

Promotion 2003
Année 3
Majeure 1
PHY551

MAJEURE DE PHYSIQUE

Optique quantique 1 : Lasers

Tome I

Alain Aspect, Claude Fabre, Gilbert Grynberg

Édition 2005

Avant-Propos

La découverte des lasers, en 1960 a complètement renouvelé l'optique. Au niveau fondamental, comme dans le domaine des applications, des possibilités radicalement nouvelles ont été offertes au physicien et à l'ingénieur, par la directivité, la monochromaticité, la puissance de la lumière laser. Certaines applications, comme les télécommunications par fibre optique, présentent des enjeux économiques considérables. Des progrès récents en recherche fondamentale, comme le refroidissement d'atomes par laser, ou le développement des « peignes optiques de fréquences », offrent des perspectives d'amélioration considérable des horloges atomiques, dont la précision et l'exactitude pourraient dépasser les 10^{-15} , soit une seconde d'incertitude en trente millions d'années, permettant alors une augmentation spectaculaire de la précision de la localisation par satellite, un meilleur contrôle des flux de données, ou de nouveaux tests de la relativité générale. Ce cours a pour but de présenter les notions de base de l'*optique quantique*, nom sous lequel on désigne traditionnellement la discipline très vaste consacrée aux *lasers et à leurs applications*. Notre premier objectif est de présenter une théorie des lasers simple mais puissante et générale. Elle donne une compréhension unifiée applicable à la grande variété des lasers dont nous disposons aujourd'hui, et sans doute à ceux qui ne manqueront pas d'apparaître dans le futur. Nous essaierons ensuite de dégager les propriétés spécifiques de la lumière laser qui la rendent apte à des applications inconcevables avec les sources traditionnelles.

Pour comprendre le principe des lasers, nous devons d'abord étudier l'interaction lumière-matière, dont l'exemple le plus simple est l'*interaction atome-rayonnement*. Les processus en jeu ne peuvent se décrire que dans le cadre quantique, et nous utiliserons donc les connaissances acquises dans le tronc commun, que nous compléterons par l'étude des transitions d'un système quantique entre deux états, ou entre un état discret et un continuum (**chapitre 1**). En fait, on utilise la mécanique quantique pour traiter les atomes, mais on décrit la lumière laser par un champ électromagnétique classique. Cette théorie « semi-classique » de l'interaction atome-rayonnement, universellement utilisée en physique des lasers, permet de rendre compte de façon quantitative d'un très grand nombre de phénomènes, tout en donnant des images simples et fécondes. Elle permet aussi de décrire en détail l'*amplification de lumière par émission stimulée*, à la base de tous les lasers (**chapitre 2**). On peut se demander pourquoi ne pas utiliser une théorie complètement quantique où la lumière aussi est quantifiée. On peut en fait montrer, en utilisant une telle théorie¹, que la description d'un *faisceau laser* par une *onde électromagnétique*

¹Voir le cours « Optique Quantique 2 » de la majeure de physique 2.

classique est particulièrement judicieuse et commode, notamment pour rendre compte des propriétés remarquables de cohérence.

Armés de ce modèle de l'interaction matière-rayonnement, nous pourrons donner une description générale du fonctionnement des lasers, et de quelques propriétés essentielles, qui seront présentées sur des exemples représentatifs de divers types de lasers très répandus (**chapitre 3**).

Dans les **chapitres 4 et 5**, nous approfondirons l'étude du fonctionnement des lasers, soit en régime stationnaire, soit au contraire en régime transitoire. Nous verrons ainsi apparaître des comportements caractéristiques de deux grandes catégories de laser : les *lasers continus* d'une part, pouvant être ultra-stables, et souvent de faible puissance ; les *lasers en impulsion* d'autre part, qui permettent de délivrer en des temps incroyablement brefs (jusqu'à 10^{-15} s, c'est-à-dire une femtoseconde, et même moins puisqu'on a atteint en 2004 l'échelle des attosecondes, 10^{-18} s) des énergies correspondant à des puissances crêtes qui se comptent en milliers voire en millions de gigawatts. Pour décrire ces différents régimes, nous partirons d'équations apparemment très simples, mais qui comportent des termes non-linéaires, à l'origine d'une très grande variété de phénomènes importants : compétition entre modes, impulsions de relaxation, déclenchement passif par absorbant saturable... L'intérêt de ces phénomènes va bien au-delà de la physique des lasers, car on ne saurait surestimer l'importance des effets non-linéaires dans la physique moderne, et plus généralement dans toute la science et la technique. Dans ce domaine, le laser nous offre des exemples concrets, relativement faciles à modéliser et à étudier expérimentalement, permettant de se familiariser avec les *phénomènes non-linéaires*.

La lumière laser est en général extraordinairement plus monochromatique que celle fournie par les sources traditionnelles. Pourtant, on ne peut se borner à la décrire par une onde monochromatique idéale. Ce problème nous conduira à introduire, au **chapitre 6**, des *éléments d'optique statistique*, nécessaires pour donner un sens à la notion de largeur de raie laser, de cohérence temporelle... Il s'agit en fait d'un exemple important de l'utilisation des concepts statistiques en physique, où nous retrouvons le nom d'Albert Einstein qui traita brillamment le problème du mouvement Brownien, avant d'introduire le concept d'émission stimulée, à la base de l'effet laser !

Ce cours est très centré sur l'interaction laser-atome. Cette interaction est à la base du fonctionnement du laser, mais la lumière laser a en retour profondément renouvelé la physique atomique, conduisant à l'invention de nouvelles méthodes qui ont débouché sur des applications spectaculaires. Ainsi, les physiciens ont appris à utiliser les lasers pour contrôler le mouvement des atomes, que l'on sait refroidir ainsi au nanokelvin, ou que l'on peut espérer focaliser pour réaliser des objets « nanoscopiques ». Cette possibilité de *manipulation et refroidissement des atomes par laser* constituera notre dernier exemple d'application de la théorie de l'interaction lumière-matière : il est présenté au **chapitre 7**. Ces méthodes ont d'ores et déjà des applications pratiques, aussi bien dans le domaine des horloges atomiques, que dans le développement de nouveaux senseurs inertiels et gravitationnels basés sur des interféromètres atomiques.

Il n'est pas insignifiant de remarquer que le contenu du **chapitre 7** correspond à un champ de recherche qui a démarré dans la deuxième moitié des années 80 et qui a conduit en 1997 à l'attribution du prix Nobel de physique à Claude Cohen-Tannoudji, Steven Chu, et Bill Phillips. Depuis, ce domaine a débouché sur une nouvelle découverte majeure, celle des « condensats de Bose-Einstein d'atomes en milieu dilué », elle aussi récompensée par un prix Nobel en 2001 (E. Cornell, C. Wieman, W. Ketterle). Ces systèmes, dont on dira quelques mots en fin du **chapitre 7**, présentent de fortes analogies avec les lasers, et on commence aujourd'hui à parler de « *Lasers à Atomes* ». On ne pouvait rêver de meilleure conclusion pour ce cours.

Le tome I du cours photocopié « Optique quantique 1 : Lasers » est basé sur un travail collectif de longue haleine avec Gilbert Grynberg et Claude Fabre, et enrichi d'un complément (II.5) rédigé par Emmanuel Rosencher, qui nous fait profiter de ses immenses connaissances en optique quantique des semi-conducteurs. Quant au tome II, rédigé plus récemment, ses imperfections sont dues à moi seul.

Je ne saurais terminer cet avant-propos sans évoquer la mémoire de Gilbert Grynberg qui nous a quittés au début de 2003, nous laissant avec une grande peine, et un grand vide. C'est lui qui a créé cet enseignement d'Optique Quantique. Il avait d'abord introduit, au sein du tronc commun de mécanique quantique, des exemples puisés dans ce domaine, à une époque où l'optique n'était pas encore redevenue une discipline incontournable. Il avait alors été assez convaincant pour qu'on lui demande de créer un cours d'optique quantique lors de la réforme ayant introduit les majeures. C'est dans ce cadre que j'ai eu la chance de travailler avec lui, découvrant sa conception originale de l'enseignement de l'optique quantique, basée sur une expérience de recherche de très haut niveau, et sur une réflexion personnelle profonde. Cette conception sous-tend le cours que vous allez recevoir. Elle consiste à vous montrer que si l'on veut comprendre en profondeur les phénomènes afin de pouvoir ensuite devenir créatif, *il ne faut pas être dogmatique, et savoir jongler avec des approches très différentes permettant de voir les phénomènes sous des angles variés, de se les représenter avec des images diverses, de les traiter avec plusieurs formalismes*. Gilbert n'avait pas son pareil pour sauter du modèle classique de l'électron lié élastiquement, au modèle complètement quantique de l'atome habillé, en passant par le modèle semi-classique de l'interaction lumière-matière. Il voulait faire partager aux polytechniciens cette expérience intellectuelle, dont il pensait qu'elle avait une valeur générale, bien au-delà de notre discipline. Son influence nous a marqués si fortement que son esprit vit toujours au travers de ce cours.

Alain Aspect
Septembre 2005

Mode d'emploi

Cet enseignement s'inscrit naturellement dans le prolongement du cours de Mécanique Quantique, et ce polycopié fera souvent référence au remarquable cours de Jean-Louis Basdevant et Jean Dalibard, désigné sous le signe « BD ». Parmi les autres ouvrages de Mécanique Quantique, nous citons sous le sigle « CDL » le livre de Claude Cohen-Tannoudji, Bernard Diu et Franck Lalœ².

Le tome I de ce polycopié est le résultat d'un travail collectif de longue haleine où Gilbert Grynberg et Claude Fabre ont joué un rôle majeur³. Son volume montre qu'il va très au-delà de ce qui sera enseigné. Le programme d'enseignement *stricto sensu* ne comprend que les chapitres, en laissant de côté certains développements – souvent en petites lettres. C'est en fait le contenu des amphes qui fait foi. Les compléments ne font pas partie du cours, mais certains seront abordés en petite classe.

Le fait que ce document aille au-delà du contenu de l'enseignement représente une petite difficulté, qui ne devrait pas être insurmontable pour un polytechnicien. L'avantage est qu'il y trouvera des informations utiles pour les modules d'enseignement plus spécialisés, les stages d'option, ou pour simplement satisfaire sa curiosité et élargir ses connaissances. A ceux que cette ambition habite, je ne saurais trop recommander d'aprendre à consulter les ouvrages dont une liste est donnée ci-contre. Ils y trouveront non seulement des informations complémentaires, mais aussi des points de vue différents, qui leur permettront de développer leur propre compréhension.

²Claude Cohen-Tannoudji, Bernard Diu et Franck Lalœ, « Mécanique Quantique », Hermann, Paris.

³G. Grynberg, A. Aspect, C. Fabre, Introduction aux lasers et à l'optique quantique, Cours de l'École polytechnique, Ellipses, 1995.

Bibliographie sommaire

A.E. SIEGMAN : *Lasers* (University Science Books).

O. SVELTO : *Principles of lasers* (Plenum).

L. MANDEL et E. WOLF : *Optical coherence and quantum optics* (Cambridge University Press).

C. COHEN-TANNOUDJI, J. DUPONT-ROC, G. GRYNBERG :

- *Introduction à l'électrodynamique quantique.*
- *Processus d'interaction entre photons et atomes* (Éditions du CNRS)⁴.

M. SARGENT, M.O. SCULLY, W.E. LAMB : *Laser Physics* (Addison Wesley).

H. HAKEN : *Laser theory* (Springer Verlag).

C. COHEN-TANNOUDJI : *Atoms in electromagnetic fields* (World Scientific).

C. FABRE et J.-P. POCHOLLE (éditeurs scientifiques) : *Les lasers et leurs applications scientifiques et médicales* (Éditions de Physique).

D. DANGOISSE, D. HENNEQUIN et V. ZEHNLE-DHAOUI : *Les lasers* (Dunod).

A. YARIV : *Quantum Electronics*, third edition (Wiley).

B.E.A. SALEH, M.C. TEICH : *Fundamentals of Photonics* (Wiley).

C.H. TOWNES : *Production of coherent radiation by atoms and molecules* in Nobel Lectures, Physics 1963-1970 (Elsevier). Disponible sur le web à

<http://www.nobel.se/physics>

On pourra également lire, dans le même ouvrage, les deux autres conférences du prix Nobel de Physique 1964 :

N.G. BASOV : *Semiconductor lasers* et A.M. PROKHOROV : *Quantum electronics* (1964).

⁴Ces deux ouvrages sont appelés CDG1 et CDG2 dans ce cours.

Table des matières

1	Probabilités de transition	15
1.1	Probabilité de transition	16
1.2	Transition entre niveaux discrets	17
1.2.1	Présentation du problème	17
1.2.2	Exemples	18
1.2.3	Développement de la fonction d'onde en série de perturbation	19
1.2.4	Théorie au premier ordre	21
1.2.5	Calculs au second ordre	29
1.2.6	Comparaison avec la solution exacte dans le cas d'un système à deux niveaux (oscillations de Rabi)	31
1.3	Cas d'un niveau discret couplé à un continuum	33
1.3.1	Exemple : auto-ionisation de l'hélium	34
1.3.2	Niveau discret couplé à un quasi-continuum. Modèle simplifié	36
1.3.3	Évolution du système	39
1.3.4	État final du processus. Largeur d'un niveau instable	41
1.3.5	Règle d'or de Fermi	43
1.3.6	Cas d'une perturbation sinusoïdale	47
1.4	Conclusion	48
	Complément I1 : Exemple d'utilisation de la règle de Fermi : autoionisation	51
1.	Fonctions d'onde du quasi-continuum	52
2.	Élément de matrice de l'hamiltonien d'interaction	52
3.	Densité des états couplés au niveau discret	53
4.	Autre approche possible du quasi-continuum : « condition aux limites périodiques »	54
	Complément I.2 : Continuum de largeur variable	57
1.	Présentation du modèle	57
2.	Evolution temporelle	58
2	Atomes en interaction avec une onde	61
2.1	Processus d'interaction atome – champ	63
2.1.1	Absorption	63
2.1.2	Émission induite	64
2.1.3	Émission spontanée	66

2.1.4	Diffusion élastique	67
2.1.5	Processus non-linéaires	69
2.2	Hamiltonien d'interaction	71
2.2.1	Électrodynamique classique : Équations de Maxwell-Lorentz	71
2.2.2	Hamiltonien d'une particule dans un champ électromagnétique clas- sique extérieur	73
2.2.3	Hamiltonien d'interaction en jauge de Coulomb	76
2.2.4	Hamiltonien dipolaire électrique	79
2.2.5	Hamiltonien dipolaire magnétique	81
2.3	Transition entre deux niveaux atomiques	82
2.3.1	Probabilité de transition au premier ordre de la théorie des pertur- bations	83
2.3.2	Oscillation de Rabi	88
2.3.3	Transitions multiphotoniques	94
2.3.4	Déplacements lumineux	97
2.4	Absorption entre deux niveaux de durée de vie finie	99
2.4.1	Présentation du modèle utilisé	99
2.4.2	Population excitée	101
2.4.3	Susceptibilité diélectrique	104
2.4.4	Propagation d'une onde électromagnétique : absorption et dispersion	108
2.4.5	Cas d'un système fermé à deux niveaux	110
2.5	Amplification laser	112
2.5.1	Alimentation du niveau supérieur : émission induite	112
2.5.2	Propagation amplifiée : effet laser	114
2.5.3	Généralisation : alimentation des deux niveaux et saturation	115
2.5.4	Gain laser et inversion de population	116
2.5.5	Bilan énergétique au cours de la propagation : absorption et émis- sion induite	117
2.5.6	Équations cinétiques pour les atomes	118
2.5.7	Interaction atomes photons. Section efficace laser	121
2.5.8	Équations cinétiques pour les photons. Gain laser	122
2.6	Conclusion : modèle semi-classique ou équations cinétiques ?	124

Complément II.1 : Polarisation du rayonnement et transitions dipolaires électriques 127

1	Règles de sélection et polarisation	127
1.1	Transition dipolaire électrique interdite	127
1.2	Lumière polarisée linéairement	129
1.3	Lumière polarisée circulairement	131
1.4	Emission spontanée	135
2	Résonance optique	137
2.1	Principe de l'expérience	137
2.2	Transfert de population dans l'état excité : collisions et effet Hanle .	138
2.3	Double résonance	140
3	Pompage optique	141

3.1	Transition $J = 1 \rightarrow J = 0$ excitée en polarisation linéaire	142
3.2	Équations de pompage	144
Complément II.2 : Matrice densité et équations de Bloch optiques		149
1	Fonction d'onde et matrice densité	150
1.1	Système isolé et système en interaction	150
1.2	Description en terme de matrice densité	150
1.3	Systèmes à deux niveaux	152
2	Traitement perturbatif	156
2.1	Résolution par itération de l'équation d'évolution de la matrice densité	156
2.2	Atome interagissant avec un champ oscillant : réponse linéaire . . .	158
3	Atome à deux niveaux : Traitement non-perturbatif	161
3.1	Introduction	161
3.2	Système fermé	163
3.3	Système ouvert	165
4	Conclusion	166
Complément II.3 : Modèle de l'électron élastiquement lié		169
1.	Notations et hypothèses simplificatrices	170
2.	Rayonnement dipolaire électrique	171
3.	Dipôle oscillant sinusoïdalement. Polarisation	173
4.	Réaction de rayonnement. Amortissement radiatif. Moment cinétique de la lumière	177
5.	Diffusion Rayleigh, diffusion Thomson, diffusion résonnante	180
6.	Susceptibilité	183
7.	Lien entre le modèle classique de l'électron élastiquement lié et le modèle quantique de l'atome à deux niveaux	184
Complément II.4 : Effet photoélectrique		187
1	Modèle utilisé	189
1.1	Etat atomique lié	189
1.2	Etats ionisés - Densité d'états	189
1.3	Hamiltonien d'interaction	192
2	Taux de photoionisation, section efficace de photoionisation	194
2.1	Taux de photoionisation	194
2.2	Section efficace de photoionisation	196
2.3	Comportement aux temps longs	196
3	Application : Photoionisation de l'atome d'hydrogène	197
Complément II.5 : Détecteurs infrarouge à multipuits quantiques		201
3 Les lasers		209
3.1	Conditions d'oscillation	212
3.1.1	Seuil d'oscillation	212
3.1.2	Régime stationnaire. Intensité et fréquence de l'onde laser	214
3.2	Description des milieux amplificateurs	219

3.2.1	Nécessité d'une inversion de population	219
3.2.2	Inversion de population dans les systèmes à quatre niveaux	221
3.2.3	Lasers à semi-conducteurs	229
3.2.4	Transition laser aboutissant sur le niveau fondamental : systèmes à trois niveaux	231
3.2.5	Lasers Raman	235
3.3	Propriétés spectrales des lasers	237
3.3.1	Modes longitudinaux	237
3.3.2	Fonctionnement monomode longitudinal	240
3.3.3	Largeur de raie laser	242
3.4	Lasers en impulsion	244
3.4.1	Laser à modes synchronisés	244
3.4.2	Laser déclenché	248
3.5	Spécificité de la lumière laser	250
Complément III.1 : Cavity résonnante Fabry-Perot		253
1.	Position du problème. Cavity linéaire	253
2.	Facteur de transmission et de réflexion de la Cavity. Résonances	255
3.	Cavity en anneau à un seul miroir de couplage	257
4.	Finesse. Facteur de qualité	258
5.	Exemple. Cavity de grande finesse	259
6.	Cavity laser linéaire	261
Complément III.2 : Faisceaux gaussiens. Modes transverses d'un laser		263
1.	Faisceau gaussien	264
2.	Mode gaussien fondamental d'une cavity stable	266
3.	Modes gaussiens	267
4.	Modes longitudinaux et transverses d'un laser	271
Complément III.3 : Le laser source d'énergie		273
1	Effets de l'irradiation laser sur la matière	274
1.1	Couplage laser-matériau	275
1.2	Transfert thermique	276
1.3	Effet mécanique	277
1.4	Effets photochimiques, photoablation	277
2	Usinage et traitement de surface des matériaux	279
2.1	Effets thermiques	279
2.2	Transfert de matière	280
2.3	Effets photochimiques, photoablation	280
3	Application médicales	280
4	Application militaires	282
5	Fusion inertielle	284
Complément III.4 : Le laser source de lumière cohérente		289
1	Les atouts du laser	289
1.1	Propriétés géométriques	289

1.2	Propriétés spectrales et temporelles	290
1.3	Manipulation du faisceau	292
1.4	Conclusion	292
2	Mesures par laser	293
2.1	Mesure de position	293
2.2	Mesure de distances	293
2.3	Mesure de vitesses	296
2.4	Contrôle non destructif	298
2.5	Analyse chimique	299
3	Le gyroscope laser ou « gyrolaser »	299
3.1	Effet Sagnac	299
3.2	Le gyrolaser	301
4	Le laser porteur d'information	303
4.1	Lecture et reproduction	304
4.2	Mémoires optiques	305
4.3	Télécommunications	307
4.4	Vers l'ordinateur optique ?	309
5	Séparation isotopique par laser	309
5.1	Effets isotopiques sur les niveaux atomiques	310
5.2	Spectre de l'Uranium	311
5.3	Méthode de séparation	311
5.4	Intérêt de la méthode	313
Complément III.5 : Spectroscopie non-linéaire		315
1	Largeur homogène et inhomogène	315
2	Spectroscopie d'absorption saturée	317
2.1	Trous dans une distribution de population	317
2.2	Absorption saturée dans un gaz	319
3	Spectroscopie d'absorption à deux photons	323
3.1	Transition à deux photons	323
3.2	Élimination de l'élargissement Doppler	323
3.3	Propriétés de la spectroscopie à deux photons sans élargissement Doppler	326
4	Spectroscopie de l'atome d'hydrogène	327
4.1	Le rôle de l'atome d'hydrogène	327
4.2	L'atome d'hydrogène	328
4.3	Mesure de la constante de Rydberg	329
Complément III.6 : Largeur spectrale des lasers : formule de Schawlow-Townes		333
Complément III.7 : Faisceau de lumière incohérente et lumière laser		339
1	Conservation de la luminance pour une source classique (incohérente) . . .	339
1.1	Étendue géométrique, luminance	339
1.2	Conservation de la luminance	341

2	Éclairement maximal d'une surface avec une source incohérente	341
3	Éclairement maximal d'une surface avec de la lumière laser	343
4	Nombre de photons par mode d'une cavité	344
4.1	Rayonnement thermique dans une cavité	344
4.2	Cavité laser	345
5	Nombre de photons par mode pour un faisceau libre	346
5.1	Mode libre propagatif	346
5.2	Pinceau de rayonnement thermique	347
5.3	Faisceau émis par un laser	348
6	Conclusion	349

Chapitre 1

Évolution des systèmes quantiques : probabilités de transition

Dans ce cours, nous allons nous intéresser à l'interaction entre la matière et la lumière. Ce sujet fait largement appel à la *mécanique quantique*. Une *description quantique de la matière* est en effet indispensable si l'on veut comprendre à *l'échelle microscopique* les divers *processus d'interaction* qui peuvent se produire entre celle-ci et le champ électromagnétique. Dans le chapitre 2, nous nous poserons par exemple la question suivante : étant donné un atome mis à l'instant initial dans un état déterminé et soumis à partir de cet instant à un rayonnement électromagnétique, que deviennent l'atome et le champ électromagnétique à tout instant ultérieur ? Pour pouvoir y répondre, il nous faut connaître les méthodes qui permettent de calculer l'évolution temporelle d'un système quantique dans un certain nombre de situations typiques. C'est l'objet de ce premier chapitre, qui rappelle quelques résultats de mécanique quantique relatifs à l'évolution des systèmes.

Cette évolution est fonction de la variation temporelle du rayonnement incident, qui peut être appliqué à partir d'un instant donné puis rester d'intensité constante, ou bien ne durer qu'un intervalle de temps bref (on parlera dans ce dernier cas de processus « impulsionnel »). L'évolution dépend aussi de la *structure du spectre énergétique* du système considéré, qui peut être un ensemble de *niveaux discrets* bien séparés, ou bien un *continuum*.

Ce chapitre débute par un rappel de mécanique quantique élémentaire qui nous permet de définir la notion de probabilité de transition (partie A). La partie B utilise une méthode perturbative pour calculer la *probabilité de transition d'un système quantique entre un état initial et un état final* donnés sous l'effet d'une interaction. Enfin, dans la partie C, nous étudions le cas où *l'état initial est couplé à un très grand nombre d'états finaux* très « serrés » en énergie (on parle de « quasi-continuum » d'états). Nous aboutissons à une expression particulièrement importante de la probabilité de transition, connue sous le nom de « *règle d'or de Fermi* », qui permet de déterminer un *taux de transition* par unité de temps. Nous récapitulons en conclusion les différents régimes d'évolution temporelle susceptibles de se produire.

Ce chapitre se prolonge par deux compléments. Le complément I.1 constitue un exemple d'application de la règle d'or de Fermi au problème de l'autoionisation (« effet Auger »), c'est-à-dire de l'éjection d'un électron par un atome doublement excité. Le complément I.2 décrit un modèle simple permettant de comprendre le passage continu entre les deux situations limites vues dans les parties B et C : transition entre deux niveaux discrets, transition entre un niveau discret et un continuum. Signalons également que les compléments II.4 et II.5 décrivant des processus de photodétection, constituent d'intéressantes applications de la règle d'or de Fermi.

1.1 Probabilité de transition

Rappelons tout d'abord quelques résultats importants relatifs à un système quantique décrit par un hamiltonien $\hat{H}_0$ indépendant du temps¹. Nous désignerons par $|n\rangle$ et E_n les vecteurs propres et valeurs propres de $\hat{H}_0$. Supposons à l'instant initial $t = 0$ le système dans l'état le plus général

$$|\psi(0)\rangle = \sum_n \gamma_n |n\rangle . \quad (1.1)$$

En utilisant l'équation de Schrödinger, on montre² que le système se trouve à un instant t ultérieur dans l'état

$$|\psi(t)\rangle = \sum_n \gamma_n e^{-iE_n t/\hbar} |n\rangle . \quad (1.2)$$

La probabilité de trouver le système dans l'état $|\varphi\rangle$ vaut

$$P_\varphi(t) = |\langle\varphi|\psi(t)\rangle|^2 . \quad (1.3a)$$

Cette quantité est la *probabilité de transition* de l'état $|\psi(0)\rangle$ vers l'état $|\varphi\rangle$ entre les instants 0 et t

$$P_{\psi(0) \rightarrow \varphi}(t) = |\langle\varphi|\psi(t)\rangle|^2 . \quad (1.3b)$$

Dans le cas particulier où le système est à l'instant initial dans un état propre $|n\rangle$ de $\hat{H}_0$, il est décrit à tout instant t par le vecteur d'état

$$|\psi(t)\rangle = e^{-iE_n t/\hbar} |n\rangle . \quad (1.4a)$$

La probabilité de le trouver ultérieurement dans un autre état propre $|m\rangle$ de $\hat{H}_0$ ($m \neq n$) est donc nulle

$$P_{n \rightarrow m}(t) = |\langle m|\psi(t)\rangle|^2 = 0 . \quad (1.4b)$$

¹Dans toute la suite de cet ouvrage, nous distinguerons les opérateurs de la théorie quantique par le symbole $\hat{}$.

²Voir par exemple BD Chapitre III, ou CDL Chapitre III.D.2.

Par exemple, un électron d'un atome d'hydrogène initialement dans un état $|n, l, m\rangle$ pourrait subsister indéfiniment dans cet état si l'atome n'était pas couplé au monde extérieur. En fait, il subit des transitions vers d'autres niveaux sous l'effet d'*interactions extérieures* de diverses origines : collisions contre des ions, électrons ou atomes présents dans l'environnement, champs électromagnétiques oscillants extérieurs, etc. Le couplage avec le champ électromagnétique quantifié est aussi responsable de transitions spontanées d'un niveau excité $|n, l, m\rangle$ vers des niveaux d'énergie inférieure avec émission de rayonnement : c'est le phénomène d'*émission spontanée*.

Dans tous ces exemples, les interactions viennent rajouter à l'hamiltonien $\hat{H}_0$ du système isolé un terme $\hat{W}$. En général, l'état $|n\rangle$ n'est pas un état propre de l'hamiltonien total $\hat{H}_0 + \hat{W}$, et l'état du système évolue de façon non triviale. La probabilité de trouver le système au bout d'un certain temps dans un autre état propre $|m\rangle$ de $\hat{H}_0$ n'est pas nulle : on dit que le système a effectué une transition de $|n\rangle$ vers $|m\rangle$.

Dans certains cas, le système évolue sous l'effet d'un *hamiltonien $\hat{W}$ dépendant du temps* : la dépendance est sinusoïdale dans le cas du champ électromagnétique, impulsionnelle dans le cas de collisions. En général, on ne peut pas calculer exactement le vecteur d'état à tout instant. Le paragraphe suivant montre qu'on peut cependant, grâce à une approche perturbative, obtenir la probabilité de transition sous la forme d'un développement en série.

Un problème formellement semblable, que nous rencontrerons également, est celui où l'hamiltonien total $\hat{H}_0 + \hat{W}$ est indépendant du temps mais où le système est préparé à l'instant $t = 0$ dans un état propre de $\hat{H}_0$, et détecté à l'instant t dans un autre état propre de $\hat{H}_0$. Les probabilités de transition se calculent à l'aide des mêmes méthodes, puisque ce problème est mathématiquement identique à celui où on appliquerait le couplage $\hat{W}$ de façon transitoire pendant un intervalle de temps $[0, t]$.

1.2 Transition entre niveaux discrets sous l'effet d'une perturbation dépendant du temps

1.2.1 Présentation du problème

Envisageons un système dont l'évolution se déduit d'un hamiltonien

$$\hat{H} = \hat{H}_0 + \hat{H}_1(t) . \quad (1.5)$$

Le premier terme $\hat{H}_0$ est indépendant du temps, ses vecteurs propres et valeurs propres sont notés $|n\rangle$ et E_n :

$$\hat{H}_0|n\rangle = E_n|n\rangle . \quad (1.6)$$

Le terme $\hat{H}_1(t)$ est un terme d'interaction dont les éléments de matrice sont supposés petits devant les écarts énergétiques $|E_n - E_m|$ entre états propres différents de $\hat{H}_0$. Il

peut dépendre explicitement du temps ou bien être constant entre les temps 0 et t et nul à l'extérieur de cet intervalle (perturbation « en créneau »).

Le couplage $\hat{H}_1(t)$ va pouvoir induire des transitions entre deux états propres $|n\rangle$ et $|m\rangle$ de $\hat{H}_0$ qui ne sont pas états propres de l'hamiltonien total $\hat{H}$. Nous nous proposons de calculer la probabilité $P_{n \rightarrow m}(t)$ de telles transitions en supposant pour simplifier que les niveaux sont non dégénérés. On trouvera des exposés plus complets de la question dans des ouvrages généraux de mécanique quantique³.

1.2.2 Exemples

Avant d'étudier les développements mathématiques, il peut être utile de présenter deux exemples de situations physiques correspondant au modèle développé. Ces exemples nous serviront dans la suite pour illustrer les résultats obtenus.

a. Interaction avec un champ électromagnétique classique

Considérons un atome décrit par un hamiltonien indépendant du temps $\hat{H}_0$ qui interagit avec une onde électromagnétique classique incidente dont le champ électrique au point où est situé l'atome est de la forme

$$\mathbf{E}(t) = \mathbf{E} \cos(\omega t + \varphi) . \quad (1.7)$$

Nous verrons dans le chapitre II que l'interaction entre l'atome et le champ peut, à une bonne approximation, s'écrire sous la forme d'un terme d'énergie dipolaire électrique

$$\hat{H}_1(t) = -\hat{\mathbf{D}} \cdot \mathbf{E}(t) , \quad (1.8)$$

où $\hat{\mathbf{D}}$ est le moment dipolaire électrique de l'atome⁴

$$\hat{\mathbf{D}} = q\hat{\mathbf{r}} \quad (1.9)$$

(q est la charge électrique de l'électron et $\mathbf{r}$ le rayon vecteur reliant le noyau à l'électron).

