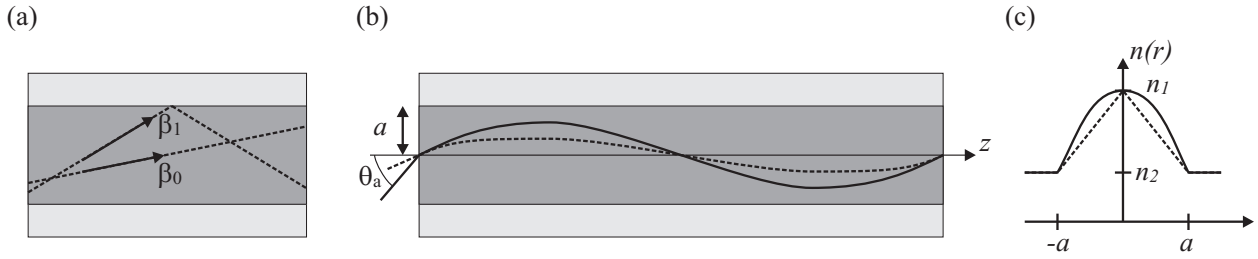


FIGURE 5.17 – (a) Deux modes se propageant dans une fibre à saut d'indice. (b) Deux modes se propageant dans une fibre à profil d'indice. (c) Différent types de profils d'indice : parabolique (trait plein) et triangulaire (traitillé).

On introduit aussi la fréquence normalisée V de la fibre (en anglais *fiber parameter* ou *V parameter*) qui permet de déterminer le nombre de modes supportés par une fibre :

$$V = \frac{2\pi a}{\lambda_0} \text{NA} = k_0 a \text{NA}, \quad (5.27)$$

où a est le rayon du coeur. Les dimensions typiques pour le diamètre du coeur d'une fibre optique sont 8, 50, 62.5, 85 et 100 μm . Les diamètres correspondants pour la gaine sont 125, 125, 125, 125 et 140 μm . La gaine n'a donc pas une épaisseur infinie. Le diamètre le plus petit correspond à une fibre monomode, Fig. 5.15(b).

La fibre est généralement protégée par une troisième couche que l'on nomme manteau et qui donne à la fibre un diamètre extérieur de l'ordre de quelques centaines de μm . Si ce manteau sert avant tout à protéger la fibre, ses propriétés optiques peuvent influencer la performance du guide. Ainsi, lorsque l'on courbe une fibre optique, il se peut qu'une partie de la lumière ne soit plus guidée dans le coeur mais dans la gaine, entre le coeur et le manteau. L'absorption est alors généralement forte et les pertes importantes. On essaie donc d'éviter ce genre de situation.

5.4.2 Fibres à profil d'indice

Considérons deux modes se propageant dans une fibre à saut d'indice, par exemple le mode fondamental β_0 se propageant avec un très faible angle de propagation et un mode β_1 plus élevé faisant de nombreux zig-zags, Fig. 5.17(a). Il est clair que ce dernier va mettre un temps plus long pour parcourir une distance de fibre donnée. En effet, si on se concentre sur la vitesse de phase $c_1 = c_0/n_1$, on remarque que chaque mode se déplace à la même vitesse puisque l'indice de réfraction est homogène dans le coeur. Après une certaine distance le mode β_1 sera "en retard" par rapport au mode β_0 . Remarquons aussi que le mode plus élevé évolue dans la région périphérique de la fibre. Ainsi, si cette région périphérique était composée d'un matériau d'indice $\tilde{n}$ plus petit que n_1 , ce mode rattraperait son retard puisque sa vitesse de phase $c = c_0/\tilde{n}$ serait plus importante que c_1 .

Pour que cette approche fonctionne, il faut évidemment qu'un rayon se propageant dans cet environnement soit guidé. Il s'avère que tel est le cas lorsqu'il existe un gradient d'indice, comme indiqué Fig. 5.17(b) : le rayon se courbe en fonction de sa propagation et est guidé le long de la fibre. Suivant l'angle d'incidence, le rayon évolue plus ou moins loin de l'axe de la fibre.

Les fibres optiques sont fabriquées à partir d'une pré-forme en silice (SiO_2) : un gros tube de verre à l'intérieur duquel on dépose des couches de SiO_2 dopé généralement avec du germanium (GeCl_4) pour augmenter l'indice ou du silicium (SiF_4) pour le diminuer par rapport à l'indice de la silice. On peut ainsi créer des profils d'indice de réfraction arbitraires dans le tube. Ce dernier est ensuite chauffé jusqu'à ce qu'il se referme et on obtient ainsi un barreau de silice. On le place dans un four vertical et on extrude la fibre optique. Cette dernière subit encore différents traitements, comme par exemple le dépôt d'un manteau protecteur.

5.4.3 Profile du champ électrique

Le calcul du champ électrique associé à un mode d'une fibre optique est simple par son principe qui est semblable à celui utilisé pour les modes d'un guide plan. Les détails du calculs dépassent cependant le cadre de ce cours et nécessitent de résoudre l'équation d'onde en coordonnées cylindriques en imposant les conditions limites associées aux équations de Maxwell à chaque interface. Le champ électrique en coordonnées cylindriques (ρ, ϕ, z) possède alors une partie à dépendance radiale (variable ρ) et une partie à dépendance azimutale (variable ϕ). La première est généralement une fonction de Bessel et la deuxième une fonction en sinus ou cosinus.

Notons que l'on ne peut pas diviser les modes d'une fibre optique en modes TE et TM comme pour un guide diélectrique planaire, l'état de polarisation lié à la symétrie cylindrique est plus compliqué. Cependant, les modes les plus utiles sont ceux qui ont une composante transverse dominante et sont donc presque polarisés linéairement dans cette direction. En anglais on

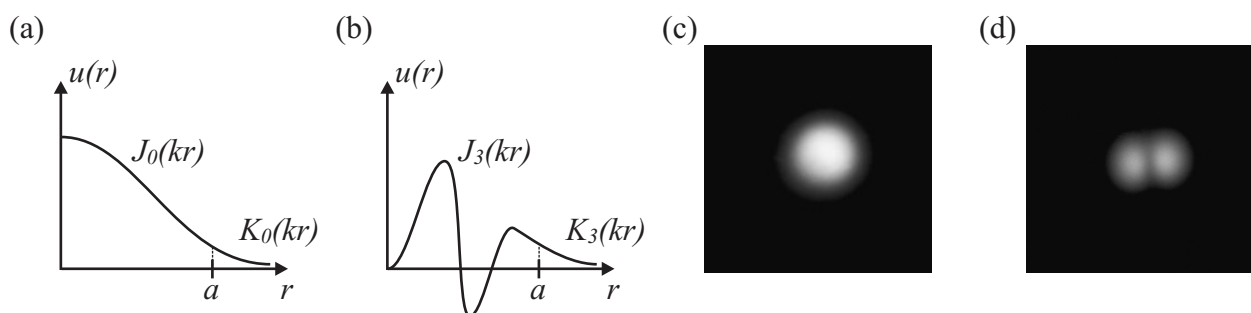


FIGURE 5.18 – Distribution radiale du champ électrique pour les modes (a) $l = 0$ et (b) $l = 3$. Intensité du champ pour les modes (c) LP_{01} et (d) LP_{11} .

les appelle *linear polarized modes* et on les note LP_{lm} , où l'indice l indique l'ordre de la fonction de Bessel radiale utilisée et m l'ordre azimutal. La figure 5.18 donne des exemples de fonctions radiales d'ordre $l = 0$ et $l = 3$. Remarquer que pour satisfaire les conditions de radiation à l'infini on utilise une fonction de Bessel de la deuxième espèce $K_l(kr)$ pour le champ à l'extérieur du coeur et une fonction de Bessel de la première espèce $J_l(kr)$ pour le champ dans le coeur.

Cette fonction radiale est ensuite multipliée par une fonction trigonométrique de périodicité $2\pi/m$ pour obtenir la composante principale du champ électrique, Fig. 5.18(c) et (d).

5.4.4 Atténuation

L'atténuation correspond à la perte de puissance subie par la lumière durant sa propagation dans une fibre optique. On la mesure généralement en dB/km. Pour calculer l'atténuation on mesure la puissance P_1 dans une section donnée de la fibre, par exemple S_1 sur Fig. 5.19 ; on mesure ensuite cette puissance dans une autre section, par exemple P_2 mesurée en S_2 Fig. 5.19. La décroissance de la puissance a une forme exponentielle,

$$P_2 = P_1 e^{-\alpha(z_2 - z_1)}, \quad (5.28)$$

où l'on a introduit le coefficient d'atténuation α qui a pour unités une distance inverse (par exemple m^{-1} ou km^{-1}). On le calcule aisément avec Eq. (5.28) :

$$\alpha = \frac{\ln(P_1/P_2)}{z_2 - z_1}. \quad (5.29)$$

Dans la pratique on exprime souvent le coefficient d'absorption en dB/m ou dB/km ; dans ce cas il faut utiliser le logarithme en base 10 et Eq. (5.29) devient,

$$\alpha_{dB} = \frac{10 \log(P_1/P_2)}{z_2 - z_1}; \quad (5.30)$$

ainsi la décroissance de la puissance optique s'exprime,

$$P_2 = P_1 10^{-\alpha_{dB}(z_2 - z_1)/10}. \quad (5.31)$$

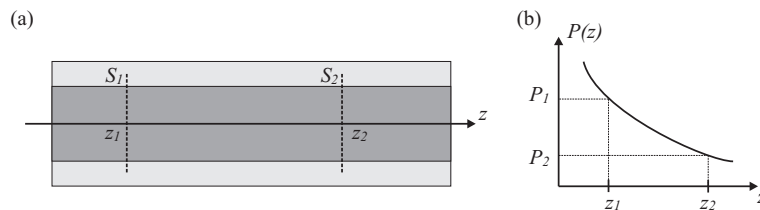


FIGURE 5.19 – (a) Mesure de la puissance en deux sections d'une fibre optique. (b) Décroissance exponentielle de la puissance avec la distance.

L'atténuation a différentes causes :

- *la diffusion Rayleigh* qui provient de la diffusion de la lumière sur des irrégularités ou des inhomogénéités de la fibre. Lorsque la lumière arrive sur de telles irrégularités elle est diffusée dans toutes les directions. Ainsi certains rayons diffusés ne satisfont plus la condition de guidage et la lumière est perdue. Il convient de noter que la diffusion Rayleigh, si elle pose problème pour la transmission sur de longues distances, permet aussi de faire des diagnostics sur une fibre en mesurant de l'extérieur la quantité de lumière se propageant dans la fibre.

La diffusion Rayleigh varie avec la longueur d'onde en sorte que l'atténuation par diffusion Rayleigh varie avec la quatrième puissance de l'inverse de la longueur d'onde dans le vide λ_0 :

$$\alpha = \frac{A_0}{\lambda_0^4}, \quad (5.32)$$

où A_0 est une constante. On peut donc minimiser les pertes par diffusion Rayleigh en utilisant une longueur d'onde dans l'infrarouge plutôt que dans le visible.

- *l'absorption par l'eau* qui existe toujours sous forme de traces dans une fibre optique. Cette absorption est associée au lien hydrogène O-H⁻ qui vibre à une fréquence correspondant à la longueur d'onde dans le vide $\lambda_0 = 2.8 \mu\text{m}$ et crée des pics d'absorption pour les longueurs d'onde $\lambda_0 = 1.38, 1.24$ et $0.95 \mu\text{m}$. Les ions OH⁻ forment aussi une liaison avec les molécules de silice SiO₂ qui absorbe la lumière à $\lambda_0 = 1.23 \mu\text{m}$.

- *l'absorption par les métaux* qui sont présents en trace dans une fibre, mais dont l'interaction avec la lumière est si efficace qu'elle donne lieu à une très forte absorption. Ainsi une concentration de quelques ppm de métal peut donner lieu une atténuation de plus de 1000 dB/km.

La figure 5.20 résume ces différentes propriétés. On y retrouve la décroissance de la diffusion Rayleigh en fonction de la longueur d'onde ainsi que les différents pics d'absorption mentionnés. On observe trois fenêtres principales pour l'utilisation des fibres optiques : $\lambda_0 \simeq 850 \text{ nm}$, $1.3 \mu\text{m}$ et $1.55 \mu\text{m}$. Cette dernière est particulièrement utilisée en raison des nombreux dispositifs semiconducteurs qui existent à cette longueur d'onde.

Figure 5.20 illustre aussi les énormes progrès qui ont été réalisés dans la fabrication de fibres optiques depuis le début des années 1970. Un contrôle des impuretés et de l'homogénéité des matériaux a permis de réduire les pertes de façon considérable.

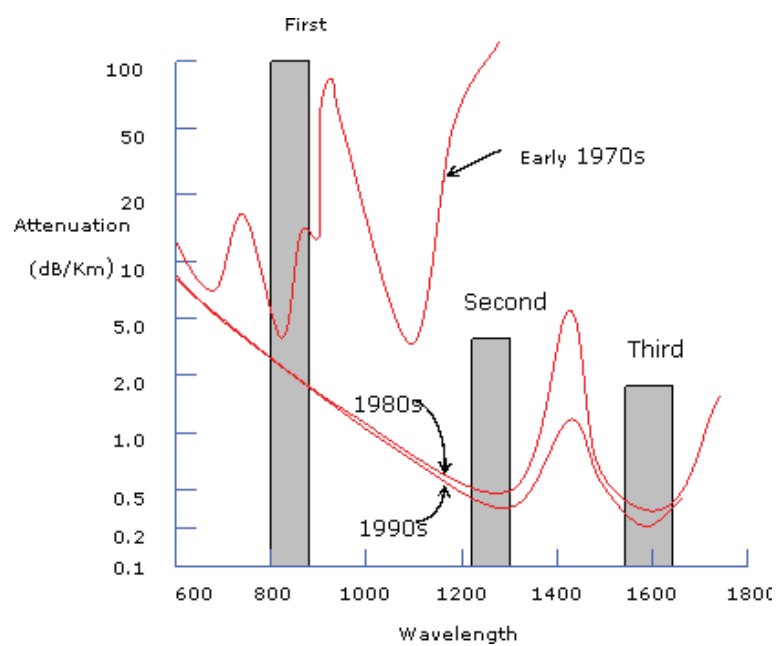


FIGURE 5.20 – Atténuation dans une fibre optique en fonction de la longueur d'onde. Trois fenêtres principales d'utilisation existent où l'atténuation est faible. L'évolution de ces performances depuis les années 1970 est aussi illustrée (d'après Electronics and Communication Engineering Department, Asansol Engineering College India).

Chapitre 6

Photons et transitions optiques

Dans ce chapitre nous apportons un éclairage nouveau sur l'optique, dans lequel nous considérons la lumière comme formée de photons. Ceux-ci permettent d'expliquer un certain nombre de phénomènes qui sont hors de portée de l'optique électromagnétique basée sur les équations de Maxwell, Fig. 6.1. La théorie sous-jacente est l'électrodynamique quantique ; il s'agit d'un formalisme compliqué qui dépasse le cadre de ce cours et nous nous contenterons dans ce chapitre d'une approche phénoménologique. Cette approche permet de comprendre nombre de phénomènes quantiques en incluant simplement dans le formalisme de Maxwell quelques concepts issus de l'électrodynamique quantique qui décrivent le caractère corpusculaire de la lumière.

6.1 Le photon

Dans la perspective de la mécanique quantique, la lumière est formée de "particules" appelées photons. Un photon est une particule indivisible qui transporte une certaine quantité

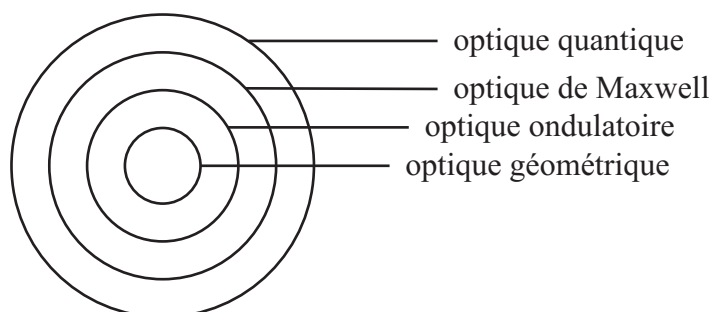


FIGURE 6.1 – L'optique quantique, plus complète que l'optique de Maxwell, permet de comprendre certains phénomènes qui ne s'expliquent pas avec cette dernière.

d'énergie électromagnétique ainsi qu'une quantité de mouvement. Le photon possède aussi un moment cinétique intrinsèque que l'on peut relier à la polarisation de la lumière. Le photon n'a par contre pas de masse. Il se déplace à la vitesse de la lumière (c_0 dans le vide, c_0/n dans un milieu d'indice de réfraction n).

L'énergie E d'un photon dépend de sa fréquence ν :

$$E = h\nu = \hbar\omega, \quad (6.1)$$

où nous avons introduit la constante de Planck $h = 6.63 \cdot 10^{-34}$ Js, ainsi que la constante de Planck réduite $\hbar = h/2\pi = 1.05 \cdot 10^{-34}$ Js. Ainsi un photon infrarouge avec une longueur d'onde $\lambda_0 = 1 \mu\text{m}$ a une fréquence $\nu = 3 \cdot 10^{14}$ Hz correspondant à une énergie $h\nu = 1.99 \cdot 10^{-19}$ J = 1.24 eV. Cette énergie correspond à l'énergie acquise par un électron lorsqu'il est accéléré à travers une différence de potentielle de 1.24 V. On a donc une relation très simple entre l'énergie en eV et la longueur d'onde en μm :

$$E(\text{eV}) = \frac{1.24}{\lambda_0(\mu\text{m})}. \quad (6.2)$$

Cette relation est particulièrement utile pour déterminer la longueur d'onde de la radiation émise par un semiconducteur dont on connaît l'énergie du bandgap. Ainsi pour le GaAs, avec un bandgap de $E = 1.43$ eV à température ambiante, l'émission se situe à environ $\lambda_0 = 870$ nm. Pour un photon micro-onde, avec une longueur d'onde de 1 cm, soit 10^4 fois plus grande que dans Eq. (6.2), l'énergie est donc 10^4 plus petite : $h\nu = 1.24 \cdot 10^{-4}$ eV. En spectroscopie, on note souvent l'énergie en unités cm^{-1} ; la valeur correspondante s'obtient en exprimant la longueur d'onde du photon en cm puis en prenant simplement l'inverse (on parle de longueur d'onde réciproque). Ainsi 1 cm^{-1} correspond à $1.24/10'000$ eV et 1 eV correspond à 8068.1 cm^{-1} . Ces correspondances sont illustrées Fig. 6.2. Remarquer que la fréquence, l'énergie et la longueur d'onde réciproque croissent dans le même sens, alors que la longueur d'onde croît en direction opposée.

A chaque photon de fréquence ν on peut associer une onde décrite par la fonction d'onde complexe $U(\mathbf{r}) \exp(j2\pi\nu t) = U(\mathbf{r}) \exp(j\omega t)$; noter la relation entre l'énergie du photon et la fréquence de la fonction d'onde. Lorsque ce photon arrive sur un détecteur de surface infinitésimale $d\mathbf{A}$ situé à la position $\mathbf{r}$ et orienté de façon normale à la direction de propagation du photon, ce photon est soit entièrement absorbé, soit pas détecté du tout. En effet le photon est indivisible et on ne peut pas en détecter seulement une partie. La probabilité $p(\mathbf{r})d\mathbf{A}$ de détecter un photon au point $\mathbf{r}$ sur une surface $d\mathbf{A}$ est proportionnelle à l'intensité de la lumière à cet endroit. Il est donc plus probable de trouver un photon dans les endroits où l'intensité est élevée. On peut utiliser Eq. (6.1) pour déterminer, à une fréquence donnée ν , le flux de photons N correspondant à une certaine intensité lumineuse :

$$N = I/(h\nu) \quad (6.3)$$

Contrairement aux ondes qui sont étendues dans l'espace et aux particules qui sont localisées dans l'espace, un photon se comporte à la fois comme quelque chose d'étendu et de localisé. On parle de *dualité onde-corpuscule* ou de *dualité onde-matière*. La localisation d'un photon

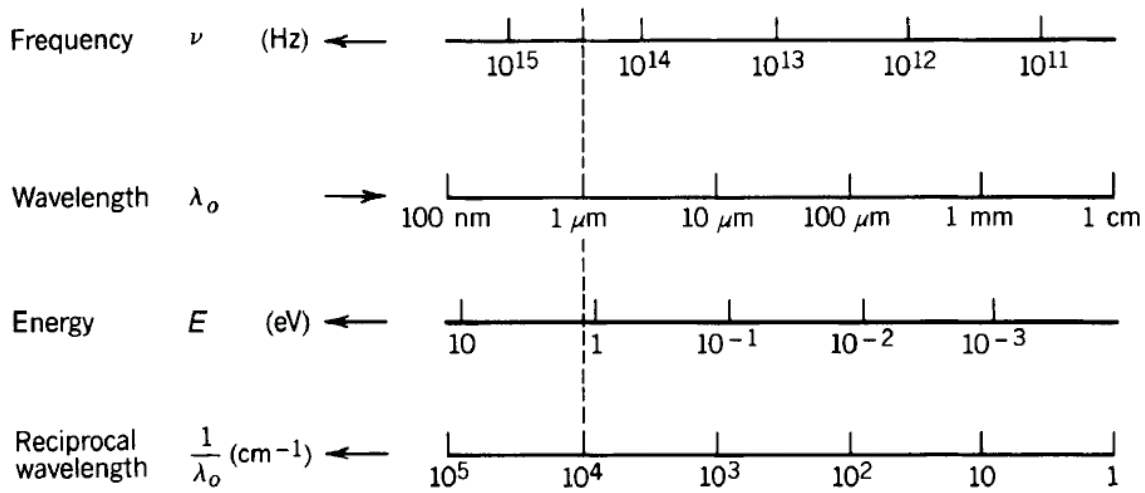


FIGURE 6.2 – Relation entre la longueur d’onde, la fréquence et l’énergie du photon, d’après B.E.A Saleh et M.C. Teich, *Fundamentals of photonics*, 2nd Ed. (Wiley, Hoboken, 2007).

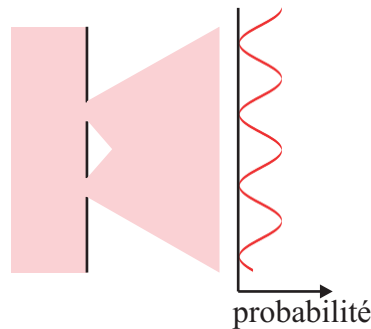


FIGURE 6.3 – Expérience des franges d’interférence de Young avec une source qui envoie les photons les uns après les autres. La statistique des photons mesurés sur l’écran reproduit la figure d’interférence connue de l’optique ondulatoire.

est mise en évidence par sa détection à un endroit bien précis. Un exemple parlant de cette dualité est l’expérience d’interférence des photons réalisée avec un nombre très limité de photons. Dans ce cas on observe que même s’ils arrivent individuellement, les photons forment une statistique qui reproduit la figure d’interférence décrite dans le Chapitre 3. Ainsi, si on réalise l’expérience des fentes d’interférence de Young avec une source de photons unique qui envoie les photons l’un après l’autre, on obtient une statistique qui reproduit parfaitement la figure d’interférence bien connue, Fig. 6.3.