Sous l'action de $\hat{H}_1(t)$, un électron initialement dans un état propre de $|n, l, m\rangle$ de $\hat{H}_0$ va pouvoir être porté dans un autre état $|n', l', m'\rangle$. Si l'énergie de ce dernier état est plus *élevée* que l'énergie initiale, l'énergie nécessaire pour exciter l'atome est prise au champ électromagnétique : il y a *absorption*. Si au contraire l'énergie de l'état final est *inférieure* à l'énergie de l'état initial, il y a transfert d'énergie de l'atome vers le champ électromagnétique : il s'agit de *l'émission induite*. Nous reviendrons dans le chapitre II sur l'ensemble de ces phénomènes et leurs conséquences.

³Voir par exemple CDL Chapitre XIII.

⁴Pour simplifier, nous prendrons un atome à un électron, et plus précisément l'atome d'hydrogène.

b. Processus collisionnels

Considérons un atome A immobile, dont les niveaux d'énergie interne sont les valeurs propres de l'hamiltonien $\hat{H}_0$, et supposons qu'une autre particule B passe au voisinage de A (figure 1.1).

Appelons $\hat{\mathcal{V}}$ le potentiel d'interaction entre l'atome et la particule incidente (qui peut être par exemple une interaction dipôle-dipôle proportionnelle à $1/R^6$ de type *Van der Waals* si la particule B est elle-même un atome, ou une interaction en $1/R^4$ si la particule incidente est chargée). Pour l'atome A, cette interaction est un opérateur agissant dans l'espace des états de A. Ses éléments de matrice sont fonction de la distance R entre A et B, et tendent rapidement vers 0 lorsque R est très grand. Puisque R varie avec le temps, *l'hamiltonien d'interaction dépend lui aussi du temps*. Si avant la collision, quand les particules A et B sont éloignées l'une de l'autre, l'atome A est dans l'état $|n\rangle$, il va pouvoir se trouver à l'issue de la collision dans un autre état $|m\rangle$. La collision est *élastique* si les énergies initiales et finales des états $|m\rangle$ et $|n\rangle$ sont *égales*, ou *inélastique* si elles sont *différentes*. Ce type de transition induite par collision est par exemple responsable de l'excitation des atomes dans une lampe à décharge (un « tube au néon » par exemple) ou, comme nous le verrons dans le chapitre III, dans certains types de laser.

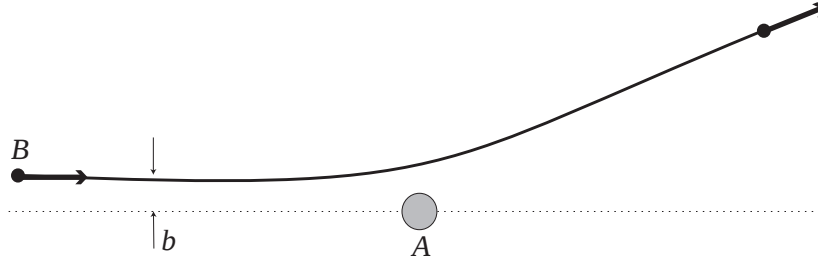


FIG. 1.1: **Interaction collisionnelle entre un atome A et une particule incidente B.** La distance b s'appelle le paramètre d'impact. La donnée de la trajectoire $\mathbf{R}(t)$ de la particule B permet de déterminer la variation temporelle du terme d'interaction $\hat{H}_1(t) = \hat{\mathcal{V}}(|\mathbf{R}(t)|)$.

1.2.3 Développement de la fonction d'onde en série de perturbation

Pour connaître l'évolution du système sous l'effet de l'hamiltonien $\hat{H}$ donné par l'équation (1.5), il faut résoudre l'équation de Schrödinger. À cette fin, nous allons utiliser une méthode de résolution approchée de type perturbatif, qui repose sur l'hypothèse que les éléments de matrice de $\hat{H}_1(t)$ sont petits⁵ devant ceux de $\hat{H}_0$. Pour mieux identifier les

⁵Plus précisément, dans la base $\{|n\rangle\}$ des états propres de $\hat{H}_0$, les éléments de matrice non diagonaux $|\langle n|\hat{H}_1|m\rangle|$ sont petits devant les différences d'énergie correspondantes $|E_n - E_m|$.

ordres successifs de perturbation, récrivons $\hat{H}_1(t)$ sous la forme

$$\hat{H}_1(t) = \lambda \hat{H}'_1(t) , \quad (1.10)$$

où $\hat{H}'_1(t)$ a des éléments de matrice du même ordre de grandeur que ceux de $\hat{H}_0$ et où λ est un paramètre réel sans dimension très petit devant 1, qui caractérise la force relative de l'interaction $\hat{H}_1(t)$. Dans le premier exemple du paragraphe (B.2), λ est proportionnel à l'amplitude E du champ incident. Dans le deuxième exemple, λ est une fonction décroissante du paramètre d'impact b . Dans chacun des cas, on pourra donc trouver des conditions expérimentales (onde électromagnétique suffisamment faible, collision à grand paramètre d'impact) pour lesquelles la méthode est valable.

Explicitons l'équation de Schrödinger

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = (\hat{H}_0 + \lambda \hat{H}'_1(t)) |\psi(t)\rangle \quad (1.11)$$

en utilisant le développement de $|\psi(t)\rangle$ sur la base des états propres de $\hat{H}_0$

$$\Rightarrow |\psi(t)\rangle = \sum_n \gamma_n(t) e^{-iE_n t/\hbar} |n\rangle . \quad (1.12)$$

Nous avons écrit le coefficient du ket $|n\rangle$ comme un produit de deux termes $\gamma_n(t)$ et $\exp(-iE_n t/\hbar)$. Cette séparation permet de tenir compte de l'évolution naturelle du système sous l'influence de $\hat{H}_0$ seul, (si $\hat{H}_1(t)$ est nul, les $\gamma_n(t)$ sont *constants* en vertu de l'équation (1.2)). Elle simplifie les développements ultérieurs.

Projetons l'équation (1.11) sur un état propre $|k\rangle$ de $\hat{H}_0$:

$$\begin{aligned} i\hbar \frac{d}{dt} \langle k | \psi(t) \rangle &= \langle k | \hat{H}_0 | \psi(t) \rangle + \lambda \langle k | \hat{H}'_1 | \psi(t) \rangle \\ &= E_k \langle k | \psi(t) \rangle + \lambda \sum_n \langle k | \hat{H}'_1(t) | n \rangle \langle n | \psi(t) \rangle . \end{aligned} \quad (1.13)$$

(On a utilisé la relation de fermeture : $\sum_n |n\rangle \langle n| = \hat{I}$).

En utilisant le développement (1.12) de $|\psi(t)\rangle$, nous récrivons (1.13) sous la forme

$$\left[E_k \gamma_k(t) + i\hbar \frac{d}{dt} \gamma_k(t) \right] e^{-iE_k t/\hbar} = E_k \gamma_k(t) e^{-iE_k t/\hbar} + \lambda \sum_n \langle k | \hat{H}'_1(t) | n \rangle \gamma_n(t) e^{-iE_n t/\hbar} . \quad (1.14)$$

Les termes proportionnels à E_k dans les membres de gauche et de droite se simplifient et nous obtenons

$$i\hbar \frac{d}{dt} \gamma_k(t) = \lambda \sum_n \langle k | \hat{H}'_1(t) | n \rangle e^{i(E_k - E_n)t/\hbar} \gamma_n(t) , \quad (1.15)$$

qui est un système (éventuellement infini) d'équations différentielles. Ce système est *exact*, aucune approximation n'ayant été faite jusqu'ici.

Les coefficients $\gamma_k(t)$ dépendent de λ . La méthode des perturbations consiste à développer $\gamma_k(t)$ en série entière de λ (qui est rappelons-le, un paramètre réel petit devant 1) :

$$\gamma_k(t) = \gamma_k^{(0)}(t) + \lambda \gamma_k^{(1)}(t) + \lambda^2 \gamma_k^{(2)}(t) + \dots \quad (1.16)$$

En reportant ces développements dans l'équation (1.15), nous pouvons identifier les termes de même puissance en λ . Nous obtenons ainsi :

– à l'ordre 0

$$i\hbar \frac{d}{dt} \gamma_k^{(0)}(t) = 0 ; \quad (1.17)$$

– à l'ordre 1

$$i\hbar \frac{d}{dt} \gamma_k^{(1)}(t) = \sum_n \langle k | \hat{H}'_1(t) | n \rangle e^{i(E_k - E_n)t/\hbar} \gamma_n^{(0)}(t) ; \quad (1.18)$$

– à l'ordre r

$$i\hbar \frac{d}{dt} \gamma_k^{(r)}(t) = \sum_n \langle k | \hat{H}'_1(t) | n \rangle e^{i(E_k - E_n)t/\hbar} \gamma_n^{(r-1)}(t) . \quad (1.19)$$

Ce système est susceptible d'être résolu par *itération*. En effet les termes d'ordre zéro $\gamma_k^{(0)}(t)$ sont connus : il s'agit de *constantes déterminées par l'état initial du système*. En portant ces termes dans (1.18) on peut trouver les termes d'ordre 1, $\gamma_k^{(1)}(t)$, qui eux mêmes permettent d'accéder aux termes d'ordre 2, $\gamma_k^{(2)}(t)$, et ainsi de suite. Il est donc possible, en principe, de déterminer successivement tous les termes du développement (1.16).

1.2.4 Théorie au premier ordre

a. Probabilité de transition

Supposons qu'à l'instant initial t_0 , le système se trouve dans un état propre $|i\rangle$ de $\hat{H}_0$. Il s'ensuit que tous les $\gamma_k(t_0)$ sont nuls à l'exception de $\gamma_i(t_0)$ qui est égal à 1. La solution de (1.17) est donc

$$\gamma_k^{(0)}(t) = \delta_{ki} . \quad (1.20)$$

Considérons à présent les transitions vers les niveaux $|k\rangle$ différents de l'état initial ($k \neq i$). En portant le résultat (1.20) dans l'équation (1.18) et en intégrant sur le temps, nous trouvons $\gamma_k^{(1)}(t)$:

$$\gamma_k^{(1)}(t) = \frac{1}{i\hbar} \int_{t_0}^t dt' \langle k | \hat{H}'_1(t') | i \rangle e^{i(E_k - E_i)t'/\hbar} . \quad (1.21)$$

L'amplitude de probabilité de trouver le système dans l'état $|k\rangle$ à l'instant t est, d'après (1.12) et (1.16), égale (à un facteur de phase près) à $(\gamma_k^{(0)}(t) + \lambda\gamma_k^{(1)}(t) + \dots)$. Pour un état $|k\rangle$ différent de $|i\rangle$, le terme d'ordre 0 est nul. Nous en déduisons, en utilisant (1.21), que l'amplitude de transition de $|i\rangle$ vers $|k\rangle$, à l'ordre 1, est, à un facteur de phase près, égale à

$$\Rightarrow S_{ki} = \lambda\gamma_k^{(1)}(t) = \frac{1}{i\hbar} \int_{t_0}^t dt' \langle k | \hat{H}_1(t') | i \rangle e^{i(E_k - E_i)t'/\hbar}, \quad (1.22)$$

puisque, d'après (1.10), $\lambda\hat{H}_1'(t')$ est égal à $\hat{H}_1(t')$. La probabilité de trouver le système dans l'état $|k\rangle$ est égale au carré du module de (1.22), c'est-à-dire :

$$\Rightarrow P_{i \rightarrow k}(t) = \left| \frac{1}{i\hbar} \int_{t_0}^t dt' \langle k | \hat{H}_1(t') | i \rangle e^{i(E_k - E_i)t'/\hbar} \right|^2. \quad (1.23)$$

Les formules (1.22) et (1.23) constituent les résultats essentiels de la théorie des perturbations dépendant du temps, au premier ordre. Ce sont ces formules que nous exploiterons par la suite.

Notons que l'approche perturbative n'est valable que si

$$P_{i \rightarrow k} \ll 1. \quad (1.24)$$

C'est en effet dans le cas où l'interaction $\hat{H}_1$ induit de faibles effets à l'ordre le plus bas du calcul perturbatif que la série de perturbation (1.16) a de bonnes chances de converger rapidement. La condition (1.24) n'est d'ailleurs qu'une condition *nécessaire*, et non pas suffisante, pour que la théorie des perturbations au premier ordre puisse être appliquée.

b. Exemple de processus collisionnel : étude qualitative du domaine d'énergie accessible

En utilisant simplement les propriétés de la *transformée de Fourier* appliquées à la formule (1.23), nous allons montrer dans ce paragraphe qu'on peut prévoir dans quel domaine d'énergie des niveaux atomiques peuvent être excités au cours d'une collision.

Supposons pour simplifier que le terme d'interaction $\hat{H}_1(t)$ soit de la forme

$$\hat{H}_1(t) = \hat{W} f(t), \quad (1.25)$$

où $\hat{W}$ est un opérateur agissant sur les variables atomiques, et $f(t)$ est une fonction réelle du temps qui tend vers 0 quand $t \rightarrow \pm\infty$ et atteint sa valeur maximale en $t = 0$ (voir figure 1.2). Nous supposons qu'avant la collision ($t_0 = -\infty$) le système est dans l'état $|i\rangle$. L'amplitude de probabilité pour trouver le système dans l'état $|k\rangle$ après la collision ($t_0 = +\infty$) a pour valeur

$$S_{ki} = \frac{W_{ki}}{i\hbar} \int_{-\infty}^{+\infty} dt f(t) e^{i(E_k - E_i)t/\hbar}, \quad (1.26)$$

où W_{ki} est l'élément de matrice

$$W_{ki} = \langle k | \hat{W} | i \rangle .$$

Introduisons la transformée de Fourier $\tilde{f}(E)$ de la fonction $f(t)$:

$$\tilde{f}(E) = \frac{1}{\sqrt{2\pi\hbar}} \int_{-\infty}^{+\infty} dt f(t) e^{iEt/\hbar} . \quad (1.27)$$

Nous obtenons alors pour la probabilité de transition $P_{i \rightarrow k}$ l'expression suivante

$$P_{i \rightarrow k} = \frac{2\pi}{\hbar} |W_{ki}|^2 |\tilde{f}(E_k - E_i)|^2 , \quad (1.28)$$

qui dépend de la valeur de $\tilde{f}(E)$ prise en $E = E_k - E_i$.

Utilisons maintenant une propriété connue des largeurs de transformées de Fourier (voir figure 1.2) : si la largeur de la fonction $f(t)$ est Δt , la largeur de sa transformée de Fourier $\tilde{f}(E)$ est de l'ordre de $\hbar/\Delta t$. L'expression (1.28) montre alors que si la durée de l'interaction est de l'ordre de Δt , seuls les niveaux d'énergie pour lesquels

$$|E_k - E_i| < \frac{\hbar}{\Delta t} \quad (1.29)$$

auront une *probabilité appréciable d'être peuplés*.

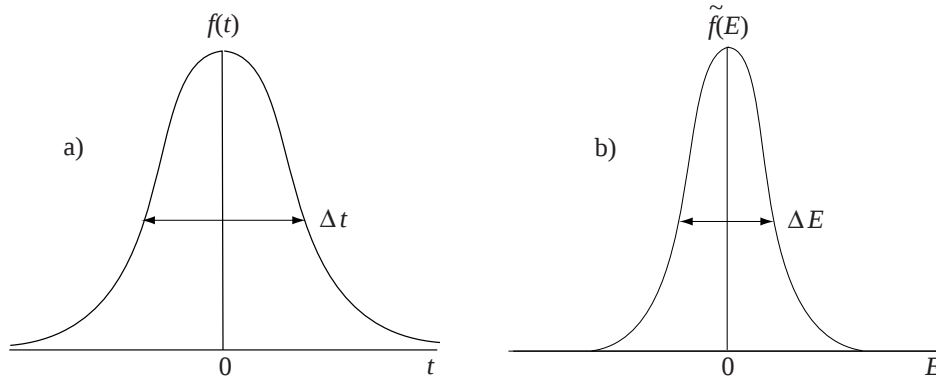


FIG. 1.2: **La transformée de Fourier de la fonction** $f(t)$, centrée à l'origine et de largeur Δt (fig. B.2.a) est une fonction de E centrée à l'origine et de largeur $\Delta E \approx \hbar/\Delta t$ (fig. B.2.b).

Prenons le cas d'une collision dans laquelle l'interaction est une fonction décroissante de la distance R entre partenaires de la collision (voir Figure 1.1). Sa durée à mi-hauteur est de l'ordre de b/V , où b est le « paramètre d'impact » (valeur minimale de R pendant la collision) et V la vitesse relative des particules. L'inégalité (1.29) implique que seuls les niveaux d'énergie $|k\rangle$ tels que

$$|E_k - E_i| \leq \frac{\hbar V}{b} \quad (1.30)$$

seront peuplés de manière appréciable à l'issue de la collision.

À titre d'exemple, considérons une collision avec un paramètre d'impact $b = 5$ nm (soit 100 fois de rayon de Bohr a_0), par un atome ayant une vitesse de 500 m/s. Une telle collision est susceptible de provoquer une transition vers des niveaux séparés de moins de $\hbar V/b \sim 10^{-22}$ J, soit encore $|E_k - E_i|/\hbar \lesssim 10^{11}$ Hz. Ce calcul simple montre que la thermalisation des populations des sous-niveaux hyperfins atomiques, dont les séparations sont au plus de quelques GHz, peut se réaliser sous l'effet de collisions à grand paramètre d'impact dues par exemple à l'interaction de Van der Waals. Ce n'est pas le cas pour les niveaux électroniques, séparés de plus de 10^{14} Hz.

En revanche, si le projectile incident est un électron d'énergie 50 eV, dont la vitesse est de l'ordre de 5×10^6 m/s, une collision avec un paramètre d'impact de 5 nm a une durée b/V de l'ordre de 10^{-15} s : elle peut provoquer une transition vers des états électroniques élevés, susceptibles de se désexciter en émettant de la lumière visible (la lumière orangée de longueur d'onde 0,6 micromètres a une fréquence de 5×10^{14} Hz). C'est ce qui se passe dans une lampe à décharge.

Remarque

Le raisonnement ci-dessus ne prend pas en compte les questions de conservation de l'énergie. Il faut vérifier que l'énergie du projectile est supérieure à l'énergie nécessaire pour exciter la cible, car sinon l'hypothèse d'une vitesse constante du projectile n'est pas valable. Dans les exemples ci-dessus il n'y a pas de problème.

c. Cas d'une perturbation constante « branchée » à partir d'un instant déterminé

Il arrive souvent qu'un système soit brusquement mis en contact à partir d'un instant initial $t = 0$ avec une perturbation $\hat{W}$ qui est ensuite constante⁶. Nous allons donner l'expression de la probabilité de transition au premier ordre dans ce cas important, expression qui nous sera utile dans la suite de ce chapitre.

Si à l'instant $t = 0$, le système est dans l'état propre $|i\rangle$ de $\hat{H}_0$, l'amplitude de probabilité de le trouver dans l'état $|k\rangle$ à l'instant T se calcule à partir de (1.22) et vaut

$$S_{ki}(T) = \frac{W_{ki}}{i\hbar} \frac{e^{i(E_k - E_i)T/\hbar} - 1}{i(E_k - E_i)/\hbar} . \quad (1.31)$$

Nous en déduisons la probabilité de transition $P_{i \rightarrow k}(T)$

$$\Rightarrow P_{i \rightarrow k}(T) = \frac{|W_{ki}|^2}{\hbar^2} g_T(E_k - E_i) , \quad (1.32)$$

où

$$g_T(E) = \frac{\sin^2(ET/2\hbar)}{(ET/2\hbar)^2} T^2 \quad (1.33)$$

est la fonction représentée sur la figure 1.3.

⁶Ce calcul s'applique aussi au cas d'un hamiltonien $\hat{H}_0 + \hat{W}$ indépendant du temps, mais où le système est préparé et détecté dans un état propre de $\hat{H}_0$.

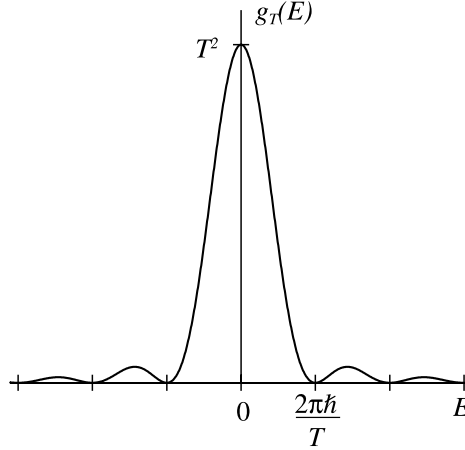


FIG. 1.3: **Allure de la fonction** $g_T(E)$. Ce graphe montre qu'une interaction de durée T permet d'atteindre préférentiellement les états dont l'énergie est égale à celle de l'état initial à $\pi\hbar/T$ près.

Les caractéristiques essentielles de cette fonction sont les suivantes :

- elle est maximum en $E = 0$, où sa valeur est T^2 ;
- sa largeur est de l'ordre de $2\pi\hbar/T$;
- sa surface est proportionnelle à T , plus précisément

$$\int_{-\infty}^{+\infty} dE g_T(E) = 2\pi\hbar T \quad (1.34)$$

(valeur qui n'est autre que le demi-produit de sa hauteur par la distance entre ses deux premiers zéros, comme si la fonction g_T était triangulaire).

Nous poserons

$$\delta_T(E) = \frac{g_T(E)}{2\pi\hbar T} = \frac{2\hbar \sin^2(ET/2\hbar)}{\pi T E^2} . \quad (1.35)$$

La fonction $\delta_T(E)$ est « piquée » en 0, de surface unité, de hauteur $T/2\pi\hbar$, de largeur $2\pi\hbar/T$. Elle constitue, lorsque T est suffisamment grand, une approximation de la fonction de Dirac, et on montre que $\lim_{T \rightarrow +\infty} \delta_T(E) = \delta(E)$. Nous pouvons écrire (1.32) sous la forme

$$\Rightarrow P_{i \rightarrow k}(T) = T \frac{2\pi}{\hbar} |W_{ki}|^2 \delta_T(E_k - E_i) . \quad (1.36)$$

Nous retrouvons dans ce cas particulier le résultat du paragraphe précédent. Un couplage de durée T n'est efficace qu'entre états d'énergies voisines. On trouve plus précisément ici

$$|E_k - E_i| \lesssim \frac{\pi\hbar}{T} . \quad (1.37)$$

Les résultats (1.36) et (1.37) seront utilisés dans la partie 1.3, où le couplage peut induire des transitions vers un ensemble de niveaux $|k\rangle$ très serrés. La distribution des états finaux

atteints après une interaction de durée T est centrée sur l'énergie initiale E_i , et sa largeur décroît avec T .

Dans le cas où on s'intéresse à la probabilité de transition entre 2 niveaux spécifiques $|i\rangle$ et $|k\rangle$ non dégénérés ($E_i \neq E_k$), les expressions (1.32) et (1.33) montrent que la probabilité de transition oscille avec le temps d'interaction

$$P_{i \rightarrow k}(T) = \frac{4|W_{ki}|^2}{|E_k - E_i|^2} \sin^2 \left(\frac{E_k - E_i}{2\hbar} T \right). \quad (1.38)$$

Comme le montre la figure (1.4), la fréquence de l'oscillation augmente avec l'écart en énergie, tandis que la valeur maximale de la probabilité de transition décroît comme $(E_k - E_i)^{-2}$. Nous retrouverons de façon plus rigoureuse ce résultat (oscillation de Rabi) par un traitement non perturbatif, au paragraphe 1.2.6.b.

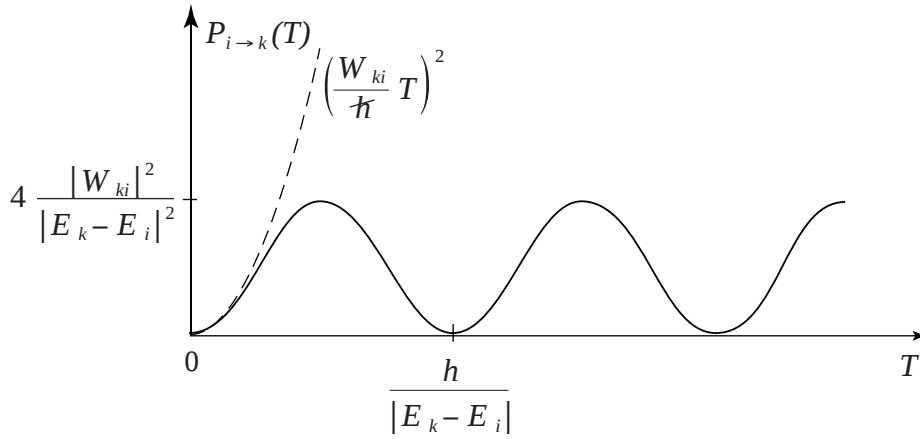


FIG. 1.4: Probabilité de transition entre 2 niveaux discrets, en fonction de temps d'interaction T . La courbe tiretée montre le démarrage parabolique indépendant de $|E_k - E_i|$.

Il est clair que le résultat de la figure (1.4) n'est valable que si $|W_{ki}| \ll |E_k - E_i|$, ce qui est précisément la condition de validité a priori du traitement perturbatif (cf. § 1.2.3, note 5). On peut pourtant se demander ce que donne $P_{i \rightarrow k}(T)$ au voisinage de $T = 0$, lorsque $|E_k - E_i|$ tend vers 0. On constate alors que la probabilité de transition démarre quadratiquement, suivant une loi indépendante de $|E_k - E_i|$:

$$P_{i \rightarrow k}(T) \underset{|E_k - E_i| \rightarrow 0}{\simeq} \left| \frac{W_{ki}}{\hbar} \right|^2 T^2. \quad (1.39)$$

La validité de cette formule est évidemment limitée aux valeurs de T plus petites que $\hbar/|W_{ki}|$. Notons que le membre de droite de (1.39) coïncide aussi avec la valeur de $P_{i \rightarrow k}(T)$ autour de $T = 0$ pour $|E_k - E_i| \neq 0$.

Remarque

Le problème que nous venons de traiter ici recouvre en fait deux situations physiques différentes, en fonction du comportement de la perturbation après l'instant T : ou bien $\hat{W}$ est interrompu à l'instant T et le système n'évolue plus ultérieurement (perturbation en « créneau », de type collisionnel), ou bien $\hat{W}$ reste constant (perturbation en « marche de trottoir »), mais on s'intéresse à l'état du système à l'instant T (par exemple en effectuant une mesure à l'instant T).

d. Cas d'une perturbation sinusoïdale

Nous avons vu dans le paragraphe 2.a que dans le cas de l'interaction atome-rayonnement (équation 1.8) on avait souvent affaire à des perturbations sinusoïdales, c'est-à-dire de la forme

$$\hat{H}_1(t) = \hat{W} \cos(\omega t + \varphi) . \quad (1.40)$$

L'amplitude de transition pour passer de $|i\rangle$ à $|k\rangle$, calculée par l'expression (1.22), vaut alors

$$S_{ki} = -\frac{W_{ki}}{2\hbar} \left[\frac{e^{i(\omega_{ki}-\omega)t-i\varphi} - e^{i(\omega_{ki}-\omega)t_0-i\varphi}}{\omega_{ki} - \omega} + \frac{e^{i(\omega_{ki}+\omega)t+i\varphi} - e^{i(\omega_{ki}+\omega)t_0+i\varphi}}{\omega_{ki} + \omega} \right] , \quad (1.41)$$

la quantité ω_{ki} étant la fréquence de Bohr associée à la transition de $|i\rangle$ à $|k\rangle$

$$E_k - E_i = \hbar\omega_{ki} . \quad (1.42)$$

L'amplitude S_{ki} apparaît comme la somme de deux termes dont les dénominateurs respectifs sont $\omega_{ki} - \omega$ et $\omega_{ki} + \omega$. Pour avoir une *amplitude de transition relativement importante*, il faut se placer dans des conditions où *l'un de ces dénominateurs devient très petit*. Le terme correspondant dans (1.41) est alors très grand devant l'autre.

Dans le cas où $\omega_{ki} > 0$ (état final d'énergie supérieure à celle de l'état initial), cette résonance se produit lorsque

$$|\omega_{ki} - \omega| \ll \omega \quad (1.43)$$

(*condition d'excitation quasi-résonnante*). Cette condition est en fait nécessaire pour avoir une probabilité d'excitation décelable. En effet, en prenant l'exemple de l'interaction d'un atome avec un champ électromagnétique du domaine visible⁷, $W_{ki}/\hbar\omega$ excède rarement 10^{-6} , même avec une source laser usuelle⁸. Pour avoir une amplitude de transition relativement importante, il faut que le dénominateur $|\omega_{ki} - \omega|$ soit bien plus petit que ω . Comme $\omega_{ki} + \omega$ est plus grand que ω , la condition (1.43) implique donc que le second terme de (1.41), appelé « anti-résonnant », peut être négligé devant le premier⁹. On a

⁷Rappelons que pour des radiations visibles, la fréquence $\omega/2\pi$ est de l'ordre de quelques 10^{14} Hz.

⁸Il faut cependant souligner qu'on peut atteindre des valeurs de ce rapport qui sont proches de 1, voir même supérieures, si l'on utilise des lasers émettant des impulsions dont la durée est de quelques dizaines de femtosecondes (voir chapitre III paragraphe D.1.).

⁹Pour des raisons historiques, cette approximation est parfois appelée « approximation du champ tournant » (en anglais « rotating wave approximation »). En effet, elle a été introduite dans le domaine des radiofréquences (Résonance Magnétique Nucléaire par exemple) où cette approximation est équivalente à remplacer un champ magnétique oscillant le long d'un axe par sa composante circulaire résonnante. En optique, l'appellation « approximation quasi-résonnante » traduit plus fidèlement la situation physique.

alors une expression simple de la probabilité de transition :

$$P_{i \rightarrow k}(T) = \frac{|W_{ki}|^2}{4\hbar^2} g_T(E_k - E_i - \hbar\omega) \quad (1.44)$$

$$\Rightarrow P_{i \rightarrow k}(T) = T \frac{|W_{ki}|^2}{4} \frac{2\pi}{\hbar} \delta_T(E_k - E_i - \hbar\omega) . \quad (1.45)$$

Dans ces formules, $g_T(E)$ et $\delta_T(E)$ sont les fonctions introduites au paragraphe précédent (équations (1.33) et (1.35)), et $T = t - t_0$ est la durée de l'interaction.

Dans le cas où ω_{ki} est négatif (état final d'énergie inférieure à celle de l'état initial), la condition d'excitation quasi-résonnante est remplie lorsque

$$|\omega_{ki} + \omega| \ll \omega . \quad (1.46)$$

C'est alors le premier terme de S_{ki} qui est négligeable, et on obtient l'expression de $P_{i \rightarrow k}$ en remplaçant ω par $-\omega$ dans (1.44) ou (1.45).