Bien qu’il n’ait pas de masse, le photon transporte une certaine quantité de mouvement $\mathbf{p}$ qui est reliée au vecteur d’onde $\mathbf{k}$:

$$\mathbf{p} = \hbar \mathbf{k}. \quad (6.4)$$

On peut développer cette équation plus avant pour l’amplitude de la quantité de mouvement

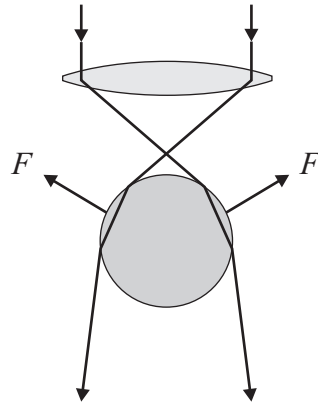


FIGURE 6.4 – Principe du transfert de la quantité de mouvement dans une pincette optique.

en utilisant $p = \hbar k = \hbar \omega / c = \hbar 2\pi / \lambda$ en sorte que

$$p = E/c = h/\lambda. \quad (6.5)$$

La quantité de mouvement associée au photon peut être transférée à un objet de masse finie, donnant lieu à une force et causant le mouvement de l'objet. On parle souvent de pression de radiation. Cet effet est exploité dans les pincettes optiques (*optical tweezer* en anglais) qui permettent de manipuler des objets dans les trois dimensions d'espace. Le principe en est illustré Fig. 6.4 qui montre une particule sphérique piégée par un faisceau gaussien focalisé fortement par une lentille. Si on considère un rayon interagissant avec la particule, on remarque qu'à chaque réfraction le rayon change de direction. Les photons correspondants échangent donc de la quantité de mouvement avec la particule, créant une force F . Globalement la particule subit une force totale verticale qui peut contrebalancer la force de gravité. Si la particule se déplace latéralement par rapport à l'axe du faisceau, il apparaît aussi une dissymétrie qui donne lieu à une force de rappel latérale permettant de manipuler la particule dans cette direction. Ainsi, en jouant sur l'orientation du faisceau dans l'espace peut-on déplacer la particule à volonté, comme avec des pincettes (*tweezer* en anglais). On parle aussi de piégeage optique ; ce dernier peut revêtir des formes très variées, par exemple dans des champs optiques périodiques obtenus par l'interférence de plusieurs faisceaux où l'on piège de nombreuses particules. Une application potentielle de tel piégeage est la création de larges entités de matière dans l'espace, maintenues par des forces optiques.

Considérons une particule sphérique d'indice de réfraction n_1 de rayon a , beaucoup plus petit que la longueur d'onde λ_0 (on parle d'approximation dipolaire), se trouvant dans un milieu homogène d'indice n_0 . La force agissant sur cette particule dépend de la variation de l'intensité du champ électrique $\mathbf{E}$ associé à l'onde optique :

$$\mathbf{F}(\mathbf{r}) = \frac{1}{2} \alpha \nabla E^2(\mathbf{r}) = 2\pi n_0^2 \varepsilon_0 a^3 \left(\frac{m^2 - 1}{m^2 + 2} \right) \nabla I(\mathbf{r}), \quad (6.6)$$

où α représente la polarisabilité de la particule, i.e. sa propension à interagir avec la lumière

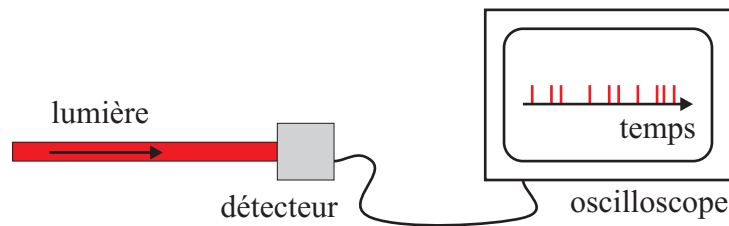


FIGURE 6.5 – Enregistrement des photons à un endroit donné en fonction du temps.

incidente, $m = n_1/n_0$ l'indice relatif de la particule par rapport au milieu environnant et I l'intensité de la lumière.

Il est intéressant de noter que si l'on dispose d'un détecteur très sensible et rapide, on peut détecter des photons uniques et étudier leur arrivée dans le temps. On remarque alors que les photons arrivent généralement de façon quasi-aléatoire, Fig. 6.5. Ceci provient du fait que les processus donnant lieu à la création de photons dans la matière sont généralement aléatoires. Cependant, il existe aussi des conditions expérimentales où un grand nombre de photons est mesuré quasi-simultanément. On parle alors de gerbes de photons (en anglais *photon bunching*). Il existe aussi un phénomène inverse (*photon anti-bunching*) dans lequel l'émission d'un photon, par exemple par la fluorescence d'une molécule, limite la probabilité d'émission des photons suivants. Tous ces phénomènes sont étudiés dans le cadre de l'optique statistique.

Terminons cette section en indiquant Tab. 6.1 les valeurs typiques de quelques flux moyens de photons pour différentes sources. Ces flux sont mesurés en prenant des temps de mesure longs en sorte que les fluctuations statistiques dont nous venons de parler n'influencent pas la mesure.

TABLE 6.1 – Densité moyenne de flux de photons pour différentes sources .

Source	Flux (photons/cm ² s)
Etoile	10^6
Pleine lune	10^8
Eclairage électrique	10^{12}
Soleil	10^{14}
Laser He-Ne (10 mW focalisé sur un spot de 20 μ m)	10^{22}

6.2 Radiation thermique

Dans cette section nous introduisons le concept de corps noir. Le spectre du rayonnement issu d'un corps noir est resté un mystère jusqu'à la fin du 19^{me} siècle et c'est son étude qui poussa

le physicien allemand Max Planck à introduire la notion de quanta. En effet, en suivant la théorie de Rayleigh, à température donnée une source devrait émettre une quantité infinie d'énergie aux courtes longueurs d'onde (on parle de "catastrophe ultraviolette", Fig. 6.7). L'introduction de la quantification du champ électromagnétique et la notion de photon ont permis de réconcilier la théorie avec l'expérience.

Comme nous l'avons vu au paragraphe précédent, tout corps ayant une température $T > 0$ absorbe et émet des photons. Le spectre de la lumière émise dépend de la température du corps. Un corps noir est un objet abstrait qui absorbe totalement les radiations électromagnétiques à toutes les fréquences, des plus petites aux plus grandes. Un tel objet n'existe pas, mais peut être approximé par une cavité aux parois imperméables, comme indiqué Fig. 6.6. La cavité est remplie par toutes sortes de radiations de différentes fréquences qui sont absorbées par les parois. Pour mesurer le spectre émis par un tel objet on pratique une petite ouverture dans la paroi, Fig. 6.6. Pour ne pas perturber la mesure, toute radiation pénétrant dans la cavité doit être absorbée ; on suppose de plus que cette radiation provenant de l'extérieur ne constitue qu'une part infime des radiations se trouvant dans la cavité. Quant au spectre émis à travers la petite ouverture vers l'extérieur il représente bien le spectre de la radiation en équilibre thermique. Figure 6.7 illustre ce spectre pour différentes températures. Une théorie classique – semblable à la diffusion Rayleigh introduite Sec. 5.4.4 pour les fibres optiques – indique que la quantité de lumière émise devrait diverger pour les courtes longueurs d'onde. L'expérience montre que tel n'est pas le cas et le spectre de radiation émis par un corps noir exhibe un maximum qui dépend de la température, Fig. 6.7.

Cette radiance spectrale d'un corps noir est définie par

$$B_{\lambda T} = \frac{2hc^2}{\lambda^5 [\exp(hc/\lambda KT) - 1]} . \quad (6.7)$$

Cette formule est issue d'une approche quantique de la radiation dans laquelle le champ électromagnétique est formé de photons qui transportent une quantité d'énergie quantifiée. La valeur de la constante de Planck h apparaissant dans Eq. (6.7) et que nous avons utilisée pour obtenir la quantité d'énergie $h\nu$ associée à un photon de fréquence ν a été déterminée par le physicien allemand en comparant Eq. (6.7) avec la mesure d'un corps noir.

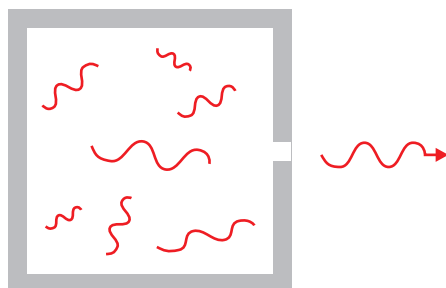


FIGURE 6.6 – Modèle d'un corps noir : toutes les radiations de toutes fréquences sont absorbées par les parois de la cavité. Une petite ouverture permet d'observer le spectre émis par le corps noir sans le perturber.

Equation (6.7) permet de déduire plusieurs lois fondamentales liées à la radiation des corps noirs. En intégrant Eq. (6.7) sur l'entier du spectre on obtient la loi de Boltzmann qui indique que l'intensité totale d'émission d'un corps noir varie avec la quatrième puissance de la température :

$$I = \sigma T^4, \quad (6.8)$$

avec la constante $\sigma = 5.669 \cdot 10^{-8} \text{ W/m}^2\text{K}^4$.

En dérivant Eq. (6.7) on obtient la longueur d'onde $\lambda_{\max}$ d'émissivité maximale pour une température donnée,

$$\lambda_{\max} T = 2.8978 \cdot 10^{-3} \text{ m K}, \quad (6.9)$$

connue sous le nom de loi de Wien. Il existe aussi deux valeurs limites d'Eq. (6.7) :

$$B_{\lambda T} = \frac{2h\nu^3}{c^2} \exp(-h\nu/KT) \quad \text{pour } h\nu \gg KT, \quad (6.10)$$

$$B_{\lambda T} = \frac{2KT\nu^2}{c^2} \quad \text{pour } h\nu \ll KT. \quad (6.11)$$

Même si un corps noir est une idéalisation qui n'existe pas en réalité, le formalisme développé par Planck et esquissé dans ce paragraphe s'applique à nombre de situations pratiques, même si le matériau ne se comporte pas comme un absorbeur parfait. En effet, il existe généralement une portion du spectre optique dans laquelle le matériau peut être considéré comme un absorbeur parfait. Un domaine d'application particulièrement important de ce formalisme est la thermographie où le spectre émis par un objet permet d'obtenir des informations sur la température ou la conformation de cet objet.

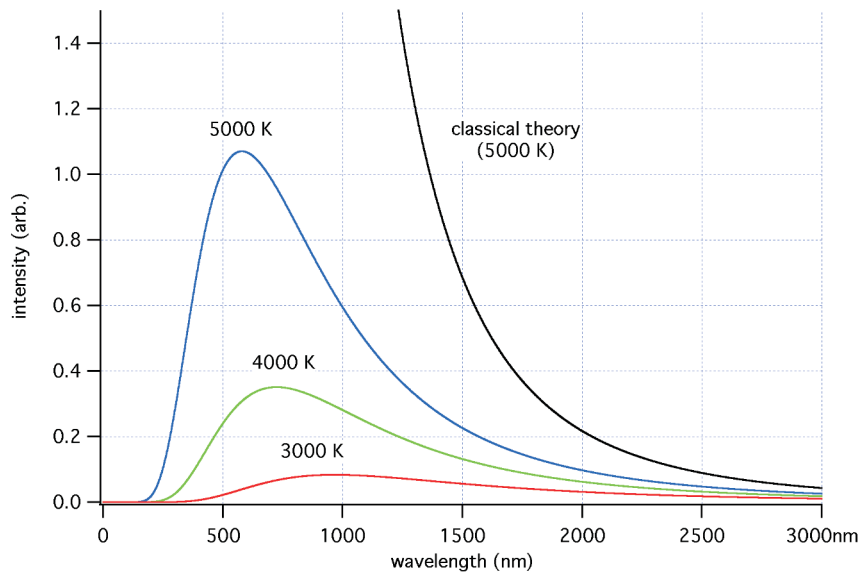


FIGURE 6.7 – Spectre de radiation d'un corps noir à différentes températures T . La courbe divergente pour $T \rightarrow 0$ est issue de la théorie classique et représente la "catastrophe ultra-violette".

6.3 Niveaux d'énergie

Les photons sont indissociablement liés aux niveaux d'énergie que l'on trouve dans la matière. En effet, lorsqu'un électron se désexcite d'un niveau énergétique élevé vers un niveau plus bas, il perd son énergie en émettant un photon. Inversement, dans un atome, un électron peut être promu à un niveau énergétique plus élevé en absorbant un photon.

6.3.1 Niveaux électroniques, modèle de Bohr

Commençons par étudier le modèle de Bohr de l'atome d'hydrogène, illustré Fig. 6.8, où l'on considère un atome avec un noyau formé d'un proton et un électron gravitant autour selon une trajectoire circulaire. Le proton a une charge e , alors que l'électron a une charge $-e$. En omettant toute notation vectorielle pour simplifier la discussion, on peut écrire l'amplitude de la force de Coulomb subie par l'électron comme,

$$F = \frac{1}{4\pi\epsilon_0} \frac{e^2}{r^2}. \quad (6.12)$$

En supposant une orbite circulaire, l'accélération vaut $a = v^2/r$, où v représente la vitesse. En utilisant la deuxième loi de Newton $F = ma$, avec m la masse de l'électron, on obtient la relation suivante :

$$\frac{mv^2}{r} = \frac{1}{4\pi\epsilon_0} \frac{e^2}{r^2}. \quad (6.13)$$

En se souvenant que l'énergie cinétique de l'électron vaut $E_c = mv^2/2$ et son énergie potentielle vaut $E_p = -e^2/4\pi\epsilon_0 r$ (potentiel de Coulomb), Eq. (6.13) indique que $E_c = E_p/2$.

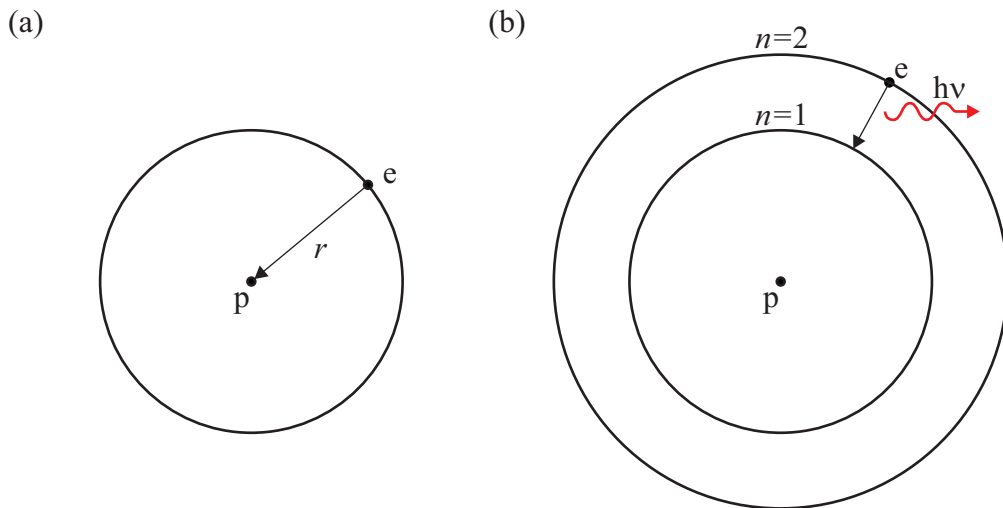


FIGURE 6.8 – (a) Modèle de Bohr pour l'atome d'hydrogène : un électron orbite circulairement autour d'un proton. (b) Emission d'un photon lorsqu'un électron se désexcite du niveau 2 vers le niveau 1.

L'énergie potentielle est négative car il s'agit d'un potentiel attractif. Ainsi, en tenant compte des signes, l'énergie totale de l'électron vaut,

$$E = E_c + E_p = -\frac{1}{8\pi\epsilon_0} \frac{e^2}{r}. \quad (6.14)$$

En combinant Eqs. (6.13) et (6.14) on obtient pour l'énergie totale

$$E = -\frac{mv^2}{2}. \quad (6.15)$$

En principe, n'importe quelle orbite est possible : l'électron peut se trouver à une distance r arbitraire du proton et l'énergie correspondante calculée à partir d'Eq. (6.14). Bohr a cependant introduit une contrainte supplémentaire sur les orbites possibles, qui est que le moment cinétique de l'électron ne peut prendre que des valeurs discrètes :

$$L = n \frac{h}{2\pi} = n\hbar \quad n = 1, 2, 3, \dots \quad (6.16)$$

En introduisant cette valeur du moment cinétique L et sa définition, $L = mvr$, dans Eq. (6.13), on peut éliminer r et obtenir une expression pour la vitesse,

$$v = \frac{1}{4\pi\epsilon_0} \frac{e^2}{L} = \frac{1}{4\pi\epsilon_0} \frac{2\pi e^2}{nh}. \quad (6.17)$$

Finalement, Eqs. (6.15) et (6.17) donnent la formule de Bohr pour les niveaux d'énergie possibles pour l'atome d'hydrogène :

$$E_n = -\left(\frac{1}{4\pi\epsilon_0}\right)^2 \frac{me^4}{2n^2\hbar^2}; \quad (6.18)$$

cette équation donne aussi les rayons possibles des orbites de l'électron,

$$r_n = 4\pi\epsilon_0 \frac{n^2\hbar^2}{me^2}. \quad (6.19)$$

On utilise cette équation pour définir le rayon de Bohr a_0 qui correspond à la valeur obtenue pour $n = 1$,

$$a_0 = r_1 = 4\pi\epsilon_0 \frac{\hbar^2}{me^2} \simeq 0.53 \text{ \AA}, \quad (6.20)$$

qui est extrêmement proche de la valeur expérimentale.

Revenons à notre objectif qui est d'établir un lien entre les niveaux énergétiques d'un atome et la fréquences ν de lumière émise lorsqu'un électron se désexcite d'un niveau n vers un niveau n' . En suivant les travaux de Planck et Einstein, Bohr a postulé que l'énergie du photon correspondait à la différence d'énergie entre les deux niveaux :

$$\Delta E_{n \rightarrow n'} = E_n - E_{n'} = h\nu. \quad (6.21)$$

En utilisant Eq. (6.18), on obtient la valeur suivante de la différence d'énergie :

$$\Delta E_{n \rightarrow n'} = \left(\frac{1}{4\pi\epsilon_0} \right)^2 \frac{me^4}{2\hbar^2} \left(\frac{1}{n'^2} - \frac{1}{n^2} \right). \quad (6.22)$$

Les premières expériences de spectroscopie mesuraient la longueur d'onde λ de la lumière plutôt que sa fréquence ν , en utilisant la relation $\lambda = c_0/\nu$, on peut obtenir d'Eq. (6.22) une expression pour les longueurs d'ondes émises par un atome d'hydrogène :

$$\lambda = \frac{hc_0}{\Delta E_{n \rightarrow n'}} = (4\pi\epsilon_0)^2 \frac{4\pi\hbar^3 c_0}{me^4} \frac{n^2 n'^2}{n^2 - n'^2}. \quad (6.23)$$

Le mathématicien suisse Johann Jakob Balmer a découvert de façon remarquable une relation entre les longueurs d'onde du spectre de l'atome d'hydrogène :

$$\lambda = \frac{bn^2}{n^2 - 4} \quad n = 3, 4, 5, \dots \quad (6.24)$$

avec $b = 3'645.6 \text{ \AA}$ et la longueur d'onde en unités d'Angstrom. Cette série de Balmer n'est autre qu'Eq. (6.23) pour le cas particulier $n' = 2$ et $n > 2$:

$$\lambda = (4\pi\epsilon_0)^2 \frac{16\pi\hbar^3 c_0}{me^4} \frac{n^2}{n^2 - 4}. \quad (6.25)$$

Les termes constants au début de cette équation correspondent exactement au paramètre $b = 3'645.6 \text{ \AA}$ trouvé par Balmer ! Cette série de lignes spectroscopiques joue un rôle très important en astronomie puisque l'hydrogène est un carburant essentiel des étoiles, dont la spectroscopie suit la série de Balmer.

Le modèle de l'atome de Bohr permet de déduire une dernière constante importante qui intervient dans la spectroscopie atomique : la constante de structure fine. Considérons la vitesse v de l'électron dans le modèle de Bohr et calculons son rapport à la vitesse de la lumière c_0 . En utilisant Eq. (6.13) on obtient la valeur suivante :

$$\left(\frac{v}{c_0} \right)^2 = \frac{1}{4\pi\epsilon_0} \frac{e^2}{rmc_0^2}. \quad (6.26)$$

En introduisant la valeur de r_n d'Eq. (6.19) on obtient finalement

$$\frac{v}{c_0} = \frac{1}{4\pi\epsilon_0} \frac{e^2}{n\hbar c_0}. \quad (6.27)$$

La plus grande vitesse pour l'électron s'obtient pour la plus petite valeur de n : $n = 1$. La valeur numérique correspondante est $v/c_0 = 0.007297$. Cette valeur de la constante de structure fine est généralement approximée par $\alpha = 1/137$. En électrodynamique quantique, cette constante joue le rôle de constante de couplage entre les électrons et les photons.

Dans le modèle de Bohr, l'état d'un électron est déterminé par le nombre quantique principal n . Une théorie quantique de l'atome d'hydrogène introduit d'autres nombres quantiques pour décrire un état électronique :

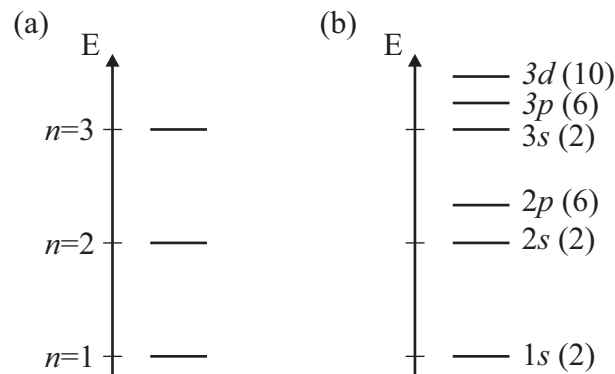


FIGURE 6.9 – (a) Modèle de Bohr pour l'atome d'hydrogène, l'énergie de chaque niveau dépend uniquement du nombre quantique principal n et (b) modèle plus précis dans lequel l'énergie d'un état électronique dépend aussi du nombre quantique azimutal l indiqué par les lettres s , p , d . Le nombre d'électrons possible est indiqué entre parenthèse pour chaque niveau.

- i) Pour chaque nombre quantique principal n , il existe les nombres quantiques azimutaux $l = 0, 1, 2, \dots, n - 1$, qui correspondent au moment angulaire orbital de l'électron. On introduit la notation suivante pour ces orbitales : $l = 0 \leftrightarrow$ orbitale s , $l = 1 \leftrightarrow$ orbitale p , $l = 2 \leftrightarrow$ orbitale d , $l = 3 \leftrightarrow$ orbitale f et $l = 4 \leftrightarrow$ orbitale g .
- ii) Pour chaque nombre quantique azimutal l , il existe $2l + 1$ nombres quantiques magnétiques $m = -l, -l + 1, \dots, -1, 0, 1, \dots, l - 1, l$.
- iii) Finalement, en plus de son moment angulaire orbital, l'électron a aussi un moment angulaire intrinsèque (le *spin* en anglais) qui peut prendre les valeurs $m_s = \pm 1/2$.

Pour une valeur n du nombre quantique principal, il existe $2n^2$ états possibles caractérisés par les valeurs des autres nombres quantiques l , m , et m_s . Equation (6.18) ne dépend que du nombre quantique principal ; on dit donc que l'on a $2n$ états dégénérés (i.e. avec la même énergie) dans ce modèle de Bohr, comme illustré Fig. 6.9(a).

Un modèle plus précis repose sur le principe d'exclusion de Pauli qui interdit à deux électrons d'occuper exactement le même état quantique (i.e. avec exactement les mêmes nombres quantiques n , l , m et m_s). Dans ce cas, l'énergie des niveaux dépend non seulement du nombre quantique principal n mais aussi du nombre quantique azimutal l que l'on indique par les lettres correspondantes s , p , d , f , g , Fig. 6.9(b). Dans cette figure, on indique entre parenthèse près de chaque niveau le nombre maximal d'électrons qu'il peut accueillir ; ces différents électrons sont caractérisés par leur nombre quantique magnétique m et par leur spin m_s ; pour une valeur de l il existe $2(2l + 1)$ états différents possibles.

La figure 6.9(b) illustre le comportement des niveaux énergétiques pour des atomes plus compliqués que l'atome d'hydrogène. Il n'existe cependant pas de formule équivalente à Eq. (6.18) pour des atomes plus compliqués. La seule règle qui demeure est le remplissage des niveaux énergétiques depuis le bas ; ainsi, pour un atome de carbone (numéro atomique $Z = 6$) avec ses 6 électrons, 2 électrons occupent l'état $1s$, deux électrons occupent l'état

2s et deux électrons occupent l'état 2p, qui n'est donc pas complet puisqu'il peut recevoir 6 électrons.

6.3.2 Etats vibrationnels et rotationnels

La relation entre émission de lumière et niveaux électroniques d'un atome, comme illustrée Fig. 6.8, est aisée à comprendre intuitivement : pour changer d'orbite, l'électron doit se débarrasser d'une certaine quantité d'énergie qui est alors transmise sous forme de photon. Ce phénomène peut aussi se produire dans une molécule avec sa structure électronique complexe : dans ce cas, un changement d'orbite pour un électron nécessite le concours d'un photon, généralement dans l'ultraviolet.

Compte tenu de son extension spatiale, une molécule possède aussi des états vibrationnels et rotationnels, en plus d'états électroniques. Ces états que l'on peut qualifier de "mécaniques" sont aussi associés à une énergie bien spécifique et sont aussi quantifiés, comme les états électroniques de l'atome de Bohr. On peut donc aussi considérer des transitions entre différents états vibrationnels ou rotationnels dans une molécule. Les énergies associées avec des transitions vibrationnelles sont extrêmement petites et correspondent à des fréquences dans l'infrarouge, alors que les transitions rotationnelles ont des énergies encore plus petites, dans les micro-ondes.

Une molécule est composée de différents atomes qui sont maintenus ensemble sous l'effet de forces moléculaires. Deux types de liaisons existent dans ce cas : ioniques et covalentes. Dans une liaison ionique les charges positives et négatives demeurent séparées spatialement, donnant lieu à la formation de dipôles. Dans une liaison covalente les atomes partagent leurs charges et la molécule résultante n'a pas de moment dipolaire permanent. On peut faire un modèle très simple d'une molécule en utilisant des masses et des ressorts, comme illustré

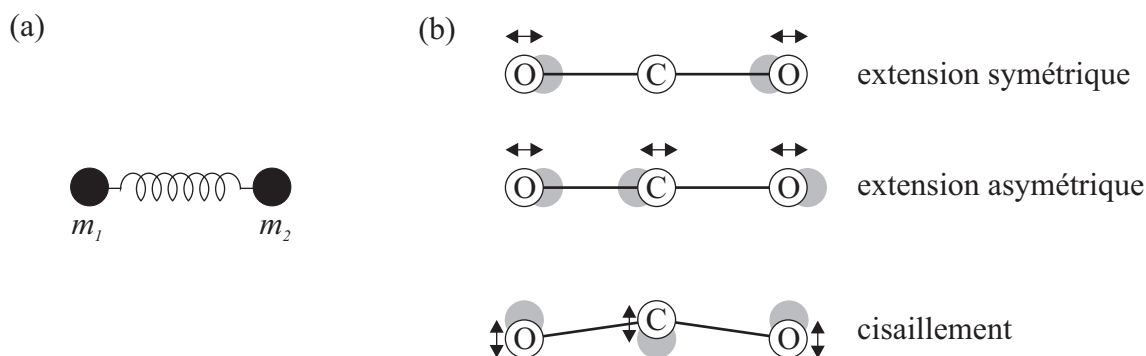


FIGURE 6.10 – (a) Modèle "mécanique" pour une molécule diatomique comme O_2 , N_2 ou CO : les deux atomes sont représentés par des masses m_i reliées par un ressort. (b) Différents modes de vibration pour la molécule CO_2 ; les énergies correspondantes sont données dans Tab. 6.2.

TABLE 6.2 – Energie des modes de vibration de la molécule CO_2 dans différentes unités : nombre d'onde $\tilde{\nu}$, fréquence ν , longueur d'onde λ et énergie E .

Mode	$\tilde{\nu}$ [cm^{-1}]	ν [THz]	λ [μm]	E [meV]
extension symétrique	1'388	41.6	7.2	172.0
extension asymétrique	2'349	70.4	4.2	291.2
cisaillement	667	20	15	82.7

Fig. 6.10(a) pour le cas d'une molécule diatomique comme O_2 , N_2 ou CO . Chaque atome est représenté par une masse m_i et le ressort représente la liaison entre ces atomes, qui peut être de deux types : ionique ou covalente. Dans une liaison ionique les charges positives et négatives demeurent séparées spatialement, donnant lieu à la formation de dipôles. Dans une liaison covalente les atomes partagent leurs charges et la molécule résultante n'a pas de moment dipolaire permanent. On comprend aisément que l'on peut associer une énergie spécifique à un état de vibration (ou mode) de la molécule, comme par exemple l'énergie cinétique correspondante. Ces mouvements mécaniques peuvent être variés, comme illustré Fig. 6.10(b) pour la molécule de CO_2 . Chaque vibration a une énergie particulière que l'on indique en général en unités spectroscopiques (nombre d'onde $\tilde{\nu}$) en cm^{-1} ; ces énergies sont très petites, comme indiqué sur Tab. 6.2, par exemple lorsqu'on considère les unités en meV. Elles correspondent donc à des longueurs d'onde très grandes, dans l'infrarouge.

On peut évidemment représenter chaque mode vibrationnel dans un diagramme d'énergie, semblable à celui utilisé pour les transitions électroniques. Le cas du CO_2 est illustré Fig. 6.11, où on a reporté les énergies par rapport à un état de référence que l'on a choisi arbitrairement. Chaque mode vibrationnel est caractérisé par des indices indiqués entre parenthèse, dont le nombre dépend du nombre de degrés de liberté de la molécule (le nombre de mouvements

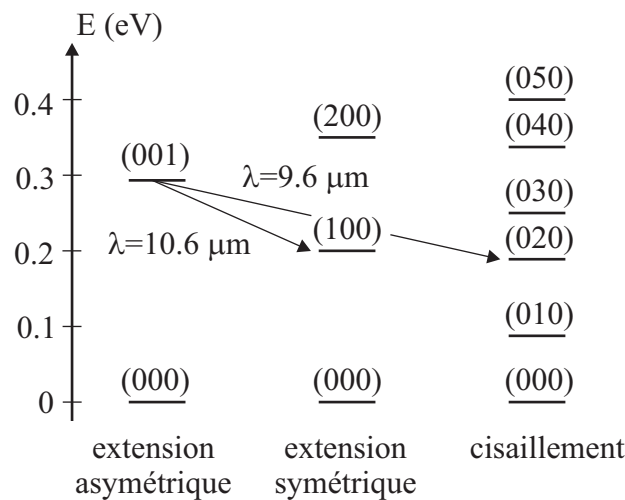


FIGURE 6.11 – Premiers niveaux vibrationnels de la molécule CO_2 ainsi que deux transitions utilisées pour les lasers CO_2 qui émettent dans l'infrarouge.

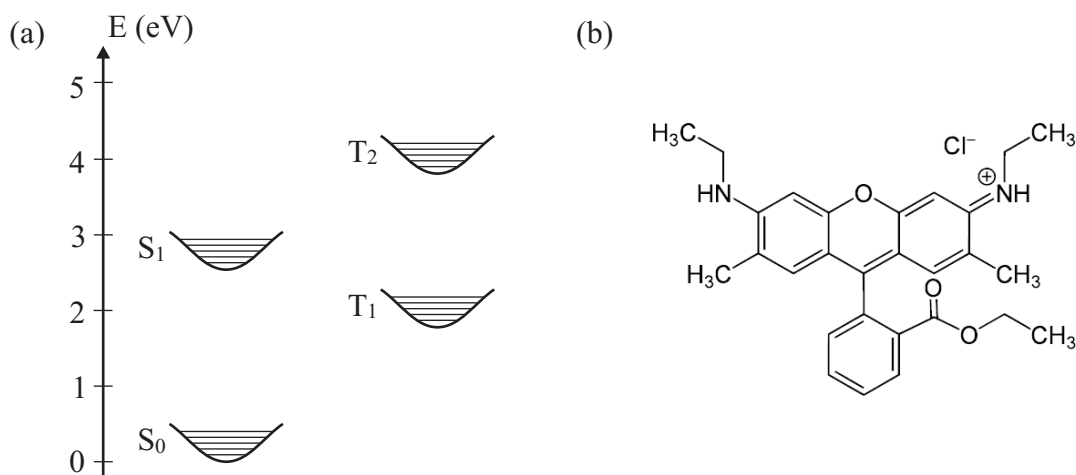


FIGURE 6.12 – (a) Niveaux énergétiques pour une molécule de Rhodamine 6G et (b) structure chimique de cette molécule d'après wikipedia.

possibles). On a aussi indiqué sur cette figure deux transitions possibles qui sont utilisées dans les lasers à CO_2 qui émettent dans l'infrarouge à de grandes longueurs d'ondes.

La spectroscopie Raman est une technique où l'on illumine une molécule à une longueur d'onde fixe et on mesure les différents modes de vibration de la molécule. On observe ainsi un spectre caractéristique de la molécule, qui permet de l'identifier de façon unique : chaque molécule ayant un spectre Raman spécifique qui dépend de ses différents modes de vibration ; on parle d'empreinte chimique de la molécule (*molecular fingerprint* en anglais). Bien que le signal Raman soit très faible, cette technique est largement utilisée pour analyser des composés chimiques inconnus.

Les colorants organiques sont des molécules compliquées dans lesquelles les atomes peuvent subir différents types de transitions (électroniques, vibrationnelles et rotationnelles). Il résulte de cette structure complexe un grand nombre de niveaux, comme indiqué Fig. 6.12(a). On distingue des niveaux singulet (S , en anglais *singlet state*) pour lesquels le spin (moment cinétique) de l'électron excité est antiparallèle aux spins du reste des électrons de la molécule et des niveaux triplet (T , en anglais *triplet state*) pour lesquels le spin de l'électron excité est parallèle. Les probabilités de transition dépendent fortement des spins relatifs des niveaux initiaux et finaux de l'électron, nous y reviendrons lorsque nous aborderons la fluorescence.

6.3.3 Solides et structure de bande

Nous avons été guidés par le principe d'exclusion de Pauli dans notre remplissage des niveaux électroniques, Fig. 6.9 : ce principe interdit que deux électrons aient exactement les mêmes nombres quantiques n , l , m et m_s . Ainsi, si on prend un atome donné et on l'associe à un deuxième atome identique, on ne peut pas avoir deux fois les mêmes niveaux électroniques

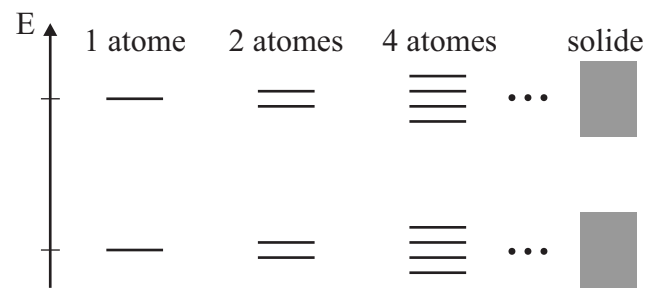


FIGURE 6.13 – Evolution des niveaux énergétiques depuis un atome isolé vers un solide. Les décalages successifs des niveaux énergétiques forment des bandes d'énergie pour le solide.

et les niveaux correspondants vont légèrement se décaler, comme illustré Fig. 6.13. Ce décalage, ou élargissement des niveaux se poursuit lorsqu'on ajoute des atomes supplémentaires. Ainsi, lorsque l'on a un solide, on n'a plus de niveaux individuels, mais des bandes d'énergie dans lesquelles les électrons se répartissent. Ces bandes d'énergie peuvent être plus ou moins occupées par des électrons, Fig. 6.14. On appelle bande de conduction la bande d'énergie la plus basse qui n'est pas occupée ou seulement partiellement occupée. On appelle bande de valence la bande complètement occupée d'énergie la plus haute. Ces deux bandes sont séparées par une bande dite interdite dont la largeur (en terme d'énergie) est nommée le bandgap. Dans un solide, les électrons remplissent toujours d'abord les bandes d'énergie les plus faibles. On distingue trois états de la matière : le métal dont la bande de conduction est partiellement occupée ; les semiconducteurs intrinsèques qui possèdent une bande de conduction inoccupée ; et les isolants qui sont semblables aux semiconducteurs mais possèdent un

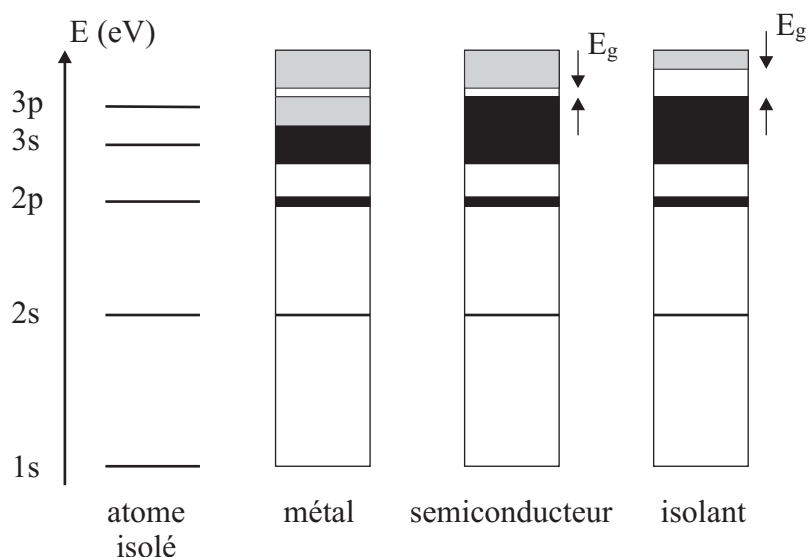


FIGURE 6.14 – Niveaux d'énergie d'un atome isolé et du solide correspondant dans différents états. Les bandes (ou parties de bandes) occupées par des électrons sont indiquées en noir ; les bandes (ou parties de bandes) inoccupées en gris.

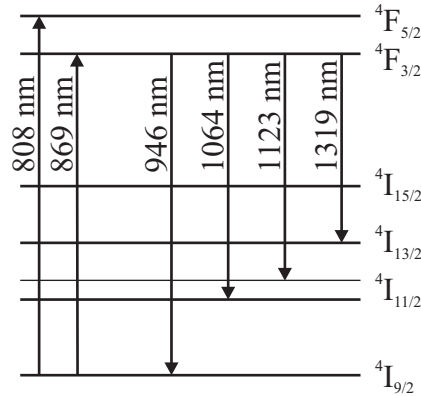


FIGURE 6.15 – Niveaux d'énergie pour un laser Nd :YAG.

bandgap important (typiquement plus de 3 eV), Fig. 6.14.

Lorsque l'on ajoute des impuretés dans un solide (par exemple des dopants métalliques dans un semiconducteur) la situation devient plus compliquée et il peut apparaître des transitions avec les niveaux électroniques des impuretés métalliques. Ainsi, pour fabriquer un laser, on utilise souvent un diélectrique transparents dans une certaine gamme de longueurs d'ondes comme matériau hôte pour recevoir le matériau actif, généralement sous forme de ions. Figure 6.15 illustre un exemple très répandu, le laser Nd :YAG. Un cristal d' $\text{Y}_3\text{Al}_5\text{O}_{12}$ (YAG) est dopé avec des ions Nd^{3+} . La structure électronique du ion permet de nombreuses transitions, Fig. 6.15. On remarque aussi sur cette figure la possibilité de pomper ce matériau dans les 800 nm ; comme l'absorption du Nd :YAG est relativement large, on peut pomper ce laser simplement avec une lampe. La longueur d'onde d'émission la plus commune pour un tel laser est $\lambda = 1064 \text{ nm}$. On peut ensuite générer des multiples de cette fréquence par doublage, triplage et quadruplage en utilisant des cristaux non-linéaires. On obtient ainsi des émissions à $\lambda = 532 \text{ nm}$, 355 nm et 266 nm .

6.4 Effets thermiques

Jusqu'à présent nous n'avons pas inclus la température dans nos considérations. Celle-ci joue cependant un rôle essentiel et détermine l'occupation des niveaux énergétiques. Lorsqu'un système possédant N niveaux d'énergie discrets E_i est à l'équilibre thermique à la température T , la probabilité qu'un niveau m soit occupé est donnée par la distribution de Boltzmann :

$$P(E_m) \propto \exp(-E_m/KT), \quad m = 1, 2, 3, \dots, \quad (6.28)$$

où $K = 1.38 \cdot 10^{-23} \text{ J/K}$ est la constante de Boltzmann. Ainsi la probabilité que les niveaux d'énergie élevée soient occupés décroît-elle comme une exponentielle, Fig. 6.16.

Si la température augmente, la distribution de Boltzmann se déplace et les niveaux d'énergie

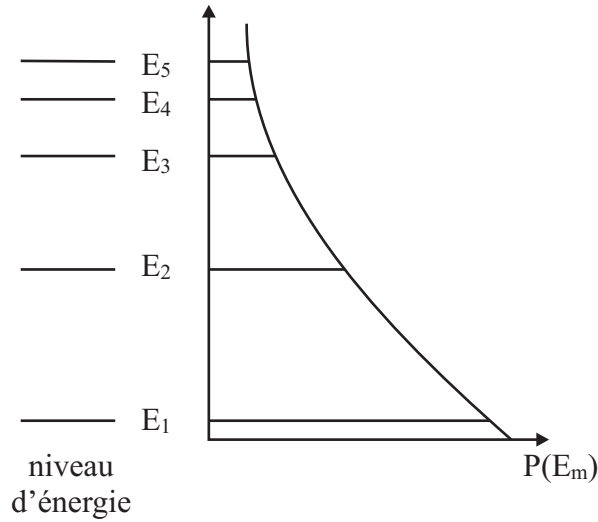


FIGURE 6.16 – La distribution de Boltzmann détermine le taux d’occupation de chaque niveau, en fonction de la température.

plus élevée commencent à se peupler. Supposons que l’on a un grand nombre d’atomes N et que N_m atomes occupent le niveau d’énergie E_m ; on a alors $N_m/N \simeq P(E_m)$ et si N_1 atomes occupent le niveau 1 et N_2 occupent le niveau 2, le rapport des populations vaut

$$\frac{N_2}{N_1} = \exp \left(-\frac{E_2 - E_1}{KT} \right). \quad (6.29)$$

Ce rapport dépend de la température : à $T = 0^\circ \text{K}$ tous les atomes sont dans le niveau d’énergie le plus bas (état fondamental). Lorsque la température augmente, la population des niveaux d’énergie supérieure augmente, mais à l’équilibre la population d’un niveau d’énergie est toujours plus grande que celle des niveaux d’énergie supérieure. Si on n’est pas à l’équilibre, il se peut que la population d’un niveau d’énergie élevé soit plus importante que la population d’un niveau d’énergie plus basse. On parle dans ce cas d’inversion de population, une situation qui forme la base du phénomène de lasage.

Dans un système quantique, comme par exemple un atome avec plusieurs électrons ou un semiconducteur, on doit aussi tenir compte du principe d’exclusion de Pauli qui empêche qu’un état énergétique soit occupé par plus d’un électron. La probabilité qu’un niveau d’énergie E soit occupé suit alors la distribution de Fermi–Dirac :

$$f(E) = \frac{1}{\exp((E - E_f)/KT) + 1}, \quad (6.30)$$

où E_f est l’énergie de Fermi. La valeur maximale de la distribution de Fermi–Dirac est $f(E) = 1$ (correspondant à un niveau totalement occupé) ; elle décroît lorsque l’énergie E croît et vaut $1/2$ lorsque $E = E_f$. Pour les valeurs de l’énergie supérieures à E_f , la distribution de Fermi–Dirac s’approche fort de celle de Boltzmann $P(E_m)$, comme indiqué Fig. 6.17.

Les transitions qui jouent un rôle en optique satisfont en général cette condition, $E \gg E_f$, et suivent donc la statistique de Boltzmann.

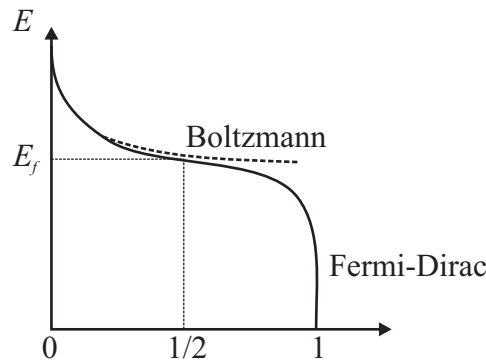


FIGURE 6.17 – La distribution de Fermi–Dirac, qui s’applique aux systèmes quantiques, a un comportement semblable à celle de Boltzmann pour $E \gg E_f$.

6.5 Luminescence, fluorescence et phosphorescence

La luminescence est l’émission de lumière par la matière ; elle trouve son origine dans la désexcitation des niveaux électroniques. Un électron se trouvant sur un niveau d’énergie élevé descend vers un niveau d’énergie plus bas. Ce faisant, il perd une certaine quantité d’énergie sous forme d’un photon. Il existe deux catégories de luminescence : la fluorescence qui est un processus rapide et la phosphorescence qui est un processus plus lent. Le taux d’émission de la fluorescence est de l’ordre de 10^8 s^{-1} alors que celui de la phosphorescence varie entre 1 et 10^3 s^{-1} .

La différence de rapidité entre ces deux processus d’émission optique s’explique par la différence des états électroniques en jeu, Fig. 6.18. Pour la fluorescence, l’électron de l’état excité 2 a un autre spin que l’électron se trouvant dans l’état fondamental 1 ; on parle d’état singulet pour cette paire d’électrons, Fig. 6.18(a). Ainsi, lorsque l’électron excité redescend sur l’état fondamental en émettant un photon, les deux électrons ont des spins différents, comme le requiert le principe d’exclusion de Pauli. Rien ne s’oppose à cette désexcitation qui peut donc se faire très rapidement. Par contre, pour la phosphorescence, les deux électrons ont le même spin. Si l’électron excité redescendait directement sur l’état fondamental les deux électrons se trouveraient dans une configuration interdite par le principe d’exclusion de Pauli. Cette désexcitation n’est donc pas favorable et prend un temps important. On note

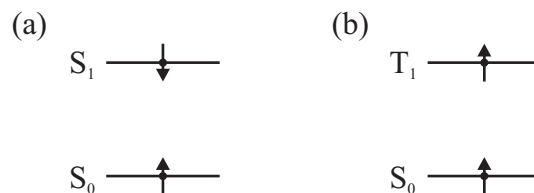


FIGURE 6.18 – Les différents états électroniques associés avec (a) la fluorescence (état singulet) et (b) la phosphorescence (état triplet).