Les expressions finales de la probabilité de transition sont donc analogues à (1.32) et (1.36) – à un facteur $1/4$ près – lorsqu'on remplace $E_k - E_i$ par $E_k - E_i \pm \hbar\omega$. Les conclusions du paragraphe (c) se transposent alors aisément au cas de la perturbation sinusoïdale : les niveaux $|k\rangle$ qui vont être peuplés de manière efficace au bout d'un temps d'interaction T sont ceux dont l'énergie est telle que

$$\begin{cases} E_i + \hbar\omega - \frac{\pi\hbar}{T} < E_k < E_i + \hbar\omega + \frac{\pi\hbar}{T} \\ E_i - \hbar\omega - \frac{\pi\hbar}{T} < E_k < E_i - \hbar\omega + \frac{\pi\hbar}{T} . \end{cases} \quad (1.47)$$

Pour des temps d'interaction suffisamment longs, il n'y a de *transition possible que vers des niveaux $|k\rangle$ distants énergétiquement de $|i\rangle$ de la quantité $\hbar\omega$: les variations d'énergie de l'atome se font uniquement par absorption ou émission de quanta d'énergie $\hbar\omega$.*

Nous retrouvons ici la règle que Bohr avait introduite empiriquement aux premiers temps de la mécanique quantique. En anticipant sur la quantification du rayonnement, et dans le cas de l'interaction matière-rayonnement, cette règle a une interprétation naturelle en termes de *photons* d'énergie $\hbar\omega$ absorbés ou émis lors de la transition. Notons cependant qu'elle a été établie sans qu'on ait eu besoin d'introduire a priori cette notion. À ce niveau, le concept de photon est une commodité, pas une nécessité.

Remarque

L'approximation quasi-résonnante implique, en plus de la condition (1.43) ou (1.46), que $W_{ki}/\hbar\omega$ soit très petit devant 1, ce qui est généralement bien vérifié dans le domaine optique. En revanche, pour les champs ayant une fréquence ω beaucoup plus petite (domaine des champs « radiofréquence », où $\omega/2\pi$ est de l'ordre de 10^9 Hz), il est possible d'atteindre des intensités pour lesquelles $W_{ki}/\hbar\omega$ est voisin de 1 (et parfois supérieur). Dans ces conditions, il n'est plus possible de négliger le terme anti-résonnant. Sa contribution première est de déplacer la position du centre de la résonance d'une quantité de l'ordre de $W_{ki}^2/\hbar^2\omega$ (effet Bloch-Siegert¹⁰).

¹⁰Voir par exemple CDG 2 Complément A_{VI}

1.2.5 Calculs au second ordre

Dans de nombreuses situations l'état initial $|i\rangle$ et l'état final $|k\rangle$ ne sont pas couplés directement par l'hamiltonien de perturbation $\hat{H}_1(t)$ (on a $\langle i|\hat{H}_1(t)|k\rangle = 0$). La probabilité de transition (1.23) est donc nulle au premier ordre. En revanche, $|i\rangle$ et $|k\rangle$ étant couplés par $\hat{H}_1(t)$ à de mêmes états $|j\rangle$ ($\langle i|\hat{H}_1(t)|j\rangle \neq 0$ et $\langle j|\hat{H}_1(t)|k\rangle \neq 0$), nous allons montrer que la probabilité de transition au second ordre est susceptible d'être non nulle. On peut se représenter cette situation en disant que le système « passe intermédiairement » de $|i\rangle$ à $|j\rangle$, puis de $|j\rangle$ à $|k\rangle$. Des exemples de telles situations se rencontrent fréquemment dans les problèmes d'interaction atome-rayonnement et nous aurons l'occasion d'utiliser les résultats de ce paragraphe dans les chapitres suivants (voir figure 2.9 du chapitre II).

Dans l'hypothèse où $\langle k|\hat{H}_1(t)|i\rangle$ est nul, il faut calculer les $\gamma_j^{(1)}(t)$ avec j différent de k et reporter leurs valeurs (données par des expressions analogues à (1.21)) dans l'expression (1.19) pour calculer $\gamma_k^{(2)}(t)$. L'amplitude de transition de $|i\rangle$ à $|k\rangle$ entre 0 et T est, à un facteur de phase près, égale à

$$\begin{aligned} S_{ki}(T) &= \lambda^2 \gamma_k^{(2)}(T) \\ &= \frac{1}{(i\hbar)^2} \int_0^T dt' \int_0^{t'} dt'' \sum_j \langle k|\hat{H}_1(t')|j\rangle \langle j|\hat{H}_1(t'')|i\rangle e^{i(E_k - E_j)t'/\hbar} e^{i(E_j - E_i)t''/\hbar}. \end{aligned} \quad (1.48)$$

(Notons que la somme porte sur les états j différents de i et k puisque nous avons supposé $\langle k|\hat{H}_1(t)|i\rangle = 0$).

Considérons à présent le cas particulier important où $\hat{H}_1(t)$ est de la forme

$$\hat{H}_1(t) = \hat{W} f(t), \quad (1.49)$$

où $\hat{W}$ est un opérateur et $f(t)$ une « fonction de branchement » de temps caractéristique θ (voir figure 1.5.a). Plus précisément nous supposons que la durée de l'interaction T est grande devant ce temps de branchement θ , et que θ est lui-même très long devant les temps d'évolution caractéristiques du système atomique libre, de la forme $\hbar/|E_i - E_j|$ où $|j\rangle$ est l'un des niveaux intermédiaires de la transition reliant $|i\rangle$ à $|k\rangle$:

$$T \gg \theta \gg \frac{\hbar}{|E_i - E_j|}. \quad (1.50)$$

Nous montrons à la fin de ce paragraphe que la probabilité de transition $P_{i \rightarrow k}$ peut se mettre sous la forme :

$$P_{i \rightarrow k} = T \left| \sum_{j \neq k, i} \frac{\langle k|\hat{W}|j\rangle \langle j|\hat{W}|i\rangle}{E_j - E_i} \right|^2 \frac{2\pi}{\hbar} \delta_T(E_k - E_i), \quad (1.51)$$

où $\delta_T(E)$ est la fonction de largeur $2\pi\hbar/T$ introduite en (1.35). La comparaison avec la formule (1.36) montre que l'on peut transposer au deuxième ordre les résultats du premier

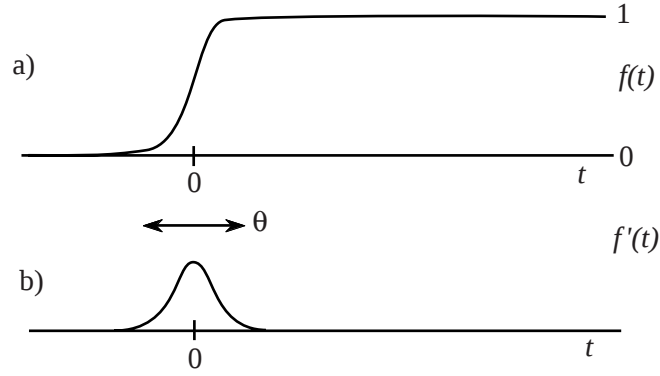


FIG. 1.5: a) **Allure de la fonction de branchement** $f(t)$, présentant une variation continue de 0 à 1 se produisant sur un temps θ ; b) allure de la dérivée de $f(t)$.

ordre en remplaçant l'hamiltonien $\hat{W}$ par un *hamiltonien effectif* $\hat{W}_{\text{eff}}$ dont l'élément de matrice entre les états $|i\rangle$ et $|k\rangle$ vaut

$$\langle k | \hat{W}_{\text{eff}} | i \rangle = \sum_{j \neq k, i} \frac{\langle k | \hat{W} | j \rangle \langle j | \hat{W} | i \rangle}{E_i - E_j}. \quad (1.52)$$

Cet élément de matrice, et donc la probabilité de transition $P_{i \rightarrow k}$, sont d'autant plus importants qu'il existe un ou plusieurs niveaux « relais » $|j\rangle$ dont l'énergie est proche de l'énergie de l'état initial $|i\rangle$.

Remarque

On peut aussi montrer que la perturbation $\hat{H}_1(t)$ provoque des déplacements des niveaux d'énergie, égaux aux éléments diagonaux du hamiltonien effectif $\hat{W}_{\text{eff}}$ (voir par exemple le paragraphe C.4 du chapitre II). Ainsi, le déplacement du niveau $|i\rangle$ vaut

$$\langle i | \hat{W}_{\text{eff}} | i \rangle = \sum_{j \neq i} \frac{\langle i | \hat{W} | j \rangle \langle j | \hat{W} | i \rangle}{E_i - E_j}. \quad (1.53)$$

À la limite des très longs temps d'interaction, il ne peut y avoir de transitions qu'entre niveaux dont les énergies *déplacées* sont identiques.

Donnons une démonstration, relativement rudimentaire, de la formule (1.51) (on trouvera dans CDG 1, Complément B_{IV} une démonstration plus élaborée). En utilisant la forme (1.49) de l'hamiltonien $\hat{H}_1(t)$, nous trouvons à partir de (1.48) l'amplitude de transition $S_{ki}(T)$ pour passer de l'état $|i\rangle$ à un instant $t_0 < 0$ quelconque avant le branchement à l'état $|k\rangle$ à l'instant T :

$$S_{ki}(T) = -\frac{1}{\hbar^2} \sum_{j \neq k, i} W_{kj} W_{ji} \int_{t_0}^T dt' \int_{t_0}^{t'} dt'' e^{i(E_k - E_j)t'/\hbar} e^{i(E_j - E_i)t''/\hbar} f(t') f(t''). \quad (1.54)$$

Modifions l'intégrale sur t'' en procédant à une intégration par parties :

$$\int_{t_0}^{t'} dt'' e^{i(E_j - E_i)t''/\hbar} f(t'') = \frac{e^{i(E_j - E_i)t'/\hbar}}{i(E_j - E_i)/\hbar} f(t') - \int_{t_0}^{t'} \frac{e^{i(E_j - E_i)t''/\hbar}}{i(E_j - E_i)/\hbar} f'(t'') dt''. \quad (1.55)$$

Les hypothèses (1.50) sur la forme de $f(t)$ nous permettent de majorer $f'(t'')$ par $1/\theta$ (voir Figure B.4.b). Le second terme du second membre de (1.55) est donc plus petit que le premier par un facteur de l'ordre de $\hbar/\theta|E_j - E_i|$, que l'on a supposé petit devant 1. En négligeant ce terme, (1.54) devient

$$S_{ki}(T) = \frac{1}{i\hbar} \sum_{j \neq i} \frac{W_{kj}W_{ji}}{E_i - E_j} \int_0^T dt' e^{i(E_k - E_i)t'/\hbar} (f(t'))^2 \quad (1.56a)$$

que l'on transforme en

$$S_{ki}(T) \approx \frac{1}{i\hbar} \sum_{j \neq i} \frac{W_{kj}W_{ji}}{E_i - E_j} \int_0^T dt' e^{i(E_k - E_i)t'/\hbar} . \quad (1.56b)$$

L'égalité entre (1.56a) et (1.56b) est vraie à des termes d'ordre θ/T près. Le carré de l'intégrale apparaissant dans (1.56b) est égal à la fonction $g_T(E_k - E_i)$ introduite en (1.33) pour le calcul au premier ordre. La formule (1.51) est donc bien établie.

1.2.6 Comparaison avec la solution exacte dans le cas d'un système à deux niveaux (oscillations de Rabi)

Il peut être intéressant de comparer la solution perturbative au premier ordre du paragraphe B.4 avec une solution exacte à tous les ordres. Cette solution exacte existe pour le système quantique le plus simple, en l'occurrence le système à deux niveaux, soumis à une interaction constante appliquée brusquement à partir de $t = 0$.

a. Quelques formules utiles

Considérons un hamiltonien $\hat{H}_0$ défini sur un espace de Hilbert de dimension 2, qui a pour états propres les états $|a\rangle$ et $|b\rangle$, d'énergie E_a et E_b . D'une manière identique à celle du paragraphe (B.4.c), on applique à partir de l'instant 0 une interaction $\hat{W}$ qui est par la suite constante. Nous supposons pour simplifier que les éléments de matrice diagonaux W_{aa} et W_{bb} sont nuls, et que W_{ab} est réel.

Supposons que le système est initialement dans l'état $|a\rangle$, et calculons la probabilité de transition $P_{a \rightarrow b}(T)$. Pour résoudre ce problème, il nous faut connaître les états propres $|\varphi_1\rangle$ et $|\varphi_2\rangle$, et les énergies E_1 et E_2 , de l'hamiltonien total $\hat{H}_0 + \hat{W}$. On montre¹¹ que les états propres sont donnés par

$$\begin{cases} |\varphi_1\rangle = \cos \theta |a\rangle + \sin \theta |b\rangle \\ |\varphi_2\rangle = -\sin \theta |a\rangle + \cos \theta |b\rangle \end{cases} , \quad (1.57a)$$

l'angle θ étant déterminé par

$$\tan 2\theta = 2W_{ab}/(E_a - E_b) . \quad (1.57b)$$

¹¹Voir CDL complément B_{IV}.

Les énergies correspondantes valent

$$\begin{cases} E_1 = \frac{1}{2}(E_a + E_b) + \frac{1}{2}\sqrt{(E_a - E_b)^2 + 4W_{ab}^2} \\ E_2 = \frac{1}{2}(E_a + E_b) - \frac{1}{2}\sqrt{(E_a - E_b)^2 + 4W_{ab}^2} \end{cases} \quad (1.58)$$

b. Évolution temporelle

L'état du système à l'instant t est obtenu en décomposant l'état initial $|a\rangle$ sur les états propres $|\varphi_1\rangle$ et $|\varphi_2\rangle$ et en multipliant les coefficients relatifs à ces deux états par les facteurs de phase $\exp(-iE_1t/\hbar)$ et $\exp(-iE_2t/\hbar)$. Un calcul sans difficulté¹² permet de calculer le vecteur d'état du système à l'instant T , et donc la probabilité de transition $P_{a \rightarrow b}(T)$, qui vaut

$$P_{a \rightarrow b}(T) = \frac{4W_{ab}^2}{(E_a - E_b)^2 + 4W_{ab}^2} \sin^2 \left\{ \sqrt{(E_a - E_b)^2 + 4W_{ab}^2} \frac{T}{2\hbar} \right\} . \quad (1.59)$$

On obtient un *comportement oscillant de la probabilité de transition*, représenté sur la figure 1.6, et appelé « *oscillation de Rabi* ». La pulsation caractéristique de cette oscillation est appelée *pulsation de Rabi*.

Lorsque les deux niveaux *ont même énergie* ($E_a = E_b$), la probabilité atteint périodiquement la valeur 1 (*transfert total de $|a\rangle$ vers $|b\rangle$*), et ce *quelle que soit la valeur du couplage W_{ab} entre les deux niveaux*. La pulsation de Rabi correspondante, $W_{ab}/\hbar$, est proportionnelle au couplage entre les deux niveaux. Lorsque W_{ab} est très petit, la probabilité maximale est toujours égale à un, mais le temps mis pour l'atteindre devient de plus en plus long.

Lorsque *les deux niveaux n'ont pas la même énergie*, l'oscillation est plus rapide, mais le transfert de $|a\rangle$ vers $|b\rangle$ n'est en revanche *jamais total*, aussi grande que soit la force de l'interaction W_{ab} .

Remarque

L'oscillation de Rabi entre les deux états $|a\rangle$ et $|b\rangle$ présente de nombreuses analogies avec ce qui se produit lorsqu'on couple deux oscillateurs en mécanique classique. On trouve en effet dans ce cas que l'énergie passe périodiquement d'un oscillateur à l'autre, avec une fréquence d'oscillation qui dépend linéairement de la valeur du couplage entre les deux oscillateurs.

c. Comparaison avec la théorie des perturbations

Envisageons d'abord *la limite des temps courts*. La probabilité de transition (1.59) a alors la valeur approchée suivante, correspondant au démarrage de la précession de Rabi :

$$P_{a \rightarrow b}(T) \approx \frac{W_{ab}^2}{\hbar^2} T^2 . \quad (1.60)$$

On a donc *une loi en T^2* , qui coïncide avec les prédictions de la théorie des perturbations au premier ordre (voir équations (1.32) et (1.33)), également à la limite des temps courts.

¹²On trouvera dans le chapitre II, paragraphe C.2, un calcul analogue, mais plus détaillé, effectué dans le cas particulier de l'interaction d'un système à deux niveaux avec un champ électromagnétique quasi-résonnant.

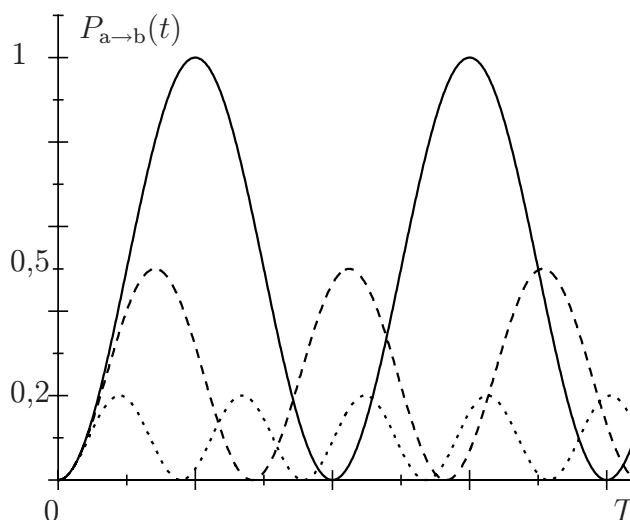


FIG. 1.6: **Évolution temporelle de la probabilité de transition** de a vers b pour des niveaux de même énergie (trait plein), ou d'énergies différentes (traits tiretés : $|E_a - E_b| = 2|W_{ab}|$; pointillés : $|E_a - E_b| = 4|W_{ab}|$).

Notons que ce résultat a été obtenu à la limite où $\sin x \approx x$, ce qui implique que $x \ll 1$, donc que $P_{a \rightarrow b}(T)$ est dans ce cas très petit devant 1.

Que donne cette comparaison pour des temps d'interaction plus longs ? Plaçons-nous dans la situation où les niveaux a et b ne sont pas dégénérés, et de plus

$$|W_{ab}| \ll |E_a - E_b|. \quad (1.61)$$

Dans ce cas, la formule (1.59) devient, à l'ordre le plus bas

$$P_{a \rightarrow b}(T) \approx \frac{4W_{ab}^2}{(E_a - E_b)^2} \sin^2 \frac{E_a - E_b}{2\hbar} T. \quad (1.62)$$

Cette expression approchée du résultat exact est identique au résultat (1.38) de la théorie perturbative. La théorie des perturbations donne donc une excellente approximation de l'amplitude et de la fréquence de l'oscillation de Rabi, dans le cas non dégénéré et lorsque le couplage est faible (équation 1.61). En revanche, la théorie des perturbations ne peut prédire l'oscillation de Rabi qui apparaît dans le cas exactement dégénéré, et où la probabilité de transition atteint la valeur de 1.

1.3 Cas d'un niveau discret couplé à un continuum : Règle d'or de Fermi

Après avoir étudié le cas d'un niveau isolé couplé à un autre niveau isolé, nous allons maintenant considérer une situation radicalement différente, où le niveau initial $|i\rangle$

est couplé à un ensemble de niveaux $|k\rangle$ faisant partie d'un *continuum*. Nous allons voir que dans ce cas, l'évolution du système présente de notables différences avec la situation précédente. Après avoir donné un exemple de telles situations, fréquentes en mécanique quantique, nous calculerons cette évolution et aboutirons à la règle d'or de Fermi. Signalons que le complément I.2 montre comment on peut passer continûment du couplage avec un niveau discret au couplage avec un continuum.

1.3.1 Exemple : auto-ionisation de l'hélium

a. Modèle des particules indépendantes

L'hélium possède deux électrons, plongés dans le potentiel coulombien du noyau doublement chargé He^{2+} . On peut donner un modèle grossier de ce système en ignorant l'interaction entre ces électrons. Chacun de ces électrons est donc considéré comme plongé dans un potentiel Coulombien. Les états propres d'un tel système hydrogénoïde sont bien connus : ils sont analogues à ceux de l'atome d'hydrogène, à un facteur près correspondant à la charge double du noyau. On a d'abord une suite discrète d'états liés. Si on ignore le spin de l'électron et le moment magnétique associé, ainsi que les effets relativistes, l'énergie de ces états $|n\rangle$ s'écrit

$$E_n = -\frac{E_I}{n^2}, \quad (1.63)$$

où n est un entier non nul, et l'énergie d'ionisation E_I vaut 4 fois celle de l'atome d'hydrogène ($E_I = 54.4$ eV). Il y a de plus les états ionisés d'énergie positive, où l'électron peut s'éloigner à l'infini. Les énergies E_k de ces états $|k\rangle$ ne sont pas quantifiées : elles forment un continuum (figure 1.7).

Dans le modèle des deux électrons indépendants, chaque électron peut être dans un quelconque des états du système hydrogénoïde dont les énergies sont représentées sur la figure 1.7. Parmi tous ces états, nous considérons d'abord la série $(1, n)$ où l'un des électrons est dans l'état $n = 1$ tandis que l'autre est dans un état quelconque. L'énergie d'un tel état vaut

$$E_{1,n} = E_1 + E_n = -E_I \left(1 + \frac{1}{n^2} \right). \quad (1.64)$$

On a représenté la suite de ces énergies sur la figure 1.8a. Le niveau fondamental $(1, 1)$ a une énergie $-2E_I$, puis on a une suite de niveaux discrets d'énergie comprise entre $-2E_I$ et $-E_I$. Pour des énergies entre $-E_I$ et 0, on a des états $(1, k)$ où le second électron est libre : les niveaux d'énergie forment un continuum.

Considérons maintenant la série $(2, n')$ où l'un des électrons est dans l'état $n = 2$ et l'autre est dans un état $n \geq 2$ (figure 1.8b). On voit apparaître la limite d'ionisation $E_{2,\infty}$ au-dessus de laquelle on a un ion He^+ dans un état excité. Mais le point qui nous intéresse ici est que le niveau discret $(2, 2)$ a une énergie supérieure à la limite d'ionisation de la série $(1, n)$. Il lui est donc énergétiquement permis à l'atome dans l'état lié $(2, 2)$ de passer

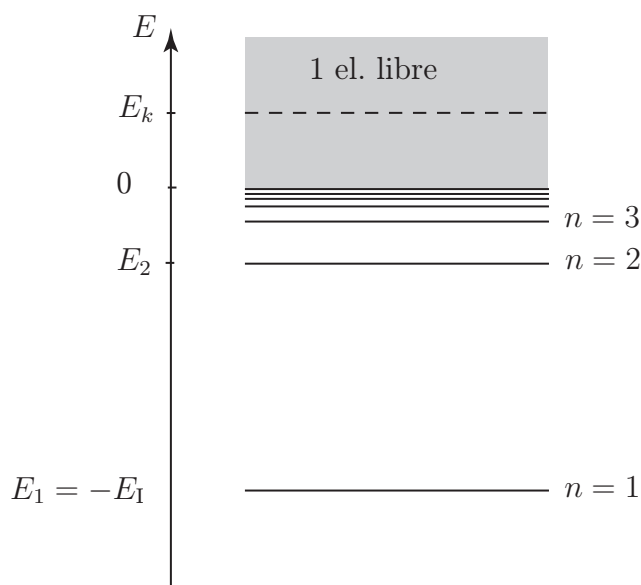


FIG. 1.7: **Niveaux d'énergie d'un électron** dans un potentiel coulombien. Les niveaux liés d'énergie négative forment une suite discrète. Les niveaux ionisés d'énergie positive (en tireté) forment un continuum.

dans l'état $(1, k_0)$ (voir la figure), et donc de s'ioniser, s'il existe un processus d'interaction permettant cette transition.

b. Autoionisation

Jusqu'à présent, nous avons ignoré l'interaction coulombienne entre les deux électrons. Nous allons maintenant la prendre en compte de façon perturbative, en rajoutant à l'hamiltonien décrivant le système à 2 électrons un terme

$$\hat{H}_1 = -\frac{1}{4\pi\epsilon_0} \frac{q_e^2}{|\mathbf{r}_1 - \mathbf{r}_2|} . \quad (1.65)$$

Ce terme a pour premier effet de modifier la valeur précise des énergies des divers états que nous avons considérés. Mais on peut continuer à distinguer les séries $(1, n)$, $(2, n)$ et le niveau discret $(2, 2)$ a toujours une énergie supérieure à la limite d'ionisation de la série $(1, n)$: il existe donc un état ionisé $(1, k_0)$ de même énergie que $(2, 2)$. Or l'hamiltonien $\hat{H}_1$ a un élément de matrice non nul entre les états $(2, 2)$ et l'état $(1, k_0)$, vers lequel il peut provoquer une transition : l'état $(2, 2)$ est un état susceptible de s'ioniser spontanément. Ce phénomène d'autoionisation s'observe effectivement lorsqu'on porte l'atome d'hélium dans l'état $(2, 2)$.

c. Couplage au continuum $(1, k)$

En fait, le terme de couplage $\hat{H}_1$ a également des éléments de matrice non nuls vers d'autres états du continuum $(1, k)$, et en particulier vers les états d'énergie voisine de $(1, k_0)$. Or on sait que la transition est possible même s'il n'y a pas égalité parfaite entre les énergies initiale et finale. Pour traiter le problème de l'autoionisation, il faut donc

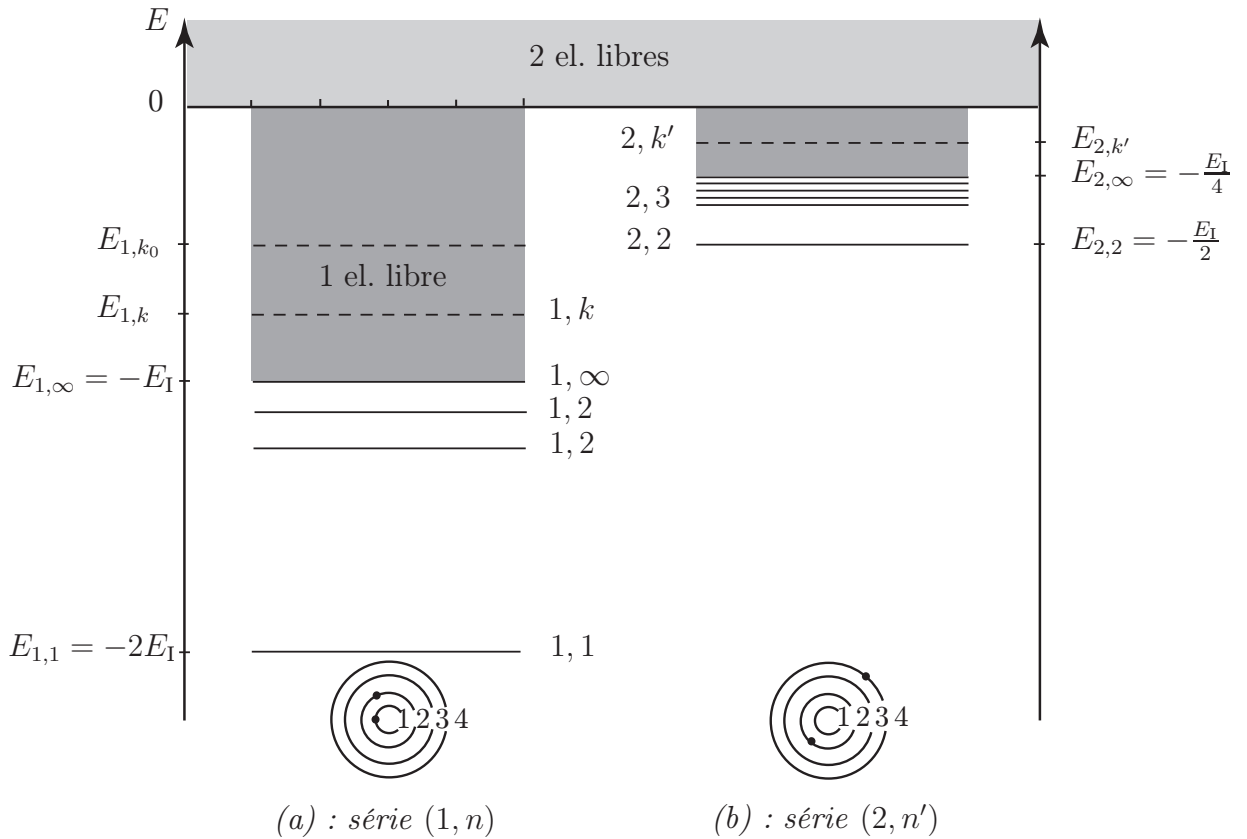


FIG. 1.8: **Niveaux d'énergie de deux électrons indépendants** dans un potentiel coulombien (comparer à la figure 1.7. a) Série (1, n) où un électron reste dans l'état $n = 1$. Pour une énergie $E \geq E_{1,\infty}$ on a un continuum d'états (1, k) où le deuxième électron est libre : l'hélium est ionisé une fois. b) Série (2, n') où un électron reste dans $n = 2$, l'autre étant dans $n' \geq 2$. On constate que le niveau discret (2, 2) a une énergie supérieure à l'énergie $E_{1,\infty}$ de la première limite d'ionisation de la série (1, n). Il est donc susceptible d'évoluer vers l'état ionisé (1, k_0) de même énergie, ou vers les états ionisés d'énergie voisine.

prendre en compte le couplage du niveau discret (2, 2) à l'ensemble des états (1, k) qui forment un continuum. C'est ce problème général du couplage d'un niveau discret à un continuum que nous allons envisager maintenant.

1.3.2 Niveau discret couplé à un quasi-continuum. Modèle simplifié

Dans l'exemple que nous venons de présenter, l'état final comporte un électron non lié qui s'éloigne indéfiniment de l'atome. On voit ainsi apparaître un comportement irréversible, qualitativement différent de l'oscillation de Rabi entre deux niveaux discrets. C'est ce comportement que nous nous proposons de traiter quantitativement maintenant.

a. Notion de quasi-continuum

Au lieu de traiter le problème du couplage entre un niveau discret et un continuum, nous allons considérer un problème proche, mais plus simple. Les états propres appartenant à un continuum sont en effet d'un traitement mathématique délicat, en particulier parce qu'ils ne sont pas normalisables. Il s'avère alors plus pratique de faire les calculs dans le cas d'un *ensemble de niveaux discrets très serrés*, appelé *quasi-continuum*, puis de passer à la limite du spectre continu sur le résultat final¹³.

Pour cela, nous allons considérer que les particules sont enfermées dans une boîte *fictive* de *grande dimension*, qui introduit de nouvelles conditions aux limites. Les états d'énergie positive sont alors des états liés, donc discrets, de la boîte fictive. Ils sont d'autant plus serrés que la boîte est plus grande. Il nous suffira de faire *tendre les dimensions de la boîte vers l'infini* à la fin du calcul pour obtenir le résultat relatif au spectre continu.

Illustrons cette démarche sur l'exemple du paragraphe 1.3.1a. L'introduction de la boîte fictive consiste à ajouter au potentiel coulombien dans lequel sont plongés les électrons un potentiel nul à l'intérieur de la boîte (par exemple un cube dont les faces sont en $x = \pm L/2$, $y = \pm L/2$, $z = \pm L/2$), et infini hors de la boîte. Les solutions stationnaires correspondant à ce potentiel sont exclusivement des états liés, formant une suite discrète même pour des valeurs positives de l'énergie. Pour illustrer la situation sans compliquer la présentation, nous considérons le problème à 1 dimension d'un potentiel attractif au milieu d'une boîte carrée (figure 1.9).

Nous admettons de plus que les états d'énergie positive sont peu différents de ceux associés à un potentiel nul à l'intérieur de la boîte : cette approximation est raisonnable lorsque la taille de la boîte L est suffisamment grande et pour des énergies E pas trop proches de 0. Les niveaux d'énergie sont alors ceux du puits de potentiel carré

$$E_k = \frac{\hbar^2 \pi^2}{2mL^2} k^2. \quad (1.66)$$

L'intervalle entre deux niveaux successifs peut s'écrire

$$E_{k+1} - E_k = \frac{2\pi\hbar}{\sqrt{2m}} \frac{\sqrt{E}}{L}. \quad (1.67)$$

On constate bien que l'écart entre niveaux varie en $1/L$ à énergie donnée.

Notons que l'amplitude des fonctions d'onde correspondantes, qui résulte d'une normalisation dans la boîte de taille L , varie en $L^{-1/2}$. Nous ferons appel à ce résultat dans la suite.

¹³C'est par commodité de calcul que les quasi-continuum sont introduits ici. Mais il existe également des situations naturelles de quasi-continuum. Par exemple les états excités de vibration-rotation d'une molécule à grand nombre d'atomes forment un ensemble dense de niveaux d'énergie ayant les propriétés d'un quasi-continuum. Cette propriété est exploitée dans la photodissociation des molécules (voir N. Bloembergen et E. Yablonovitch, *Physics Today* **31**, 5, 23 (1978)), ou dans les lasers accordables à colorants (voir § 3.2.2.d).

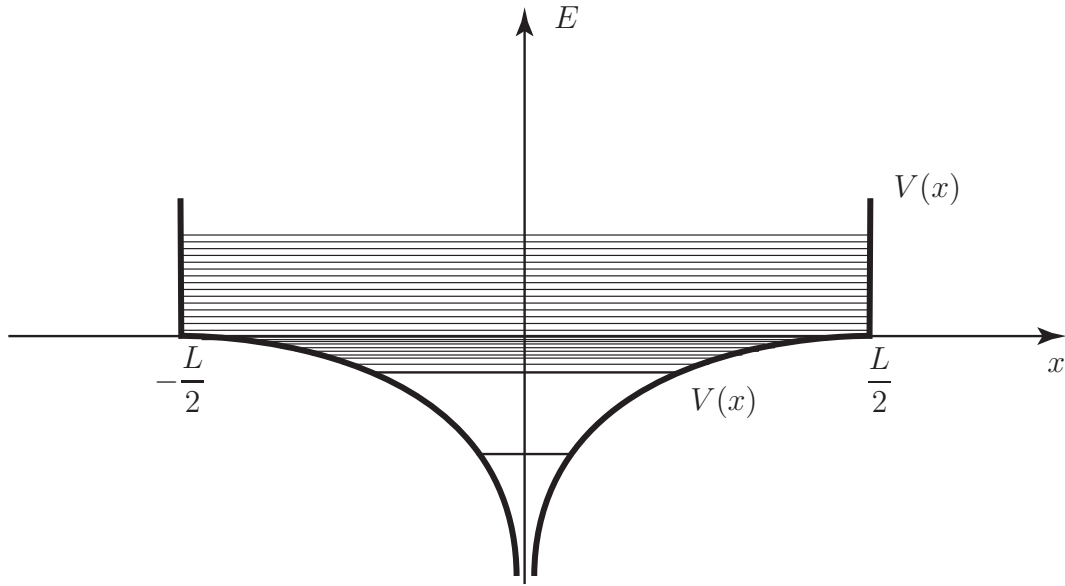


FIG. 1.9: **Électron dans un potentiel attractif** plongé dans une boîte carrée de taille $\frac{L}{2}$. On a des niveaux discrets même pour les états d'énergie positive, mais si L est suffisamment grand ils sont très serrés : on a un quasi-continuum.

b. Modèle simple de quasi-continuum

Revenons au problème général d'un niveau discret couplé à un quasi-continuum. Nous allons le traiter tout d'abord dans un cadre extrêmement simplifié, qui nous permettra d'obtenir les résultats essentiels en évitant des calculs compliqués. Considérons un hamiltonien $\hat{H}_0$ dont les états propres sont d'une part le niveau initial $|i\rangle$ (nous prendrons $E_i = 0$), et d'autre part un ensemble d'états $|k\rangle$ que nous supposons régulièrement espacés, d'énergie $E_k = k\varepsilon$ (k entier variant de $-\infty$ à $+\infty$). La distance ε entre deux niveaux consécutifs sera supposée très petite (nous préciserons cette hypothèse par la suite), de sorte que les niveaux $|k\rangle$ forment un quasi-continuum (Figure 1.10).

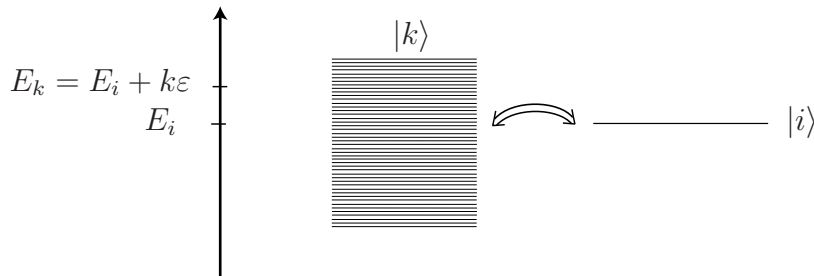


FIG. 1.10: **Niveaux d'énergie** d'un système comprenant un niveau isolé $|i\rangle$ couplé à une infinité d'états $|k\rangle$ régulièrement espacés formant un quasi-continuum. Si le système est initialement dans $|i\rangle$, le couplage le fait passer vers le quasi-continuum, de façon irréversible à la limite $\varepsilon \rightarrow 0$.

La deuxième simplification consiste à supposer que le niveau $|i\rangle$ est couplé aux niveaux

$|k\rangle$ par un terme de couplage $\hat{W}$ dont les éléments de matrice s'écrivent :

$$\begin{cases} \langle k|\hat{W}|i\rangle = w \\ \langle k|\hat{W}|k'\rangle = 0 \\ \langle i|\hat{W}|i\rangle = 0 . \end{cases} \quad (1.68)$$

L'interaction $\hat{W}$ ne couple donc pas deux états différents du continuum. L'élément de matrice w est pris indépendant de k et réel pour simplifier.

1.3.3 Évolution du système

a. Comportement aux temps courts : probabilité de transition par unité de temps

Calculons la probabilité $P_i(T)$ pour qu'un système initialement dans l'état $|i\rangle$ soit encore dans cet état à l'instant T . Les divers états du continuum sont distincts (orthogonaux entre eux), et on doit donc sommer les probabilités relatives aux divers états finaux possibles. La probabilité $P_i(T)$ s'écrit alors :

$$P_i(T) = 1 - \sum_{k=-\infty}^{+\infty} P_{i \rightarrow k}(T) , \quad (1.69)$$

où $P_{i \rightarrow k}(T)$, probabilité de transition entre deux niveaux discrets sous l'effet d'une perturbation constante branchée entre $t = 0$ et $t = T$, est donnée au premier ordre par l'expression (1.36) :

$$P_{i \rightarrow k}(T) = T \frac{2\pi}{\hbar} w^2 \delta_T(k\varepsilon) . \quad (1.70)$$

Plaçons-nous dans le cas où l'écart entre deux niveaux ε est très petit devant la largeur $2\pi\hbar/T$ de δ_T . La fonction $\delta_T(k\varepsilon)$ varie alors lentement avec k , et on peut remplacer la somme discrète dans (1.69) par l'intégrale correspondante. On trouve

$$\begin{aligned} P_i(T) &= 1 - T \frac{2\pi}{\hbar} w^2 \int_{-\infty}^{+\infty} \frac{dE}{\varepsilon} \delta_T(E) \\ &= 1 - T \frac{2\pi w^2}{\hbar \varepsilon} , \end{aligned} \quad (1.71)$$

car l'aire totale de la fonction δ_T vaut 1. Nous obtenons donc dans le cadre du présent modèle de niveaux très serrés une *évolution linéaire en T* , et non pas quadratique comme dans le cas de deux niveaux discrets (équation 1.60). On peut donc définir une *probabilité de départ par unité de temps* Γ

$$\Gamma = \frac{1 - P_i(T)}{T} , \quad (1.72)$$

qui vaut

$$\Gamma = \frac{2\pi}{\hbar} w^2 \frac{1}{\varepsilon} . \quad (1.73)$$

La quantité Γ , homogène à l'inverse d'un temps, est aussi appelée *probabilité de transition* vers le quasi-continuum *par unité de temps* (ou encore taux de transition).

La quantité w^2/ε ne dépend pas de L pour l'exemple présenté ci-dessus : en effet, ε varie comme L^{-1} , et w est proportionnel à l'amplitude de la fonction d'onde du quasi-continuum, en $L^{-1/2}$. C'est cette propriété, à vérifier systématiquement à la fin des calculs, qui valide le recours à un quasi-continuum.

Remarque

Deux conditions doivent être satisfaites pour pouvoir utiliser les formules (1.72) et (1.73). Elles imposent deux limites supérieures au temps d'interaction T :

$$T \ll 2\pi\hbar/\varepsilon \quad \text{et} \quad T \ll \Gamma^{-1} = \hbar\varepsilon/2\pi w^2 .$$

La première est nécessaire pour pouvoir remplacer la somme par une intégrale dans (1.69) et elle peut toujours être satisfaite en prenant L suffisamment grand. La seconde assure la validité du calcul de perturbation au premier ordre ($P_i(T) \ll 1$). D'autre part, dans les cas réalistes, le quasi-continuum ne s'étend pas sur un domaine d'énergie infini, mais a une largeur Δ (cf. complément I.2). Il faut alors rajouter une condition supplémentaire sur T , nécessaire pour donner la valeur 1 à l'intégrale $\int dE \delta_T(E)$ figurant dans (1.71), calculée maintenant sur un domaine d'intégration de largeur Δ :

$$T \gg \hbar/\Delta .$$

Ces conditions sont compatibles à condition que $\Gamma \ll \Delta$.

b. Comportement aux temps longs : décroissance exponentielle

Pour étudier l'évolution sur des temps pouvant être longs devant Γ^{-1} , il faut renoncer à l'approche perturbative. La méthode de résolution de ce type de problème a été introduite en 1930 par Weisskopf et Wigner. Ce paragraphe en donne une présentation simplifiée.

Le vecteur d'état du système, donné par (1.12) s'écrit dans notre cas

$$|\psi(t)\rangle = \gamma_i(t)|i\rangle + \sum_{k=-\infty}^{+\infty} \gamma_k(t) e^{-ik\epsilon t/\hbar} |k\rangle , \quad (1.74)$$

et l'équation de Schrödinger (1.15) est équivalente au système suivant d'équations différentielles :

$$\begin{cases} i\hbar \frac{d}{dt} \gamma_i(t) = w \sum_{k=-\infty}^{+\infty} \gamma_k(t) e^{-ik\epsilon t/\hbar} \\ i\hbar \frac{d}{dt} \gamma_k(t) = w e^{ik\epsilon t/\hbar} \gamma_i(t) . \end{cases} \quad (1.75)$$

Le système étant initialement dans l'état $|i\rangle$, on a $\gamma_k(0) = 0$ et donc $\gamma_k(t)$ peut s'écrire

$$\gamma_k(t) = \frac{w}{i\hbar} \int_0^t dt' \gamma_i(t') e^{ik\epsilon t'/\hbar} . \quad (1.76)$$

En portant ce résultat dans la première équation (1.75) et en utilisant l'expression (1.73) de Γ , on obtient l'équation exacte suivante :

$$\frac{d}{dt}\gamma_i(t) = -\frac{\Gamma}{2\pi\hbar} \int_0^t dt' \gamma_i(t') \left[\sum_{k=-\infty}^{+\infty} \varepsilon e^{ik\varepsilon(t'-t)/\hbar} \right]. \quad (1.77)$$

Nous supposons toujours le continuum suffisamment serré pour que $\varepsilon \ll \hbar/T$. On peut alors remplacer la somme sur k entre crochets par l'intégrale suivante :

$$\int_{-\infty}^{+\infty} dE e^{iE(t'-t)/\hbar} = 2\pi\hbar\delta(t' - t). \quad (1.78)$$

Posons $\tau = t' - t$. On peut alors récrire l'équation (1.77) sous la forme

$$\frac{d}{dt}\gamma_i(t) = -\Gamma \int_{-t}^0 d\tau \delta(\tau) \gamma_i(t + \tau). \quad (1.79)$$

L'intégrale figurant dans (1.79) vaut $\gamma_i(t)/2$ puisque la fonction δ est paire et que, pour tout $t > 0$

$$\int_{-t}^t d\tau \delta(\tau) \gamma_i(t + \tau) = \gamma_i(t) \int_{-t}^t d\tau \delta(\tau) = \gamma_i(t). \quad (1.80)$$

L'équation (1.79) est alors très simple et sa solution est

$$\gamma_i(t) = \exp(-\Gamma t/2). \quad (1.81)$$

On en déduit la probabilité $P_i(T) = |\gamma_i(T)|^2$:

$$P_i(T) = \exp(-\Gamma T). \quad (1.82)$$

La probabilité de trouver le système dans l'état initial *décroît exponentiellement* avec le temps et tend vers 0 aux temps très longs comme dans un processus de désintégration radioactive. En d'autres termes, (1.82) exprime qu'un système préparé dans l'état $|i\rangle$ a une *durée de vie* égale à Γ^{-1} .

Notons enfin qu'aux temps courts, l'expression (1.82) redonne bien l'expression (1.71) obtenue grâce à la méthode perturbative.

1.3.4 État final du processus. Largeur d'un niveau instable

À la différence du problème du couplage entre niveaux discrets étudié précédemment (expression 1.59), le système atteint aux temps longs un état stationnaire, dans lequel la probabilité de présence dans l'état initial est nulle. La question se pose alors de savoir quels sont les états finaux possibles pour le système à $t = +\infty$. Il nous faut donc calculer les coefficients $\gamma_k(t = +\infty)$. Pour cela, reportons la solution (1.81) dans l'équation (1.76). Le calcul de l'intégrale est élémentaire et donne

$$\gamma_k(t) = w \frac{1 - e^{(ik\varepsilon/\hbar - \Gamma/2)t}}{k\varepsilon + i\hbar\Gamma/2}. \quad (1.83)$$

La probabilité de présence du système dans l'état $|k\rangle$ à l'issue du processus ($t \rightarrow +\infty$) vaut donc

$$P_k = \frac{w^2}{(k\varepsilon)^2 + \hbar^2\Gamma^2/4} . \quad (1.84)$$

La probabilité $dP(E)$ de trouver le système dans l'un quelconque des états $|k\rangle$ du quasi-continuum dont l'énergie est comprise entre E et $E + dE$ est égale à

$$dP(E) = P_k \frac{dE}{\varepsilon} , \quad (1.85)$$

puisque dE/ε représente le nombre d'états du quasi-continuum dans l'intervalle d'énergie dE . En utilisant (1.73) et (1.84), nous trouvons

$$\frac{dP}{dE} = \frac{\hbar\Gamma}{2\pi} \frac{1}{E^2 + \hbar^2\Gamma^2/4} . \quad (1.86)$$

La distribution en énergie des états finaux du processus dans le quasi-continuum est donc une Lorentzienne centrée sur l'énergie $E_i = 0$ de l'état discret initial; sa largeur à mi-hauteur est égale à $\hbar\Gamma$ (voir figure 1.11).

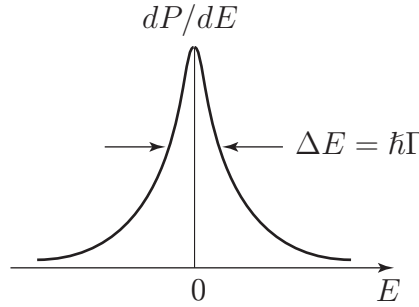


FIG. 1.11: **Distribution en énergie des états finaux** à l'issue d'un processus couplant un niveau discret à un quasi-continuum par une perturbation constante. La largeur de cette distribution vaut $\Delta E = \hbar\Gamma$.

Le niveau initial discret se « vide » donc dans des états du quasi-continuum distribués sur une largeur $\hbar\Gamma$ autour de l'énergie initiale. Comme a priori il y a conservation de l'énergie dans le processus de transition, on peut interpréter ce résultat en considérant que l'état initial a une indétermination en énergie $\hbar\Gamma$ liée à sa durée de vie finie Γ^{-1} .

Ce résultat a une portée extrêmement générale. Par exemple, il existe des particules élémentaires dont la durée de vie est extrêmement brève (de l'ordre de 10^{-23} s pour le méson rho). Leur masse (proportionnelle à leur énergie) n'est alors pas déterminée à mieux que $\Delta m \sim \hbar/\tau c^2$ (150 MeV/c² pour le méson rho, ce qui est beaucoup plus que la masse de l'électron qui vaut environ 1 MeV/c²).

Remarque

Une caractéristique frappante de cette évolution en $e^{-\Gamma t}$ (comparée à l'oscillation de Rabi) est son caractère *monotone* : la probabilité de présence dans l'état initial décroît toujours. On pourrait être tenté d'y voir un comportement irréversible. Le problème est en fait plus subtil. À n'importe quel moment t_0 de son évolution, le système se trouve en effet dans un état parfaitement déterminé $|\psi(t_0)\rangle$, qui s'exprime à partir des coefficients $\gamma_k(t_0)$ donnés par les relations (1.83). Faisons agir sur cet état l'opérateur de renversement du sens du temps $\hat{K}$, qui revient à prendre le complexe conjugué de sa fonction d'onde¹⁴ (cet opérateur est l'équivalent quantique de l'inversion des vitesses instantanées en Mécanique classique). Si le système évolue à partir de ce nouvel état initial $\hat{K}|\psi(t_0)\rangle$ sous l'effet du même hamiltonien, il est facile de montrer qu'au bout du temps t_0 , il se retrouvera exactement dans l'état initial $|i\rangle$. C'est la manifestation dans ce cas particulier de la *réversibilité* de la dynamique résultant de l'équation de Schrödinger.

Il faut cependant noter que, d'un point de vue pratique, il est infiniment plus facile de préparer le système dans l'état propre $|i\rangle$ de $\hat{H}_0$ que dans l'état $\hat{K}|\psi(t_0)\rangle$ avec les *valeurs exactes des amplitudes et les phases relatives* nécessaires pour que le système « se concentre » ultérieurement dans l'état $|i\rangle$. L'évolution inverse de (1.81) est donc possible, mais très peu probable. C'est à ce niveau statistique que s'introduit l'irréversibilité dans ce type de processus.

1.3.5 Règle d'or de Fermi

Le cas considéré ci-dessus est très schématique. Dans la pratique, l'intensité du couplage W_{ik} est une fonction du niveau $|k\rangle$, et la différence d'énergie $E_k - E_{k-1}$ entre niveaux consécutifs dans le quasi-continuum dépend aussi du niveau $|k\rangle$ considéré. Nous ne traiterons pas ici le cas général¹⁵, et nous nous bornerons à donner sans démonstration le résultat exact en mentionnant les similitudes et les différences avec le cas simple que nous venons d'étudier.

a. Présentation d'un quasi-continuum plus réaliste : notion de densité d'états

Considérons donc la situation générale où les niveaux $|k\rangle$ du quasi-continuum ne sont pas répartis de façon régulière : la différence d'énergie $E_k - E_{k-1}$ est une fonction de k , et de plus les niveaux du quasi-continuum ne s'étendent pas forcément de $-\infty$ à $+\infty$ et peuvent être limités à une certaine plage d'énergie (Figure 1.12). Nous pouvons alors définir une *densité d'états* $\rho(E)$, qui est égale au nombre de niveaux dN du quasi-continuum dans l'intervalle d'énergie comprise entre E et $E + dE$ divisé par la largeur dE de cet intervalle :

$$\rho(E) = \frac{dN(E)}{dE} . \quad (1.87)$$

La valeur de $\rho(E)$ dépend du système quantique considéré. Dans le modèle simplifié précédent, $\rho(E)$ valait $1/\varepsilon$. Dans le cas du quasi-continuum à une dimension introduit au paragraphe (1.3.2), les niveaux d'énergie sont donnés par la formule (1.66) et la densité

¹⁴Voir par exemple Messiah Tome 2, Chapitre XV.

¹⁵Voir CDLII, Chapitre XIII.

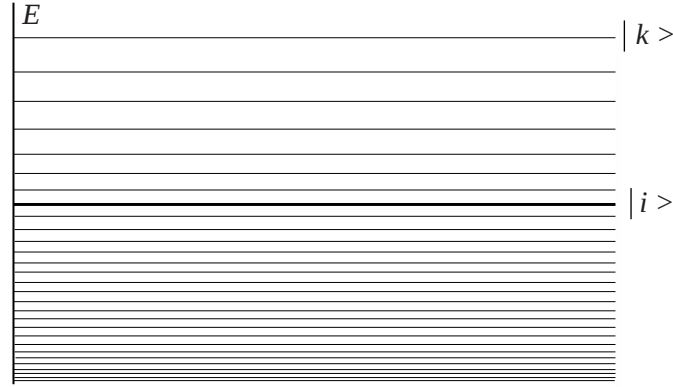


FIG. 1.12: **Couplage entre un niveau discret et un quasi-continuum** dans lequel l'espacement entre niveaux est une fonction de l'énergie.

$\rho(E)$ vaut

$$\rho(E) = \left(\frac{dE}{dn} \right)^{-1} = \frac{L}{\pi \hbar} \left(\frac{m}{2E} \right)^{1/2} . \quad (1.88)$$

Par ailleurs, dans un modèle réaliste, l'élément de matrice n'est pas nécessairement constant. Nous supposons néanmoins que $\langle k | \hat{W} | i \rangle$ est une *fonction lentement variable* de k .

b. Règle d'or de Fermi pour un continuum non dégénéré

Supposons le système initialement dans le niveau $|i\rangle$. On peut généraliser les résultats du paragraphe 1.3.3 et calculer par la théorie des perturbations au premier ordre la probabilité $P_i(T)$ de trouver le système dans l'état $|i\rangle$ pour des temps T suffisamment courts. On trouve une dépendance linéaire en T

$$P_i(T) = 1 - \Gamma T . \quad (1.89)$$

La probabilité de transition par unité de temps Γ est donnée par la formule importante suivante souvent utilisée dans les applications de la mécanique quantique (et qualifiée par Fermi de « règle d'or de la physique quantique », d'où son nom consacré de « *règle d'or de Fermi* »)

$$\Rightarrow \quad \Gamma = \frac{2\pi}{\hbar} |W_{fi}|^2 \rho(E_f = E_i) . \quad (1.90)$$

Dans cette formule, $|f\rangle$ est le niveau du quasi-continuum qui a même énergie que l'état discret $|i\rangle$, et W_{fi} et $\rho(E_f)$ sont respectivement l'élément de matrice de couplage et la densité d'états calculés pour cet état.

On montre dans le complément I.1 de ce chapitre comment on peut appliquer cette formule à un cas pratique schématisant le problème de l'autoionisation.

La formule (1.90) est la généralisation naturelle de (1.73) puisque nous savons qu'au bout d'un temps T , seuls les niveaux situés dans un domaine d'énergie de largeur $\hbar/T$

autour de E_i peuvent être peuplés : si $\hbar/T$ est petit devant l'échelle de variation de $\rho(E)$ et de W_{ik} , la distribution des niveaux d'énergie du quasi-continuum couplés à l'état initial coïncide localement avec celle du modèle simple étudié précédemment, qui peut alors être utilisé si $\rho(E)$ varie lentement au voisinage de E_i .

La formule (1.90) se démontre à partir de l'expression (1.36). Pour trouver $P_i(T)$, il suffit en effet de connaître les probabilités de transition $P_{i \rightarrow k}(T)$ vers tous les niveaux $|k\rangle$ possibles :

$$P_i(T) = 1 - \sum_k T \frac{2\pi}{\hbar} |W_{ki}|^2 \delta_T(E_k - E_i) . \quad (1.91)$$

Si W_{ki} est une fonction lentement variable de k sur la largeur $\hbar/T$ de la fonction δ_T , on peut remplacer la somme discrète par une intégrale en introduisant la densité d'états $\rho(E)$:

$$P_i(T) = 1 - T \frac{2\pi}{\hbar} \int dE_k \rho(E_k) |W_{ki}|^2 \delta_T(E_k - E_i) . \quad (1.92)$$

Si à son tour $\rho(E_k)$ varie lentement, on peut assimiler δ_T à une distribution de Dirac, et aboutir aux formules (1.89) et (1.90).

c. Règle d'or de Fermi pour un continuum dégénéré

La formule (1.90) suppose implicitement que les états du continuum ne sont pas dégénérés, c'est-à-dire que la donnée de E suffit à définir entièrement un état du quasi-continuum. On rencontre souvent des problèmes plus complexes où l'on a besoin d'autres paramètres, que nous regroupons sous la notation condensée β , pour définir les états du quasi-continuum, notés $|E_k, \beta\rangle$. Dans l'exemple de l'autoionisation introduit dans le paragraphe (C.1.b), le quasi-continuum est celui d'un électron libre enfermé dans une boîte de dimensions grandes devant celles de l'atome. Un état du quasi-continuum $|E_k, \theta, \varphi\rangle$ est par exemple défini par la donnée de l'énergie cinétique E_k de l'électron, mais aussi par les deux angles (θ, φ) précisant la direction dans laquelle se déplace l'électron éjecté¹⁶. Cherchons alors la probabilité différentielle $dP(T)/d\beta$ de trouver le système dans un état du continuum $|E, \beta\rangle$ dont le paramètre β est compris entre β et $\beta + d\beta$. Un raisonnement analogue à celui de la remarque précédente permet de montrer que celle-ci croît linéairement avec T . On peut donc définir une probabilité de transition différentielle par unité de temps

$$\Rightarrow \quad \frac{d\Gamma}{d\beta} = \frac{1}{T} \frac{dP(T)}{d\beta} = \frac{2\pi}{\hbar} |\langle E_f, \beta | \hat{W} | i \rangle|^2 \rho(\beta, E_f = E_i) . \quad (1.93)$$

La probabilité $P_i(T)$ de trouver le système dans l'état initial à l'instant T a toujours la dépendance linéaire de la formule (1.89), avec Γ donné par :

$$\Rightarrow \quad \Gamma = \int d\beta \frac{d\Gamma}{d\beta} . \quad (1.94)$$

¹⁶On utilise ici les notations habituelles pour les coordonnées sphériques (voir Complément II.3, Fig. 2).

Cette formule s'interprète facilement : on obtient Γ en sommant les probabilités sur toutes les transitions possibles, puisqu'elles aboutissent à des états finaux distincts.

Dans le cas du modèle simple d'autoionisation, la formule (1.93) permet de calculer la probabilité $d\Gamma/d\Omega$ de trouver l'électron éjecté dans une direction $\beta = (\theta, \varphi)$ donnée. Il faut ensuite sommer sur toutes les directions possibles pour calculer la probabilité globale d'autoionisation Γ , ce qui donne

$$\Gamma = \int \frac{d\Gamma}{d\Omega} d\Omega = \int \frac{d\Gamma}{d\Omega} \sin \theta d\theta d\varphi. \quad (1.95)$$

Le calcul effectif de Γ conduit à un résultat indépendant de la taille L de la boîte fictive (voir compléments I.1 et II.3).

Remarques

(i) On rencontre également la situation où les états du continuum s'écrivent $|E, m\rangle$, où m est un indice prenant des valeurs discrètes et non pas continues. Une formule analogue à (1.93) permet de déterminer la probabilité partielle de transition Γ_m vers les états $|E, m\rangle$ pour m fixé. Il suffit ensuite de sommer Γ_m sur toutes les valeurs possibles de m pour trouver la probabilité de transition totale Γ .

(ii) Donnons l'expression de la densité d'états du quasi-continuum des états d'un électron libre se propageant dans la direction (θ, φ) avec l'énergie E , à l'intérieur d'une boîte cubique de volume L^3 (voir complément II.3) :

$$\rho_{el}(E, \theta, \varphi) = \left(\frac{L}{2\pi\hbar} \right)^3 (2m^3 E)^{1/2}. \quad (1.96)$$

On constate sur les expressions (1.88) et (1.96) que la densité d'états croît proportionnellement au « volume » (L à une dimension, et L^3 à 3 dimensions) de la portion d'espace utilisée pour définir le quasi-continuum.

En raison de la symétrie sphérique du problème, la densité est dans ce cas particulier indépendante des paramètres supplémentaires θ et φ , et donc de la direction de propagation de l'électron. Dans le cas général, la densité d'états du quasi-continuum est une fonction des variables supplémentaires β .

d. Comportement aux temps longs

La solution non perturbative aux temps longs calculée par la méthode de Wigner-Weisskopf (paragraphe C.2.d) se généralise à un quasi-continuum quelconque. On obtient une décroissance exponentielle de la population du niveau $|i\rangle$ (formule 1.82) avec Γ donné maintenant par la formule (1.90) ou (1.94). Cette formule n'est toutefois valable que si le continuum s'étend sur un domaine d'énergie suffisamment grand. Plus précisément, $\hbar\Gamma$ étant la « largeur » du niveau initial induite par le couplage avec le continuum, on aura une décroissance exponentielle de la probabilité lorsque la quantité $|W_{fi}|^2 \rho(E)$ varie peu dans un intervalle de largeur $\hbar\Gamma$ autour de E_i . Nous verrons dans le complément I.2 ce qui peut se produire lorsque cette condition n'est pas réalisée.

Remarque

La distribution en énergie des états finaux du processus est, ici aussi, donnée par une courbe lorentzienne de largeur $\hbar\Gamma$, Γ étant déterminé par (1.90) ou (1.94). Toutefois, en toute

rigueur, le couplage avec les niveaux du quasi-continuum entraîne généralement un déplacement des niveaux d'énergie, et le centre de la lorentzienne est généralement légèrement déplacé par rapport à la position du niveau discret E_i . Dans le cas du modèle simple du §C.2, ce déplacement est nul parce que les contributions de deux niveaux $|k\rangle$ symétriques par rapport à $|i\rangle$ ont des signes opposés.

1.3.6 Cas d'une perturbation sinusoïdale

Nous pouvons utiliser les résultats du paragraphe (B.4.d) pour généraliser la règle d'or de Fermi au cas d'une perturbation sinusoïdale

$$\hat{H}_1(t) = \hat{W} \cos(\omega t + \varphi) . \quad (1.97)$$

Il suffit pour cela d'utiliser la formule (1.45) pour la probabilité $P_{i \rightarrow k}(T)$ au lieu de la formule (1.36). On obtiendra de la même manière une variation linéaire en T de $P_i(T)$, permettant de définir une probabilité de transition par unité de temps Γ qui vaut maintenant

$$\Gamma = \frac{2\pi}{\hbar} \frac{1}{4} [|W_{f'i}|^2 \rho(E_{f'} = E_i + \hbar\omega) + |W_{f''i}|^2 \rho(E_{f''} = E_i - \hbar\omega)] \quad (1.98)$$

les états $|f'\rangle$ et $|f''\rangle$ du quasi-continuum ayant respectivement les énergies $E_i + \hbar\omega$ et $E_i - \hbar\omega$, et $W_{f'i}$ et $W_{f''i}$ étant les éléments de matrice correspondants¹⁷. On voit ainsi apparaître les taux de transition Γ' et Γ'' .

À l'issue du processus, seront peuplés par le couplage sinusoïdal les niveaux du quasi-continuum dont l'énergie vaut $E_i \pm \hbar\omega$, à $\hbar\Gamma'$ ou $\hbar\Gamma''$ près (voir Figure 1.13). Comme dans le paragraphe (B.4.d), l'interprétation de la formule (1.98) en termes d'absorption ou d'émission d'un quantum d'énergie $\hbar\omega$ s'introduit ici naturellement, bien que nous n'ayons pas introduit la quantification du champ qui provoque la transition.

Remarque

Après l'étude du couplage entre deux niveaux discrets, puis entre un niveau discret et un continuum, on pourrait continuer l'étude de situations plus compliquées dans lesquelles le système est initialement dans un ensemble de niveaux discrets, ou dans un continuum d'états.

Si le système part d'un niveau du continuum pour aboutir à un autre niveau du continuum (c'est, par exemple, le cas pour la diffusion d'un électron libre par un potentiel), on peut encore appliquer la formule (1.93) donnant la probabilité de transition par unité de temps. Cependant, l'état $|i\rangle$ appartenant à un continuum, sa fonction d'onde et donc la formule (1.93) dépendent de la dimension de la boîte fictive dans lequel est enfermé le système. La quantité physiquement significative, parce qu'indépendante de la dimension de la boîte fictive, est la *section efficace différentielle* qui est le rapport de la probabilité de transition par unité de temps au flux des particules.