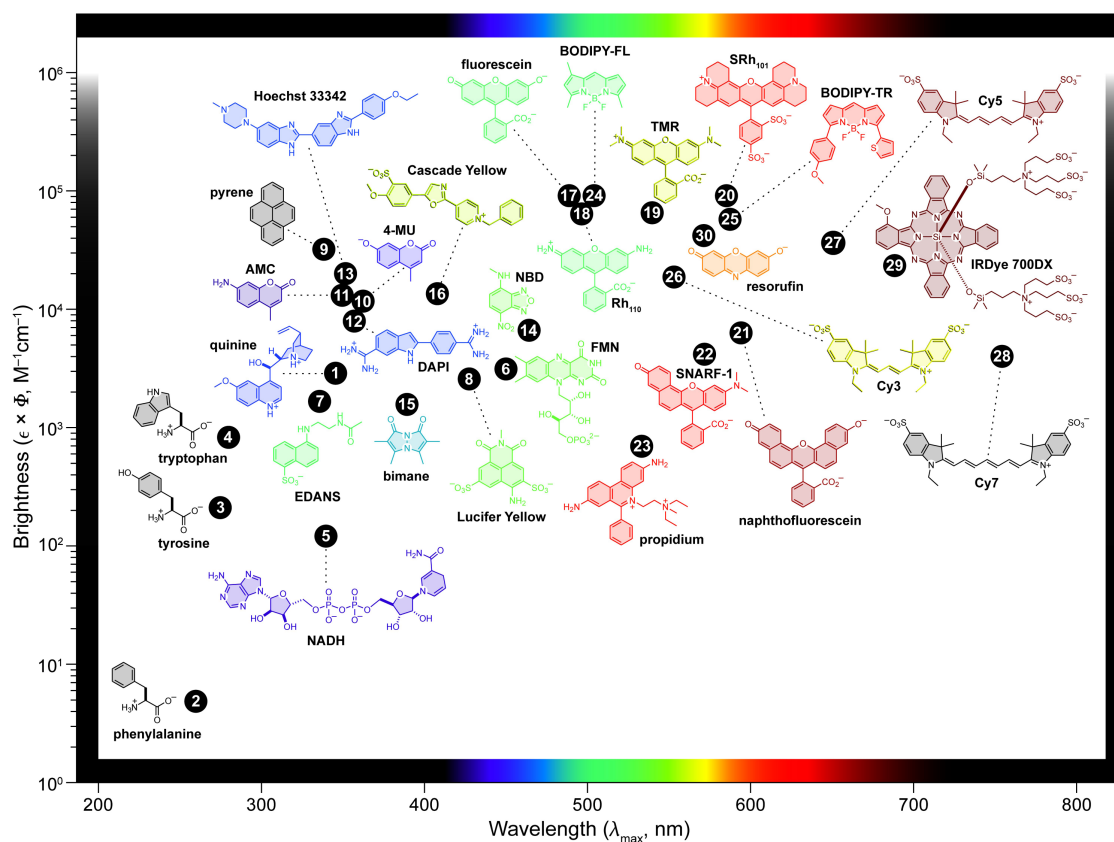


FIGURE 6.19 – Différentes molécules fluorescentes avec leur longueur d'onde d'absorption et leur luminosité ; reproduit de <http://pubs.acs.org/doi/full/10.1021/cb700248m>.

dans les diagrammes les états singulet avec la lettre majuscule S et les états triplets avec la lettre majuscule T. On ajoute souvent un indice qui indique l'ordre énergétique, par exemple 0 pour l'état fondamental et 1 pour l'état excité, Fig. 6.18. En général, on utilise la lettre T uniquement pour les états excités.

L'inverse du taux de fluorescence permet de définir un temps caractéristique que l'on nomme la durée de vie. Pour la fluorescence, cette durée de vie est de l'ordre de $\tau = 10$ ns. Ce temps correspond à la durée moyenne que met un électron excité pour redescendre à son niveau fondamental. Pour la phosphorescence, la durée de vie varie entre une milliseconde et une seconde, voire des temps plus longs pour les matériaux qui émettent une lueur dans le noir après avoir été exposés à la lumière.

Les molécules fluorescentes sont souvent des molécules aromatiques composées de cycles de carbone. En modifiant leur structure, il est possible d'ajuster la longueur d'onde de la fluorescence et d'obtenir des molécules fluorescentes sur l'entier du spectre visible, Fig. 6.19

Dans une molécule, la fluorescence ne peut se décrire avec un simple modèle ne comportant que deux niveaux énergétiques. Il faut prendre en compte la structure plus complexe de

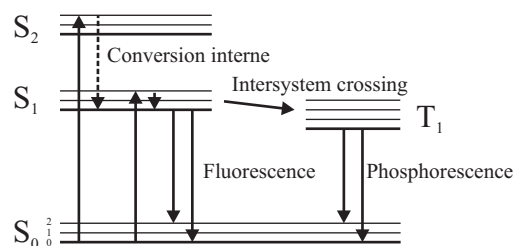


FIGURE 6.20 – Diagramme de Jablonski indiquant la fluorescence et la phosphorescence.

la molécule dans un diagramme de Jablonski, du nom du physicien polonais Aleksander Jablonski qui a fait des contributions essentielles à l'étude de l'influence de l'environnement (pression, polarisation) sur la fluorescence. Figure 6.20 illustre un diagramme de Jablonski typique. On y remarque l'état fondamental S_0 ainsi que les états excités S_1 et S_2 . On nomme ces états singulets puisque le spin de l'électron excité est opposé à celui de l'électron resté dans l'état fondamental. Pour un niveau électronique donné, il peut exister plusieurs niveaux vibrationnels avec des énergies légèrement différentes. C'est cette famille de niveaux proches les uns des autres qui donne lieu à la largeur spectrale observée pour l'absorption de la molécule. La fluorescence est émise depuis le niveau S_1 vers le niveau S_0 . Le photon qui a excité la molécule peut être absorbé par l'un des niveaux associé au niveau S_1 ou par l'un des niveaux associés au niveau S_2 . Dans ce dernier cas, une partie de l'énergie sera perdue par conversion interne au sein de la molécule, c'est à dire qu'elle génèrera des vibrations de la molécule (phonons) qui dissiperont de l'énergie. Lorsque l'électron redescend sur le niveau fondamental, il peut soit atterrir directement sur le niveau S_0 , soit atterrir sur l'un des niveaux vibrationnels 1 ou 2 ; son surplus d'énergie sera alors dissipé par les vibrations (phonons) de la molécule.

Figure 6.20 indique qu'un autre chemin est possible pour la désexcitation de l'électron. Ce chemin passe par le niveau T_1 , un niveau triplet dans lequel l'électron excité a le même spin que l'électron se trouvant à l'état fondamental. Nous avons vu que la transition vers l'état fondamental n'est alors pas très favorable et cette transition prend plus de temps. La transition entre le niveau excité S_1 et le niveau T_1 se nomme *intersystem crossing* car on passe d'un système singulet à un système triplet. Finalement, l'électron se désexcite vers le niveau fondamental S_0 ou un des niveaux vibrationnels associés en émettant de la phosphorescence.

Trois observations importantes peuvent être faites en étudiant Fig. 6.21. D'une part, l'énergie absorbée par la molécule est plus grande que l'énergie émise par fluorescence ou phosphorescence. Cela indique que la longueur d'onde d'absorption sera plus courte que la longueur d'onde d'émission, comme illustré Fig. 6.21 qui représente les spectres d'absorption et d'émission pour la molécule Rhodamine 6G. On nomme *Stokes shift* la différence spectrale entre ces deux spectres, du nom du physicien irlandais George Stokes qui a été le premier à mettre en évidence cette différence de longueur d'onde dans la fluorescence de certains minerais. D'autre part, les différents niveaux électroniques et vibrationnels existants dans la molécule produisent un élargissement des spectres d'absorption et d'émission qui peut être assez important (plus d'une centaine de nanomètres). Finalement, on remarque que les spectres

d'absorption et d'émission sont très sensibles au milieu dans lequel la mesure est faite. Ainsi, le solvant utilisé peut influencer de façon importante ces spectres.

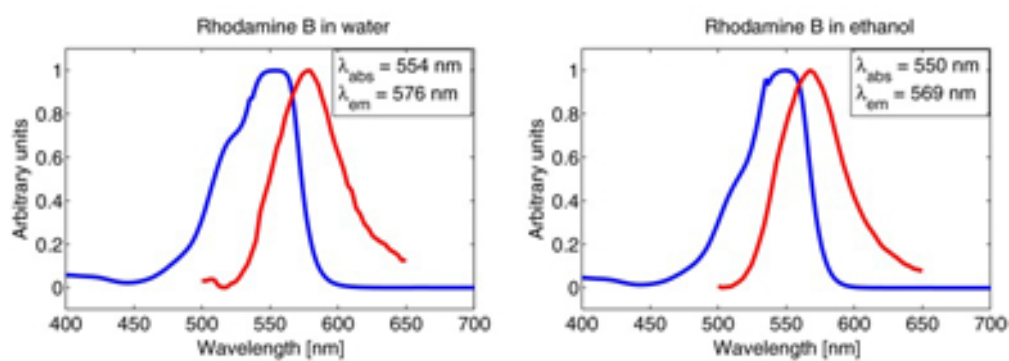


FIGURE 6.21 – Spectre d'absorption et d'émission pour la molécule Rhodamine 6G dans de l'eau (à gauche) et dans de l'éthanol (à droite).

6.6 Transitions optiques et coefficients d'Einstein

Dans l'interaction entre un atome (ou molécule ou solide) et un champ électromagnétique, un atome peut gagner de l'énergie en absorbant un photon ou perdre de l'énergie en émettant un photon. Nous considérerons deux niveaux et désignerons par 1 le niveau dans lequel l'atome est au repos et 2 l'état dans lequel l'atome est excité. On définit la différence d'énergie entre ces deux niveaux :

$$\Delta E = E_2 - E_1 = h\nu = \hbar\omega. \quad (6.31)$$

Il existe alors trois types d'interactions :

1) L'émission spontanée où un atome se trouvant dans l'état excité 2 émet un photon d'énergie ν et retombe dans l'état 1, Fig. 6.22(b). Le nombre de transitions $2 \rightarrow 1$ par unité de volume V et de temps s'écrit

$$p_{\text{sp}} = \frac{c}{V} \sigma(\nu), \quad (6.32)$$

où $\sigma(\nu)$ représente la section efficace de transition qui a pour unités cm^2 puisque p_{sp} doit avoir comme unités s^{-1} . La probabilité d'émettre un photon de façon spontanée pendant l'intervalle de temps Δt s'écrit $p_{\text{sp}} \Delta t$. Si on a un nombre N d'atomes excités, une fraction $\Delta N = N p_{\text{sp}} \Delta t$ se désexcitera pendant l'intervalle Δt . On peut donc écrire $dN/dt = -p_{\text{sp}} N$ et la population d'atomes excités évolue selon

$$N(t) = N(0) e^{-p_{\text{sp}} t}, \quad (6.33)$$

comme indiqué Fig. 6.23. Cette décroissance exponentielle se fait avec une constante de temps $1/p_{\text{sp}}$.

2) L'absorption où un photon est incident sur un système se trouvant dans l'état 1 et le fait passer dans l'état 2, Fig. 6.22(a). Le nombre de transitions $1 \rightarrow 2$ par unité de volume V et de temps suit une loi similaire à l'émission spontanée :

$$p_{\text{ab}} = \frac{c}{V} \sigma(\nu). \quad (6.34)$$

Si nous avons n photons, la probabilité d'absorption d'un photon est multipliée par autant :

$$P_{\text{ab}} = n \frac{c}{V} \sigma(\nu). \quad (6.35)$$

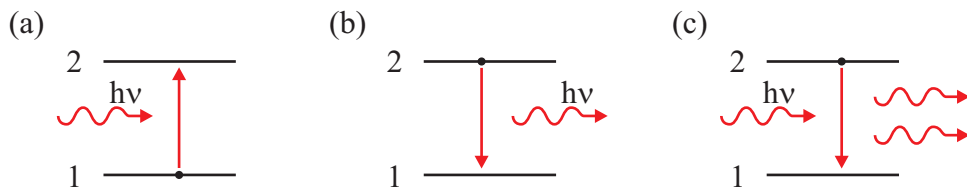


FIGURE 6.22 – Les trois types de transitions optiques : (a) absorption, (b) émission spontanée et (c) émission stimulée.

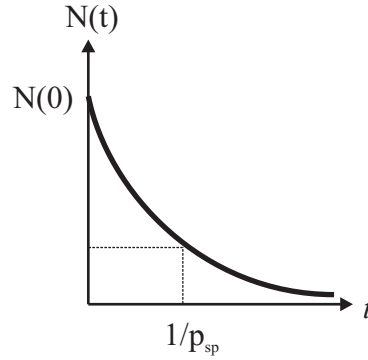


FIGURE 6.23 – L'émission spontanée résulte en une décroissance exponentielle du nombre d'atomes excités avec une constante de temps $1/p_{sp}$.

3) L'émission stimulée où un photon incident sur un atome excité le fait revenir au niveau fondamental en émettant un photon, Fig. 6.22(c). Ce deuxième photon a la même énergie et est en phase avec le photon incident ; on a donc une émission cohérente. On a une relation semblable pour la densité de probabilité d'émission stimulée,

$$p_{st} = \frac{c}{V} \sigma(\nu) ; \quad (6.36)$$

et si l'on a n photons, la probabilité d'émission stimulée est aussi multipliée par autant :

$$P_{st} = n \frac{c}{V} \sigma(\nu) . \quad (6.37)$$

Comme $P_{st} = P_{ab}$, on introduit $W_i = P_{st} = P_{ab}$ pour représenter la densité de probabilité d'émission stimulée ou d'absorption.

Puisqu'à la fois l'émission spontanée et l'émission stimulée peuvent exister, la probabilité totale pour un atome d'émettre un photon s'écrit,

$$p_{sp} + P_{st} = (n + 1) \frac{c}{V} \sigma(\nu) . \quad (6.38)$$

La section efficace de transition $\sigma(\nu)$ caractérise l'interaction de l'atome avec la lumière. Lorsqu'on l'intègre sur l'entier du spectre, on obtient la force de l'oscillateur correspondant,

$$S = \int_0^\infty d\nu \sigma(\nu) , \quad (6.39)$$

qui a pour unités cm^2Hz . Cette fonction a généralement la forme d'une Lorentzienne, comme indiqué Fig. 6.24(a). Sa surface est proportionnelle à l'interaction entre l'atome et la lumière, parfois on définit une fonction normalisée $g(\nu) = \sigma(\nu)/S$, la ligne spectrale dont la surface vaut 1, comme indiqué Fig. 6.24(b). Cette dernière fonction permet de définir une grandeur importante, sa largeur à mi-hauteur (en anglais *Full width at half maximum* FWHM) :

$$\Delta\nu = 1/g(\nu_0) , \quad (6.40)$$

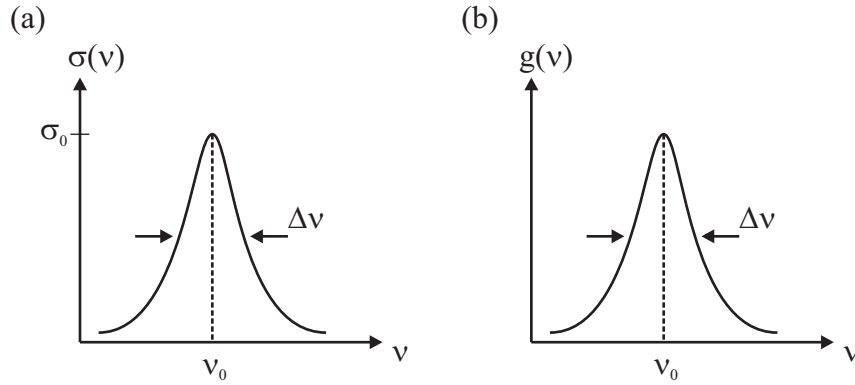


FIGURE 6.24 – (a) Section efficace de transition (la surface sous la courbe vaut S , Eq. (6.39)) et (b) ligne spectrale (la surface sous la courbe vaut 1).

comme indiqué Fig. 6.24(b).

Considérons maintenant un atome illuminé par de la lumière *monochromatique* de fréquence ν et d'intensité I . Le flux de photons,

$$\phi = \frac{I}{h\nu}, \quad (6.41)$$

indique le nombre de photons par unité de temps et de surface. Nous souhaitons déterminer les densités de probabilités $W_i = P_{st} = P_{ab}$ pour l'absorption et l'émission stimulée. On détermine le nombre de photons n interagissant avec l'atome en considérant un cylindre de base A et volume $V = cA$ parallèle à la direction de propagation de la lumière. Ce cylindre contient $n = \phi A = \phi V/c$ photons. Ainsi,

$$\phi = n \frac{c}{V}; \quad (6.42)$$

et W_i s'obtient en introduisant cette valeur dans Eq. (6.35) :

$$W_i = \phi \sigma(\nu). \quad (6.43)$$

On remarque que $\sigma(\nu)$ est la constante de proportionnalité entre le flux de photon et la probabilité d'une transition induite par ces photons, d'où son nom de section efficace.

Si l'illumination est polychromatique avec une densité d'énergie spectrale $\rho(\nu)$ distribuée de part et d'autre de la transition optique centrée à la fréquence ν_0 , on doit définir un nombre moyen de photons $\bar{n}$ à la fréquence de transition,

$$\bar{n} = \frac{c^3}{8\pi h\nu_0^3} \rho(\nu_0), \quad (6.44)$$

qui nous permet d'écrire

$$W_i = \frac{\bar{n}}{t_{sp}}, \quad (6.45)$$

où on a introduit la durée de vie t_{sp} pour l'émission spontanée. Les coefficients qui apparaissent dans Eq. (6.44) sont liés au caractère dipolaire de la transition optique associée avec un système à deux niveaux.

Einstein a étudié l'interaction des atomes avec de la lumière polychromatique et obtenu deux coefficients A et B qui portent son nom et décrivent la probabilité d'émission spontanée et d'émission stimulée :

$$A = \frac{1}{t_{\text{sp}}}, \quad (6.46)$$

$$B = \frac{\lambda^3}{8\pi h t_{\text{sp}}}, \quad (6.47)$$

où on a introduit la longueur d'onde dans le milieu correspondant à la fréquence de la transition optique : $\lambda = c/\nu_0$. On a la proportion suivante entre les coefficients d'Einstein :

$$\frac{B}{A} = \frac{\lambda^3}{8\pi h}. \quad (6.48)$$

On remarque que le taux d'émission spontanée pour un atome excité dans l'état 2 est constant et vaut $A = 1/t_{\text{sp}}$ alors que le taux d'émission stimulée occasionné par une source polychromatique vaut $W_i = B\rho(\nu_0)$ et est proportionnel à la densité d'énergie lumineuse à la fréquence ν_0 de la transition atomique considérée. Ces deux taux sont égaux lorsque $\rho(\nu_0) = A/B = 8\pi h/\lambda^3$.

Si l'on observe le spectre de la radiation émise par un ensemble d'atomes tous semblables, on s'attend à n'observer qu'un pic très mince de fréquence ν_0 correspondant à la transition entre les deux niveaux concernés. Dans la réalité on n'observe pas un tel rayonnement monochromatique, mais plutôt un pic $g(\nu)$ d'une certaine largeur, comme indiqué Fig. 6.25(a). Ceci indique que le niveau énergétique d'émission 2 n'est pas infiniment mince mais a une certaine largeur : on parle de bande d'émission, Fig. 6.25(b) et on appelle $g(\nu)$ la largeur de la ligne d'émission.

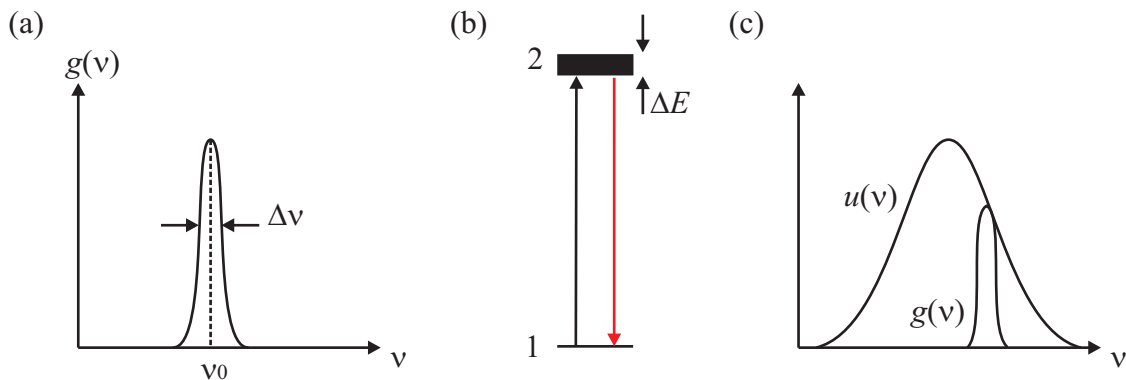


FIGURE 6.25 – Pour un ensemble d'atomes ou de molécules, la largeur de ligne $\Delta\nu$ est finie (a) et correspond à la largeur du niveau énergétique 2 (b). Un tel système peut absorber une bande d'énergie du champ incident $u(\nu)$ (c).

6.7 Colorimétrie

Nous terminons ce chapitre par une digression sur la colorimétrie, l'art de rendre les couleurs par la lumière. Ce sujet est particulièrement important pour le développement de nouvelles sources d'éclairage qui produisent un rendu des couleurs satisfaisant pour l'oeil humain. Celui-ci possède des récepteurs (cônes) sensibles respectivement aux longueurs d'ondes longues (X), moyennes (Y) et courtes (Z) du spectre visible. La perception d'une couleur peut donc s'exprimer par la combinaison de ces trois composantes que l'on appelle les valeurs tristimulus.

Afin de déterminer la perception d'une couleur donnée, la Commission Internationale d'Eclairage (CIE) a défini en 1931 une relation entre cette couleur et ses valeurs tristimulus. Cette relation définit un espace de couleurs appelé CIE 1931. Dans cet espace, chaque couleur est représenté par un triple XYZ , où X correspond à la composante rouge, Y à la composante verte et Z à la composante bleue. Il faut remarquer que chaque paramètre ne représente pas réellement une couleur, mais plutôt permet de ramener n'importe quelle couleur à la perception que s'en fait l'oeil. Ainsi, deux sources lumineuses peuvent sembler avoir la même couleur à l'oeil, bien que leurs composantes spectrales soient différentes. Par exemple on peut fabriquer une diode organique blanche (en anglais *organic light emitting diode*, OLED) en combinant trois sources (rouge, vert et bleu) ou en combinant deux sources (orange et bleu). Si la lumière produite par chaque diode apparaît la même à l'oeil humain, leurs valeurs XYZ seront les mêmes.

La détermination des paramètres XYZ est relativement complexe puisqu'elle doit tenir compte de la sensibilité de l'oeil. La CIE a cependant établi une procédure permettant de calculer les valeurs XYZ associées à un spectre d'intensité quelconque $I(\lambda)$ en fonction de la longueur d'onde λ . Cette procédure est basée sur la décomposition du spectre en les trois courbes $\bar{x}(\lambda)$, $\bar{y}(\lambda)$ et $\bar{z}(\lambda)$ illustrées Fig. 6.26. On dit que ces courbes correspondent à un observateur standard. Les valeurs X , Y et Z s'obtiennent en projetant le spectre étudié $I(\lambda)$ sur chaque courbe :

$$\begin{aligned} X &= \int_0^\infty d\lambda I(\lambda) \bar{x}(\lambda), \\ Y &= \int_0^\infty d\lambda I(\lambda) \bar{y}(\lambda), \\ Z &= \int_0^\infty d\lambda I(\lambda) \bar{z}(\lambda). \end{aligned} \tag{6.49}$$

Il existe évidemment d'autres espaces colorimétriques, comme le RGB (rouge, vert et bleu) ou CYMK (cyan, jaune, magenta et noir), pour représenter une couleur par des nombres. L'avantage de la norme CIE 1931 est son adéquation avec la perception de l'oeil humain.