¹⁷Si le continuum ne s'étend pas aux régions d'énergie voisines de $E_i + \hbar\omega$ ou $E_i - \hbar\omega$, la probabilité de transition correspondante est évidemment nulle, puisque la densité d'état correspondante $\rho(E_{f'})$ ou $\rho(E_{f''})$ est nulle.

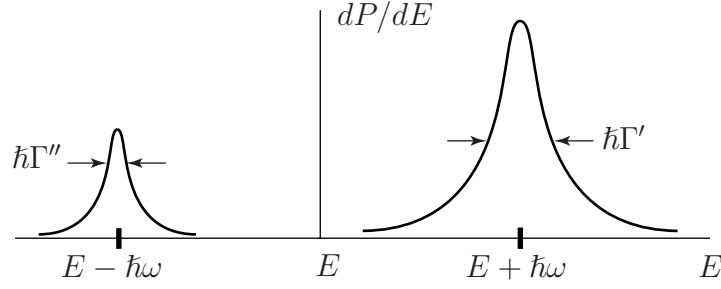


FIG. 1.13: **Distribution en énergie des états finaux** à l'issue du processus de couplage entre le niveau discret et le quasi-continuum par une perturbation sinusoïdale. Les hauteurs des deux pics sont proportionnelles aux premier et deuxième termes de la relation (1.98). On peut interpréter ce résultat en termes d'émission ou d'absorption de quanta d'énergie $\hbar\omega$ par le système. Les largeurs à mi-hauteur $\hbar\Gamma'$ et $\hbar\Gamma''$ de ces deux Lorentziennes sont données par les deux termes de la formule (1.98).

Si le système est initialement dans une *superposition incohérente* (c'est-à-dire un mélange statistique) d'états discrets ou appartenant à un continuum, il suffit de moyenner les probabilités $P_{i \rightarrow k}$ entre niveaux discrets (expression B.19) sur l'ensemble des niveaux initiaux possibles et de sommer le résultat sur l'ensemble des états finaux possibles.

1.4 Conclusion

Ce chapitre nous a permis de passer en revue quelques aspects essentiels du comportement dynamique des systèmes non-dissipatifs¹⁸ en mécanique quantique (systèmes obéissant à une équation de Schrödinger). Dans le cas de couplages « simples » (constants ou sinusoïdaux), on a pu ainsi dégager deux grands types de variation de la population $P_i(t)$ d'un état initialement peuplé (voir figure 1.14) :

- (a) une *évolution en t^2 aux temps courts, sinusoïdale aux temps longs*, dans le cas où le niveau initial est couplé à un *niveau final isolé* ;
- (b) une *évolution décroissante, monotone, linéaire en t aux temps courts, exponentielle aux temps longs*, dans le cas où le niveau initial est couplé à un ensemble de niveaux très serrés, formant un *quasi-continuum* suffisamment large.

Dans l'interaction atome rayonnement, on rencontre des situations correspondant à ces deux cas : couplage entre un atome à deux niveaux et une onde monochromatique pour le cas (a) ; désexcitation spontanée d'un atome isolé préparé dans un état excité pour le cas (b). Ces deux comportements sont fondamentalement différents puisque dans le cas (a) le système se retrouve périodiquement dans un état identique à l'état initial, alors que dans le cas (b) il atteint aux temps longs un état stationnaire différent de l'état

¹⁸Le comportement dynamique des systèmes dissipatifs, couplés à l'environnement extérieur donne lieu à des comportements beaucoup plus riches (voir par exemple le complément II.2). On y retrouve cependant les deux comportements typiques vus dans ce chapitre, dans les situations limites.

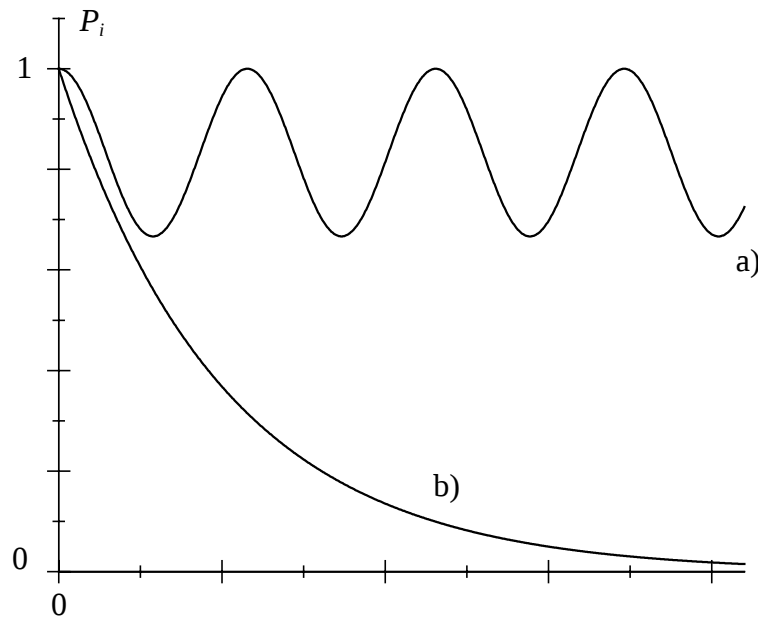


FIG. 1.14: **Variation de la population P_i d'un état discret initialement peuplé et couplé par une perturbation constante a) à un autre niveau discret b) à un quasi-continuum.**

initial. Ils correspondent en fait à deux cas limites extrêmes. Le complément I.2 de ce chapitre montre le passage continu entre ces deux régimes extrêmes. On s'y intéresse à l'évolution d'un niveau discret couplé à un quasi-continuum de largeur Δ variable. On trouve le régime d'oscillations (a) si la largeur Δ est très petite devant Γ , et le régime d'évolution monotone (b) pour les très grandes valeurs de Δ .

Remarque

Nous avons considéré jusqu'ici uniquement des cas où la perturbation $\hat{W}(t)$ était parfaitement déterminée. Il arrive aussi dans des cas concrets que celle-ci présente un aspect aléatoire. C'est le cas de l'exemple du paragraphe (B.2.b) relatif aux collisions, où l'instant de collision, la vitesse et le paramètre d'impact présentent des fluctuations statistiques. Il faut alors moyenner les expressions des probabilités de transition sur toutes les réalisations possibles. Dans ce cas, il est facile de voir que les *oscillations de Rabi* (1.59), qui ont des phases différentes selon les réalisations, *vont avoir tendance à se brouiller* au bout d'un temps plus ou moins long. La probabilité de transition moyennée ne présentera plus d'oscillations et atteint alors un régime stationnaire. En revanche, le comportement exponentiel se révèle beaucoup plus résistant à ce type de moyenne statistique. Nous en verrons un exemple dans le chapitre II paragraphe D.2 lorsque nous étudierons l'absorption et l'émission induite : le pompage et la relaxation des deux niveaux intéressés par la transition introduisent une moyenne sur les oscillations de Rabi commençant à des instants différents et entraînent le brouillage de ces oscillations.

Complément I1

Exemple d'utilisation de la règle de Fermi dans le problème de l'autoionisation

Dans ce complément, nous allons détailler le calcul d'une probabilité de transition à l'aide de la règle d'or de Fermi dans le cadre d'un problème donné, en l'occurrence celui de l'autoionisation d'un système à deux particules, présenté dans le paragraphe 1.3.1. Il s'agit ici uniquement de comprendre comment « appliquer » cette règle, et non pas d'étudier en détail le problème physique réel de l'autoionisation d'un atome, ce calcul techniquement délicat sortant du cadre de cet ouvrage. Nous nous contenterons donc de calculer la probabilité de transition pour le modèle simplifié à une dimension avec un puits de potentiel en créneau introduit dans le paragraphe 3.1.a. Nous verrons en particulier comment on calcule concrètement la densité d'états $\rho(E)$, et introduirons les « conditions aux limites périodiques » qui permettent souvent de simplifier les calculs de densité d'états.

La probabilité de désexcitation par unité de temps de l'état doublement excité $|i\rangle = |1 : \beta ; 2 : \beta\rangle$ vers les états $|k\rangle = |1 : E ; 2 : \alpha\rangle$ est donnée par la règle d'or de Fermi (formule (1.90)) :

$$\Gamma_1 = \frac{2\pi}{\hbar} |\langle 1 : E ; 2 : \alpha | \hat{W} | 1 : \beta ; 2 : \beta \rangle|^2 \rho(E_f = E_i) \quad (1)$$

où $\hat{W}$ est le terme de couplage entre les deux particules, ρ la densité d'états finaux prise pour une énergie égale à celle de l'état initial $2E_\beta$, et où l'énergie E est égale à $2E_\beta - E_\alpha$. Pour calculer Γ_1 , nous avons besoin de connaître la forme des fonctions d'onde du quasi-continuum, l'élément de matrice d'interaction, et la densité d'état ρ correspondante. Nous allons les calculer en détail dans les trois paragraphes suivants.

1. Fonctions d'onde du quasi-continuum

Le système atomique à une dimension est enfermé dans une boîte fictive à une dimension de longueur L (voir paragraphe C.2.a.). Si nous supposons pour simplifier que les états du quasi-continuum auxquels nous nous intéressons ont une énergie suffisamment élevée pour que le puits de potentiel V_0 n'exerce qu'une faible influence sur la forme de la fonction d'onde, les fonctions d'onde $\psi_n(x)$ du quasi-continuum sont celles d'une particule dans un puits infini¹, et s'écrivent donc :

$$\psi_n(x) = \sqrt{\frac{2}{L}} \sin \frac{n\pi}{L} \left(x + \frac{L}{2} \right) \quad (2)$$

Elles ont pour énergie :

$$E_n = \frac{\hbar^2 \pi^2}{2mL^2} n^2 \quad (n \text{ entier strictement positif}) \quad (3)$$

Notons que les fonctions d'onde $\psi_n(x)$ ont pour parité $(-1)^{n+1}$. En particulier, celles correspondant à $n = 2p$ sont impaires et s'annulent à l'origine.

2. Élément de matrice de l'hamiltonien d'interaction

Appelons $\psi_\alpha(x)$ et $\psi_\beta(x)$ les fonctions d'onde des états liés du puits de potentiel de la figure C.1 du chapitre I, de profondeur V_0 . L'élément de matrice que nous devons évaluer s'écrit alors :

$$\langle k | \hat{W} | i \rangle = \int_{-L/2}^{L/2} dx_1 \int_{-L/2}^{L/2} dx_2 \psi_n^*(x_1) \psi_\alpha^*(x_2) W(|x_1 - x_2|) \psi_\beta(x_1) \psi_\beta(x_2) \quad (4)$$

où W est une fonction de la distance $|x_1 - x_2|$ entre les deux particules. Les fonctions d'onde des deux premiers états liés $\psi_\alpha(x)$ et $\psi_\beta(x)$ et du puits de potentiel de profondeur V_0 sont réelles et respectivement paires et impaires vis à vis du changement de signe de la coordonnée x . On en déduit que l'expression sous le signe somme dans l'expression (4) a la parité de $\psi_n(x)$. L'élément de matrice n'est donc différent de 0 que pour les valeurs impaires de n , $n = 2p + 1$. On peut l'écrire dans ce cas :

$$\langle k | \hat{W} | i \rangle = \sqrt{\frac{2}{L}} \int_{-L/2}^{L/2} dx_1 \int_{-L/2}^{L/2} dx_2 \cos(k_p x_1) \psi_\alpha(x_2) W(|x_1 - x_2|) \psi_\beta(x_1) \psi_\beta(x_2) \quad (5)$$

où k_p est égal à $(2p + 1)\pi/L$. La valeur exacte de $\langle k | \hat{W} | i \rangle$ dépend alors de la forme exacte du terme d'interaction $\hat{W}$, que nous ne précisons pas ici.

Enfin, nous pouvons étendre dans (5) les bornes d'intégration à $-\infty$ et à $+\infty$. En effet, les fonctions $\psi_\alpha(x)$ et $\psi_\beta(x)$ sont localisées essentiellement dans l'intervalle $[-a, a]$ (elles décroissent exponentiellement pour $|x| \geq a$). Elles sont donc négligeables au delà de $\pm L/2$ pour des valeurs de L suffisamment grandes.

¹Voir par exemple BD Chapitre III, ou CDL Complément $H_1(c)$.

3. Densité des états couplés au niveau discret

D'après l'expression (3), deux niveaux consécutifs du quasi-continuum sont séparés par un intervalle énergétique ε valant :

$$\varepsilon = \frac{dE_n}{dn} = \frac{\hbar^2 \pi^2}{mL^2} n = \frac{\pi \hbar}{L} \sqrt{\frac{2E_n}{m}} \quad (6)$$

Dans l'intervalle d'énergie dE autour de E_n , il y a $dn = dE/\varepsilon$ états. Nous en déduisons que la densité d'états $\rho_c(E) = dn/dE_n$ n'est autre que l'inverse de l'expression (6) :

$$\rho_c(E) = \varepsilon^{-1} = \frac{1}{\pi \hbar} \sqrt{\frac{m}{2E}} \quad (7)$$

Il faut cependant prendre garde que la densité d'états totale ρ_c n'est pas forcément celle qu'il faut utiliser dans l'expression (1) de la règle d'or de Fermi. Pour établir cette formule, nous avons en effet supposé que les éléments de matrice de couplage variaient lentement en fonction de l'énergie. Ce n'est pas le cas ici, puisque, pour les raisons de symétrie expliquées dans le précédent paragraphe, un élément de matrice sur deux est nul. Il faut en fait considérer *la densité des états couplés à l'état initial $|i\rangle$* , qui ont pour énergie :

$$E_p = \frac{\hbar^2 \pi^2}{2mL^2} (2p+1)^2 \quad (8)$$

Pour ce sous-ensemble d'états, les éléments de matrice (5) varient lentement avec l'énergie, et on peut utiliser la règle d'or de Fermi. Leur densité $\rho(E)$ est égale à la moitié de la densité totale $\rho_c(E)$ donnée en (7). Nous aboutissons ainsi à l'expression finale de la probabilité d'autoionisation par unité de temps Γ_1 dans le cadre de ce modèle simple :

$$\Gamma_1 = \frac{1}{\hbar^2} \sqrt{\frac{2m}{E_f}} \left| \int_{-\infty}^{+\infty} dx_1 \int_{-\infty}^{+\infty} dx_2 \cos(k_f x_1) \psi_\alpha(x_2) W(|x_1 - x_2|) \psi_\beta(x_1) \psi_\beta(x_2) \right|^2 \quad (9)$$

où l'énergie finale $E_f = \hbar^2 k_f^2 / 2m$ est déterminée par la conservation de l'énergie ($E_f = 2E_\beta - E_\alpha$). Il est important de noter que Γ_1 ne dépend pas de la dimension L de la boîte, le terme linéaire en L de la densité d'états étant compensé par celui provenant du carré de l'élément de matrice.

Remarques

(i) Nous n'avons envisagé ici que le cas où la particule numérotée 1 est libre à l'issue du processus. Pour obtenir la probabilité totale de désexcitation, il faut rajouter à Γ_1 le terme Γ_2 obtenu en échangeant les rôles des indices 1 et 2. W étant invariant dans cet échange, Γ_1 et Γ_2 sont égaux. Il faut donc doubler le résultat (9) pour obtenir la probabilité totale d'autoionisation Γ .

(ii) Les particules 1 et 2 étant identiques, elles sont indiscernables et il faut en toute rigueur symétriser (dans le cas de bosons) les fonctions d'onde à utiliser². Il est facile de se convaincre qu'on aboutit dans ce cas à la même valeur pour la probabilité d'autoionisation. Dans le cas de fermions, il n'est pas possible de mettre deux particules dans l'état β (sauf à rajouter un degré de liberté supplémentaire comme le spin).

²Voir par exemple BD Chapitre XIV, ou CDL 2, Chapitre XIV.

4. Autre approche possible du quasi-continuum : « condition aux limites périodiques »

Reprenons le potentiel représenté sur la figure C.5 du chapitre I. Pour transformer le continuum de niveaux d'énergie en un ensemble discret, donc plus facile à manipuler qu'un ensemble continu, nous avons rajouté des « parois » matérielles infranchissables en $x = \pm L/2$. On peut utiliser une deuxième approche : elle consiste à conserver le potentiel initial considéré sur la figure C.1, mais à *restreindre l'espace des fonctions d'onde solutions*. Plus précisément, on ne considère que les solutions définies dans un intervalle $[-L/2, +L/2]$ autour de l'origine, avec L suffisamment grand. Or on sait qu'il est possible de décomposer toute fonction $f(x)$ de cet intervalle en série de Fourier :

$$f(x) = \sum_{n=-\infty}^{+\infty} c_n \varphi_n(x) \quad (10)$$

avec :

$$\varphi_n(x) = \frac{1}{\sqrt{L}} e^{i\left(\frac{2n\pi}{L}\right)x} \quad (11)$$

Ces fonctions φ_n , forment donc une *base* (qui est de plus orthonormée) de l'espace de Hilbert correspondant. Elles correspondent à des ondes progressives (se propageant dans le sens positif ou négatif de l'axe des x selon le signe de n) de la particule libre, à la différence des ondes stationnaires $\psi_n(x)$ introduites précédemment (équation (2)). Ce sont des fonctions propres de l'hamiltonien d'une particule libre, avec pour énergie :

$$E_n = \frac{\hbar^2}{2m} \left(\frac{2n\pi}{L} \right)^2 \quad (12)$$

Ces fonctions n'ont pas une parité définie. Elles sont donc *toutes* couplées à l'état initial par le terme d'interaction $\hat{W}$, à la différence des fonctions $\psi_n(x)$. L'intervalle entre deux valeurs de E_n consécutives vaut alors :

$$\frac{dE_n}{dn} = \frac{4\hbar^2\pi^2}{mL^2} n = \frac{2\pi\hbar}{L} \sqrt{\frac{2E_n}{m}} \quad (13)$$

Il faut maintenant tenir compte du fait qu'il y a maintenant *deux états* pour chaque valeur de l'énergie E_n , correspondant à n positif ou négatif, c'est-à-dire à des ondes se propageant vers les x positifs ou négatifs. La densité d'états $\rho(E)$ est donc *égale au double* de celle que l'on déduit de (13). Elle est donc identique à la densité totale $\rho_c(E)$ des états du puits de potentiel avec « parois », donnée par l'expression (7).

On déduit de (11) et (13) la nouvelle expression de la probabilité de désintégration par unité de temps, obtenue à la limite L infini :

$$\Gamma_1 = \frac{2}{\hbar^2} \sqrt{\frac{m}{2E_f}} \left| \int_{-\infty}^{+\infty} dx_1 \int_{-\infty}^{+\infty} dx_2 \exp(-ik_f x_1) \psi_\alpha(x_2) W(|x_1 - x_2|) \psi_\beta(x_1) \psi_\beta(x_2) \right|^2 \quad (14)$$

Si l'on décompose l'exponentielle en sinus et cosinus, et qu'on tient compte des propriétés de symétrie de $\psi_\alpha(x)$, $\psi_\beta(x)$ et W dans la réflexion d'espace, on montre facilement que cette expression coïncide exactement avec celle trouvée en (9).

On retrouvera cette technique de décomposition en série de Fourier dans de nombreux problèmes. La relation (10) impliquant que l'extension de f à tout l'axe réel est périodique de période L (mais pas nécessairement continue en $x = \pm L/2$), cette technique est connue sous le nom de « *condition aux limites périodiques* ».

Complément I.2

Continuum de largeur variable

Dans le chapitre I, nous avons vu deux types bien différents d'évolution temporelle, selon que le niveau discret initial est couplé à un niveau final isolé ou à un ensemble de niveaux très serrés formant un quasi-continuum. Nous allons dans ce complément préciser, dans le cadre d'un modèle simple, comment s'effectue la transition entre ces régimes extrêmes, en étudiant le comportement d'un niveau discret $|i\rangle$ couplé à un quasi-continuum ayant une largeur variable Δ .

1. Présentation du modèle

Nous reprenons le modèle simplifié de quasi-continuum introduit dans le paragraphe (C.2.b). Il s'agit d'états $|k\rangle$ équidistants d'énergie $E_k = k\varepsilon$, avec k entier variant de $-\infty$ à $+\infty$: nous supposons maintenant que les éléments de matrice de l'hamiltonien $\hat{W}$ qui couple l'état discret au continuum ne sont pas constants. Ils sont égaux à :

$$\begin{aligned}\langle k|\hat{W}|i\rangle &= w_k = \frac{w}{\sqrt{1 + \left(\frac{k\varepsilon}{\Delta}\right)^2}} \\ \langle k|\hat{W}|k'\rangle &= 0 \\ \langle i|\hat{W}|i\rangle &= 0\end{aligned}\tag{1}$$

Le carré de l'élément de matrice de couplage a donc une dépendance lorentzienne en fonction de l'énergie : « Vu du niveau $|i\rangle$ initial », le quasi-continuum a ainsi une largeur en énergie de l'ordre de Δ , puisque les niveaux $|k\rangle$ d'énergie notablement supérieure (en valeur absolue) à cette valeur ne sont que très faiblement couplés à l'état initial. Nous supposons de plus que w et Δ sont grands devant ε . L'intérêt de la forme (1) du couplage est qu'il permet la résolution des équations d'évolution pour toutes les valeurs de la largeur du continuum Δ .

2. Evolution temporelle

Pour obtenir l'évolution temporelle du système, nous suivrons une procédure analogue à celle du paragraphe (C.2.d), en écrivant la fonction d'onde solution sous la forme :

$$|\psi(t)\rangle = \gamma_i(t)|i\rangle + \sum_{k=-\infty}^{+\infty} \gamma_k(t)e^{-ik\varepsilon t/\hbar}|k\rangle \quad (2)$$

L'équation de Schrödinger est alors équivalente au système suivant d'équations différentielles :

$$\begin{cases} \frac{d}{dt}\gamma_i(t) = \frac{1}{i\hbar} \sum_{k=-\infty}^{+\infty} w_k \gamma_k(t) e^{-ik\varepsilon t/\hbar} \\ \frac{d}{dt}\gamma_k(t) = \frac{w_k}{i\hbar} e^{ik\varepsilon t/\hbar} \gamma_i(t) \end{cases} \quad (3)$$

L'intégration formelle de la deuxième équation (voir équation (1.76)) permet d'obtenir l'équation suivante pour $\gamma_i(t)$:

$$\frac{d}{dt}\gamma_i(t) = - \int_0^t dt' \gamma_i(t') \left[\sum_{k=-\infty}^{+\infty} \frac{w_k^2}{\hbar^2} e^{ik\varepsilon(t'-t)/\hbar} \right] \quad (4)$$

L'intervalle ε étant très petit, on peut, après avoir utilisé (1), remplacer la somme entre crochets par l'intégrale :

$$\int_{-\infty}^{+\infty} \frac{dE}{\varepsilon} \frac{w^2}{\hbar^2} \frac{1}{1 \left(\frac{E}{\Delta}\right)^2} e^{iE(t'-t)/\hbar} = \frac{w^2}{\hbar^2} \frac{\pi\Delta}{\varepsilon} e^{-\Delta|t'-t|/\hbar} \quad (5)$$

Introduisant la notation habituelle $\Gamma = (2\pi/\hbar)(w^2/\varepsilon)$ (probabilité de transition par unité de temps donnée par la règle d'or de Fermi), on aboutit ainsi à l'équation intégrodifférentielle suivante pour γ_i :

$$\frac{d}{dt}\gamma_i(t) = -\frac{\Gamma\Delta}{2\hbar} \int_0^t dt' \gamma_i(t') e^{\Delta(t'-t)/\hbar} \quad (6)$$

ou, de manière équivalente, en dérivant cette équation par rapport à t , à l'équation différentielle du deuxième ordre :

$$\frac{d^2}{dt^2}\gamma_i(t) + \frac{\Delta}{\hbar} \frac{d}{dt}\gamma_i(t) + \frac{\Gamma\Delta}{2\hbar} \gamma_i = 0 \quad (7)$$

Cette équation différentielle peut se résoudre simplement à cause de la forme particulière de couplage au continuum que nous avons choisie au moyen de l'équation (1).

Supposons que le système se trouve à l'instant initial dans le niveau $|i\rangle$. Les conditions initiales sont donc : $\gamma_i(0) = 1$, $\gamma_k(0) = 0$ d'où $\dot{\gamma}_i(0) = 0$ en vertu de (3). La solution correspondante de (7) s'écrit alors :

- Dans le cas où $\Delta > 2\hbar\Gamma$:

$$\gamma_i(t) = e^{-\Delta t/2\hbar} \left(\cosh \frac{\Delta' t}{2\hbar} + \frac{\Delta}{\Delta'} \sinh \frac{\Delta' t}{2\hbar} \right), \text{ avec } \Delta' = \Delta \sqrt{1 - \frac{2\hbar\Gamma}{\Delta}} \quad (8)$$

- Dans le cas où $\Delta < 2\hbar\Gamma$:

$$\gamma_i(t) = e^{-\Delta t/2\hbar} \left(\cos \frac{\Delta'' t}{2\hbar} + \frac{\Delta}{\Delta''} \sin \frac{\Delta'' t}{2\hbar} \right), \text{ avec } \Delta'' = \Delta \sqrt{\frac{2\hbar}{\Delta} - 1} \quad (9)$$

Considérons alors les deux cas extrêmes pour la largeur du continuum :

- Si $\hbar\Gamma/\Delta \ll 1$ (continuum large voir la courbe a) sur la figure 1), $\gamma_i(t)$ s'écrit au premier ordre de perturbation en $\hbar\Gamma/\Delta$:

$$\gamma_i(t) = e^{-\Gamma t/2} \left(1 + \frac{\hbar\Gamma}{2\Delta} \right) - \frac{\hbar\Gamma}{2\Delta} e^{-\Delta t/\hbar} \quad (10)$$

Pour des temps longs devant l'inverse de la largeur du continuum $\Delta/\hbar$, la contribution de la deuxième exponentielle est négligeable et on retrouve la décroissance exponentielle habituelle (équation (C.17)). Le deuxième terme joue par contre un rôle important aux temps courts ($t < \hbar/\Delta$), où il assure le démarrage en t^2 , et non pas en t , de $\gamma_i(t)$. En effet, au voisinage de $t = 0$, et au premier ordre en $\hbar\Gamma/\Delta$, $\gamma_i(t)$, s'écrit :

$$\gamma_i(t) = 1 - \frac{\Gamma\Delta}{4\hbar} t^2 \quad (11)$$

- Si $\hbar\Gamma/\Delta \gg 1$ (continuum étroit courbe c)), la solution approchée de (7) s'écrit, à l'ordre le plus bas :

$$\gamma_i(t) = e^{-\Delta t/2\hbar} \cos \sqrt{\frac{\Gamma\Delta}{2\hbar}} t \quad (12)$$

On obtient des oscillations de pulsation :

$$\left(\frac{\Gamma\Delta}{2\hbar} \right)^{1/2} = \frac{w}{\hbar} \left(\frac{\pi\Delta}{\varepsilon} \right)^{1/2} \quad (13)$$

analogues aux oscillations de Rabi, mais qui s'amortissent sur un temps de l'ordre de $(\hbar/\Delta)$. Aux temps courts, lorsque cet amortissement est négligeable, tout se passe donc comme si le niveau initial était couplé à un niveau final unique, avec un élément de matrice égal à $w(\pi\Delta/\varepsilon)^{1/2}$.

Dans ce modèle, la valeur particulière $\Delta = 2\hbar\Gamma$ marque la frontière entre le régime à décroissance monotone et le régime pour lequel des oscillations analogues aux oscillations de Rabi commencent à apparaître. La figure (1) donne l'évolution de $P_i(t) = |\gamma_i(t)|^2$ pour diverses valeurs du paramètre $p = 2\hbar\Gamma/\Delta$. On y voit comment s'opère le passage continu du régime oscillant au régime exponentiel.

Si la forme précise de $P_i(t)$ dépend évidemment du modèle particulier choisi ici, son allure générale n'en dépend pas : il faut en particulier retenir que *la transition du régime exponentiel au régime oscillant se fait lorsque la probabilité de transition par unité de temps Γ calculée par la règle d'or de Fermi est de l'ordre de grandeur de la largeur du continuum*. Il faut aussi noter que même dans le régime oscillant, on aura toujours une décroissance globale de l'amplitude de l'oscillation sur des temps de l'ordre de $\hbar$ divisé par la largeur du continuum.

Remarques

(i) Notons que, dans le modèle étudié dans ce complément, la largeur du continuum est en toute rigueur infinie, puisque la densité d'états est constante et égal à $1/\varepsilon$. Ce qui compte en définitive est une « largeur effective », associée à la décroissance avec k de l'élément de matrice de couplage (voir équation 1). D'une manière générale, la largeur effective du continuum sera associée à la variation avec k de la quantité $\rho(E)|\langle k|\hat{W}|i\rangle|^2$.

(ii) Pour le cas de continums plus complexes, on pourra consulter CDG 2, Complément C_{III}.

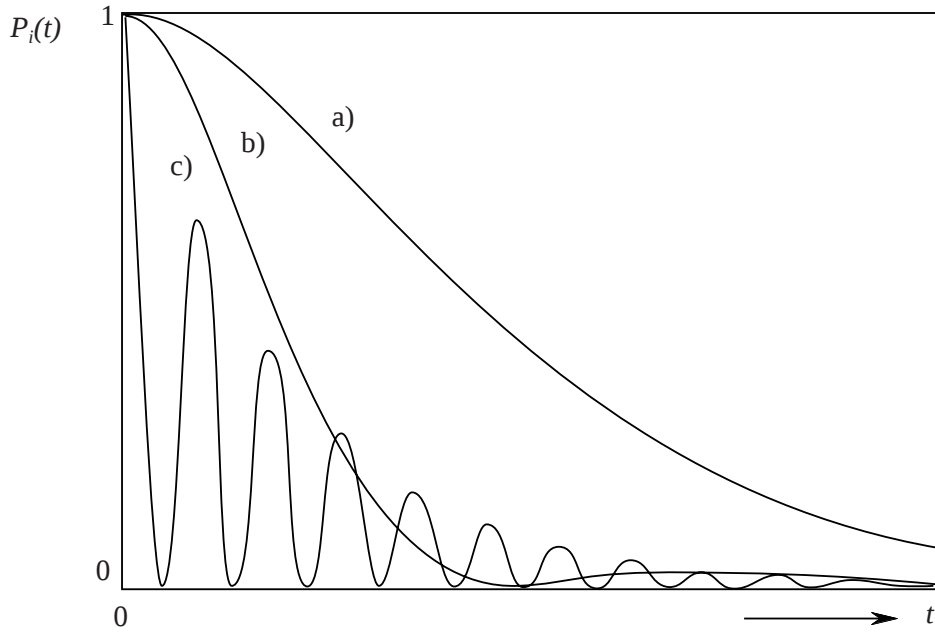


FIG. 1: **Évolution temporelle** de la probabilité de rester dans le niveau initial pour différentes valeurs du paramètre $p = \hbar\Gamma/\Delta$ correspondant à des continums de plus en plus étroits (Δ étant maintenu constant) : a) $p = 0,5$; b) $p = 2$; c) $p = 100$. Le temps t varie de 0 à $5 \hbar/\Delta$.

Chapitre 2

Atomes en interaction avec une onde électromagnétique classique. Amplification laser

Nous abordons dans ce chapitre le problème général de l'interaction entre un atome (ou une molécule) et un champ électromagnétique. Son importance tient d'abord au fait qu'une grande partie de notre connaissance des atomes est obtenue par l'étude du rayonnement électromagnétique absorbé et émis par les atomes (nous considérons aussi bien la lumière que les champs radiofréquence ou le rayonnement X). Réciproquement, la matière modifie la propagation des ondes électromagnétiques, notamment par les effets d'absorption, de réfraction, et de diffusion. Le domaine concerné est immense, et il est hors de question de le couvrir en un seul chapitre. Nous nous proposons ici de présenter quelques éléments de base sur l'interaction entre un *atome*, traité *quantiquement*, et un champ *électromagnétique classique*, c'est-à-dire une onde électromagnétique constituée de champs électriques et magnétiques réels obéissant aux équations de Maxwell.

Une description rigoureuse devrait prendre en compte la nature quantique de la lumière. Il s'avère néanmoins que de nombreux résultats importants peuvent être obtenus dans le point de vue « semi-classique » adopté ici (qu'il serait sans doute plus logique d'appeler semi-quantique). On peut d'ailleurs démontrer, dans le cadre de l'optique quantique, que ce traitement semi-classique est pour l'essentiel équivalent au traitement quantique si l'onde lumineuse interagissant avec l'atome a été émise par un laser très au dessus du seuil, ou par une source classique (lampe à décharge ou à incandescence). Il en est de même pour un champ radiofréquence produit par les techniques habituelles de l'électronique (oscillateurs, klystrons...).

Il faut néanmoins savoir que l'approche semi-classique est incapable de traiter de façon rigoureuse un phénomène aussi important que l'émission spontanée, que l'on ne peut déduire des premiers principes que dans le cadre d'une théorie quantique du rayonnement (il en est de même de certains processus de diffusion). Cependant, il est possible d'intro-

duire de façon phénoménologique la durée de vie des niveaux atomiques, et le formalisme semi-classique est alors capable de rendre compte de façon simple d'un très grand nombre d'aspects de l'interaction atome-rayonnement, y compris ceux à la base de l'effet laser. Ce chapitre est donc très important, même si certains effets associés à l'émission spontanée, ainsi que les domaines nouveaux des « états non-classiques de la lumière » et de « l'électrodynamique quantique en cavité » échappent au traitement semi-classique.

Après avoir présenté qualitativement dans la partie 2.1 divers processus d'interaction entre matière et rayonnement – absorption, émission spontanée, émission induite, diffusion, processus multiphotoniques – nous mettons le formalisme en place dans la partie B, en justifiant le choix de l'*hamiltonien dipolaire électrique* pour décrire l'interaction entre un atome et la lumière considérée comme un champ électromagnétique classique. Nous introduisons également à la fin de cette partie l'interaction dipolaire magnétique qui est généralement le terme dominant de l'interaction entre un atome et une onde radiofréquence. Dans la partie C, nous calculons les probabilités de transition, sous l'effet du rayonnement, entre deux états atomiques discrets de durée de vie infinie. Après une première approche perturbative, nous étudions le très important phénomène de l'*oscillation de Rabi*. Nous traitons également dans cette partie les *transitions multiphotoniques*, et les *déplacements lumineux*. Dans la partie D, nous reprenons la question d'une transition quasi-résonnante entre deux niveaux, en introduisant phénoménologiquement la durée de vie des niveaux concernés. Afin de conserver un traitement simple, nous nous limitons au cas particulier d'une transition entre deux niveaux ayant des durées de vie égales. Nous obtenons alors une expression réaliste de la réponse d'une assemblée d'atomes soumis à une onde lumineuse, réponse caractérisée par une susceptibilité diélectrique. Dans la partie 2.5, nous exploitons ce modèle pour montrer comment un tel milieu, convenablement excité, peut donner une amplification de l'onde incidente : on a alors un milieu susceptible de donner lieu à l'**amplification laser**.

Ce chapitre est suivi de cinq compléments. Le **complément II.1** est consacré au **pompage optique**, méthode inventée et développée par A. Kastler et J. Brossel, grâce à laquelle il est possible de préparer une assemblée d'atomes dans des sous-niveaux Zeeman précis, en utilisant la sélectivité des transitions provoquées par de la lumière polarisée (règles de sélection). Le **complément II.2** présente le formalisme de la matrice densité qui permet d'écrire les **équations de Bloch optiques**, grâce auxquelles il est possible de traiter de façon générale le problème de l'interaction d'une onde électromagnétique et d'un atome, pour une transition entre niveaux de durée de vie quelconque. Ce formalisme, difficile à justifier au niveau de cet ouvrage, n'est généralement abordé que dans des cours de troisième cycle ; mais il est d'une telle importance pour l'interaction lumière-matière qu'il nous a paru difficile de le passer sous silence. Le **complément II.3** présente le **modèle de l'électron élastiquement lié**. Dans ce modèle complètement classique, ni la matière ni le rayonnement ne sont quantifiés. Ce modèle permet pourtant d'obtenir de nombreux résultats intéressants, qu'il sera instructif de comparer aux résultats quantiques. De plus, il fournit des images simples, notamment en ce qui concerne la polarisation du rayonnement émis par un atome. Le **complément II.4**, portant sur l'**effet photoélectrique**, peut être considéré comme un exercice corrigé présentant un effet physique intéressant, à l'origine de l'optique quantique puisque c'est pour l'interpréter qu'Einstein introduisit la notion de photon : il s'agit de l'éjection d'un électron hors de la matière sous l'effet d'un rayonnement incident.

Cet effet est le mécanisme à la base du fonctionnement de certains détecteurs photoélectriques comme les photomultiplicateurs. Le **complément II.5** présente un effet analogue dans des **hétérostructures semi-conductrices**, à la base de détecteurs de rayonnement infrarouge.

2.1 Processus d'interaction atome – champ électromagnétique

Cette première partie est destinée à présenter qualitativement les phénomènes les plus importants qui peuvent se produire dans l'interaction entre une onde électromagnétique et un atome. Il s'agit ici de permettre au lecteur de se familiariser avec ces processus qui reviendront tout au long de l'ouvrage, et d'insister sur un certain nombre de leurs caractéristiques essentielles. Cette partie dépasse donc le cadre du chapitre 2, et il ne faut pas s'étonner d'y voir mentionnées des notions comme celle de photon, qui font partie de la culture de base en optique moderne.

2.1.1 Absorption

Considérons un atome dont le centre de masse est situé à l'origine des coordonnées, et dont l'hamiltonien interne $\hat{H}_0$ admet des états propres atomiques $|a\rangle, |b\rangle, |c\rangle \dots$ d'énergies $E_a, E_b, E_c \dots$ (le niveau $|a\rangle$, correspondant à la plus petite des énergies, est le niveau fondamental). Cet atome est soumis à une onde électromagnétique incidente monochromatique, dont le champ électrique s'écrit :

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0 \cos(\omega t + \varphi(\mathbf{r})) . \quad (2.1)$$

Sous l'action de cette onde électromagnétique, l'atome, initialement dans l'état $|a\rangle$, peut être porté vers un état $|b\rangle$ ou $|c\rangle$ d'énergie plus élevée, *tout en prélevant cette énergie sur le champ* qui est donc atténué (Figure 2.1). Comme on l'a vu au chapitre 1 (§ 1.2.4.d), ce processus n'est important que si l'onde électromagnétique est quasi-résonnante, c'est-à-dire si sa fréquence est très proche d'une fréquence de Bohr atomique, par exemple

$$\omega = \omega_{ba} = \frac{E_b - E_a}{\hbar} . \quad (2.2)$$

Pour ω voisin de ω_{ba} , on pourra en première approximation négliger les autres niveaux atomiques, et utiliser le modèle simplifié de l'**atome à deux niveaux** (a et b). C'est ce modèle que nous utiliserons dans tout ce chapitre, chaque fois que nous aurons affaire à des **processus quasi-résonnants**. Nous noterons alors la fréquence de résonance ω_0 .

La condition de résonance suggère naturellement d'interpréter le processus ci-dessus comme l'absorption dans l'onde incidente d'un *photon* d'énergie $\hbar\omega$ faisant passer l'atome du niveau d'énergie E_a au niveau d'énergie E_b . Il s'agit d'une *image* commode et souvent

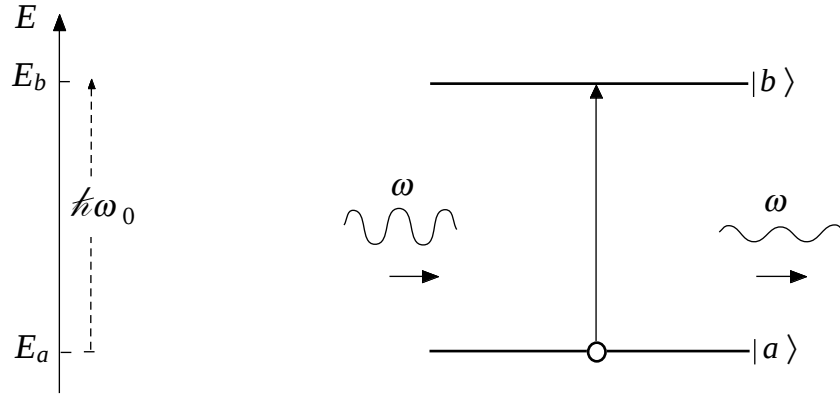


FIG. 2.1: **Processus d'absorption.** L'excitation de l'atome de l'état fondamental $|a\rangle$ vers l'état excité $|b\rangle$ s'accompagne d'une diminution de l'amplitude de l'onde incidente de pulsation ω . L'effet ne se produit que si l'onde est quasi-résonnante avec la fréquence de Bohr ω_{ba} , que nous noterons ω_0 pour le modèle de l'atome à deux niveaux.

fructueuse. Il convient néanmoins de remarquer qu'une telle interprétation n'est en rien justifiée par le formalisme de ce chapitre, où la lumière n'est pas quantifiée. Comme on l'a montré au chapitre 1 (et comme nous le reverrons au paragraphe 2.3), cette condition de résonance provient de l'application de la théorie des perturbations au cas d'une interaction modulée sinusoïdalement. Nous nous permettons néanmoins d'utiliser ce type d'image quand nous le jugerons utile à la compréhension des phénomènes.

Remarques

(i) Le modèle de l'atome à deux niveaux pourrait sembler sans intérêt lorsque plusieurs niveaux atomiques sont dégénérés en énergie : la condition de résonance n'est alors plus suffisante pour isoler deux niveaux particuliers. En fait, si on utilise de la lumière polarisée, le rayonnement n'interagit qu'avec certains sous-niveaux (règles de sélection, cf. Complément II.1). La portée du modèle est donc beaucoup plus grande qu'on ne pourrait le penser a priori.

(ii) C'est pour simplifier les notations que nous avons considéré un atome à l'origine des coordonnées, mais cette hypothèse n'a rien d'essentiel. Il est facile, et *parfois nécessaire*, d'introduire dans le formalisme la position $\mathbf{r}_0$ du centre de masse. Par exemple, dans le cas d'un atome en mouvement uniforme dans une onde progressive, on introduit explicitement une position $\mathbf{r}_0 = \mathbf{v}_{\text{at}} t$ dépendant du temps, et le champ électrique sur l'atome s'écrit

$$\mathbf{E} = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r}_0) = \mathbf{E}_0 \cos(\omega - \mathbf{k} \cdot \mathbf{v}_{\text{at}}) t. \quad (2.3)$$

On constate que la fréquence vue par l'atome est déplacée par l'effet Doppler $\Delta\omega_D = -\mathbf{k} \cdot \mathbf{v}_{\text{at}}$. Ce phénomène joue un rôle important en spectroscopie (cf. Complément III.5).

2.1.2 Émission induite

Si nous considérons maintenant un atome dans le niveau excité $|b\rangle$, irradié par une onde électromagnétique de fréquence ω , il peut se désexciter vers le niveau $|a\rangle$ sous l'effet de l'onde, qui sera alors **amplifiée** (Figure 2.2). Ce processus, appelé émission induite

– ou *émission stimulée* – a été introduit par Einstein en 1916 pour des raisons théoriques. Comme l'absorption dont il est le symétrique, ce processus n'est important que si l'onde est quasi-résonnante avec la fréquence atomique ω_0 , ce qui suggère ici l'image de l'émission induite d'un photon d'énergie $\hbar\omega$ venant s'ajouter à l'onde incidente. Pour rendre compte du fait que le rayonnement incident voit son *amplitude* augmentée, on peut dire que le rayonnement induit possède la *même fréquence*, la même *direction* de propagation, et une phase bien définie par rapport au rayonnement incident, avec lequel il interfère constructivement.

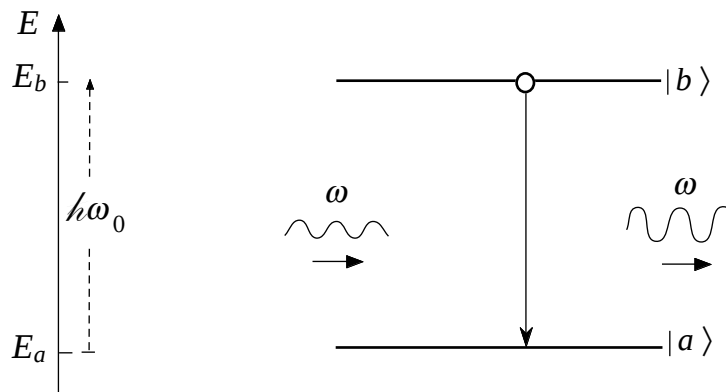


FIG. 2.2: **Émission induite.** L'atome se désexcite sous l'effet de l'onde incidente, dont l'amplitude augmente. L'effet n'est important que si la fréquence de l'onde incidente est quasi-résonnante avec la fréquence de Bohr ω_0 de la transition.

Longtemps considérée comme une curiosité théorique, car elle jouait peu de rôle dans les sources de lumière traditionnelles, l'émission induite est devenue un phénomène très important qui permet l'amplification des ondes lumineuses. Nous verrons qu'elle est à la base du fonctionnement des *lasers*.

Remarques

(i) On peut se demander avec quelle précision la fréquence de l'onde incidente doit coïncider avec la fréquence atomique ω_0 pour que l'*émission induite* – ou l'*absorption* – soit un phénomène important. Nous verrons dans les parties C et D qu'à faible intensité la *résonance* suit une *loi Lorentzienne* dont la largeur est de l'ordre soit de l'inverse du temps d'interaction, soit de l'inverse de la durée de vie des niveaux atomiques (c'est la plus grande des deux largeurs qui compte). À forte intensité, la largeur de la résonance devient plus grande que la largeur précédente, et elle croît comme l'amplitude du champ, c'est-à-dire comme la racine carrée de l'éclairement de l'onde (puissance par unité de surface).

(ii) On considère souvent que l'apport majeur de l'article d'Einstein de 1916 est la prédiction du phénomène d'émission induite, ce qui est exact sur le plan historique. Sur un plan conceptuel, et avec un recul de plus de cinquante ans, on peut avoir un point de vue différent, dans lequel absorption et émission induite sont deux processus totalement symétriques, descriptibles avec des champs classiques. L'originalité de la contribution d'Einstein est alors la nécessité de l'existence de l'émission spontanée, et la prédiction quantitative de ses propriétés, par exemple sa variation avec la fréquence.

2.1.3 Émission spontanée

Si l'atome est initialement dans l'état excité $|b\rangle$, il peut se désexciter vers l'état fondamental $|a\rangle$ même en l'absence de tout rayonnement incident : ce phénomène est l'émission spontanée. Le rayonnement émis a une fréquence égale à la fréquence de Bohr ω_0 . Sa direction de propagation, aussi bien que sa phase, ont un caractère aléatoire.

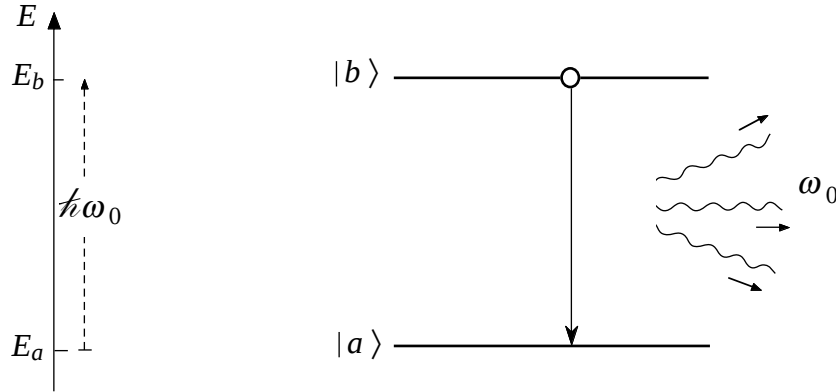


FIG. 2.3: **Émission spontanée.** L'atome, initialement dans l'état excité $|b\rangle$, se désexcite en émettant un rayonnement de fréquence égale à la fréquence de Bohr atomique ω_0 . Ce processus, qui se produit même en l'absence de rayonnement incident, crée un rayonnement dont la direction de propagation, aussi bien que la phase, ont un caractère aléatoire.

Ici encore, on peut donner une image suggestive de ce processus en terme de photon émis : une émission spontanée donne lieu à l'émission d'un photon d'énergie $\hbar\omega_0$, de direction, de phase, et de polarisation aléatoires. En fait, c'est même *la seule image simple possible*, car l'émission spontanée est un phénomène qui est fondamentalement lié à la quantification du champ électromagnétique. Si nous le présentons dans ce chapitre, c'est à cause de son importance physique. Mais à la différence de l'absorption et de l'émission induite, *l'émission spontanée ne peut être déduite du formalisme semi-classique* utilisé ici.

Remarquons en effet qu'un atome isolé dans l'état $|b\rangle$ est dans un état stationnaire (état propre de l'hamiltonien atomique). Les théorèmes élémentaires de la mécanique quantique nous indiquent alors que l'atome reste indéfiniment dans cet état. Ce n'est qu'en considérant un système quantique plus grand – l'atome en interaction avec tous les modes du champ électromagnétique quantifié – que l'on peut voir apparaître naturellement l'émission spontanée.

Dans ce chapitre, nous nous contenterons de savoir que le phénomène d'émission spontanée existe, et si nécessaire nous pourrions l'introduire de façon phénoménologique en utilisant la notion de durée de vie radiative d'un état. Nous admettrons que la probabilité par unité de temps de désexcitation par émission spontanée à partir de l'état $|b\rangle$ vaut Γ . Si l'atome est dans l'état $|b\rangle$ à l'instant $t = 0$, la probabilité P_b de le trouver encore dans $|b\rangle$ à l'instant t est donc :

$$P_b = e^{-\Gamma t} = e^{-\frac{t}{\tau}} . \quad (2.4)$$

La quantité $\tau = \Gamma^{-1}$ est la *durée de vie radiative* de l'état excité $|b\rangle$.

Remarques

(i) La caractérisation de l'émission spontanée par un taux de transition (cf. § 1.3 du chapitre 1) est liée au fait qu'il s'agit, en théorie quantique, d'une transition vers un continuum (l'ensemble des états ayant toutes les directions, polarisations, et fréquences possibles pour le photon émis).

(ii) Si l'on voulait donner une représentation du rayonnement émis dans une direction donnée en terme de champ classique, on pourrait le décrire approximativement comme un train d'onde amorti, d'amplitude

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0(\mathbf{r}) H(t - r/c) \exp\left(-\frac{\Gamma}{2}(t - r/c)\right) \cos(\omega_0(t - r/c) + \phi), \quad (2.5)$$

formule dans laquelle $H(u)$ est la fonction de Heaviside (qui vaut zéro pour u négatif et 1 pour u positif), et ϕ est une phase aléatoire. En prenant le carré du module de la transformée de Fourier de cette expression, on obtient la densité spectrale de puissance du rayonnement émis, qui a une forme Lorentzienne de largeur à mi-hauteur Γ :

$$\rho(\omega) = \frac{\rho(\omega_0)}{1 + \frac{4(\omega - \omega_0)^2}{\Gamma^2}}. \quad (2.6)$$

Le point à retenir ici est que le rayonnement émis a une répartition spectrale de forme *Lorentzienne*, de largeur égale à l'inverse de la durée de vie du niveau excité :

$$\Delta\omega = \Gamma = \frac{1}{\tau}. \quad (2.7)$$

(iii) La loi Lorentzienne (2.6) décrit également le comportement résonnant des sections efficaces d'absorption et d'émission induite lorsque la seule cause d'élargissement de la transition est l'instabilité radiative de son niveau supérieur.

(iv) On connaît en électrodynamique classique le modèle de l'électron élastiquement lié, qui permet de donner une image complètement classique des processus évoqués dans ce paragraphe. Ce modèle donne en particulier les équations (2.5) et (2.6).

(v) Il existe une autre cause d'élargissement homogène (c'est-à-dire identique pour tous les atomes) des raies spectrales : il s'agit des collisions avec d'autres atomes, qui provoquent des sauts brutaux et aléatoires de la phase ϕ de l'onde émise (cf. Équation (2.5)). On peut montrer que si T_2 est l'intervalle de temps moyen entre deux collisions déphasantes, la forme de raie reste Lorentzienne, mais avec une largeur de raie :

$$\Delta\omega = \frac{1}{\tau} + \frac{1}{T_2}. \quad (2.8)$$

2.1.4 Diffusion élastique

Considérons à nouveau un atome à deux niveaux initialement dans son état fondamental $|a\rangle$, et soumis à une onde incidente de fréquence ω pouvant être très différente de la fréquence de résonance atomique ω_0 . Dans cette situation, une fraction de l'onde incidente est *diffusée* : l'amplitude de l'onde transmise est inférieure à celle de l'onde incidente, et

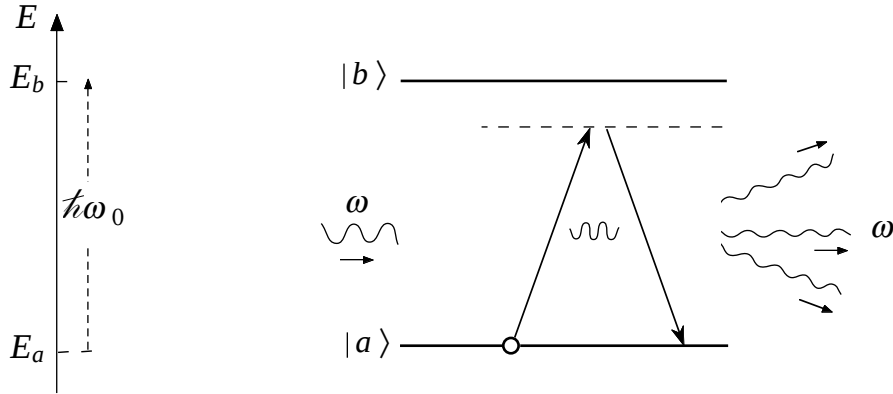


FIG. 2.4: **Processus de diffusion élastique.** L'amplitude de l'onde incidente diminue, tandis qu'il y a apparition de lumière à la fréquence de l'onde incidente dans des directions différentes de celle de l'onde incidente. On peut considérer que l'atome passe virtuellement dans le niveau excité.

l'atome émet une onde sphérique, de fréquence exactement égale à la fréquence ω de l'onde incidente, avec une phase bien définie par rapport à celle de l'onde incidente.

L'interprétation la plus simple du processus de diffusion élastique repose sur l'excitation forcée, sous l'effet du champ incident, d'un dipôle atomique induit oscillant à la fréquence ω (cf. § 2.4.3.a). Ce dipôle rayonne lui-même dans tout l'espace un champ à fréquence ω que l'on peut calculer par les méthodes de l'électrodynamique classique. Un tel calcul permet de préciser la répartition angulaire et l'efficacité du processus de diffusion.

Anticipant sur la description quantique du rayonnement, on peut également proposer une interprétation en deux étapes : (i) absorption d'un photon d'énergie $\hbar\omega$, portant « virtuellement » l'atome dans le niveau excité pendant un intervalle de temps très court de l'ordre de $1/|\omega - \omega_0|$, compatible avec la relation de dispersion temps-énergie $\Delta E \Delta t \geq \hbar$; (ii) désexcitation de l'atome et émission d'un photon d'énergie $\hbar\omega$ dans une direction différente de celle de l'onde incidente. L'égalité entre les fréquences des ondes incidente et diffusée s'interprète alors simplement comme la conservation de l'énergie.

L'efficacité de la diffusion élastique varie fortement avec la fréquence ω de l'onde incidente. Si ω est petit devant la fréquence de résonance ω_0 , la puissance diffusée varie suivant une loi en ω^4 , et on parle alors de *diffusion Rayleigh* (la diffusion de la lumière solaire par les molécules atmosphériques, dont les résonances électroniques sont dans l'ultraviolet, suit une loi de ce type qui favorise les courtes longueurs d'onde, ce qui explique la couleur bleue du ciel). Si au contraire ω est très grand devant la fréquence de résonance ω_0 (diffusion de rayons X), l'efficacité du processus est constante : il s'agit de la *diffusion Thomson*. Enfin, autour de ω_0 le comportement est résonnant, suivant une loi Lorentzienne analogue à celle décrivant l'absorption ou l'émission induite.

Remarques

(i) Lorsque la seule cause d'élargissement de la transition est l'émission spontanée depuis l'état excité, la largeur de la résonance est l'inverse Γ de la durée de vie de cet état (cf. § 2.1.3). Cette dépendance spectrale pourrait inciter à interpréter la diffusion élastique résonnante comme résultant de deux processus indépendants non corrélés : une absorption portant l'atome dans l'état excité, puis une émission spontanée à partir du niveau excité (Figure 2.3). Une telle image risque de conduire à la conclusion erronée que la lumière diffusée a une largeur spectrale Γ et une phase aléatoire, alors qu'elle est monochromatique à la fréquence de l'onde incidente avec laquelle elle a une relation de phase constante. Même si en théorie quantique certains processus de diffusion peuvent être décrits en terme d'absorption suivie d'émission spontanée, on ne peut pas séparer les deux étapes et faire le calcul comme s'il s'agissait de deux processus successifs indépendants.

(ii) Il est intéressant de noter qu'ici aussi le modèle complètement classique de l'électron élastiquement lié prédit correctement la forme des lois de diffusion. En particulier, on trouve ainsi les valeurs des sections efficaces de diffusion Rayleigh et Thomson.

2.1.5 Processus non-linéaires

Les effets d'absorption et d'émission induite décrits ci-dessus apparaissent au premier ordre du traitement perturbatif de l'interaction atome-champ. À des ordres plus élevés, on prévoit l'existence de transitions lorsque la fréquence du champ est un sous-multiple d'une fréquence de Bohr atomique

$$\omega = \frac{\omega_0}{p}, \quad (2.9)$$

où p est un entier. Ces transitions, associées à des termes non-linéaires de degré p du champ, dans l'hamiltonien effectif d'interaction (voir § 2.3.3), peuvent naturellement s'interpréter comme une absorption de p photons d'énergie $\hbar\omega$ (Figure 2.5). On les appelle aussi *processus multiphotoniques*.

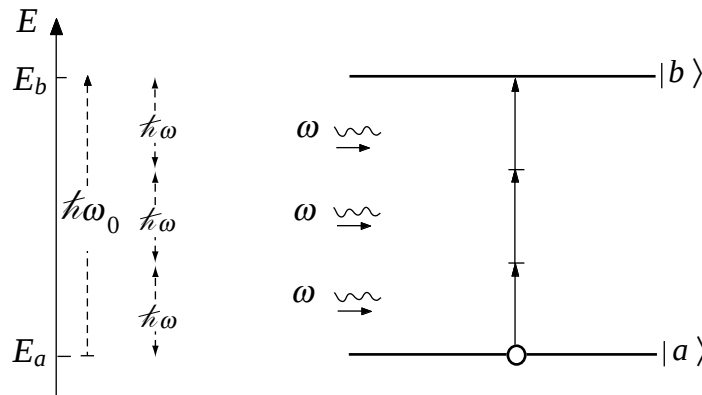


FIG. 2.5: **Processus non-linéaire d'ordre 3**, résonnant lorsque $\omega = \omega_0/3$, et qui peut être interprété comme l'absorption de trois photons d'énergie $\hbar\omega$.

Dans le cas où le champ électromagnétique incident comporte deux fréquences ω_1 et ω_2 , les termes non-linéaires font apparaître des combinaisons linéaires de ces deux fréquences, et il existe des effets résonnants lorsque

$$p_1\omega_1 + p_2\omega_2 = \omega_0 , \quad (2.10)$$

p_1 et p_2 étant des entiers *éventuellement négatifs*. On voit apparaître ici le domaine très riche de la spectroscopie non-linéaire (cf. Complément III.5). À titre d'exemple, la figure 2.6 illustre l'*effet Raman stimulé*, qui est résonnant pour un champ incident comportant deux fréquences ω_1 et ω_2 telles que

$$\omega_1 - \omega_2 = \omega_0 . \quad (2.11)$$

Ce processus donne lieu à l'atténuation du champ à la fréquence ω_1 et à l'amplification de celui à fréquence ω_2 . Il peut s'interpréter comme une transition de a vers b accompagnée de l'absorption d'un photon $\hbar\omega_1$ immédiatement suivie de l'émission stimulée d'un photon $\hbar\omega_2$.

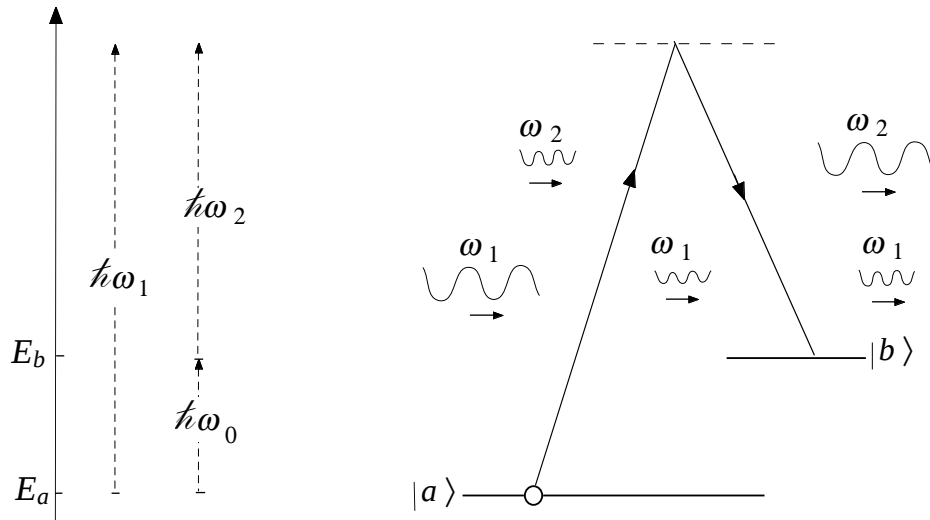


FIG. 2.6: **Effet Raman stimulé**. On peut avoir une transition de $|a\rangle$ vers $|b\rangle$ sous l'effet des ondes de fréquences ω_1 et ω_2 à condition que $\omega_1 - \omega_2 = \omega_0$. Si l'atome est initialement dans l'état $|b\rangle$, le processus inverse peut se produire.

Remarque

Même si l'onde à la fréquence ω_2 est absente, un processus analogue à celui de la Figure 2.6 peut se produire : un photon $\hbar\omega_1$ est absorbé dans l'onde incidente, et un photon $\hbar\omega_2$ est émis spontanément, avec une phase et une direction quelconques. Ce phénomène s'appelle *diffusion Raman spontanée*. Il ne peut être décrit quantitativement que dans le cadre d'une théorie quantique du rayonnement.