Cette représentation à trois paramètres peut être réduite à deux paramètres seulement x et

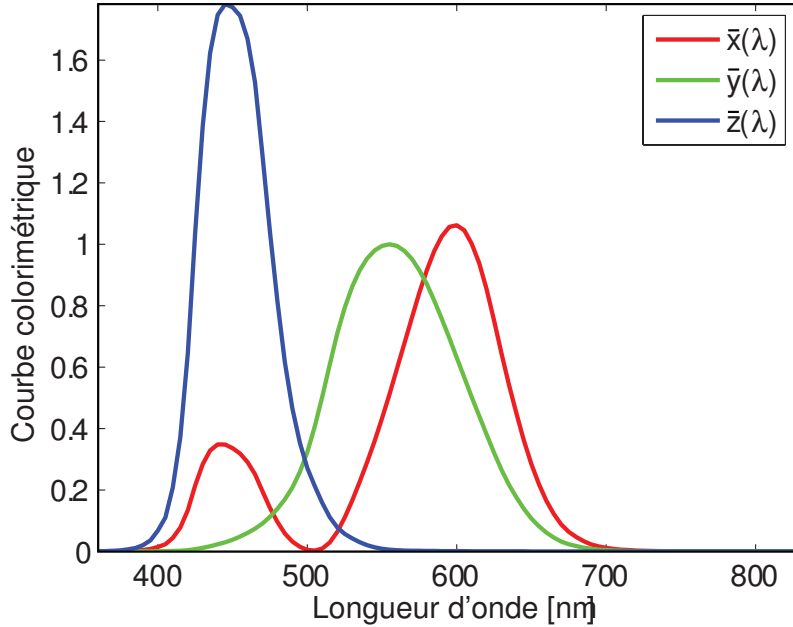


FIGURE 6.26 – Spectres $\bar{x}(\lambda)$, $\bar{y}(\lambda)$ et $\bar{z}(\lambda)$ correspondant à l'observateur standard déterminé par la norme CIE 1931.

y en utilisant les relations suivantes :

$$\begin{aligned} x &= \frac{X}{X + Y + Z}, \\ y &= \frac{Y}{X + Y + Z}, \\ z &= \frac{Z}{X + Y + Z} = 1 - x - y. \end{aligned} \tag{6.50}$$

Les paramètres x et y prennent des valeurs entre 0 et 1. Ils permettent de représenter de façon unique une couleur dans le diagramme de chromaticité, Fig. 6.27. Le pourtour du diagramme s'appelle le spectrum locus et correspond au spectre formé par les couleurs monochromatiques pures, dont les longueurs d'ondes sont indiquées en nm. Les points intérieurs au diagramme correspondent à des couleurs obtenues par mélange de couleurs pures. L'ensemble du diagramme de chromaticité s'appelle le gamut de couleur et correspond à l'ensemble des couleurs perceptibles par l'œil humain.

Si on choisit deux points A et B dans le diagramme de chromaticité, toutes les couleurs qui peuvent être obtenues en mélangeant ces deux couleurs en proportions diverses se trouvent sur la ligne AB reliant ces deux points. Si on prend trois points A , B et C dans le diagramme, toutes les couleurs possibles que l'on peut obtenir par mélange de ces couleurs se trouvent à l'intérieur du triangle ABC . Ceci se généralise au mélange de $N > 3$ couleurs. Cependant,

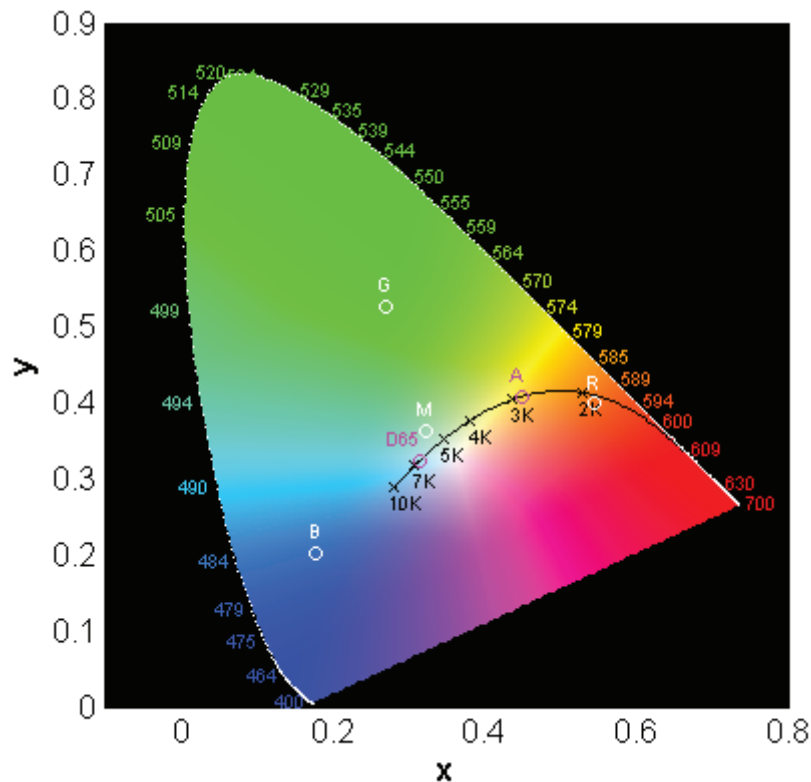


FIGURE 6.27 – Diagramme de chromaticité permettant de représenter n'importe quelle couleur par le couple xy .

si on mélange deux couleurs A et B en proportions égales, la couleur résultante n'est en général pas au milieu du segment AB . Une distance dans le diagramme de chromaticité ne correspond donc pas au niveau de différence entre deux couleurs.

Figure 6.27 indique aussi que trois sources distinctes ne permettent pas de reproduire l'entier du gamut de couleur (géométriquement, le bord du gamut est incurvé et ne peut donc pas se représenter par un triangle). On a aussi noté sur le diagramme les 3 couleurs standard rouge, vert et bleu, dont le mélange permet de créer le blanc. La ligne noire partant du coin inférieur droit représente des émetteurs thermiques de température donnée T . Une source blanche idéale (i.e. donnant un bon rendu des couleurs) correspond dans ce diagramme aux points se trouvant sur cette ligne entre les températures 3'000 K et 6'000 K.

Chapitre 7

Lasers

Le terme laser (de l'anglais *light amplification by stimulated emission of radiation*) décrit bien les mécanismes de fonctionnement de ce dispositif dans lequel la lumière est amplifiée par émission stimulée. Comme nous l'avons vu, une telle émission permet d'une part de multiplier le nombre de photons (un photon incident, deux photons en sortie) et de créer des photons cohérents, i.e. des photons qui ont tous la même phase.

Si le phénomène d'émission stimulée a été décrit la première fois par Einstein en 1917, un dispositif utilisant l'émission stimulée n'a été fabriqué qu'en 1954 par Townes, Gordon et Zeiger. Ce dispositif ne fonctionnait pas aux longueurs d'ondes optiques, mais aux fréquences micro-ondes et portait le nom de maser (*microwave amplification by stimulated emission of radiation*). En 1958, Schalow et Townes ont postulé théorétiquement la possibilité d'étendre ce dispositif aux fréquences optiques et le premier laser a été construit en 1960 par Maiman à l'aide d'un cristal de rubis. La même année Javan et Benett ont construit le premier laser à gaz. En 1962, trois groupes de recherche de General Electric, IBM et MIT ont publié quasi simultanément les premiers résultats sur les diodes laser.

Pour réaliser un laser, il faut réunir trois conditions comme illustré Fig. 7.1 :

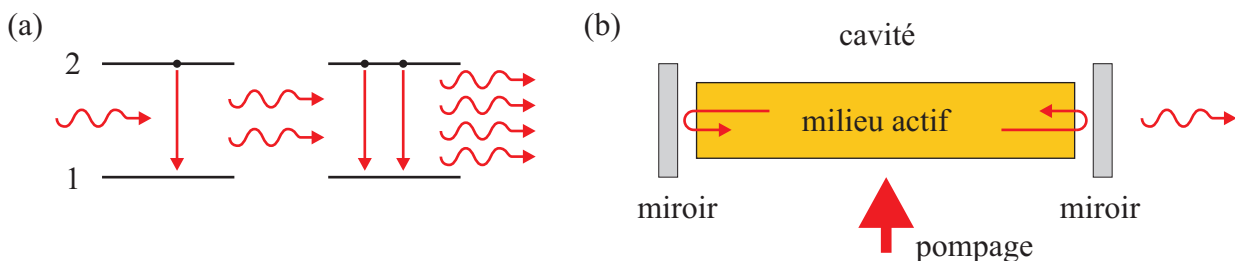


FIGURE 7.1 – (a) Principe de l'amplification laser par émissions stimulées successives. (b) Un laser se compose d'un matériau actif, d'un système de feedback (cavité) et d'un moyen d'exciter le matériau actif.

- Un matériau actif, dans lequel on peut créer des transitions stimulées.
- Un moyen d'exciter ce matériau, en sorte que les électrons se trouvent dans le niveau excité.
- Un système de feedback, permettant de garder les photons dans le laser afin qu'ils génèrent d'autres photons par émission stimulée.

D'un point de vue pratique, un laser se compose d'une cavité, d'un milieu actif et d'un moyen de pomper le milieu actif pour créer une inversion de population (une majorité de porteurs de charges se trouve dans un état excité).

7.1 Cavité laser

Nous considérons ici les modes d'une cavité formée par deux miroirs parallèles parfaits (sans pertes) séparés par une distance d , Fig. 7.2. Cette cavité se nomme un étalon Fabry-Perot, du nom de ses deux inventeurs. Nous cherchons les modes de ce résonateur, i.e. des solutions $U(z)$ de l'équation d'Helmholtz en sorte que l'amplitude du champ électrique soit nulle en $z = 0$ et $z = d$. En utilisant Fig. 7.2 on remarque qu'il y a dans la cavité une onde harmonique se propageant vers la droite et une se propageant vers la gauche ; leur superposition donne lieu à une onde stationnaire,

$$U(z) = \sin(kz), \quad (7.1)$$

où on doit imposer $kd = q\pi$ ($q = 1, 2, \dots$) pour satisfaire la condition que l'amplitude du champ électrique s'annule en $z = d$ (la condition que la champ s'annule en $z = 0$ est déjà satisfaite par Eq. 7.1). Seules les valeurs,

$$k_q = q \frac{\pi}{d}, \quad (7.2)$$

sont donc possibles pour le nombre d'onde k , avec q un nombre entier. On a donc dans le résonateurs des modes d'amplitude complexe,

$$U(z) = A_q \sin(k_q z), \quad (7.3)$$

où A_q est une constante et q est le numéro du mode. Nous avons jusqu'à présent négligé la partie temporelle du mode. En l'introduisant, le mode s'écrit

$$u(z, t) = \Re \{ U(z) \exp(j2\pi\nu t) \} = \Re \{ A_q \sin(k_q z) \exp(j2\pi\nu_q t) \}, \quad (7.4)$$

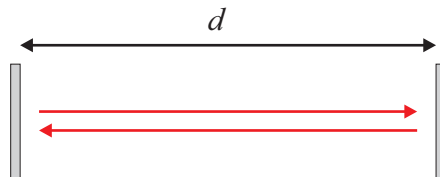


FIGURE 7.2 – Cavité laser : dans la cavité il y a superposition d'ondes se propageant dans les deux directions et donnant lieu à une onde stationnaire.

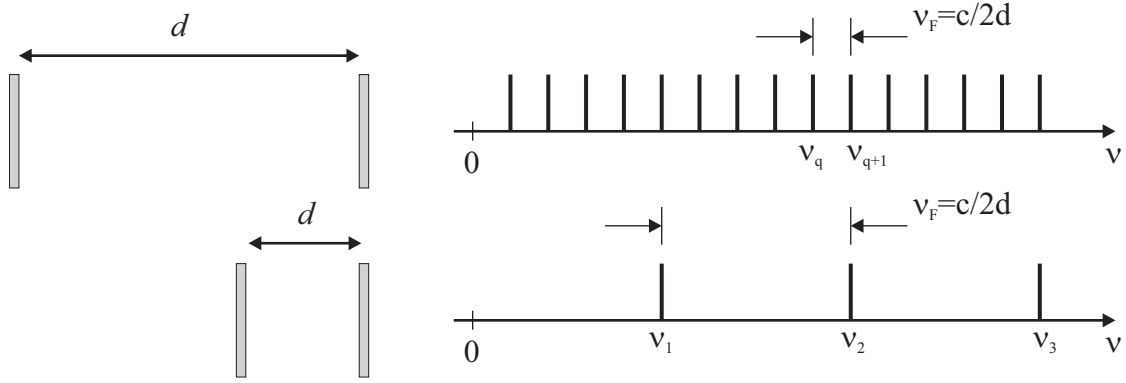


FIGURE 7.3 – Les différentes fréquences de résonance pour une cavité Fabry–Perot. L’espacement ν_F entre les modes successifs dépend de la largeur d de la cavité.

où nous avons introduit la fréquence ν_q du mode q :

$$\nu_q = \frac{ck_q}{2\pi} = q \frac{c}{2d}, \quad q = 1, 2, \dots \quad (7.5)$$

Un résonateur possède donc une série de fréquences de résonances ν_q , comme illustré Fig. 7.3. Deux fréquences voisines sont séparées par une différence de fréquence constante

$$\nu_F = \frac{c}{2d}. \quad (7.6)$$

On remarquera que c est la vitesse de la lumière dans le matériau considéré (ici le matériau dont est rempli le résonateur).

Un résonateur construit avec des miroirs parfaits n’a qu’une utilité limitée pour réaliser un laser, puisque la lumière ne peut pas en sortir ! Nous devons donc maintenant considérer un résonateur avec des miroirs ayant une réflectivité plus petite que unité. D’une façon générale, on introduit un facteur de perte r qui tient compte non seulement de la réflexion non-parfaite sur les miroirs, mais aussi des pertes qui peuvent se produire par absorption dans le matériau se trouvant dans la cavité ; si celles-ci sont nulles r représente la réflectivité des miroirs. Pendant un aller-retour dans la cavité, l’intensité lumineuse décroît d’un facteur $|r|^2$ avec $|r| < 1$.

On introduit la finesse $\mathcal{F}$ de la cavité,

$$\mathcal{F} = \frac{\pi\sqrt{|r|}}{1 - |r|}. \quad (7.7)$$

On remarque que pour une cavité avec des miroirs parfaits $\mathcal{F} = \infty$ puisque $|r| = 1$. Alors que pour une cavité avec des miroirs parfaits on avait un peigne de fréquences possibles, l’intensité lumineuse dans une cavité avec pertes suit l’équation,

$$I = \frac{I_{\max}}{1 + (2\mathcal{F}/\pi)^2 \sin^2(\pi\nu/\nu_F)}, \quad I_{\max} = \frac{I_0}{(1 - |r|)^2}, \quad (7.8)$$

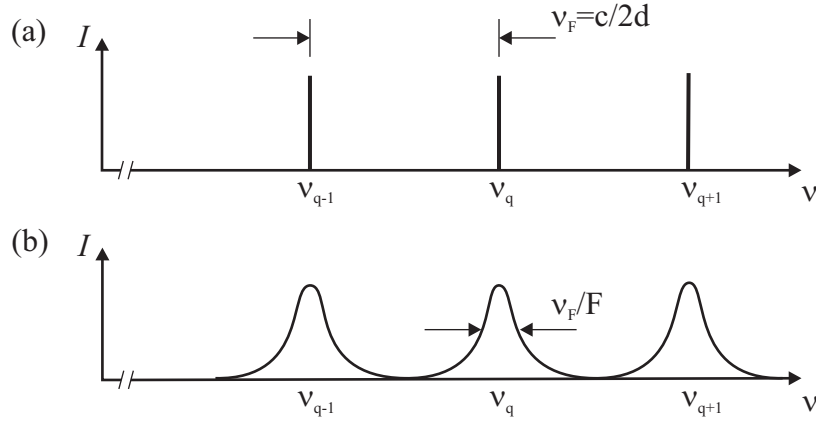


FIGURE 7.4 – Comparaison des intensités dans un résonateur (a) sans aucune perte (miroirs parfaits et pas d'absorption dans le matériau) et (b) réaliste avec des pertes.

où I_0 est l'intensité de l'onde incidente dans la cavité.

Figure 7.4 illustre les réponses spectrales pour une cavité sans perte ($\mathcal{F} = \infty$) et pour une cavité avec perte. Lorsque la finesse est grande ($\mathcal{F} \gg 1$), la réponse du résonateur est très piquée autour de multiples de la fréquence des modes. En plus de l'espacement entre les modes ν_F , on introduit la largeur spectrale des modes,

$$\delta\nu = \frac{\nu_F}{\mathcal{F}}. \quad (7.9)$$

On remarque que la largeur spectrale est inversement proportionnelle à la finesse de la cavité. Lorsque les pertes dans la cavité augmentent, la finesse diminue et la largeur spectrale augmente.

On introduit aussi le facteur de qualité Q d'un résonateur qui correspond à l'énergie emmagasinée divisée par l'énergie perdue par cycle :

$$Q = 2\pi \frac{\text{Energie emmagasinée}}{\text{Energie perdue par cycle}}. \quad (7.10)$$

Pour un résonateur avec une fréquence de résonance ν_0 , on peut définir le facteur de qualité comme

$$Q = \frac{\nu_0}{\nu_F} \mathcal{F}. \quad (7.11)$$

7.2 Amplification laser

Dans cette section nous nous intéressons à une onde $E(z) \exp(j2\pi\nu t)$ d'amplitude complexe $E(z)$ et d'intensité $I(z) = |E(z)|^2/2\eta$ se propageant dans la direction z dans un milieu atomique à deux niveaux dont l'énergie entre les niveaux correspond à l'énergie $h\nu$ des photons.

La densité du flux de photons est $\phi(z) = I(z)/h\nu$ et $\eta = \sqrt{(\mu_0/\epsilon_0\epsilon_r)}$ est l'impédance du milieu que nous supposons non-magnétique. Nous supposons qu'il y a N_2 atomes dans l'état excité et N_1 atomes dans l'état fondamental. Alors qu'elle se propage, l'onde est amplifiée avec un coefficient de gain $\gamma(\nu)$ par unité de longueur et elle subit un changement de phase $\varphi(\nu)$ (aussi par unité de longueur), suite à son interaction avec les atomes du milieu (cette variation de phase est liée au caractère dispersif du matériau). Nous cherchons à établir des formules pour $\varphi(\nu)$ et $\gamma(\nu)$. Notons que si ce dernier paramètre est positif il y a amplification, alors que s'il est négatif il y a atténuation de l'onde.

Les trois types d'interactions étudiées Sec. 6.6 peuvent avoir lieu dans ce système : si un électron se trouve dans l'état fondamental, un photon peut être absorbé. Si un électron se trouve dans un état excité, un photon peut être émis par émission stimulée. Ces deux processus donnent lieu à l'atténuation, respectivement l'amplification, de l'onde incidente. Le troisième processus – l'émission spontanée – donne lieu à du bruit, puisque le photon ainsi émis n'a pas de relation de phase avec l'onde incidente. La densité de probabilité qu'un atome à l'état fondamental absorbe un photon est

$$W_i = \phi\sigma(\nu), \quad (7.12)$$

avec $\sigma(\nu) = (\lambda^2/8\pi t_{sp})g(\nu)$ la section efficace de transition à la fréquence ν , $g(\nu)$ la largeur de ligne normalisée, t_{sp} la durée de vie pour l'émission spontanée et λ la longueur d'onde dans le milieu. Comme nous l'avons vu Sec. 6.6, la densité de probabilité pour l'émission stimulée est la même que pour l'absorption et a comme unité s^{-1} ; il s'agit donc d'un taux de transitions.

La densité moyenne de photons absorbés par unité de temps et de volume est N_1W_i et la densité de photons générés par émission stimulée est N_2W_i . Par unité de temps et de volume, on a donc un gain net de photons $(N_2 - N_1)W_i = NW_i$, où on a introduit la différence de population N . Si N est positif, on a une inversion de population (plus d'atomes sont dans l'état excité que dans l'état fondamental). Si N est négatif, on a de l'absorption et le flux de photons diminue. Si $N = 0$ il ne peut pas y avoir d'interaction et le milieu est transparent.

Si à l'aide d'un moyen de pompage externe on crée une inversion de population, le flux de photons $\phi(z)$ augmente au fur et à mesure de la propagation dans la direction z . Comme chaque photon émis par émission stimulée peut contribuer à son tour à stimuler l'émission d'autres photons, le flux de photons augmente de façon exponentielle. Ce processus est illustré Fig. 7.5 : l'augmentation de photons $d\phi(z)$ sur une tranche dz est simplement

$$d\phi = NW_idz, \quad (7.13)$$

que nous pouvons mettre sous forme d'équation différentielle,

$$\frac{d\phi(z)}{dz} = \gamma(\nu)\phi(z), \quad (7.14)$$

avec le coefficient de gain

$$\gamma(\nu) = N\sigma(\nu) = N\frac{\lambda^2}{8\pi t_{sp}}g(\nu). \quad (7.15)$$

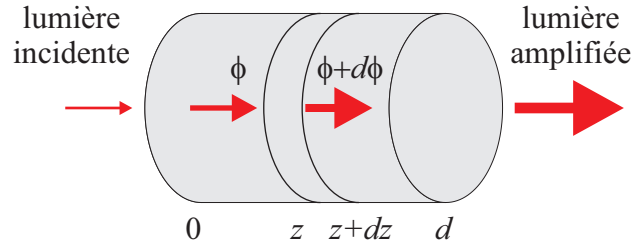


FIGURE 7.5 – Variation du flux de photons alors qu'il se propage dans un milieu actif présentant du gain ou de l'absorption.

La solution d'Eq. (7.14) est une fonction exponentielle,

$$\phi(z) = \phi(0) \exp(\gamma(\nu)z). \quad (7.16)$$

Comme l'intensité de photons vaut $I(z) = h\nu\phi(z)$, on peut récrire Eq. (7.16) comme

$$I(z) = I(0) \exp(\gamma(\nu)z); \quad (7.17)$$

et $\gamma(\nu)$ représente le gain en intensité par unité de longueur dans le milieu. Ce gain est proportionnel à la différence de population $N = N_2 - N_1$; en absence d'une inversion de population on a $N_2 < N_1$ et N est négatif et le gain est aussi négatif. La lumière sera alors atténuée exponentiellement par le milieu. A l'équilibre thermique, on a toujours $N_2 < N_1$ et le milieu ne peut produire une amplification. Pour une longueur de milieu d , le gain total $G(\nu)$ vaut

$$G(\nu) = \exp(\gamma(\nu)d). \quad (7.18)$$

Nous avons vu au chapitre précédent que la section efficace d'une transition atomique dépend de la fréquence, avec en général une forme Lorentzienne (cf. par exemple Fig. 6.24). En conséquence, le gain $\gamma(\nu)$ aura une dépendance en fréquence similaire. Ainsi, si la ligne spectrale de la transition a la forme d'une Lorentzienne

$$g(\nu) = \frac{\Delta\nu/2\pi}{(\nu - \nu_0)^2 + (\Delta\nu/2)^2}, \quad (7.19)$$

le coefficient de gain a aussi la forme d'une lorentzienne :

$$\gamma(\nu) = \gamma(\nu_0) \frac{(\Delta\nu/2)^2}{(\nu - \nu_0)^2 + (\Delta\nu/2)^2}, \quad (7.20)$$

comme illustré Fig. 7.6(a). Le coefficient $\gamma(\nu_0) = N(\lambda^2/4\pi^2 t_{sp} \Delta\nu)$ représente le gain pour la fréquence centrale ν_0 .

Au début de cette section, nous avons indiqué qu'en plus d'être amplifiée par le facteur $\gamma(\nu)$, l'onde de propageant dans le milieu atomique subit un changement de phase $\varphi(\nu)$. L'origine de ce changement de phase se trouve dans le développement que nous venons de faire : le gain dépend de la fréquence, donc l'indice de réfraction dépend de la fréquence et le milieu

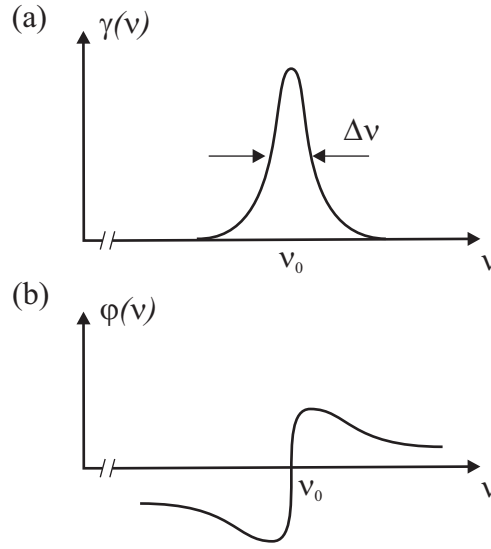


FIGURE 7.6 – (a) Coefficient de gain $\gamma(\nu)$ et (b) variation de phase $\varphi(\nu)$ pour un milieu actif décrit par une transition Lorentzienne.

est dispersif. L'onde subit donc un changement de phase qui dépend de la fréquence. Pour un milieu dans lequel la transition atomique et donc le gain sont de forme Lorentzienne, Eq. (7.19), le changement de phase est donné par,

$$\varphi(\nu) = \frac{\nu - \nu_0}{\Delta\nu} \gamma(\nu), \quad (7.21)$$

comme illustré Fig. 7.6(b).

Comme tout amplificateur, un laser nécessite une source externe d'énergie pour permettre l'amplification du signal incident. Cette source externe doit exciter les électrons en sorte de créer une inversion de population avec $N = N_2 - N_1 > 0$. Ce pompage peut prendre la forme optique d'un flash, la forme électrique d'une décharge, d'un faisceau d'électrons ou de l'injection directe de porteurs de charges ; il peut aussi prendre la forme d'une réaction chimique.

Pour maintenir le phénomène de lasage, il faut trouver un équilibre entre les populations N_1 et N_2 qui sont influencées par le pompage ainsi que les transitions radiatives et non-radiatives. On établit ainsi des équations de taux (en anglais *rate equations*) qui donnent l'évolution des populations dans le temps. On considère pour cela le schéma de Fig. 7.7(a) où nous nous concentrons sur les niveaux 1 et 2 qui ont τ_1 et τ_2 comme durées de vies des transitions vers des niveaux énergétiques plus bas. Comme le taux de désexcitation est inversement proportionnel à la durée de vie, la durée de vie globale du niveau 2 vaut $\tau_2^{-1} = \tau_{21}^{-1} + \tau_{20}^{-1}$. Ainsi, s'il existe des canaux multiples de désexcitation pour un niveau donné, la durée de vie de ce niveau diminue (i.e. la décroissance de la population de ce niveau est plus rapide). Pour la transition $2 \rightarrow 1$ il existe deux canaux possibles : une émission spontanée avec durée de vie t_{sp} et une émission non radiative avec durée de vie τ_{nr} ; on a donc $\tau_{21}^{-1} = t_{sp}^{-1} + \tau_{nr}^{-1}$.

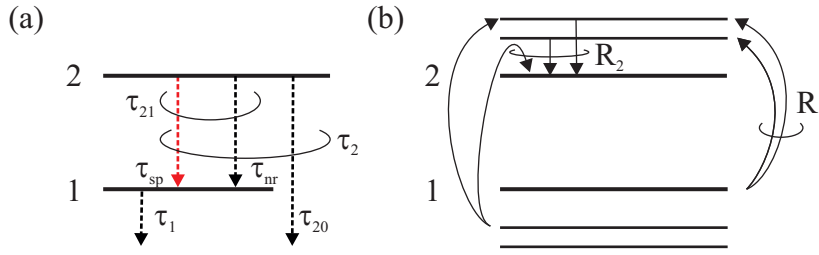


FIGURE 7.7 – (a) Niveaux énergétiques 1 et 2 ainsi que leurs durées de vie. (b) Evolution des populations des niveaux 1 et 2 à l'aide du pompage.

Si le système décrit Fig. 7.7(a) se trouve dans un état initial avec N_2 électrons sur le niveau 2 et N_1 électrons sur le niveau 1, il va se désexciter progressivement en suivant les durées de vie τ_2 et τ_1 . Un état d'équilibre pour les populations des niveaux 1 et 2 peut cependant être établi en amenant des électrons vers le niveau 2 avec un taux R_2 et en enlevant des électrons du niveau 1 avec le taux R_1 , comme illustré Fig. 7.7(b). Ce pompage permet de diminuer la population du niveau 1 et d'augmenter la population du niveau 2 comme illustré Fig. 7.8(a).

Considérons d'abord la situation où il n'y a pas de lumière dans le système : ni absorption directe entre $1 \rightarrow 2$ ni émission stimulée $2 \rightarrow 1$; les densités de populations pour les niveaux 1 et 2 sont alors décrites par le système d'équations différentielles suivant :

$$\frac{dN_2}{dt} = R_2 - \frac{N_2}{\tau_2} \quad (7.22)$$

$$\frac{dN_1}{dt} = -R_1 - \frac{N_1}{\tau_1} + \frac{N_2}{\tau_{21}}. \quad (7.23)$$

A l'équilibre, lorsque $dN_1/dt = dN_2/dt = 0$, on peut résoudre ce système d'équation pour obtenir la différence de population $N_0 = N_{\text{equilibre}} = N_2 - N_1$,

$$N_0 = R_2\tau_2 \left(1 - \frac{\tau_1}{\tau_{21}}\right) + R_1\tau_1. \quad (7.24)$$

On observe que l'on peut atteindre une grande différence de population à l'équilibre si le niveau 2 est pompé de façon intense (R_1 et R_2 grands) et décroît doucement (grande durée de vie τ_2). Si $R_1 = 0$ Eq. (7.24) devient $N_0 \approx R_2\tau_{sp}$.

Considérons maintenant le cas où il y a de la lumière dans le système et donc il y a émission stimulée $2 \rightarrow 1$ et absorption $1 \rightarrow 2$, deux processus caractérisés par le taux $W_i = \phi\sigma(\nu)$ (i.e. par la durée de vie W_i^{-1}), Fig. 7.8(b). Equations (7.22) et (7.23) doivent être étendues pour inclure ces phénomènes :

$$\frac{dN_2}{dt} = R_2 - \frac{N_2}{\tau_2} - N_2W_i + N_1W_i \quad (7.25)$$

$$\frac{dN_1}{dt} = -R_1 - \frac{N_1}{\tau_1} + \frac{N_2}{\tau_{21}} + N_2W_i - N_1W_i. \quad (7.26)$$

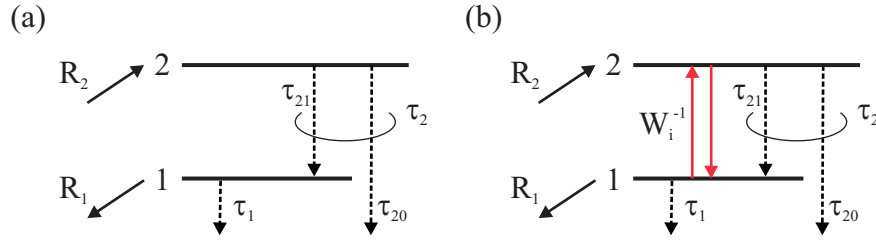


FIGURE 7.8 – Pompage des niveaux 1 et 2 ainsi que leur décroissance avec les durées de vie τ_1 et τ_2 . La population du niveau 2 est augmentée avec un taux R_2 et celle du niveau 1 est diminuée avec un taux R_1 . (a) En absence de lumière et (b) en présence de lumière donnant lieu à une absorption et une émission stimulée, toutes deux avec un taux W_i (i.e. une durée de vie W_i^{-1}).

Maintenant, la densité de population du niveau 2 peut décroître par émission stimulée $2 \rightarrow 1$ et augmenter par absorption $1 \rightarrow 2$. L'émission spontanée est incluse dans τ_{21} . À l'équilibre, la différence de population s'écrit maintenant

$$N = \left(\frac{N_0}{1 + \tau_s W_i} \right). \quad (7.27)$$

On introduit aussi une constante de temps de saturation,

$$\tau_s = \tau_2 + \tau_1 \left(1 - \frac{\tau_2}{\tau_{21}} \right), \quad (7.28)$$

qui est positive puisque $\tau_2 \leq \tau_{21}$. Ainsi, la différence de population en présence de radiation (i.e. avec émission stimulée et absorption) est plus petite en valeur absolue que sans ces phénomènes : $|N| \leq |N_0|$. Si la radiation est faible ($\tau_s W_i \ll 1$), on peut considérer $N \simeq N_0$. Lorsque l'interaction entre les deux niveaux augmente (le taux W_i d'absorption et d'émission stimulée devient grand) $N \rightarrow 0$ comme illustré Fig. 7.9.

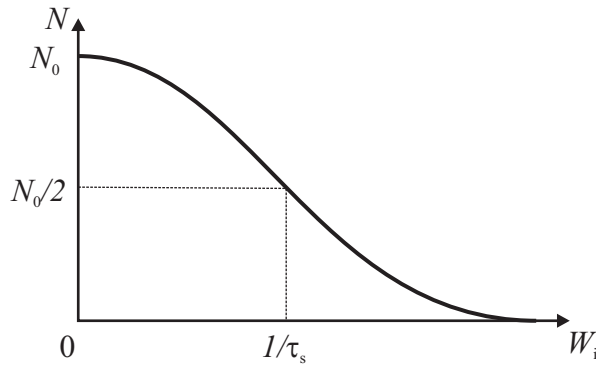


FIGURE 7.9 – Evolution de la différence de population $N = N_2 - N_1$ en fonction du taux d'absorption et d'émission stimulée W_i .

Nous avons vu que le milieu actif d'un laser a généralement de nombreux niveaux atomiques qui peuvent contribuer à la dynamique du système en recevant et en perdant des charges. Le phénomène du lasage peut cependant bien se décrire avec un système à seulement quatre niveaux, qui correspond par exemple au laser Nd :YAG (grenat d'yttrium-aluminium dopé au néodyme, Nd : Y₃Al₅O₁₂) qui est pompé optiquement au moyen de lampes flash ou de diodes laser et constitue un des types les plus communs de lasers. Figure 7.10 illustre un tel système à 4 niveaux avec son niveau fondamental 0 et les trois niveaux 1, 2 et 3. Noter que le lasage s'effectue entre les niveaux 1 et 2. A l'équilibre thermique, le niveau 1 n'a pratiquement pas de population pour autant que $E_1 \gg KT$. Le pompage s'effectue en amenant des charges dans le niveau 3 qui se situe à une énergie supérieure au niveau 2. La transition $3 \rightarrow 2$ a une durée de vie courte et il n'y a pas de population qui s'accumule au niveau 3. Le niveau 2 a par contre une grande durée de vie en sorte que la population s'y accumule. Le niveau 1 a une durée de vie courte. Il s'établit donc une inversion de population entre les niveaux 2 et 1.

A l'aide d'une source d'énergie extérieure, par exemple des photons avec une fréquence E_3/h , on pompe les atomes du niveau 0 vers le niveau 3 avec un taux R . Si le taux de transition $3 \rightarrow 2$ est suffisamment rapide, on peut le considérer comme instantané et ce pompage vers le niveau 3 est équivalent à pomper directement le niveau 2 avec un taux $R_2 = R$. On se retrouve dans la même situation que Fig. 7.8(b) décrite par Eqs. (7.25)–(7.26), avec $R_1 = 0$ puisque les charges ne sont pas extraites du niveau 1. La différence de population est donnée par Eq. (7.24) avec $R_1 = 0$:

$$N_0 = R_2 \tau_2 \left(1 - \frac{\tau_1}{\tau_{21}} \right). \quad (7.29)$$

De plus, dans la plupart des systèmes à quatre niveaux utilisés comme laser, la décroissance non-radiative $2 \rightarrow 1$ est négligeable ($t_{sp} \ll \tau_{nr}$) et $\tau_{20} \gg t_{sp} \gg \tau_1$. Il en résulte,

$$N_0 \approx R t_{sp}, \quad (7.30)$$

et

$$\tau_s \approx t_{sp}. \quad (7.31)$$

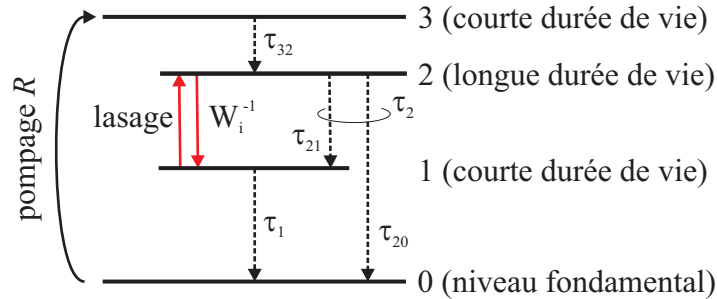


FIGURE 7.10 – Système de pompage à 4 niveaux. Le pompage augmente la population du niveau 3 et le lasage s'effectue entre les niveaux 2 et 1. Les transitions $3 \rightarrow 2$ et $1 \rightarrow 0$ sont très rapides.

On en déduit

$$N \approx \frac{Rt_{sp}}{1 + t_{sp}W_i}. \quad (7.32)$$

Notons que nous avons supposé que le taux de pompage R est indépendant de la différence de population $N = N_2 - N_1$, ce qui n'est pas toujours le cas comme nous allons le voir. On a cependant la relation suivante entre les différentes populations,

$$N_g + N_1 + N_2 + N_3 = N_a, \quad (7.33)$$

où on a introduit la population de l'état fondamental N_g (*ground state* en anglais) et la densité atomique totale du système N_a qui est constante.

Si le pompage utilise une transition $0 \rightarrow 3$ avec une densité de probabilité W , alors $R = (N_g - N_3)W$ et si les niveaux 1 et 3 ont une durée de vie courte, $N_1 \approx N_3 \approx 0$. Equation (7.33) donne alors $N_g + N_2 \approx N_a$ et $N_g \approx N_a - N_2 \approx N_a - N$. On en déduit pour le taux de pompage,

$$R \approx (N_a - N)W, \quad (7.34)$$

qui indique que le taux de pompage décroît linéairement avec la différence de population N et n'est donc pas constant. La raison en est que l'inversion de population établie entre les niveaux 2 et 1 réduit le nombre d'atomes pouvant être pompés. En introduisant Eq. (7.34) dans Eq. (7.32), on obtient

$$N \approx \frac{t_{sp}N_aW}{1 + t_{sp}W + t_{sp}W_i}. \quad (7.35)$$

Finalement, la différence de population d'écrit

$$N \approx \frac{N_0}{1 + \tau_s W_i}, \quad (7.36)$$

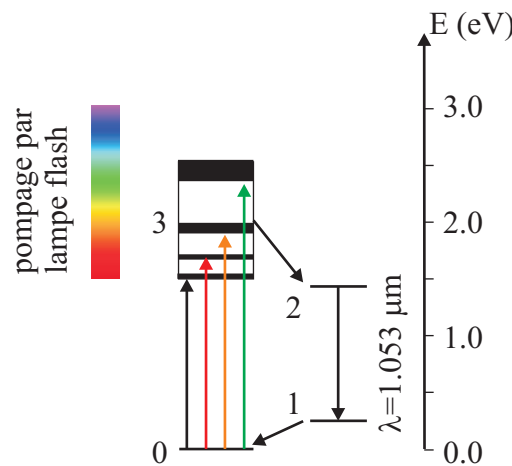


FIGURE 7.11 – Le système à quatre niveaux d'un laser Nd :YAG. Le pompage s'effectue à l'aide d'une lampe flash. Le niveau 3 est composé de quatre bandes d'absorption. L'émission se fait entre les niveaux 2 et 1 ; elle correspond à une longueur d'onde $\lambda_0 = 1.053 \mu\text{m}$.

avec les coefficients N_0 et τ_s donnés maintenant par

$$N_0 \approx \frac{t_{sp} N_a W}{1 + t_{sp} W}, \quad (7.37)$$

et

$$\tau_s \approx \frac{t_{sp}}{1 + t_{sp} W}. \quad (7.38)$$

En cas de pompage faible ($W \ll 1/t_{sp}$), $N_0 \approx t_{sp} N_a W$ est proportionnel à la densité de probabilité du pompage W et $\tau_s \approx t_{sp}$ et Eqs. (7.37)–(7.38) redonnent Eqs. (7.30)–(7.31). Au fur et à mesure que le pompage augmente, N_0 et τ_s diminuent.

Figure 7.11 illustre un tel système de pompage à 4 niveaux pour Nd :YAG. La transition associée au ion Nd^{3+} utilisée dans ce matériau a une énergie $\lambda_0 = 1.053 \mu\text{m}$ et le laser est pompé en utilisant une lampe flash puissante. Le niveau 1 a une énergie 0.24 eV au dessus de l'état fondamental ; cette valeur est bien plus grande que l'énergie thermique à température ambiante $KT \approx 0.026 \text{ eV}$ et la population du niveau 1 est donc négligeable. Le niveau 3 est en réalité composé de quatre bandes d'absorption entre 805 nm et 520 nm. Le ion excité décroît rapidement du niveau 3 vers le niveau 2 et reste à ce dernier pour un temps important ($t_{sp} = 375 \mu\text{s}$). Comme $\tau_1 \approx 300 \text{ ps}$ est très court, on retrouve bien la structure du système à 4 niveaux illustrée Fig. 7.10.

7.3 Système laser

Un système laser combine les deux éléments précédents – une cavité et un milieu d'amplification – pour réaliser l'amplification cohérente de la lumière sur une bande de fréquences très étroite. Dans un laser, on a un gain optique distribué sur toute la longueur de la cavité ; ce gain par unité de longueur $\gamma(\nu)$ détermine le taux d'amplification de la densité de flux de photons ϕ . Lorsque le flux de photons est petit (donc avant que le système ne lase), on introduit le coefficient de gain $\gamma_0(\nu)$,

$$\gamma_0(\nu) = N_0 \sigma(\nu) = N_0 \frac{\lambda^2}{8\pi t_{sp}} g(\nu), \quad (7.39)$$

avec N_0 la différence de population à l'équilibre, qui augmente avec le pompage ; $\sigma(\nu) = (\lambda^2/8\pi t_{sp})g(\nu)$ représente la section efficace de transition, t_{sp} la durée de vie de l'émission spontanée, $g(\nu)$ la largeur de ligne normalisée et $\lambda = \lambda_0/n$ la longueur d'onde dans le milieu.

Au fur et à mesure que la densité de flux de photons augmente, le laser entre dans un régime d'opération non-linéaire dans lequel il sature et le gain diminue. Le processus d'amplification fait décroître la différence de population initial N_0 qui devient

$$N = \frac{N_0}{1 + \phi/\phi_s(\nu)}, \quad (7.40)$$

avec $\phi_s(\nu) = [\tau_s \sigma(\nu)]^{-1}$ la saturation de la densité de flux de photons et τ_s la constante de saturation que nous avons illustrée Fig. 7.9 et qui vaut $\tau_s \approx t_{sp}$ pour un système de pompage à 4 niveaux.

Le coefficient de gain dans cette situation de saturation devient

$$\gamma(\nu) = N\sigma(\nu) = \frac{\gamma_0(\nu)}{1 + \phi/\phi_s(\nu)}. \quad (7.41)$$

Le processus d'amplification laser introduit aussi une phase à l'onde se propageant dans la cavité :

$$\varphi(\nu) = \frac{\nu - \nu_0}{\Delta\nu} \gamma(\nu); \quad (7.42)$$

où cette expression est valable pour un profil de gain Lorentzien.

Lorsque le milieu amplificateur est placé dans une cavité Fabry–Perot comme illustré sur Fig. 7.1(b), la lumière s'y propage avec un nombre d'onde $k = 2\pi\nu/c$ et peut subir une atténuation par unité de longueur α_s . Si les deux miroirs qui forment la cavité ont pour réflectances $\mathcal{R}_1$ et $\mathcal{R}_2$, la densité de flux de photons diminue, pendant un aller-retour dans la cavité, d'une valeur

$$\mathcal{R}_1 \mathcal{R}_2 \exp(-2\alpha_s d) = \exp(-2\alpha_r d), \quad (7.43)$$

où nous avons introduit le coefficient de perte par unité de longueur α_r qui inclut les pertes associées aux miroirs. On définit ainsi les trois coefficients de pertes α_r , α_{m1} et α_{m2} :

$$\alpha_r = \alpha_s + \alpha_{m1} + \alpha_{m2} \quad (7.44)$$

$$\alpha_{m1} = \frac{1}{2d} \ln \frac{1}{\mathcal{R}_1} \quad (7.45)$$

$$\alpha_{m2} = \frac{1}{2d} \ln \frac{1}{\mathcal{R}_2}; \quad (7.46)$$

les deux derniers coefficients sont associés aux pertes dues à chacun des deux miroirs. Le paramètre $\tau_p = 1/(\alpha_r c)$ correspond à la durée de vie des photons dans la cavité laser. Il exprime le temps passé en moyenne par un photon dans la cavité laser avant de s'en échapper.

Comme nous l'avons vu au début du chapitre, la cavité possède des modes dont la fréquence vaut $\nu_q = q\nu_F$ avec $q = 1, 2, \dots$ et $\nu_F = c/2d$. Comme indiqué Fig. 7.4 la largeur spectrale des modes vaut $\delta\nu \approx \nu_F/\mathcal{F}$. Lorsque les pertes de la cavité sont petites et la finesse grande, on a de plus

$$\mathcal{F} \approx \frac{\pi}{\alpha_r d} = 2\pi\tau_p\nu_F. \quad (7.47)$$

Pour obtenir le lasage et l'émission de lumière, il est évident qu'il faut que le gain soit supérieur aux pertes de la cavité :

$$\gamma_0(\nu) > \alpha_r. \quad (7.48)$$

Comme $\gamma_0(\nu)$ est proportionnel à la différence de population à l'équilibre N_0 , le lasage requiert,

$$N_0 > N_t, \quad (7.49)$$

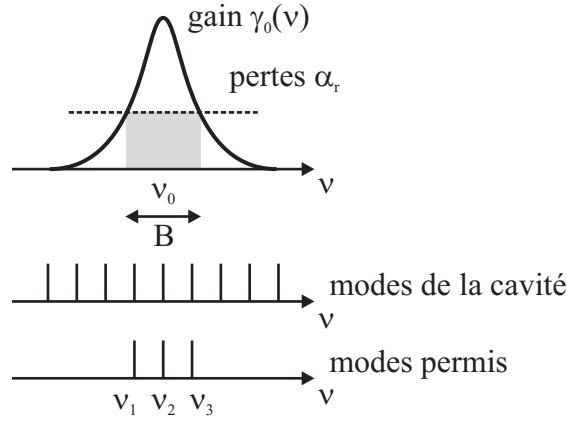


FIGURE 7.12 – Le lasage procède de l’interaction entre la courbe de gain $\gamma(\nu)$ et les modes possibles dans la cavité. Seuls les modes qui possèdent un gain supérieur aux pertes α_r contribuent au lasage.

avec $N_t = \alpha_r/\sigma(\nu)$ la différence de population au seuil de lasage (en anglais le seuil se nomme *threshold*). On peut récrire cette différence de population $N_t = 1/(c\tau_p\sigma(\nu))$ qui montre qu’elle est inversement proportionnelle à τ_p : des pertes importantes (i.e. une durée de vie des photons courte) nécessitent un pompage plus vigoureux pour obtenir le lasage. En utilisant $\sigma(\nu) = (\lambda^2/8\pi t_{sp})g(\nu)$ pour la section efficace de transition, on peut récrire la différence de population au seuil,

$$N_t = \frac{8\pi}{\lambda^2 c} \frac{t_{sp}}{\tau_p} \frac{1}{g(\nu)}. \quad (7.50)$$

On remarque qu’Eq. (7.50) indique que N_t dépend de la fréquence et que le seuil de lasage est le plus petit lorsque $g(\nu)$ est maximum, i.e. à la fréquence $\nu = \nu_0$. En effet, c’est à cette fréquence que l’interaction entre la lumière et le milieu atomique est la plus forte.