2.2 Hamiltonien d'interaction

Nous allons mettre en place dans cette partie le formalisme permettant de décrire l'interaction entre un champ électromagnétique classique et un atome traité comme un objet quantique. L'élément central du formalisme est l'hamiltonien d'interaction, dont les éléments de matrice non-diagonaux entre deux états atomiques sont responsables des transitions entre ces deux états. Contrairement à ce que l'on pourrait penser a priori, cet hamiltonien n'est pas déterminé de façon univoque. Il s'exprime en effet en fonction des potentiels électromagnétiques, et on peut, par des *transformations de jauge*, en donner plusieurs formes équivalentes. Cette liberté au niveau du choix de la jauge – et donc de l'hamiltonien d'interaction – ne remet évidemment pas en cause **l'unicité des prédictions physiques**. Il est donc possible de choisir la jauge qui conduit à l'hamiltonien d'interaction le plus commode pour l'étude du processus physique considéré. En particulier, nous montrerons que la transformation de Goppert-Mayer permet, au prix d'approximations que nous préciserons, d'aboutir à un hamiltonien d'interaction ne mettant en jeu que le champ électrique : c'est l'**hamiltonien dipolaire électrique**, dont la forme très naturelle rappelle l'expression de l'énergie d'interaction entre un dipôle électrique classique et un champ électrique.

Dans le cas où l'hamiltonien dipolaire électrique ne peut donner lieu à des transitions (parce que les éléments de matrice correspondant sont nuls), il existe d'autres termes d'interaction que l'on trouve en conservant des termes d'ordre supérieur dans les approximations. Le plus important pour nous est l'hamiltonien d'interaction *dipolaire magnétique*, qui décrit en particulier les transitions entre sous-niveaux d'un même niveau électronique atomique (par exemple sous-niveaux Zeeman, ou niveaux de structure fine, ou hyperfine...). Le couplage dipolaire électrique est alors nul, mais le couplage entre le champ magnétique d'une onde radiofréquence et le dipôle magnétique atomique va permettre aux transitions de se produire. On aborde ici le domaine de la spectroscopie des radiofréquences, dont les applications sont extrêmement importantes, telle l'horloge atomique à Césium qui constitue l'étalon de temps actuel.

Cette partie 2.2 fait appel à des notions de base en électromagnétisme qui sont supposées connues et que l'on rappelle simplement pour mémoire dans le paragraphe 2.2.1. Occasionnellement, on fera référence à des notions plus avancées, mais elles ne sont pas indispensables pour comprendre l'essentiel de cette partie.

2.2.1 Électrodynamique classique : Équations de Maxwell-Lorentz

Les équations fondamentales de l'électrodynamique classique décrivent l'évolution d'un système de particules chargées et de champs en interaction. Les équations de Maxwell relient les champs électrique $\mathbf{E}(\mathbf{r}, t)$ et magnétique $\mathbf{B}(\mathbf{r}, t)$ aux densités de charge $\rho(\mathbf{r}, t)$

et de courant $\mathbf{j}(\mathbf{r}, t)$:

$$\nabla \cdot \mathbf{E}(\mathbf{r}, t) = \frac{1}{\varepsilon_0} \rho(\mathbf{r}, t) , \quad (2.12a)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}, t) = 0 , \quad (2.12b)$$

$\Rightarrow$

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{B}(\mathbf{r}, t) , \quad (2.12c)$$

$$\nabla \times \mathbf{B}(\mathbf{r}, t) = \frac{1}{c^2} \frac{\partial}{\partial t} \mathbf{E}(\mathbf{r}, t) + \frac{1}{\varepsilon_0 c^2} \mathbf{j}(\mathbf{r}, t) . \quad (2.12d)$$

On sait que les deux équations (2.12b) et (2.12c) entraînent l'existence de potentiels vecteur et scalaire, $\mathbf{A}(\mathbf{r}, t)$ et $U(\mathbf{r}, t)$, qui caractérisent complètement les champs. Ces derniers s'en déduisent en effet de façon univoque par les relations :

$$\mathbf{B}(\mathbf{r}, t) = \nabla \times \mathbf{A}(\mathbf{r}, t) , \quad (2.13a)$$

$\Rightarrow$

$$\mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{A}(\mathbf{r}, t) - \nabla U(\mathbf{r}, t) . \quad (2.13b)$$

Il existe en revanche une infinité de couples $\{\mathbf{A}(\mathbf{r}, t), U(\mathbf{r}, t)\}$ associés au même champ électromagnétique $\{\mathbf{E}(\mathbf{r}, t), \mathbf{B}(\mathbf{r}, t)\}$. On passe de l'un de ces couples à l'autre par une *transformation de jauge* :

$$\mathbf{A}'(\mathbf{r}, t) = \mathbf{A}(\mathbf{r}, t) + \nabla F(\mathbf{r}, t) , \quad (2.14a)$$

$\Rightarrow$

$$U'(\mathbf{r}, t) = U(\mathbf{r}, t) - \frac{\partial}{\partial t} F(\mathbf{r}, t) , \quad (2.14b)$$

$F(\mathbf{r}, t)$ étant un champ scalaire arbitraire. Il est possible de tirer parti de cet arbitraire pour fixer des potentiels adaptés au problème considéré, en imposant une condition supplémentaire, la *condition de jauge*.

Pour calculer les champs créés par un ensemble de *particules ponctuelles* de charges q_α , localisées aux points $\mathbf{r}_\alpha$, et de vitesses $\mathbf{v}_\alpha$, on exprime les densités de charge et de courant qui apparaissent dans les équations de Maxwell (2.12a) et (2.12d) à l'aide de « fonctions δ de Dirac » :

$$\rho(\mathbf{r}, t) = \sum_{\alpha} q_{\alpha} \delta(\mathbf{r} - \mathbf{r}_{\alpha}(t)) , \quad (2.15a)$$

$\Rightarrow$

$$\mathbf{j}(\mathbf{r}, t) = \sum_{\alpha} q_{\alpha} \mathbf{v}_{\alpha} \delta(\mathbf{r} - \mathbf{r}_{\alpha}(t)) . \quad (2.15b)$$

Par ailleurs, les équations classiques de Newton-Lorentz décrivent la dynamique de chaque particule (de masse m_α) sous l'effet des forces électriques et magnétiques exercées par les champs :

$$\Rightarrow m_\alpha \frac{d\mathbf{v}_\alpha}{dt} = q_\alpha [\mathbf{E}(\mathbf{r}_\alpha(t), t) + \mathbf{v}_\alpha \times \mathbf{B}(\mathbf{r}_\alpha(t), t)] . \quad (2.16)$$

Lorsqu'on a un système isolé de particules chargées en interactions, l'ensemble des équations (2.12) (2.15) et (2.16) constitue un système fermé d'équations couplées. Les champs dépendent du mouvement des particules, et le mouvement des particules dépend en retour des champs. À partir de ces équations, on peut obtenir la *totalité des résultats de l'électrodynamique classique*, théorie dans laquelle ni les champs ni les mouvements des particules ne sont quantifiés. Nous verrons au chapitre V comment il est possible de quantifier un tel problème, et d'obtenir ainsi l'électrodynamique quantique, théorie dans laquelle les champs et les mouvements des particules sont quantifiés.

Il existe néanmoins des situations plus simples, dans lequel les champs sont appliqués de l'extérieur et ne dépendent pas du mouvement des particules chargées du problème. Dans ce cas, les seules variables dynamiques du problème sont celles qui décrivent le mouvement des particules¹, et il suffit de quantifier ces variables pour obtenir une description quantique du problème. Une situation de ce type est précisément celle qui se présente lorsqu'un atome est soumis à un rayonnement électromagnétique classique émis par une source extérieure. Il suffira alors de considérer le mouvement des électrons atomiques placés dans le champ électromagnétique extérieur constitué d'une part du champ coulombien du noyau, et d'autre part du rayonnement électromagnétique appliqué. C'est ce cas que nous traitons dans la suite.

2.2.2 Hamiltonien d'une particule dans un champ électromagnétique classique extérieur

a. Forme de l'hamiltonien

Nous nous proposons de décrire l'interaction entre un champ électromagnétique classique et l'atome le plus simple, constitué d'un électron dans le champ coulombien d'un noyau supposé immobile. Nous nous intéressons donc à la dynamique d'un électron, dont le mouvement est traité par la mécanique quantique, et qui est plongé dans un champ électromagnétique classique. Ce champ est complètement caractérisé par les potentiels $\mathbf{A}(\mathbf{r}, t)$ et $U(\mathbf{r}, t)$, qui prennent en compte aussi bien le *champ coulombien du noyau* que les *champs extérieurs* qui interagissent avec l'atome.

Nous admettrons que la dynamique de cet électron est déterminée par l'hamiltonien :

$$\Rightarrow \hat{H} = \frac{1}{2m} (\hat{\mathbf{p}} - q\mathbf{A}(\hat{\mathbf{r}}, t))^2 + q U(\hat{\mathbf{r}}, t) , \quad (2.17)$$

¹Le cas échéant, on prend aussi en compte l'évolution du spin des particules, qui est par définition une variable quantique.

où $\hat{\mathbf{r}}$ est l'opérateur position de l'électron, tandis que $\hat{\mathbf{p}}$ est l'opérateur $-i\hbar\nabla$ (notons que $\mathbf{A}(\hat{\mathbf{r}}, t)$ et $U(\hat{\mathbf{r}}, t)$ sont des opérateurs en tant que fonctions d'opérateurs). Des justifications rigoureuses de cet hamiltonien, basées sur le formalisme lagrangien, peuvent être trouvées dans des ouvrages plus avancés². Ici, nous nous contenterons de prouver que cet hamiltonien est plausible, en montrant qu'il conduit aux équations classiques du mouvement rappelées au § 2.2.1, lorsque l'on s'intéresse à l'évolution des *valeurs moyennes quantiques des opérateurs position et vitesse*.

b. Opérateur vitesse

Pour justifier l'emploi de l'hamiltonien (2.17), nous allons établir les équations d'évolution des valeurs moyennes de la position et de la vitesse de l'électron atomique, en utilisant le *théorème d'Ehrenfest*³. Nous devons donc d'abord préciser quel est l'opérateur qui représente la vitesse lorsqu'on utilise l'hamiltonien (2.17). A priori, la valeur moyenne de cet opérateur $\hat{\mathbf{v}}$ doit être telle que

$$\Rightarrow \quad \langle \hat{\mathbf{v}} \rangle = \frac{d}{dt} \langle \hat{\mathbf{r}} \rangle . \quad (2.18)$$

Or d'après le théorème d'Ehrenfest

$$\frac{d}{dt} \langle \hat{x} \rangle = \frac{1}{i\hbar} \langle [\hat{x}, \hat{H}] \rangle . \quad (2.19)$$

Comme $\hat{x}$ commute avec $U(\hat{\mathbf{r}}, t)$, on obtient

$$\langle [\hat{x}, \hat{H}] \rangle = i\hbar \left\langle \frac{\hat{p}_x - qA_x(\hat{\mathbf{r}}, t)}{m} \right\rangle . \quad (2.20)$$

De (2.19) et (2.20), nous déduisons

$$\frac{d}{dt} \langle \hat{x} \rangle = \left\langle \frac{\hat{p}_x - qA_x(\hat{\mathbf{r}}, t)}{m} \right\rangle , \quad (2.21)$$

ce qui d'après (2.18) suggère l'expression suivante pour l'*opérateur vitesse* :

$$\Rightarrow \quad \hat{\mathbf{v}} = \frac{\hat{\mathbf{p}} - q\mathbf{A}(\hat{\mathbf{r}}, t)}{m} . \quad (2.22)$$

Notons que $\hat{\mathbf{v}}$ diffère de $\frac{\hat{\mathbf{p}}}{m}$.

c. Équations du mouvement

Pour trouver les équations du mouvement, nous aurons besoin des *relations de commutation entre les composantes de l'opérateur vitesse*. Considérons par exemple le commutateur $[v_x, v_y]$. En partant de (2.22), on trouve :

$$[\hat{v}_x, \hat{v}_y] = -\frac{q}{m^2} [\hat{p}_x, A_y(\hat{\mathbf{r}}, t)] - \frac{q}{m^2} [A_x(\hat{\mathbf{r}}, t), \hat{p}_y] . \quad (2.23)$$

²Voir par exemple CDG1, Chapitre IV.

³Voir BD, Chapitre VIII, § 3.

Or, quelle que soit la fonction $f(\hat{\mathbf{r}})$, le commutateur $[\hat{p}_x, f(\hat{\mathbf{r}})]$ est égal à $-i\hbar \frac{\partial f}{\partial x}$, de sorte que

$$[\hat{v}_x, \hat{v}_y] = i\hbar \frac{q}{m^2} \left(\frac{\partial}{\partial x} A_y(\hat{\mathbf{r}}, t) - \frac{\partial}{\partial y} A_x(\hat{\mathbf{r}}, t) \right). \quad (2.24)$$

En utilisant la relation (2.13a) entre le champ magnétique et le potentiel vecteur, nous obtenons les relations dont nous aurons besoin :

$$[\hat{v}_x, \hat{v}_y] = i\hbar \frac{q}{m^2} B_z(\hat{\mathbf{r}}, t), \quad (2.25a)$$

$$[\hat{v}_y, \hat{v}_z] = i\hbar \frac{q}{m^2} B_x(\hat{\mathbf{r}}, t), \quad (2.25b)$$

$$[\hat{v}_z, \hat{v}_x] = i\hbar \frac{q}{m^2} B_y(\hat{\mathbf{r}}, t). \quad (2.25c)$$

Nous nous proposons maintenant de calculer l'évolution de la valeur moyenne de l'opérateur vitesse. Appliquons à nouveau le théorème d'Ehrenfest :

$$\frac{d}{dt} \langle \hat{v}_x \rangle = \frac{1}{i\hbar} \langle [\hat{v}_x, \hat{H}] \rangle + \left\langle \frac{\partial}{\partial t} \hat{v}_x \right\rangle. \quad (2.26)$$

En utilisant la définition (2.22) de l'opérateur vitesse, qui dépend explicitement du temps par l'intermédiaire de $\mathbf{A}(\hat{\mathbf{r}}, t)$, nous pouvons exprimer le dernier terme de (2.26) :

$$\frac{\partial}{\partial t} \hat{v}_x = -\frac{q}{m} \frac{\partial}{\partial t} A_x(\hat{\mathbf{r}}, t). \quad (2.27)$$

L'équation (2.22) nous permet aussi d'exprimer l'hamiltonien (2.17) en fonction de l'opérateur vitesse

$$\Rightarrow \hat{H} = \frac{1}{2} m \hat{\mathbf{v}}^2 + qU(\hat{\mathbf{r}}, t), \quad (2.28)$$

et le commutateur de (2.26) se transforme en

$$[\hat{v}_x, \hat{H}] = \frac{m}{2} [\hat{v}_x, \hat{v}_y^2] + \frac{m}{2} [\hat{v}_x, \hat{v}_z^2] + q[\hat{v}_x, U(\hat{\mathbf{r}}, t)]. \quad (2.29)$$

Développons les deux premiers termes de (2.29) en tenant compte de (2.25) et en prenant garde que $\hat{\mathbf{v}}$ et $\mathbf{B}(\hat{\mathbf{r}}, t)$ ne commutent pas. On trouve

$$[\hat{v}_x, \hat{v}_y^2] = \hat{v}_y [\hat{v}_x, \hat{v}_y] + [\hat{v}_x, \hat{v}_y] \hat{v}_y = i\hbar \frac{q}{m^2} (\hat{v}_y B_z(\hat{\mathbf{r}}, t) + B_z(\hat{\mathbf{r}}, t) \hat{v}_y), \quad (2.30a)$$

et de même

$$[\hat{v}_x, \hat{v}_z^2] = -i\hbar \frac{q}{m^2} (\hat{v}_z B_y(\hat{\mathbf{r}}, t) + B_y(\hat{\mathbf{r}}, t) \hat{v}_z). \quad (2.30b)$$

Le dernier terme de (2.29) se développe aisément en remarquant que $U(\hat{\mathbf{r}}, t)$ et $\mathbf{A}(\hat{\mathbf{r}}, t)$ commutent :

$$[\hat{v}_x, U(\hat{\mathbf{r}}, t)] = \frac{1}{m} [\hat{p}_x, U(\hat{\mathbf{r}}, t)] = -\frac{i\hbar}{m} \frac{\partial}{\partial x} U(\hat{\mathbf{r}}, t). \quad (2.31)$$

Finalement, en reportant les résultats intermédiaires (2.27) à (2.31) dans (2.26), on obtient

$$\frac{d}{dt}\langle\hat{v}_x\rangle = \frac{q}{2m}\langle\hat{v}_y\hat{B}_z - \hat{v}_z\hat{B}_y + \hat{B}_z\hat{v}_y - \hat{B}_y\hat{v}_z\rangle - \frac{q}{m}\left\langle\frac{\partial U(\hat{\mathbf{r}},t)}{\partial x} + \frac{\partial A_x(\hat{\mathbf{r}},t)}{\partial t}\right\rangle. \quad (2.32)$$

Nous reconnaissons dans le premier terme la composante suivant x de l'opérateur symétrisé associé à la force de Lorentz ($q\mathbf{v} \times \mathbf{B}$). Le deuxième terme fait apparaître la force électrique, puisqu'on reconnaît le champ électrique (Éq. 2.13b). En définitive, on obtient l'équation vectorielle

$$\Rightarrow \quad m\frac{d}{dt}\langle\hat{\mathbf{v}}\rangle = q\left\langle\frac{\hat{\mathbf{v}} \times \mathbf{B}(\hat{\mathbf{r}},t) - \mathbf{B}(\hat{\mathbf{r}},t) \times \hat{\mathbf{v}}}{2}\right\rangle + q\langle\mathbf{E}(\hat{\mathbf{r}},t)\rangle, \quad (2.33)$$

qui est l'analogie quantique de l'équation classique de Newton-Lorentz (2.16).

Nous avons donc montré que le choix de l'hamiltonien (2.17) conduit à la définition (2.22) de l'opérateur vitesse, et à une équation d'évolution (2.33) de la valeur moyenne de cet opérateur vitesse qui est la généralisation naturelle de l'équation classique (2.16) correspondante. Répétons qu'il ne s'agit pas d'une démonstration rigoureuse de la validité de l'hamiltonien (2.17), mais plutôt d'un argument de plausibilité, montrant la cohérence de ce choix.

2.2.3 Hamiltonien d'interaction en jauge de Coulomb

a. Jauge de Coulomb

Nous savons qu'il existe une infinité de couples de potentiels $\{\mathbf{A}(\mathbf{r},t), U(\mathbf{r},t)\}$ correspondant au même champ électromagnétique. On passe d'un couple à l'autre par la transformation de jauge (B.3). On peut mettre à profit cet arbitraire pour imposer une condition supplémentaire sur les potentiels, ce qui revient à choisir une jauge particulière. Parmi les diverses jauges possibles, l'une est bien adaptée aux problèmes de l'optique quantique, c'est la jauge de Coulomb, définie par la condition de jauge

$$\nabla \cdot \mathbf{A}(\mathbf{r},t) = 0. \quad (2.34)$$

Ce choix correspond à une valeur particulière du potentiel vecteur, notée $\mathbf{A}_\perp(\mathbf{r},t)$, et appelée « *A transverse* ». Donnons l'exemple, utile pour la suite, d'une onde plane électromagnétique :

$$\mathbf{E} = \mathbf{E}_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r}), \quad (2.35a)$$

$$\mathbf{B} = \frac{\mathbf{k} \times \mathbf{E}_0}{\omega} \cos(\omega t - \mathbf{k} \cdot \mathbf{r}), \quad (2.35b)$$

$$\mathbf{E}_0 \cdot \mathbf{k} = 0. \quad (2.35c)$$

Le potentiel vecteur correspondant en jauge de Coulomb est

$$\mathbf{A}_\perp = -\frac{\mathbf{E}_0}{\omega} \sin(\omega t - \mathbf{k} \cdot \mathbf{r}), \quad (2.36a)$$

tandis que le potentiel scalaire associé est nul :

$$U(\mathbf{r}, t) = 0 . \quad (2.36b)$$

Il est facile de vérifier que l'équation (2.34) est satisfaite, et que les potentiels (B.25) redonnent bien les champs (2.35) en utilisant les formules (B.2). En particulier, le champ électrique (2.35a) de l'onde électromagnétique libre s'écrit :

$$\mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{A}_\perp(\mathbf{r}, t) . \quad (2.36c)$$

b. Hamiltonien atomique et hamiltonien d'interaction

Plaçons-nous donc en jauge de Coulomb pour traiter le problème, décrit par l'hamiltonien (2.17), de l'interaction entre un *champ électromagnétique extérieur et un atome constitué d'un électron soumis au potentiel coulombien du noyau*. Un grand avantage de cette jauge, pour ce problème, est de permettre une séparation claire entre le champ coulombien statique créé par le noyau atomique, et le rayonnement électromagnétique extérieur appliqué sur le système. Prenons en effet comme champ extérieur l'onde plane monochromatique caractérisée par les potentiels (B.25) en jauge de Coulomb. En ce qui concerne le champ électrostatique créé par le noyau sur l'électron, le potentiel vecteur associé est nul, tandis que le potentiel scalaire est tout simplement le *potentiel coulombien* habituel $U_{\text{coul}}(\mathbf{r})$.

L'hamiltonien total (2.17) s'écrit alors, en posant $V_{\text{coul}}(\mathbf{r}) = qU_{\text{coul}}(\mathbf{r})$

$$\Rightarrow \hat{H} = \frac{1}{2m} (\hat{\mathbf{p}} - q\mathbf{A}_\perp(\hat{\mathbf{r}}, t))^2 + V_{\text{coul}}(\hat{\mathbf{r}}) , \quad (2.37)$$

soit encore, en développant :

$$\hat{H} = \frac{\hat{\mathbf{p}}^2}{2m} - \frac{q}{2m} (\hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\hat{\mathbf{r}}, t) + \mathbf{A}_\perp(\hat{\mathbf{r}}, t) \cdot \hat{\mathbf{p}}) + \frac{q^2 \mathbf{A}_\perp^2(\hat{\mathbf{r}}, t)}{2m} + V_{\text{coul}}(\hat{\mathbf{r}}) . \quad (2.38)$$

Les deux termes $\hat{\mathbf{p}} \cdot \mathbf{A}_\perp$ et $\mathbf{A}_\perp \cdot \hat{\mathbf{p}}$ sont égaux, car

$$\hat{\mathbf{p}} \cdot \mathbf{A}_\perp - \mathbf{A}_\perp \cdot \hat{\mathbf{p}} = [\hat{\mathbf{p}}, \mathbf{A}_\perp(\hat{\mathbf{r}}, t)] = -i\hbar \nabla \cdot \mathbf{A}_\perp(\hat{\mathbf{r}}, t) = 0 \quad (2.39)$$

et $\mathbf{A}_\perp(\hat{\mathbf{r}}, t)$ obéit à (2.34). En définitive, l'hamiltonien de l'électron s'écrit

$$\Rightarrow \hat{H} = \hat{H}_0 + \hat{H}_I , \quad (2.40a)$$

avec

$$\Rightarrow \hat{H}_0 = \frac{\hat{\mathbf{p}}^2}{2m} + V_{\text{coul}}(\hat{\mathbf{r}}) \quad (2.40b)$$

et

$$\Rightarrow \hat{H}_I = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\hat{\mathbf{r}}, t) + \frac{q^2 \mathbf{A}_\perp^2(\hat{\mathbf{r}}, t)}{2m} . \quad (2.40c)$$

L'écriture ci-dessus fait d'abord apparaître l'hamiltonien atomique habituel $\hat{H}_0$: c'est l'hamiltonien décrivant une particule soumise à un potentiel coulombien, c'est-à-dire l'hamiltonien de *l'atome d'hydrogène*. Le terme d'interaction $\hat{H}_I$ fait intervenir à la fois les variables quantifiées $\hat{\mathbf{p}}$ et $\hat{\mathbf{r}}$, décrivant le mouvement de l'électron, et le champ appliqué extérieur complètement caractérisé par le potentiel vecteur en jauge de Coulomb $\mathbf{A}_\perp(\hat{\mathbf{r}}, t)$. Il décrit l'interaction entre l'atome d'hydrogène et le champ extérieur appliqué.

L'hamiltonien d'interaction $\hat{H}_I$ apparaît comme la somme de deux contributions $\hat{H}_{I1}$ et $\hat{H}_{I2}$, respectivement linéaire et quadratique en fonction du potentiel vecteur :

$$\Rightarrow \hat{H}_{I1} = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\hat{\mathbf{r}}, t), \quad (2.40d)$$

$$\Rightarrow \hat{H}_{I2} = \frac{q^2 \mathbf{A}_\perp^2(\hat{\mathbf{r}}, t)}{2m}. \quad (2.40e)$$

Les processus linéaires en fonction du champ incident pourront être décrits à l'aide de l'hamiltonien $\hat{H}_{I1}$ seul.

c. Approximation des grandes longueurs d'onde

En optique quantique, les phénomènes d'interaction entre atome et rayonnement correspondent souvent à des situations où la longueur d'onde du rayonnement λ est *très grande* devant les dimensions atomiques. Par exemple, pour l'atome d'hydrogène, les raies d'émission ou d'absorption ont des longueurs d'onde de l'ordre de 100 nm ou plus, tandis que les dimensions atomiques sont de l'ordre du rayon de Bohr ($a_0 = 0,053$ nm). Dans ces conditions, le champ extérieur est quasiment constant sur l'étendue de l'atome, et on remplacera $\mathbf{A}_\perp(\hat{\mathbf{r}}, t)$ par sa valeur $\mathbf{A}_\perp(\mathbf{r}_0, t)$ à la position $\mathbf{r}_0$ du noyau, effectuant ainsi *l'approximation des grandes longueurs d'onde*.

L'hamiltonien d'interaction (2.40c) s'écrit alors

$$\hat{H}_I = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\mathbf{r}_0, t) + \frac{q^2}{2m} \mathbf{A}_\perp^2(\mathbf{r}_0, t) = \hat{H}_{I1} + \hat{H}_{I2}. \quad (2.41)$$

Cette expression est beaucoup plus simple que (2.40c), puisque l'opérateur de position $\hat{\mathbf{r}}$ relatif à l'électron ne figure pas dans (2.41). Notons en particulier que le deuxième terme de (2.41) – correspondant à $\hat{H}_{I2}$ – est un *scalaire* dont les éléments de matrice entre deux états atomiques différents *sont nuls* : il ne peut donc pas provoquer de transitions.

En définitive, à *l'approximation des grandes longueurs d'onde*, on pourra calculer les transitions induites par un champ extérieur entre deux niveaux atomiques, en utilisant *l'hamiltonien d'interaction*

$$\Rightarrow H_{I1} = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\mathbf{r}_0, t). \quad (2.42)$$

Dans cet hamiltonien souvent appelé « hamiltonien $\mathbf{A} \cdot \mathbf{p}$ » le champ extérieur intervient par le potentiel vecteur en jauge de Coulomb pris à la position du noyau.

2.2.4 Hamiltonien dipolaire électrique

a. Jauge de Göppert-Mayer

Nous introduisons maintenant une deuxième jauge très utilisée pour la description de l'interaction atome-rayonnement, car elle conduit à une forme de l'hamiltonien d'interaction particulièrement suggestive. Il s'agit de la jauge de Göppert-Mayer, qui se déduit de la jauge de Coulomb par la transformation de jauge (2.14a) et (2.14b) dans laquelle on prend

$$F(\mathbf{r}, t) = -(\mathbf{r} - \mathbf{r}_0) \cdot \mathbf{A}_\perp(\mathbf{r}_0, t) . \quad (2.43)$$

Cette transformation privilégie la position $\mathbf{r}_0$ du noyau. Elle donne les potentiels de Göppert-Mayer :

$$\mathbf{A}'(\mathbf{r}, t) = \mathbf{A}_\perp(\mathbf{r}, t) - \mathbf{A}_\perp(\mathbf{r}_0, t) , \quad (2.44a)$$

$$U'(\mathbf{r}, t) = U_{\text{coul}}(\mathbf{r}) + (\mathbf{r} - \mathbf{r}_0) \cdot \frac{\partial}{\partial t} \mathbf{A}_\perp(\mathbf{r}_0, t) . \quad (2.44b)$$

Avec ces potentiels, l'hamiltonien (2.17) s'écrit

$$\hat{H} = \frac{1}{2m}(\hat{\mathbf{p}} - q\mathbf{A}'(\hat{\mathbf{r}}, t))^2 + V_{\text{coul}}(\hat{\mathbf{r}}) + q(\hat{\mathbf{r}} - \mathbf{r}_0) \cdot \frac{\partial}{\partial t} \mathbf{A}_\perp(\mathbf{r}_0, t) . \quad (2.45)$$

En se souvenant que le champ électrique associé au rayonnement appliqué vaut, d'après l'équation (B.25)

$$\mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{A}_\perp(\mathbf{r}, t)$$

et en introduisant l'opérateur *moment dipolaire électrique* de l'atome

$$\hat{\mathbf{D}} = q(\hat{\mathbf{r}} - \mathbf{r}_0) , \quad (2.46)$$

l'hamiltonien (2.45) s'écrit finalement

$$\hat{H} = \frac{1}{2m}(\hat{\mathbf{p}} - q\mathbf{A}'(\hat{\mathbf{r}}, t))^2 + V_{\text{coul}}(\hat{\mathbf{r}}) - \hat{\mathbf{D}} \cdot \hat{\mathbf{E}}(\mathbf{r}_0, t) . \quad (2.47)$$

b. Approximation des grandes longueurs d'onde

Faisons maintenant l'approximation des grandes longueurs d'onde, qui nous permet comme plus haut de remplacer les potentiels associés au rayonnement extérieur par leur valeur à la position $\mathbf{r}_0$ du noyau. On remplace donc, dans (2.47), $\mathbf{A}'(\hat{\mathbf{r}}, t)$ par $\mathbf{A}'(\mathbf{r}_0, t)$. Mais d'après (2.44a)

$$\mathbf{A}'(\mathbf{r}_0, t) = 0 , \quad (2.48)$$

de sorte que (2.47) prend la forme :

$$\Rightarrow \quad \hat{H} = \hat{H}_0 + \hat{H}'_1 , \quad (2.49a)$$

où l'on a, comme en (B.29)

$$\Rightarrow \hat{H}_0 = \frac{\hat{\mathbf{p}}^2}{2m} + V_{\text{coul}}(\hat{\mathbf{r}}) \quad (2.49b)$$

qui est l'hamiltonien atomique habituel. Quant à l'hamiltonien d'interaction

$$\Rightarrow \hat{H}'_I = -\hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{r}_0, t) \quad (2.49c)$$

il porte le nom d'**hamiltonien dipolaire électrique**. Cet hamiltonien rappelle l'expression classique de l'énergie d'interaction entre un champ électrique extérieur $\mathbf{E}(\mathbf{r}, t)$ et un dipôle électrique classique $\mathbf{D}$ situé au point $\mathbf{r}_0$.

Remarques

(i) Nous avons considéré ici le cas simple de l'atome d'hydrogène. Pour un atome quelconque, on peut toujours définir un opérateur dipôle électrique, et l'hamiltonien d'interaction ci-dessus garde la même forme.

(ii) Nous aurons besoin par la suite de connaître la valeur numérique des éléments de matrice de l'opérateur dipôle électrique, pour des transitions optiques. On peut les calculer si on connaît les fonctions d'onde des états atomiques entre lesquels se produit la transition. Pour l'atome d'hydrogène, un bon ordre de grandeur est donné par le produit de la charge électrique de l'électron ($1,6 \times 10^{-19}$ C) par le rayon de Bohr ($a_0 = 0,53 \times 10^{-10}$ m) qui fixe l'échelle de la distance entre le proton et l'électron.

c. Discussion

Nous avons été amenés à décrire le même système dynamique (un atome en interaction avec un champ électromagnétique extérieur) par deux hamiltoniens différents. Cela peut sembler inquiétant quant à *l'unicité* des résultats physiques. En fait, il n'y a aucun problème *tant que l'on utilise les formes exactes* et que l'on effectue des calculs *sans approximation* (les transformations de jauge laissent les champs invariants, et les résultats physiques ne dépendent que des champs). On peut ainsi affirmer⁴ que les deux hamiltoniens d'interaction « $\mathbf{A} \cdot \mathbf{p}$ » et « $\mathbf{D} \cdot \mathbf{E}$ » conduisent aux mêmes *probabilités de transition*. Ceci reste vrai à l'approximation des grandes longueurs d'onde.