En étudiant le phénomène de lasage dans l’espace des fréquences, on remarque qu’il est le produit de deux effets, la courbe de gain du matériau et les modes possibles de la cavité, comme illustré Fig. 7.12. Pour avoir lasage, il faut avant tout que le coefficient de gain soit supérieur aux pertes : $\gamma_0(\nu) > \alpha_r$. Cette condition est satisfaite pour une bande de fréquences de largeur B autour de la fréquence de résonance ν_0 . Cette bande de fréquence augmente avec la largeur spectrale de la résonance atomique utilisée dans le laser. Seul un nombre limité M de modes contribue au lasage (on parle de modes permis) :

$$M \approx \frac{B}{\nu_F}. \quad (7.51)$$

Dès que le laser est enclenché la densité de flux de photons augmente dans ces différents modes permis. Compte tenu de la forme de la courbe de gain, les modes dont l’énergie est la plus proche de ν_0 croissent plus rapidement que les autres. En parallèle, tous ces nouveaux photons interagissent avec le milieu et réduisent le gain de façon homogène, comme illustré

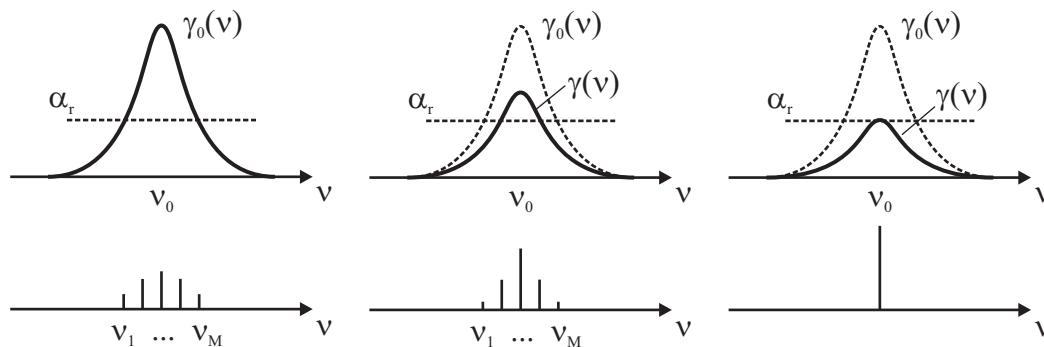


FIGURE 7.13 – Au moment où le laser se met en marche, tous les modes de la cavité ayant un gain supérieur aux pertes peuvent contribuer au lasage. Après un certain temps, seuls les modes proches de la fréquence de résonance ν_0 continuent à être amplifiés et, finalement il ne reste – dans la situation idéale – plus qu’un mode dans la cavité.

Fig. 7.13. Ainsi, après un certain temps, les modes périphériques disparaissent et il ne reste plus que le mode central à la fréquence ν_0 . Ainsi, la combinaison de la cavité avec le milieu à gain permet d’obtenir une émission de lumière quasi mono-chromatique ; qui plus est, avec des photons cohérents ayant tous la même phase.

7.4 Différents types de lasers

7.4.1 Laser He–Ne

Un état gazeux est caractérisé par une faible interaction entre les atomes ou entre les molécules. Ainsi les transitions optiques forment-elles des lignes étroites. Ceci vaut aussi pour l’absorption et rend une excitation optique par lampe flash inefficace. Il faut donc avoir recours à des mécanismes d’excitation plus sélectifs et dans la plupart des cas on utilise une excitation par décharge électrique.

Le premier laser avec une émission continue était un laser He–Ne construit en 1961 par A.W. Javan, W.R. Bennett et D.R. Herriott. Le milieu gazeux est constitué d’un mélange à basse pression, typiquement de l’hélium à une pression de 1 mmHg et du néon à une pression de 0.1 mmHg. L’excitation se fait par décharge électrique en mode DC ou AC. L’énergie des états excités de l’hélium est très proche de l’énergie des états excités du néon, Fig. 6.8, et l’énergie des premiers est transférée aux seconds par collisions. Ce mécanisme de transfert est très efficace et permet de créer une inversion de population dans les niveaux des atomes de néon.

Figure 7.14 montre un laser He–Ne : le tube de verre qui contient le mélange de gaz est fermé par des fenêtres de Brewster qui assurent que la radiation du laser est polarisée. La

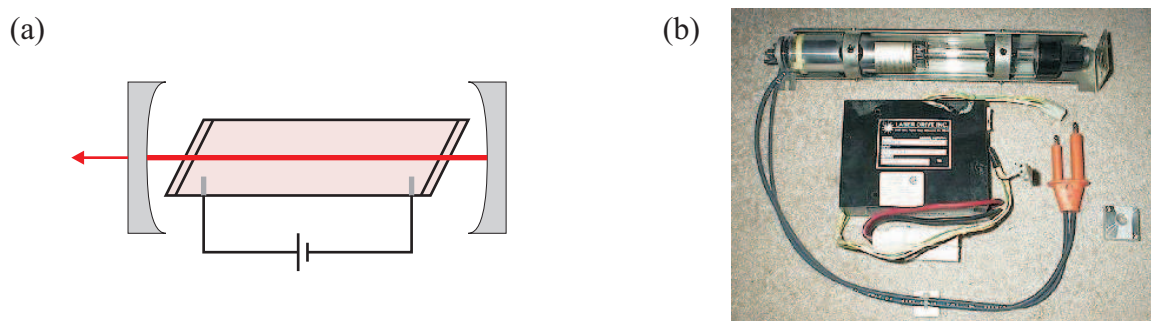


FIGURE 7.14 – Laser He-Ne : (a) principe de fonctionnement et (b) vue éclatée.

longueur du tube est de l'ordre de 10 à 50 cm et le diamètre intérieur de 1.5 mm, la pression intérieure de l'ordre de quelques mbar et la rapport de pression He :Ne de 5 :1. Les électrodes permettent de créer un champ électrique de l'ordre de 90 V/cm. Ceci nécessite des tensions de 1.5–3 kV et un courant de 10–20 mA. Pour la ligne rouge à $\lambda_0 = 632.8 \text{ nm}$ la puissance émise est de l'ordre de 1–5 mW et le rendement de conversion électrique–optique d'environ 0.1%.

7.4.2 Laser à semiconducteur

Le premier laser à rubis fut réalisé par Theodore Maiman en 1960. Cependant, le laser est resté une curiosité de laboratoire durant plusieurs années. Coûteux et difficiles à réaliser, les lasers à gaz et à corps solide ont connu un développement relativement lent et demeurèrent réservés à des applications bien spécifiques. Les lasers à colorant ont eu leur période de gloire, mais leur maintenance difficile les a presque fait disparaître. L'utilisation de semiconducteurs a permis par contre de développer toute une série de lasers particulièrement fiables et simples d'utilisation. La possibilité de les produire en masse à faible coût les a rendus omniprésents dans quantités de dispositifs, du stockage de l'information (CD, DVD) aux télécommunications par fibre optique. L'utilisation de plusieurs lasers semiconducteur en parallèle permet aussi d'atteindre de très fortes puissances optiques, par exemple pour des applications en chirurgie ou en photothérapie.

Le matériel actif d'un tel laser est un semiconducteur avec un bandgap direct : les porteurs de charge d'énergie minimale dans la bande de conduction ont la même quantité de mouvement que les porteurs de charge d'énergie maximale dans la bande de valence. L'arséniure de gallium est un exemple de tel semiconducteur à bandgap direct, Fig. 7.15(a). On a représenté sur cette figure deux transitions directes, ainsi que leur énergie. Ces transitions se trouvent dans un point particulier du réseau cristallin, identifié par la lettre Γ et qui correspond à une direction de propagation spécifique dans le cristal.

Pour un semiconducteur sans bandgap direct, comme par exemple le silicium Fig. 7.15(b), une transition optique ne peut pas avoir lieu car il faudrait que les porteurs de charge

TABLE 7.1 – Semiconducteurs avec bandgap direct (indirect), à $T = 300$ K

Semiconducteur	GaN	AlGaN	AlGaAs	GaInAsP	InP
$E_g(\text{eV})$	3.4	3.4 – (5.9)	1.4 – 2.2	0.7 – 2	1.28
$\lambda_0(\mu\text{m})$	0.36	(0.2) – 0.4	(0.6) – 0.9	0.6 – 1.7	1

Semiconducteur	InAs	InSb	PbSe	SnTe	PbSnTe
$E_g(\text{eV})$	0.36	0.24	0.26	0.18	0.04 – 0.2
$\lambda_0(\mu\text{m})$	3.4	5.2	4.8	6.9	6.5 – 3

participant à cette transition acquièrent (ou perdent) une certaine quantité de mouvement, correspondant au déplacement latéral indiqué par une flèche en traitillé.

La longueur d'onde d'émission d'un laser à semiconducteur dépend directement de l'énergie du bandgap E_g par la relation $\lambda_0 = hc_0/E_g$ (on se souviendra ici de la relation Eq. (6.2) $E(\text{eV}) = 1.24/\lambda_0(\mu\text{m})$ qui permet de connaître facilement la longueur d'onde λ_0 émise par un semiconducteur dont on connaît l'énergie de bandgap E_g). Dans un semiconducteur, le bandgap varie avec la température et donc la longueur d'onde d'émission dépend aussi de la température. La Table 7.1 donne quelques énergies et longueurs d'onde pour les semiconducteurs utilisés couramment pour réaliser des lasers. On remarque que la plupart des matériaux disponibles émettent dans le proche infrarouge. Ainsi, certains lasers semiconducteur émettant apparemment dans le visible sont en réalité des lasers émettant dans l'infrarouge dont on a doublé la fréquence (i.e. dont on a réduit la longueur d'onde de moitié). Un pointeur laser "vert" est généralement réalisé avec une diode AlGaAs émettant dans l'infrarouge à $\lambda_0 = 808\text{ nm}$ qui pompe un cristal d'yttrium aluminium vanadate dopé avec du néodyum

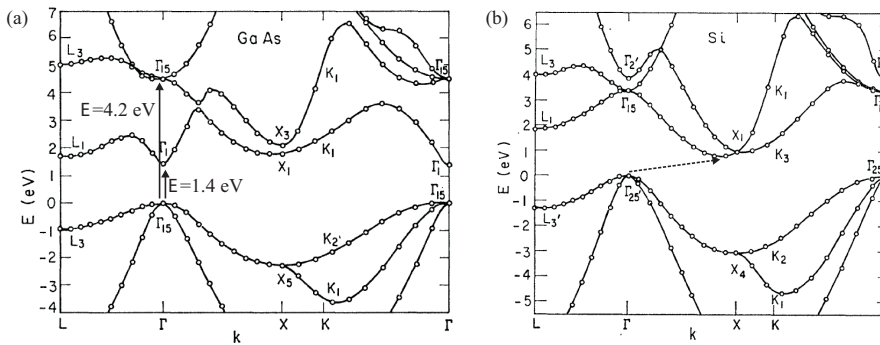


FIGURE 7.15 – Structure de bande du (a) GaAs et (b) Si. Seul le GaAs possède un bandgap direct permettant des transitions optiques ; pour le Si il n'existe pas de bandgap direct : les porteurs de charge ayant une énergie minimale dans la bande de conduction n'ont pas la même quantité de mouvement que les porteurs de charge d'énergie maximale dans la bande de valence. D'après M.L. Cohen et T.K. Bergstresser, *Phys. Rev.* **141**, 789 (1996).

(Nd :YVO₄) qui lase à $\lambda_0 = 1064\text{nm}$ et dont la fréquence est doublée en passant à travers un cristal de potassium titanyl phosphate (KTP) en sorte que la longueur d'onde finale d'émission est $\lambda_0 = 532\text{nm}$ dans le vert.

Pour qu'un semiconducteur lase, il faut créer une inversion de population. Plusieurs méthodes existent pour exciter ainsi un semiconducteur à bandgap direct : pompage optique avec des photons d'énergie supérieur au bandgap, pompage par faisceau d'électrons, etc... . La méthode la plus simple pour parvenir à cette inversion de population consiste cependant à injecter des porteurs de charge directement dans une jonction $p-n$, Fig. 7.16. On utilise donc une diode comme structure laser (d'où le terme "diode laser"). Si le biais est positif, i.e. si la polarisation est positive côté p , la diode devient conductrice, sa résistance est faible et un courant passe, Fig. 7.16(b). Ce courant est généré à l'interface $p-n$ par des courants de diffusion des électrons vers le côté p et des trous vers le côté n . Ces deux types de porteurs de charge se combinent à l'interface où leur énergie est émise sous forme de radiation. L'énergie du photon ainsi généré correspond donc (en première approximation) à l'énergie du gap : $h\nu = E_g$. Dans une telle diode, le courant se transforme directement en lumière, avec un rendement proche de 100%. Le gain optique peut aussi être très élevé, avec des coefficients de gain de l'ordre de 100cm^{-1} .

Le principe que nous venons de décrire permet de générer de la lumière à partir d'un courant électrique, par recombinaison de porteurs de charge. Dans sa version la plus simple, un tel dispositif n'émet pas de radiation cohérente, on parle de diode lumineuse. Pour obtenir un laser, il faut ajouter un dispositif de feedback; le plus simple consiste à cliver (casser sous tension) deux faces parallèles du semiconducteur pour réaliser un cavité, Fig. 7.17. La plupart des semiconducteurs ont un indice de réfraction élevé (par exemple $n = 3.6$ pour le GaAs à $\lambda_0 = 0.9\mu\text{m}$) et un simple interface entre le semiconducteur et l'air produit déjà une réflectivité importante, de l'ordre de $R = 0.32$ pour le GaAs. Comme le semiconducteur est un cristal, il est possible d'obtenir des facettes extrêmement nettes par ce procédé de clivage. On peut améliorer la réflectivité en déposant des couches diélectriques (souvent des

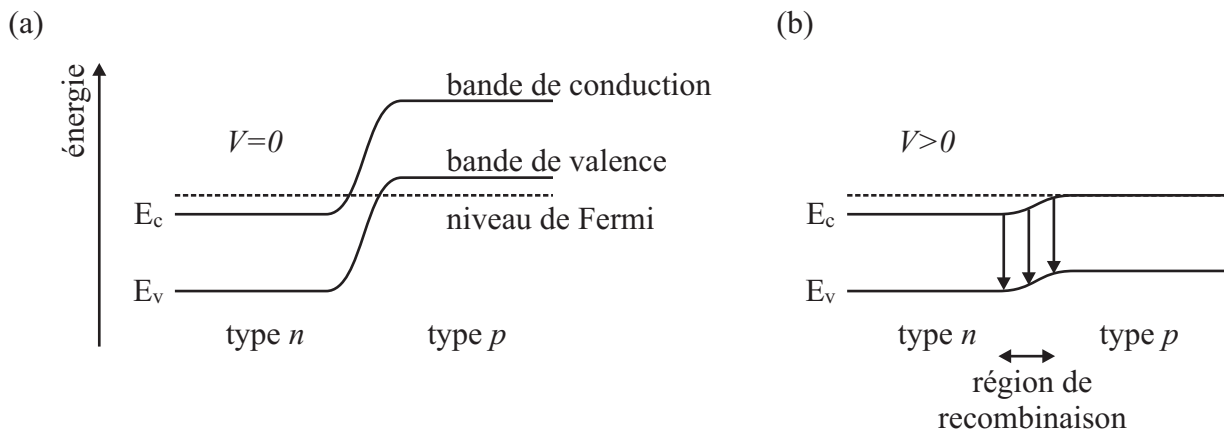


FIGURE 7.16 – Diagramme d'énergie pour une jonction $p-n$ (a) en équilibre thermique ($V = 0$) et (b) polarisée positivement ($V > 0$).

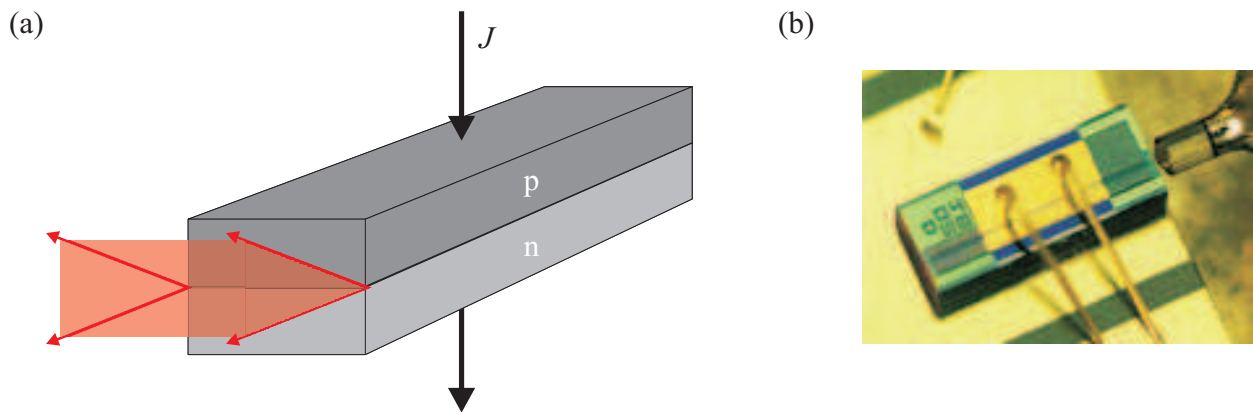
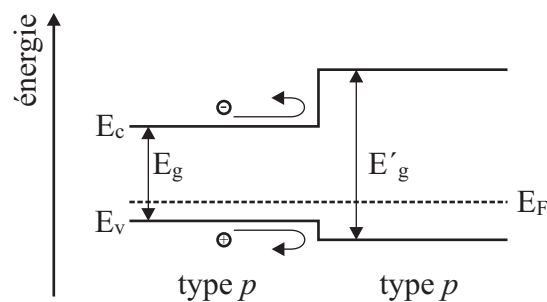


FIGURE 7.17 – (a) Schéma et (b) photographie d'une diode laser.

oxydes) sur la facette du laser. Malheureusement, lors de l'utilisation du laser ces facettes peuvent s'échauffer très fortement, donnant lieu à des dégâts catastrophiques qui mettent le laser hors d'usage. Un savoir faire extrêmement pointu a donc été développé pour passiver les facettes d'un laser semiconducteur en y déposant des couches diélectriques particulières.

Les diodes laser constituées d'une jonction simple, Fig. 7.16, ne sont pas des plus efficaces : une grande quantité de porteur de charges est perdue par diffusion à l'interface $p-n$. Deux couches semiconductrices de même structure cristalline mais avec des bandgaps différents permettent de fabriquer des lasers plus efficaces. Une telle hétérostructure peut être réalisée avec des semiconducteurs de même type : $p-p$, $n-n$, ou avec des semiconducteurs différents : $n-p$. La figure 7.18 illustre le cas d'une hétérostructure $p-p$. Des barrières de potentiel apparaissent à l'interface pour les bandes de conduction et de valence. L'amplitude de ces barrières de potentiel correspond à la différence d'énergie des deux bandgaps E_g et E'_g . La diffusion des porteurs de charge du matériau à faible bandgap vers le matériau à grand bandgap est empêchée par ces barrières de potentiel.

De telles structures sont réalisées par croissance épitaxiale (i.e. atome par atome) en phase liquide (en anglais *liquid phase epitaxy*, *LPE*), en phase gazeuse (en anglais *metalorganic*

FIGURE 7.18 – Hétérostructure formée entre deux semiconducteurs de type p avec des bandgaps différents.

vapour phase epitaxy, MOVPE) ou par jet moléculaire (en anglais *molecular beam epitaxy, MBE*). Il existe plusieurs systèmes cristallins qui permettent de créer de telles hétérostructures (i.e. des matériaux ayant des bandgaps différents mais une structure cristalline compatible, en sorte que l'on puisse les faire croître l'un sur l'autre). Citons par exemple $\text{Al}_x\text{Ga}_{1-x}\text{As}$ et GaAs ; $\text{Ga}_x\text{In}_{1-x}\text{P}_y\text{As}_{1-y}$ et InP ; $\text{Al}_x\text{Ga}_{1-x}\text{N}$ et GaN. Chaque système émet dans une région particulière du spectre : le premier à $\lambda_0 = 0.7\text{--}0.9\ \mu\text{m}$; le deuxième dans l'infrarouge entre $\lambda_0 = 6.5$ et $30\ \mu\text{m}$; le dernier dans le bleu, $\lambda_0 \simeq 0.4\ \mu\text{m}$

Afin de confiner encore mieux l'émission de la lumière et d'augmenter l'efficacité d'une diode laser, les diodes actuelles sont construites sur une double hétérostructure, comme illustré Fig. 7.19(a). Il s'agit d'une première hétérostructure $n\text{--}p$ suivie d'une seconde $p\text{--}p$. La structure du bandgap est illustrée Fig. 7.19(b). La couche intermédiaire avec le plus petit bandgap correspond à la zone active. Cette zone de recombinaison n'est plus déterminée par la longueur de diffusion des porteurs de charge, mais par la distance entre les deux hétérostructures. Les deux barrières de chaque côté de la zone active évitent la diffusion des porteurs de charge qui sont injectés à travers la jonction $p\text{--}n$. Ceci limite fortement les pertes des porteurs de charge par diffusion.

L'indice de réfraction d'un semiconducteur est relié à son bandgap. Une règle empirique indique que le produit de l'indice de réfraction n par la puissance quatrième du bandgap est constant. Ainsi n est plus grand dans la région où le gap est petit, Fig. 7.19(c). Cette figure indique que la région active fait aussi office de guide d'onde, puisque l'indice de réfraction y

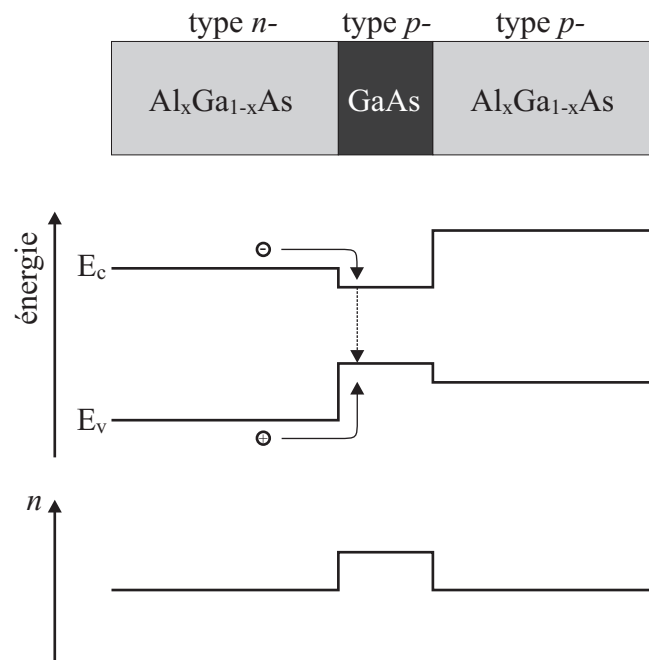


FIGURE 7.19 – Double hétérostructure formée entre trois semiconducteurs avec des bandgaps différents. La recombinaison a lieu dans la région où le bandgap est le plus petit. Cette région exhibe aussi un indice de réfraction n plus élevé et produit un guide d'onde optique.

est plus élevé. Ainsi les photons générés dans cette région par la recombinaison des porteurs de charge sont guidés à l'intérieur de cette couche d'indice élevé. Les pertes par diffraction dans le cristal en sont considérablement réduites.

La réduction des pertes par diffusion des porteurs de charge et par diffraction des photons émis rend le laser beaucoup plus efficace et permet de réduire le seuil de lasage de façon significative.

Chapitre 8

Photodétecteurs

Un photodétecteur est un dispositif qui mesure une puissance optique (un flux de photons) et le convertit en une forme mesurable, soit sous forme de courant (effet photoélectrique), soit sous forme de chaleur (effet photothermique). Dans ce chapitre, nous nous limitons à l'effet photoélectrique.

8.1 Effet photoélectrique

Lorsqu'un photon est absorbé par un métal, son énergie peut être absorbée par un électron. Si l'énergie du photon est suffisamment grande, l'électron peut s'échapper de la barrière de potentiel associée à la surface du métal et se retrouver comme électron libre dans l'espace libre (en anglais *vacuum*). Cette émission photoélectrique dans un métal est illustrée Fig. 8.1(a). Pour passer dans l'espace libre, un électron doit passer à travers une barrière de potentiel de hauteur W , correspondant à la différence d'énergie entre le niveau de Fermi et le *vacuum*. L'énergie cinétique maximale de l'électron ainsi libéré est

$$E_{max} = h\nu - W . \quad (8.1)$$

Equation (8.1) est connue sous le nom d'équation de photoémission d'Einstein. Notons que si l'électron se trouve à un niveau énergétique plus bas que le niveau de Fermi, il acquerra moins d'énergie cinétique, puisqu'il lui faudra plus d'énergie pour atteindre le *vacuum*. Dans un métal, W est généralement supérieur à 2 eV, si bien que l'effet photoélectrique dans un métal ne permet de détecter que des photons dans le spectre visible et dans les UV.

L'effet photoélectrique existe aussi dans les semiconducteurs intrinsèques, Fig. 8.1(b). Dans ce cas, le photon absorbé libère un électron de la bande de valence et l'énergie cinétique acquise par l'électron vaut

$$E_{max} = h\nu - W = h\nu - (E_g + \chi) , \quad (8.2)$$

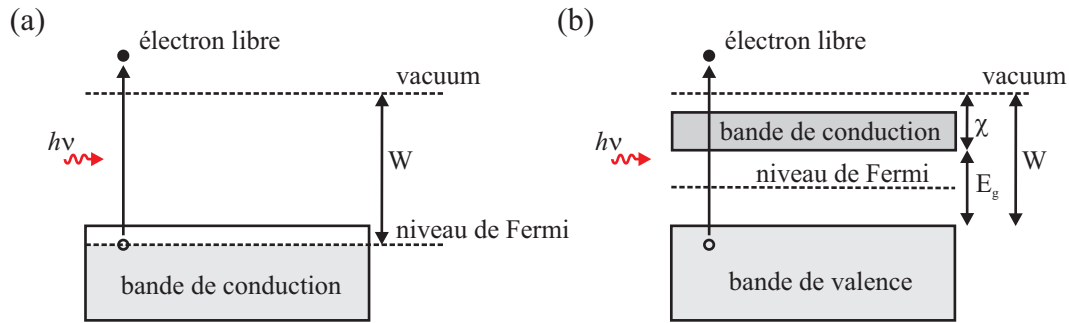


FIGURE 8.1 – Effet photoélectrique pour (a) un métal et (b) un semiconducteur intrinsèque.

où E_g est l'énergie du bandgap et χ l'affinité du matériau qui mesure l'énergie entre le bas de la bande de conduction et le *vacuum*. La valeur de $E_g + \chi$ pour un semiconducteur ($\approx 1.4 \text{ eV}$) est plus petite que la valeur de W pour un métal et permet donc de détecter des photons dans l'infrarouge proche et dans le visible (on se souviendra une fois de plus d'Eq. (6.2) pour traduire une énergie en longueur d'onde). Il existe aussi des semiconducteurs avec une affinité négative, en sorte que la bande de conduction se trouve au dessus du *vacuum*, si bien qu'il suffit que l'énergie du photon soit supérieure à E_g pour que l'effet photoélectrique ait lieu.

D'un point de vue pratique, un photodétecteur prend la forme d'un tube à vide, comme illustré Fig. 8.2(a). Les électrons sont émis de la surface d'une photocatode et vont vers une contre-électrode (anode) maintenue à un potentiel élevé. Si la photocatode travaille en transmission, les électrons émis peuvent être amplifiés par un processus de cascade sur des éléments semiconducteurs successifs maintenus à des potentiels de plus en plus élevés, Fig. 8.2(b). Chaque élément émet des électrons secondaires et cet effet photomultiplicateur peut augmenter le courant jusqu'à un facteur de 10^8 . Un tel tube photomultiplicateur est très efficace mais malheureusement assez encombrant. Une galette de microcanaux (en anglais *microchannel plate*) fonctionne sur le même principe, mais en utilisant de petits canaux parallèles d'un diamètre de $\approx 10 \mu\text{m}$ percés dans une plaque de verre et dont l'intérieur est recouvert d'un matériau qui émet des électrons secondaires. On obtient ainsi un effet

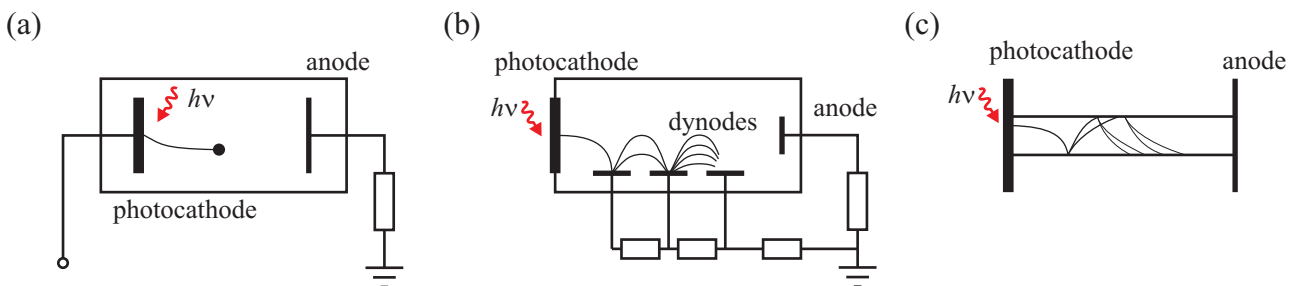


FIGURE 8.2 – Différentes réalisations de photodétecteurs : (a) tube à vide, (b) photomultiplicateur et (c) galette de microcanaux.

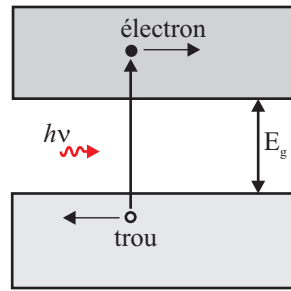


FIGURE 8.3 – Effet photoélectrique interne : un photon génère une paire électron–trou qui reste dans le semiconducteur.

d'amplification en continu, comme illustré Fig. 8.2(c). Comme il est formé de nombreux canneaux parallèles, ce système permet même de former une image entre la photocatode et l'anode. En mettant une couche phosphorescente sur l'anode, on obtient un dispositif d'amplification d'image.

Dans les deux dispositifs présentés ci-dessus, l'électron est éjecté vers le *vacuum*; on parle d'effet photoélectrique externe. La technologie actuelle utilise plutôt l'effet photoélectrique interne, dans lequel le photon incident génère une paire électron–trou qui reste dans le semiconducteur, Fig. 8.3. Ainsi, l'absorption d'un photon par un semiconducteur intrinsèque génère un électron de la bande de valence vers la bande de conduction et un trou dans la bande de valence. Ces deux porteurs de charges peuvent être transportés dans le dispositif dès lors que l'on y applique un champ électrique. Ce transport de charges correspond à un courant qui peut être mesuré.

On définit l'efficacité quantique η d'un photodétecteur comme la probabilité qu'un photon incident génère une paire électron–trou qui contribuera au courant mesuré. On a $0 \leq \eta \leq 1$ et η se définit comme le rapport entre le flux de paires électron–trou contribuant au courant divisé par le flux de photons incidents. En utilisant Fig. 8.4 on remarque que tous les photons incidents ne contribuent pas à la génération d'un photocourant : tout d'abord une partie des photons incidents est réfléchiée par la surface du dispositif (réflectivité $\mathcal{R}$);

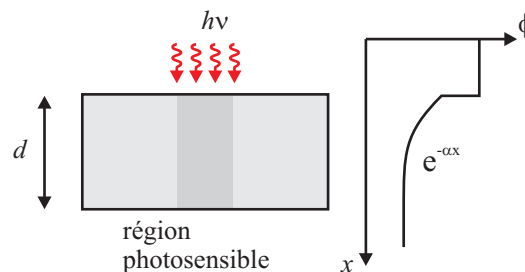


FIGURE 8.4 – Variation du flux de photons ϕ durant son interaction avec un photodétecteur. Une partie du flux est réfléchiée, une partie est absorbée progressivement et une partie n'est pas absorbée du tout.

ensuite l'absorption est progressive et nécessite une certaine épaisseur de matériau (on a là un phénomène semblable à celui du gain ou des pertes dans un laser) ; finalement certaines paires électron-trou peuvent se recombiner et ne pas participer à la génération du photocourant. Ainsi, l'efficacité quantique s'écrit

$$\eta = (1 - \mathcal{R})\varsigma[1 - \exp(-\alpha d)], \quad (8.3)$$

où ς représente la fraction de paires électron-trou qui contribuent au photocourant, α l'absorption du matériau et d l'épaisseur du photodétecteur.

L'efficacité quantique d'un semiconducteur dépend de la longueur d'onde λ_0 car l'absorption α dépend de la longueur d'onde. En particulier, lorsque la longueur d'onde est plus grande que celle correspondant au gap du semiconducteur, $\lambda_0 > \lambda_g = hc_0/E_g$, la probabilité que le photon soit absorbé est extrêmement faible. Malheureusement, lorsque la longueur d'onde est très petite, les photons sont essentiellement absorbés près de la surface du dispositif, où ils risquent de se recombiner très rapidement, sans contribuer au photocourant. D'une façon générale, pour un coefficient d'absorption α la plupart des photons sont absorbés sur une distance $1/\alpha$ (par exemple pour une valeur $\alpha = 10^4 \text{ cm}^{-1}$ la plupart de la lumière est absorbée sur une distance de $1 \mu\text{m}$). On peut augmenter l'efficacité d'un photodétecteur en l'incluant dans une cavité qui permet de maintenir les photons un temps plus long et d'en absorber davantage.

On définit la réponsivité d'un photodétecteur $\mathbf{R}$ comme le rapport entre le photocourant généré et la puissance optique illuminant le photodétecteur. Si chaque photon génère une paire électron-trou, un flux de photons ϕ (photons par seconde) produit un flux d'électrons ϕ (électrons par seconde), soit un courant $i_p = e\phi$, où e est la charge de l'électron. Ainsi, une puissance optique $P = h\nu\phi$ à une fréquence ν correspondrait à un courant $i_p = eP/h\nu$. Malheureusement, seul un taux η de photons participe à la production de courant :

$$i_p = \eta e\phi = \frac{\eta e P}{h\nu} = \mathbf{R}P. \quad (8.4)$$

On en déduit la réponsivité,

$$\mathbf{R} = \frac{\eta e}{h\nu} = \eta \frac{\lambda_0}{1.24}. \quad (8.5)$$

On fera attention que la dernière égalité suppose que λ_0 est exprimée en μm (voir Eq. (6.4)). On remarque que la réponsivité dépend linéairement de l'efficacité quantique et de la longueur d'onde. L'ordre de grandeur de la réponsivité est 1 A/W .

Dans ce qui précède nous avons supposé que chaque paire électron-trou générerait une charge e dans le circuit du photodétecteur ; dans la réalité, la charge générée q est différente et peut être supérieure ou inférieure à e . On définit ainsi le gain G du photodétecteur,

$$G = q/e. \quad (8.6)$$

En présence de gain, Eq. (8.4) devient,

$$i_p = \eta q\phi = \eta G e\phi = \frac{\eta G e P}{h\nu}. \quad (8.7)$$

De même pour Eq. (8.5),

$$\mathbf{R} = \frac{\eta G e}{h\nu} = \eta G \frac{\lambda_0}{1.24}. \quad (8.8)$$

Il ne faut pas confondre le gain avec l'efficacité quantique η d'un photodétecteur qui représente la probabilité qu'un photon incident génère une paire électron-trou qui contribuera au courant mesuré.

8.2 Photoconducteurs

Lorsqu'un photon est absorbé dans un semiconducteur, des paires électron-trou sont générées ; ainsi la conductivité σ du matériau augmente en proportion du flux de photons ϕ . En appliquant un champ électrique sur la matériel (par exemple à l'aide d'une source de tension), les électrons et les trous sont transportés dans le matériau et donnent lieu à un courant mesurable dans le circuit du photodétecteur, Fig. 8.5(a). Un tel dispositif fonctionne soit en mesurant directement ce photocourant i_p qui est proportionnel au flux de photons ϕ , soit en mesurant la tension qu'il produit sur une résistance en série avec le photodétecteur.

Dans un semiconducteur intrinsèque, le photon qui a une énergie suffisante est absorbé par une transition électronique de la bande de valence vers la bande de conduction. Un tel dispositif prend généralement la forme d'un film plus ou moins épais, avec des contacts (anode et cathode) qui sont souvent entrelacés sur la même face du semiconducteur pour d'une part maximiser la surface d'absorption et d'autre part assurer que les porteurs de charge sont extraits rapidement du dispositif et ne s'y recombinent pas, Fig. 8.5(b). Si le semiconducteur est placé sur un substrat transparent, les photons peuvent aussi arriver par le dessous du dispositif.

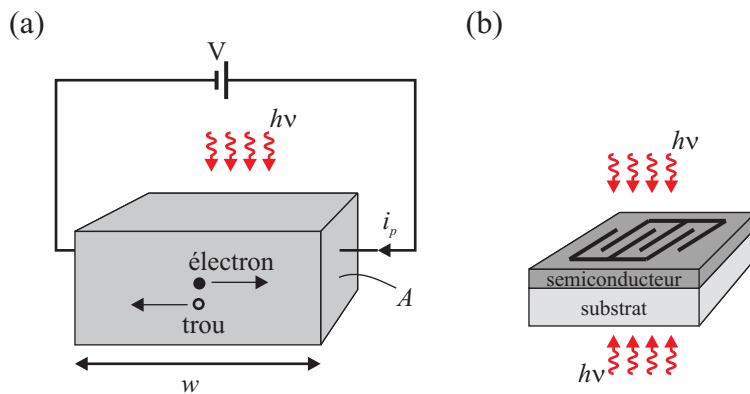


FIGURE 8.5 – (a) Dans un photdétecteur, les paires électron-trou diffusent dans le champ électrique créé par la tension appliquée V ; le photocourant i_p est proportionnel au flux de photons ϕ . (b) La structure d'un photodétecteur doit être optimisée pour maximiser les photons absorbées, par exemple en utilisant des contacts interdigités ou même transparents.

Calculons l'augmentation de conductivité produite par un flux de photons ϕ illuminant un cube de semiconducteur de volume wA comme indiqué Fig. 8.5(a). Une fraction η des photons incidents est absorbée et produit des paires électron-trou à un taux $R = \eta\phi/wA$. Cet excès de charges a une durée de vie τ et les charges ainsi créées disparaissent avec un taux $\Delta n/\tau$, où Δn est la concentration d'électrons dans le semiconducteur. A l'équilibre, ces deux taux doivent être égaux, $R = \Delta n/\tau$, en sorte que $\Delta n = \eta\tau\phi/wA$, qui nous donne une valeur pour l'augmentation des charges dans le semiconducteur. Cette augmentation se traduit en une augmentation de la conductivité $\Delta\sigma$ du semiconducteur,

$$\Delta\sigma = \frac{\eta e\tau(\mu_e + \mu_h)}{wA}\phi, \quad (8.9)$$

où l'on a introduit la mobilité μ_e des électrons et la mobilité μ_h des trous (*holes* en anglais). Equation (8.9) indique clairement que la conductivité augmente proportionnellement au flux de photons ϕ .

Un conducteur intrinsèque ne peut qu'absorber des photons qui ont une énergie égale ou supérieure au bandgap. En utilisant des semiconducteurs dopés, on peut absorber des photons dont l'énergie correspond aux impuretés et se trouve dans le bandgap. Ainsi, on étend vers l'infrarouge la plage des fréquences qui peuvent être mesurées avec un photodétecteur. On introduit une énergie d'activation E_A qui correspond à l'énergie nécessaire pour créer une paire électron-trou dans un tel semiconducteur. Quelques exemples sont donnés Tab. 8.1 avec les valeurs de E_A ainsi que les longueurs d'onde correspondantes pour les photons absorbés. Ils indiquent qu'il est possible de détecter ainsi des photons dans l'infrarouge lointain.

Jusqu'à présent nous avons considéré des semiconducteurs homogènes ; une hétérostructure peut absorber de façon efficace les photons et générer un photocourant. Figure 8.6 illustre le cas d'un photodétecteur infrarouge à puits quantiques (en anglais *quantum-well infrared photodetector (QWIP)*), dans lequel le photon incident libère vers la bande de conduction l'électron se trouvant dans le puit quantique, augmentant par la même la conductivité du matériau. Ces dispositifs sont généralement fabriqués avec des semiconducteurs III-V et sont sensibles dans l'infrarouge ($\lambda_0 = 4 \dots 20 \mu\text{m}$).

TABLE 8.1 – Quelques semiconducteurs dopés avec leur énergie d'activation et la longueur d'onde correspondante

Semiconducteur :Dopant	E_A (eV)	λ_A (μm)
Ge :Hg	0.088	14
Si :B	0.044	23
Ge :Cu	0.041	30
Ge :Zn	0.033	38
Ge :Ga	0.010	115

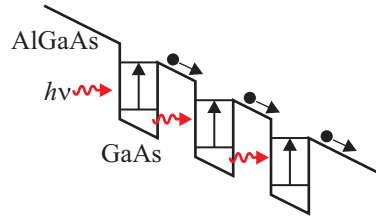


FIGURE 8.6 – Génération de porteurs de charges dans un photodétecteur à puits quantiques ; chaque puits quantique formé entre deux semiconducteurs (par exemple du GaAs et de l'AlGaAs) peut absorber des photons pour générer un photocourant.

8.3 Photodiodes

Une photodiode utilise une jonction p - n entre deux semiconducteurs ; son avantage par rapport à un photoconducteur est la rapidité de réponse. En absence de polarisation (circuit ouvert ou mode photovoltaïque) elle crée une tension. En polarisation inverse par une alimentation externe (mode photoampérique), elle crée un courant. Généralement, on sépare les régions p et n par une région intrinsèque (non-dopée), donnant la photodiode p - i - n illustrée Fig. 8.7(a). Cette région intrinsèque permet d'augmenter le volume dans lequel la lumière est absorbée et les porteurs de charge sont générés. Cette région intrinsèque diminue aussi la capacité électrique de la jonction et permet une fréquence d'opération plus élevée. La courbe courant-tension d'une photodiode est illustrée Fig. 8.7(b) ; elle prend la forme,

$$i = i_s \left[\exp \left(\frac{eV}{KT} \right) - 1 \right] - i_p, \quad (8.10)$$

qui correspond à la réponse d'une diode avec un terme supplémentaire i_p lié au photocourant.

Figure 8.8(a) illustre le mode de fonctionnement en circuit ouvert (mode photovoltaïque) :

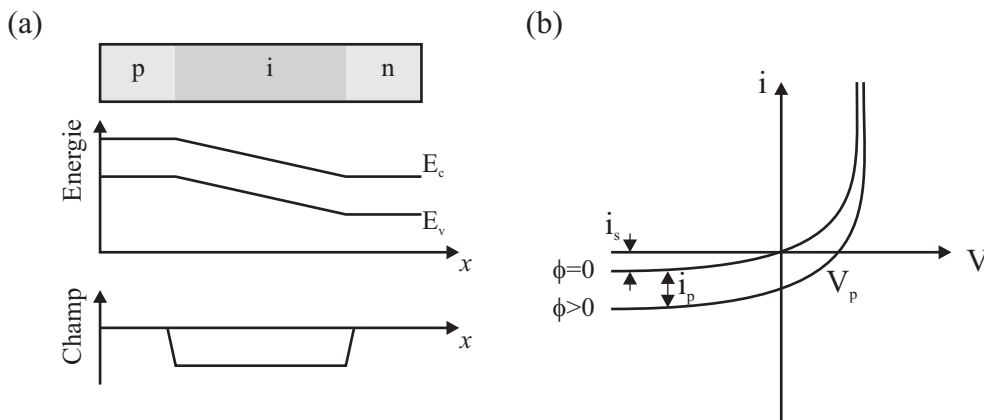


FIGURE 8.7 – Photodiode p - i - n : (a) géométrie, structure de bande et distribution du champ électrique qui permet d'extraire le photocourant ; (b) relation courant-tension.

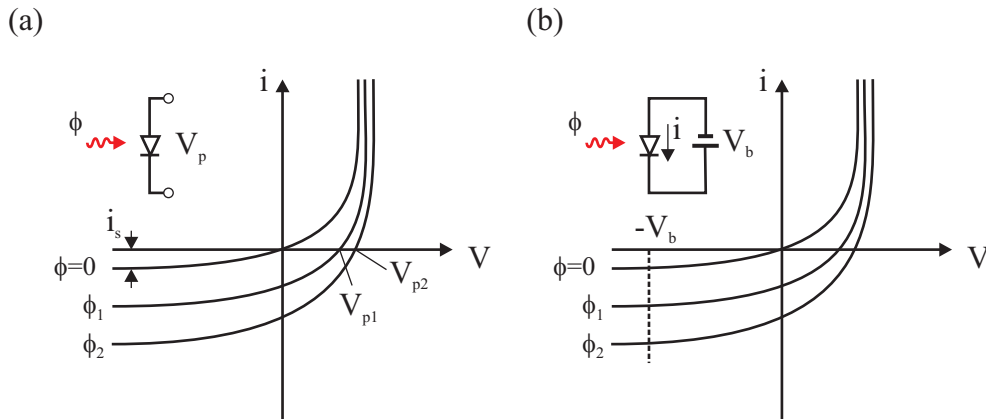


FIGURE 8.8 – Opération d’une photodiode (a) en mode photovoltaïque et (b) en mode photoampérique (polarisation inverse).

dans ce mode, la lumière génère une paire électron-trou dans la région intrinsèque. Cette paire augmente le champ électrique dans le matériau et produit une phototension V_p qui augmente avec le flux de photons. C’est le mode d’opération que l’on utilise dans les cellules solaires.

Le mode en mode polarisation inverse (mode photoampérique) permet d’augmenter la réponse de la photodiode, Fig. 8.8(b). Dans ce mode, le biais appliqué crée d’une part un fort champ électrique dans la jonction qui augmente la vitesse d’extraction des porteurs de charge et d’autre part il diminue la capacité électrique de la jonction et augmente son temps de réponse. D’une façon générale, le biais appliqué dans ce cas permet d’augmenter le nombre de photons absorbés dans la jonction.