En revanche, lorsque l'on fait des approximations plus fortes, les résultats peuvent dépendre du type d'hamiltonien utilisé, et ce d'une manière subtile. Par exemple, si on remplace la fonction d'onde atomique initiale (ou finale) par le même développement approché, on peut aboutir à des résultats différents suivant l'hamiltonien choisi. On montre ainsi⁵ qu'un calcul approché donne généralement des résultats plus précis avec l'hamiltonien dipolaire électrique lorsque l'on traite une transition entre deux niveaux discrets, alors que l'hamiltonien « $\mathbf{A} \cdot \mathbf{p}$ » donne de meilleurs résultats dans le cas d'une transition vers un continuum d'états.

Dans le cadre de l'approximation des grandes longueurs d'onde, un avantage de la jauge de Göppert-Mayer est que les diverses quantités mathématiques y ont généralement une interprétation physique simple. Par exemple, la *vitesse de l'électron* $\hat{\mathbf{v}}$ coïncide avec

⁴Voir CDG 1, Chapitre IV, § B.2.a, § D, et Complément B_{IV}.

⁵CDG 1, Complément E_{IV}, Exercice 2.

l'opérateur $\hat{\mathbf{p}}/m$ (équation (2.22) en tenant compte de (2.48)). Ceci implique que l'hamiltonien $\hat{H}_0$ (2.49b) coïncide *dans cette jauge* avec la somme de l'énergie cinétique et de l'énergie coulombienne de l'électron. De même, nous avons déjà indiqué que l'hamiltonien d'interaction coïncide avec l'énergie dipolaire électrique habituelle.

Lorsque l'approximation des grandes longueurs d'onde n'est plus valable, il n'est pas possible d'utiliser l'hamiltonien dipolaire électrique. Il faut alors soit revenir à l'hamiltonien exact en jauge de Coulomb, soit améliorer le changement de jauge⁶ de façon à obtenir les termes suivants du développement multipolaire (interaction dipolaire magnétique, interaction quadrupolaire électrique...). Ces termes jouent un rôle particulièrement important quand, pour des raisons de symétrie, le terme dipolaire électrique est nul (cf. Complément II.1).

2.2.5 Hamiltonien dipolaire magnétique

Si l'élément de matrice de l'opérateur dipolaire électrique est nul entre deux états, l'hamiltonien d'interaction dipolaire électrique ne peut provoquer de transition entre les deux états. Ceci est le cas si les deux états atomiques considérés ont la même parité, puisque l'opérateur dipolaire électrique est impair, c'est-à-dire qu'il est transformé en son opposé dans une symétrie par rapport à l'origine (Complément II.1). Cette situation se rencontre par exemple lorsqu'on considère deux sous-niveaux hyperfins d'un même niveau atomique qui ont les mêmes nombres quantiques principal et azimuthal n et l . Or on sait qu'il existe pourtant des transitions – tombant dans la bande des radiofréquences ou des hyperfréquences – entre de tels niveaux. Citons par exemple la raie à 1420 MHz de l'hydrogène (« raie de 21 cm » entre les sous-niveaux $F = 0$ et $F = 1$ du niveau fondamental) importante en radio-astronomie, ou encore la transition à 9192 MHz entre les sous-niveaux $F = 3$ et $F = 4$ du niveau fondamental du césium 133, « transition d'horloge » qui sert à définir la seconde.

Pour trouver le terme d'interaction responsable des telles transitions, il faut pousser le développement multipolaire à l'ordre suivant, et aussi tenir compte du couplage entre le champ électromagnétique incident et les moments magnétiques de spin des particules composant l'atome⁷. On obtient alors un terme d'interaction qui s'interprète naturellement comme le couplage dipolaire magnétique entre le champ magnétique $\mathbf{B}$ de l'onde et le moment dipolaire magnétique $\hat{\mathbf{M}}$ de l'atome :

$$\Rightarrow \quad \hat{H}_I'' = -\hat{\mathbf{M}} \cdot \mathbf{B}(\mathbf{r}_0, t) . \quad (2.50)$$

Pour l'électron de l'atome d'hydrogène, $\hat{\mathbf{M}}$ a pour origine non seulement le moment cinétique $\hat{\mathbf{L}}$ associé au mouvement orbital de l'électron, mais aussi le spin $\hat{\mathbf{S}}$ de l'électron :

$$\Rightarrow \quad \hat{\mathbf{M}} = \frac{q}{2m} (\hat{\mathbf{L}} + 2\hat{\mathbf{S}}) . \quad (2.51)$$

⁶CDG 1, Complément C_{IV}.

⁷CDL 2, Complément A_{XIII}

Il est facile de comparer l'ordre de grandeur de ce couplage à celui du couplage dipolaire électrique. Pour un électron atomique, les éléments de matrice de $\hat{\mathbf{M}}$ sont de l'ordre du magnéton de Bohr

$$\mu_B = \frac{\hbar q}{2m}, \quad (2.52)$$

q étant la charge de l'électron. Par ailleurs, pour une onde électromagnétique progressive dans le vide, B est égal à E/c . En adoptant l'ordre de grandeur qa_0 pour un dipole électrique (cf. Remarque (ii), § 2.2.4.b), on a donc :

$$\Rightarrow \frac{\langle \hat{H}_I'' \rangle}{\langle \hat{H}_I' \rangle} \approx \frac{\hbar}{mca_0} = \alpha, \quad (2.53)$$

α étant la constante de structure fine ($\alpha \cong 1/137$). On voit que l'amplitude d'un couplage dipolaire magnétique typique est deux ordres de grandeur plus faible qu'un couplage dipolaire électrique typique. Les couplages dipolaires magnétiques donnent néanmoins lieu à de nombreux effets observables.

Remarque

Au même ordre de développement de l'hamiltonien (2.17), on a un terme d'interaction entre le moment quadropolaire électrique de l'atome et le gradient de champ électrique de l'onde. Ce terme peut également induire des transitions à la fréquence de l'onde électromagnétique. Le poids respectif de ces divers termes est en définitive contrôlé par les symétries des états atomiques et par l'amplitude des éléments de matrice non-nuls.

2.3 Transition entre deux niveaux atomiques sous l'effet d'un champ électromagnétique oscillant

Munis d'un hamiltonien d'interaction, nous pouvons maintenant étudier comment un atome passe d'un état $|i\rangle$ à un état $|k\rangle$ sous l'effet d'une onde électromagnétique sinusoïdale. Nous le ferons d'abord dans l'hypothèse d'un couplage faible, en utilisant le premier ordre de la théorie des perturbations. Pour aborder les cas où la probabilité de transition n'est pas très petite devant 1, nous présenterons au paragraphe 2.3.3 un calcul plus exact, dans le cadre du modèle de *l'atome à deux niveaux* et de *l'approximation résonnante*. On trouve alors **l'oscillation de Rabi**, phénomène physique d'une importance considérable. Dans les paragraphes 2.3.3 et 2.3.4, nous aborderons la question des *transitions multiphotoniques*, et des *déplacements lumineux*.

2.3.1 Probabilité de transition au premier ordre de la théorie des perturbations

a. Absorption et émission induite

On considère un système atomique décrit par un hamiltonien $\hat{H}_0$, dont les états propres $|n\rangle$ ont des énergies E_n . À l'instant initial, l'atome est dans l'état propre $|i\rangle$ de $\hat{H}_0$, et on le soumet à un champ électromagnétique incident de fréquence ω . Nous cherchons la probabilité de trouver l'atome dans un autre état propre $|k\rangle$ de $\hat{H}_0$ à un instant ultérieur t . Nous admettons que les résultats établis dans la partie B en ce qui concerne l'hamiltonien d'interaction d'un atome d'hydrogène restent vrais quel que soit l'atome. En particulier, à l'approximation des grandes longueurs d'ondes, on pourra utiliser un hamiltonien d'interaction

$$\hat{H}'_I(t) = -\hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{r}_0, t) , \quad (2.54)$$

où le champ électrique $\mathbf{E}(\mathbf{r}_0, t)$ est évalué à la position $\mathbf{r}_0$ du noyau de l'atome, et où $\hat{\mathbf{D}}$ est un opérateur atomique (dipôle électrique) ayant des éléments de matrice non nuls entre les états propres de $\hat{H}_0$ considérés.

Écrivant le champ électrique

$$\mathbf{E}(\mathbf{r}_0, t) = \mathbf{E}(\mathbf{r}_0) \cos(\omega t + \varphi(\mathbf{r}_0)) , \quad (2.55)$$

nous mettons l'hamiltonien d'interaction sous la forme rencontrée au chapitre 1 (Éq.1.40 du chapitre 1) :

$$\hat{H}_I(t) = \hat{W} \cos(\omega t + \varphi) , \quad (2.56a)$$

avec

$$\hat{W} = -\hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{r}_0) . \quad (2.56b)$$

Remarque

Comme nous l'avons vu plus haut, dans le cas où $\hat{\mathbf{D}}$ a des éléments de matrice nuls entre les états $|i\rangle$ et $|k\rangle$, il peut y avoir couplage dipolaire magnétique (§ 2.2.5) entre le champ magnétique de l'onde et le moment dipolaire magnétique de l'atome. L'hamiltonien d'interaction se met alors sous la forme (2.56a), mais avec un terme de couplage

$$\hat{W} = -\hat{\mathbf{M}} \cdot \mathbf{B}_0(\mathbf{r}_0) . \quad (2.56c)$$

Les résultats obtenus dans la suite de cette partie se transposent aisément à une telle situation.

Nous pouvons utiliser les résultats obtenus au § I.1.2.4.d dans le cadre d'un calcul perturbatif au premier ordre. On sait que la probabilité de transition n'est notable que pour des situations quasi-résonnantes c'est-à-dire dans le cas où l'énergie de l'état final E_k est égale soit à

$$\implies E_k \simeq E_i + \hbar\omega , \quad (2.57a)$$

soit à

$$\implies E_k \simeq E_i - \hbar\omega, \quad (2.57b)$$

l'égalité étant vérifiée à $\Delta\omega$ près, avec

$$\Delta\omega = \frac{\pi}{T} \quad (2.57c)$$

où T est la durée de l'interaction.

L'atome est donc susceptible de passer du niveau $|i\rangle$ au niveau $|k\rangle$, à condition que la fréquence de l'onde soit résonnante avec la fréquence de Bohr de la transition. Dans le cas (2.57a) où le niveau $|k\rangle$ a son énergie plus élevée que celle du niveau $|i\rangle$, on a *absorption* (§ 2.1.1). Dans le cas opposé (2.57b) on a *émission induite* (§ 2.1.2). Nous traitons maintenant ces phénomènes de façon plus quantitative.

b. Probabilité de transition

Dans le cadre d'un calcul perturbatif au premier ordre, la probabilité de transition de $|i\rangle$ vers $|k\rangle$, après une durée d'interaction T , est donnée par l'équation (I.B.40)

$$P_{i \rightarrow k}(T) = \frac{|W_{ki}|^2}{4\hbar^2} g_T(\hbar\delta), \quad (2.58a)$$

dans laquelle

$$g_T(\hbar\delta) = T^2 \left(\frac{\sin(\delta T/2)}{\delta T/2} \right)^2 \quad (2.58b)$$

est une fonction étroite de δ de hauteur T^2 et de demi-largeur π/T , piquée autour de zéro (Figure 3 du chapitre 1). La quantité

$$W_{ki} = \langle k | \hat{W} | i \rangle \quad (2.58c)$$

est l'élément de matrice permettant la transition, et

$$\delta = \omega - \frac{|E_i - E_k|}{\hbar} \quad (2.58d)$$

est le désaccord du champ électromagnétique par rapport à la résonance atomique (la valeur absolue dans (2.58d) permet de traiter simultanément l'absorption et l'émission induite).

Il est utile d'explicitier l'élément de matrice W_{ki} . Introduisons le vecteur unitaire $\vec{\epsilon}$, parallèle au champ électrique $\mathbf{E}(\mathbf{r}_0)$, qui décrit donc la polarisation linéaire de l'onde :

$$\mathbf{E}(\mathbf{r}_0) = \vec{\epsilon} E(\mathbf{r}_0). \quad (2.59a)$$

L'hamiltonien d'interaction (2.56b) conduit alors à

$$W_{ki} = -\langle k | \hat{\mathbf{D}} \cdot \vec{\epsilon} | i \rangle E(\mathbf{r}_0) = -d E(\mathbf{r}_0), \quad (2.59b)$$

formule dans laquelle on a introduit l'élément de matrice d de la composante de $\hat{\mathbf{D}}$ suivant la polarisation $\vec{\varepsilon}$ du champ électrique. Notons qu'il est toujours possible, en choisissant judicieusement la phase relative des états $|i\rangle$ et $|k\rangle$, de rendre cet élément de matrice réel négatif. Il est alors habituel de poser

$$\Rightarrow W_{ki} = -d E(\mathbf{r}_0) = \hbar\Omega_1, \quad (2.59c)$$

ce qui définit la *pulsation de Rabi* Ω_1 (la justification de ce nom apparaîtra clairement au paragraphe 2.3.2).

Avec ces notations, la probabilité de transition s'écrit

$$\Rightarrow P_{i \rightarrow k}(T) = \left(\frac{\Omega_1 T}{2}\right)^2 \left(\frac{\sin(\delta T/2)}{\delta T/2}\right)^2 = \left(\frac{\Omega_1}{\delta}\right)^2 \sin^2\left(\frac{\delta T}{2}\right). \quad (2.60)$$

Remarques

- (i) Le choix Ω_1 positif a pour simple but d'éviter d'alourdir les expressions. Le lecteur vérifiera sans peine que les résultats établis ici se généralisent de façon naturelle si Ω_1 est complexe. Par exemple, dans (2.60), on remplacera Ω_1^2 par $|\Omega_1|^2$.
- (ii) Dans le cas d'une polarisation circulaire, on peut introduire un vecteur polarisation $\vec{\varepsilon}$ complexe (voir Complément II.1, § 3). L'équation (2.55) est alors remplacée par

$$\mathbf{E}(\mathbf{r}_0, t) = \frac{1}{2} \vec{\varepsilon} E(\mathbf{r}_0) \exp\{-i(\omega t + \varphi(\mathbf{r}_0))\} + c.c. .$$

La suite des calculs se mène sans difficulté, en faisant l'approximation résonnante. On peut en particulier écrire (2.59b) sans modification, et définir la pulsation de Rabi par (2.59c).

L'expression (2.60) montre que la probabilité de transition est proportionnelle au *carré du module de la pulsation de Rabi*, ou encore au carré de l'amplitude du champ électrique. Pour caractériser cette grandeur, nous définirons l'*intensité* $I(\mathbf{r}_0)$ de l'onde électromagnétique au point $\mathbf{r}_0$ comme la *valeur moyenne du carré du champ électrique*

$$I(\mathbf{r}_0) = \frac{1}{\theta} \int_t^{t+\theta} (\mathbf{E}(\mathbf{r}_0, t'))^2 dt' = \frac{\mathbf{E}^2(\mathbf{r}_0)}{2} \quad (2.61a)$$

l'intervalle θ étant long devant la période d'oscillation $2\pi/\omega$. On écrira alors

$$\Omega_1^2 = 2 \left(\frac{d}{\hbar}\right)^2 I(\mathbf{r}_0). \quad (2.61b)$$

Remarques

- (i) Dans le cas d'une onde progressive, l'*intensité* est uniforme, et proportionnelle à la puissance incidente par unité de surface (aussi appelée *éclairement* en optique) qui est égale à la norme de la valeur moyenne Π du vecteur de Poynting :

$$\Pi = \frac{\varepsilon_0 c}{2} |\mathbf{E}_0|^2 = \varepsilon_0 c I. \quad (2.61c)$$

C'est pourquoi on exprime souvent l'intensité en W/m^2 , par un abus de langage qui se révèle commode car faisant référence à une grandeur mesurable expérimentalement. Il est alors entendu que le nombre en question doit être divisé par $\varepsilon_0 c$. Si l'intensité n'est pas uniforme, on pourra continuer à l'exprimer dans la même unité, avec la convention ci-dessus.

- (ii) On trouve dans d'autres ouvrages des définitions différentes de l'intensité. Le point important est sa proportionnalité au carré de la pulsation de Rabi.
- (iii) Lorsque le champ électromagnétique n'est pas monochromatique, on applique la définition (2.61a) en prenant la moyenne temporelle sur une durée θ grande devant les périodes lumineuses, mais courte devant les temps de réponse des instruments d'observation. Pour de la lumière visible, on prendra par exemple θ de l'ordre de 10^{-11} s. On a alors une intensité susceptible de dépendre du temps.
- (iv) Il est intéressant de connaître un ordre de grandeur typique de Ω_1 . Pour un faisceau laser de puissance 1 milliwatt et de section 1 mm^2 , l'éclairement Π vaut 10^3 W/m^2 . Prenant d de l'ordre de $qa_0 (10^{-29} \text{ Cm})$, on trouve Ω_1 de l'ordre de 10^8 s^{-1} . Nous constatons alors qu'à résonance la probabilité de transition atteint des valeurs notables en quelques nanosecondes.

c. Discussion

La formule (2.60) met en évidence plusieurs caractéristiques importantes des processus d'absorption et d'émission induite, bien qu'elle ne soit valable que si la probabilité de transition $P_{i \rightarrow k}(T)$ reste petite devant 1. Dans le cas où le désaccord δ est grand devant la pulsation de Rabi Ω_1 , cette condition est vérifiée quelle que soit la durée d'interaction T , et la probabilité de transition oscille en fonction de T , à la fréquence δ : il s'agit du cas perturbatif de l'oscillation de Rabi, que nous verrons de façon plus exacte au paragraphe 2.3.2, dans le cas d'un atome à 2 niveaux. La valeur maximale $(\Omega_1/\delta)^2$ de la probabilité de transition présente le caractère résonnant autour de $\delta = 0$ déjà mentionné, mais le traitement perturbatif ne nous permet pas de connaître les caractéristiques de l'oscillation de Rabi lorsque δ devient égal ou inférieur à Ω_1 . La formule (2.60) nous montre cependant que le démarrage (quadratique en T) à partir de $T = 0$ présente le même caractère résonnant autour de $\delta = 0$. Nous donnerons au paragraphe 2.3.2 un traitement valable quelle que soit la valeur du rapport Ω_1/δ .

Une autre limitation de ce calcul est la non prise en compte de l'émission spontanée (§ 2.1.3) ou, plus généralement, de la possibilité pour les niveaux atomiques d'être instables, avec une durée de vie Γ^{-1} . Nous admettrons que le résultat (2.60) est qualitativement valable pour des temps d'interaction T plus petits que Γ^{-1} .

Pour des transitions dans le visible, la durée de vie radiative (c'est-à-dire due à l'émission spontanée) des niveaux excités peut valoir de 1 nanoseconde à 1 microseconde. Avec des *lasers*, la pulsation de Rabi Ω_1 atteint couramment des valeurs plus grandes que ces valeurs de Γ (cf. Remarque (iv) du paragraphe 2.3.1.b), et la probabilité de transition (2.60) peut évoluer en un temps assez petit pour respecter la condition ci-dessus. En revanche, si l'onde lumineuse est produite par une *source traditionnelle* (lampe à incandescence, lampe à décharge), la largeur spectrale $\Delta\omega$ de la lumière, qui joue un rôle analogue à Γ , est typiquement plus grande que Γ par plusieurs ordres de grandeurs, alors que les pulsations de Rabi sont plus faibles, au mieux inchangées. Cette situation « d'excitation par spectre large » n'est donc pas correctement décrite par l'équation (2.60), et on n'observe pas l'oscillation de Rabi.

À l'opposé, les transitions radiofréquence peuvent se produire entre des niveaux stables de durées de vie quasiment infinies (sous-niveaux atomiques fondamentaux par exemple, ou niveaux moléculaires non excités électroniquement). L'équation (2.60) décrit alors très bien la situation, tant que la probabilité de transition reste petite devant 1.

En conclusion, le calcul perturbatif a fait apparaître quelques caractéristiques essentielles de la probabilité de transition sous l'effet d'une onde électromagnétique : caractère résonnant, proportionnalité à l'intensité de l'onde, possibilité d'une oscillation de Rabi. Ce calcul reste néanmoins d'une portée relativement limitée, et nous établirons des résultats plus généraux au § 2.3.2 dans le cas non perturbatif pour des niveaux de durée de vie infinie, et dans la partie *D* lorsqu'on prend en compte la durée de vie finie des niveaux.

d. Équivalence des points de vue A.p et D.E

Considérons le cas d'une transition dipolaire électrique. Si au lieu de l'hamiltonien dipolaire électrique (C.3) nous avons utilisé l'hamiltonien $\mathbf{A} \cdot \mathbf{p}$, nous aurions obtenu un résultat analogue en remplaçant $\hat{W}$ par

$$\hat{W}' = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}_\perp(\mathbf{r}_0, t) . \quad (2.62)$$

Nous nous proposons de démontrer qu'à résonance les probabilités de transition sont les mêmes, c'est-à-dire que les éléments de matrice $\hat{W}'_{ki}$ et $\hat{W}_{ki}$ ont le même module.

Rappelons d'abord la relation entre l'amplitude du champ électrique et le potentiel vecteur en jauge de Coulomb (cf. Éq. (2.36a)).

$$\mathbf{A}(\mathbf{r}_0, t) = \frac{\mathbf{E}(\mathbf{r}_0)}{\omega} \cos \left(\omega t + \varphi(\mathbf{r}_0) + \frac{\pi}{2} \right) . \quad (2.63)$$

Il nous faut par ailleurs comparer $\langle k | \hat{\mathbf{D}} | i \rangle$ et $\langle k | \hat{\mathbf{p}} | i \rangle$. Intéressons nous par exemple à la composante suivant *Oz* de $\hat{\mathbf{p}}$. En partant de la forme (2.40b) de l'hamiltonien atomique $\hat{H}_0$, et en utilisant la relation de commutation

$$[\hat{z}, \hat{p}_z] = i\hbar , \quad (2.64)$$

on obtient sans difficulté

$$[\hat{z}, \hat{H}_0] = i\hbar \frac{\hat{p}_z}{m} . \quad (2.65)$$

En projetant l'équation (2.65) à gauche sur $\langle k |$ et à droite sur $|i\rangle$, on trouve

$$\langle k | \hat{z} | i \rangle (E_i - E_k) = \frac{i\hbar}{m} \langle k | \hat{p}_z | i \rangle . \quad (2.66)$$

Finalement, en se souvenant que $\hat{D}_z = q\hat{z}$ pour un atome à un électron, et en se restreignant aux termes résonnants on trouve

$$\left| \frac{W'_{ki}}{W_{ki}} \right| = \left| \frac{\omega_{ki}}{\omega} \right| , \quad (2.67)$$

ce qui prouve l'égalité des probabilités de transition à résonance.

Lorsqu'on n'est pas exactement à résonance, il semble que les probabilités de transition ne sont pas rigoureusement égales. En fait, la différence est du même ordre de grandeur que les termes négligés dans l'approximation résonnante, et elle n'est donc pas significative dans le contexte de ce paragraphe. Un calcul plus précis permet de montrer qu'en fait il y a égalité parfaite entre les probabilités de transition évaluées dans les deux points de vue même quand l'excitation n'est pas exactement résonnante.

2.3.2 Oscillation de Rabi

a. Solution non perturbative de l'équation d'évolution

Nous nous proposons maintenant d'aller au-delà du traitement perturbatif présenté ci-dessus, et de calculer exactement la probabilité de transition d'un niveau atomique $|i\rangle$ vers un niveau atomique $|k\rangle$ sous l'effet d'une onde quasi-résonnante, de pulsation ω proche de la pulsation de Bohr $|E_k - E_i|/\hbar$, dans le cas où cette probabilité n'est pas petite devant 1. Dans cette situation résonnante, nous ferons *l'approximation de l'atome à deux niveaux*, et nous noterons respectivement $|a\rangle$ et $|b\rangle$ les états d'énergie inférieur et supérieur. Nous prendrons nulle l'énergie E_a de l'état inférieur, et nous noterons ω_0 la pulsation de Bohr

$$E_b - E_a = \hbar\omega_0, \quad (2.68)$$

de sorte que la restriction de l'hamiltonien atomique au sous-espace $\{|a\rangle, |b\rangle\}$ s'écrit

$$\hat{H}_0 = \hbar \begin{pmatrix} 0 & 0 \\ 0 & \omega_0 \end{pmatrix}. \quad (2.69)$$

Nous prenons l'hamiltonien d'interaction dipolaire électrique (C.3), et nous écrivons l'élément de matrice W_{ba} (cf. Éq. 2.59b) sous la forme

$$W_{ba} = -\langle b | \hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{r}_0) | a \rangle = \hbar\Omega_1 \quad (2.70)$$

qui définit la pulsation de Rabi Ω_1 . (Rappelons que les termes de phase arbitraires sur $|a\rangle$ et $|b\rangle$ ont été choisis de sorte à rendre Ω_1 réel et positif.) L'hamiltonien atomique total est donc

$$\hat{H} = \hat{H}_0 + \hat{H}_1 = \hbar \begin{pmatrix} 0 & \Omega_1 \cos(\omega t + \varphi) \\ \Omega_1 \cos(\omega t + \varphi) & \omega_0 \end{pmatrix}. \quad (2.71)$$

Pour décrire l'évolution de l'atome, nous développons comme au chapitre 1 (partie 1.2) son état sur la base $\{|a\rangle, |b\rangle\}$ sous la forme

$$|\psi(t)\rangle = \gamma_a(t)|a\rangle + \gamma_b(t)e^{-i\omega_0 t}|b\rangle \quad (2.72)$$

qui permet, grâce au facteur $e^{-i\omega_0 t}$, de séparer l'évolution libre. L'équation de Schrödinger conduit alors aux équations d'évolution

$$i \frac{d}{dt} \gamma_a = \frac{\Omega_1 e^{i\varphi}}{2} e^{i(\omega - \omega_0)t} \gamma_b + \frac{\Omega_1 e^{-i\varphi}}{2} e^{-i(\omega_0 + \omega)t} \gamma_b, \quad (2.73a)$$

$$i \frac{d}{dt} \gamma_b = \frac{\Omega_1 e^{-i\varphi}}{2} e^{-i(\omega - \omega_0)t} \gamma_a + \frac{\Omega_1 e^{i\varphi}}{2} e^{i(\omega_0 + \omega)t} \gamma_a. \quad (2.73b)$$

Ce système se simplifie considérablement si on fait *l'approximation résonnante* : à condition que $|\omega - \omega_0|$ soit petit devant $|\omega + \omega_0|$, les termes en $\exp(\pm i(\omega + \omega_0)t)$ qui oscillent

très vite donnent des contributions négligeables, et on peut les supprimer comme nous l'avons déjà vu au paragraphe 1.2.4.d dans le cas perturbatif. On obtient alors

$$i \frac{d}{dt} \gamma_a = \frac{\Omega_1 e^{i\varphi}}{2} e^{i(\omega - \omega_0)t} \gamma_b, \quad (2.74a)$$

$$i \frac{d}{dt} \gamma_b = \frac{\Omega_1 e^{-i\varphi}}{2} e^{-i(\omega - \omega_0)t} \gamma_a, \quad (2.74b)$$

pour l'évolution des coefficients introduits en (2.72).

Pour résoudre ce système, introduisons le désaccord à résonance $\delta = \omega - \omega_0$ (cf. Éq. 2.58d), et effectuons le changement de variable

$$\gamma_a = \tilde{\gamma}_a \exp\left(i \frac{\delta}{2} t\right), \quad (2.75a)$$

$$\gamma_b = \tilde{\gamma}_b \exp\left(-i \frac{\delta}{2} t\right). \quad (2.75b)$$

Le système (C.21) se transforme en un système d'équations couplées à coefficients constants

$$i \frac{d}{dt} \tilde{\gamma}_a = \frac{\delta}{2} \tilde{\gamma}_a + \frac{\Omega_1 e^{i\varphi}}{2} \tilde{\gamma}_b, \quad (2.76a)$$

$$i \frac{d}{dt} \tilde{\gamma}_b = \frac{\Omega_1 e^{-i\varphi}}{2} \tilde{\gamma}_a - \frac{\delta}{2} \tilde{\gamma}_b. \quad (2.76b)$$

Un tel système admet deux solutions propres oscillantes, de la forme

$$\begin{bmatrix} \tilde{\gamma}_a \\ \tilde{\gamma}_b \end{bmatrix} = \begin{bmatrix} \alpha \\ \beta \end{bmatrix} \exp\left(-i \frac{\lambda}{2} t\right), \quad (2.77a)$$

λ prenant l'une des deux valeurs

$$\lambda_{\pm} = \pm \Omega = \pm \sqrt{\Omega_1^2 + \delta^2}, \quad (2.77b)$$

auxquelles correspondent deux solutions pour le rapport $\frac{\alpha}{\beta}$

$$\left(\frac{\alpha}{\beta}\right)_{\pm} = -\frac{\Omega_1 e^{i\varphi}}{\delta \mp \Omega}. \quad (2.77c)$$

La grandeur Ω définie en (2.77b) s'appelle *pulsation de Rabi généralisée*.

La solution générale du système (C.23) est donc de la forme

$$\tilde{\gamma}_a = K \frac{\Omega_1 e^{i\varphi}}{\Omega - \delta} \exp\left(-i \frac{\Omega}{2} t\right) + L \frac{\Omega_1 e^{i\varphi}}{\Omega + \delta} \exp\left(i \frac{\Omega}{2} t\right). \quad (2.78)$$

Cherchons la solution qui correspond aux conditions initiales

$$\tilde{\gamma}_a(t_0) = 1 , \quad (2.79a)$$

$$\tilde{\gamma}_b(t_0) = 0 , \quad (2.79b)$$

c'est-à-dire à un atome dans l'état inférieur $|a\rangle$ à l'instant $t = t_0$. La solution de (C.23) est dans ce cas :

$$\tilde{\gamma}_a(t) = \cos \frac{\Omega}{2}(t - t_0) - i \frac{\delta}{\Omega} \sin \frac{\Omega}{2}(t - t_0) , \quad (2.80a)$$

$$\tilde{\gamma}_b(t) = -i \frac{\Omega_1 e^{-i\varphi}}{\Omega} \sin \frac{\Omega}{2}(t - t_0) . \quad (2.80b)$$

En remplaçant Ω par sa valeur $\sqrt{\Omega_1^2 + \delta^2}$ définie en (2.77b), les formules (2.80) fournissent directement les probabilités de présence dans les états $|a\rangle$ ou $|b\rangle$ de l'atome. Elles permettent également, en revenant aux coefficients $\gamma_a(t)$ et $\gamma_b(t)$ (Éq. C.22), de calculer la *valeur moyenne de n'importe quelle observable atomique*. C'est ainsi qu'au paragraphe 2.4.3 nous calculerons la valeur moyenne $\langle \hat{\mathbf{D}} \rangle(t)$ du *dipole atomique* (la notation rappelle que la valeur moyenne quantique $\langle \hat{\mathbf{D}} \rangle = \langle \psi(t) | \hat{\mathbf{D}} | \psi(t) \rangle$ évolue avec le temps, cf. (2.87)).

Remarque

La solution (C.27) du système (C.23) est la solution particulière associée à la condition initiale (C.26). Pour chaque condition initiale, il faut recalculer la solution correspondante. Il existe des situations où cette solution est très sensible aux conditions initiales, par exemple dans ce que l'on appelle la méthode des champs séparés de Ramsey qui est utilisée en spectroscopie de haute résolution, et en particulier dans les horloges atomiques à Césium.

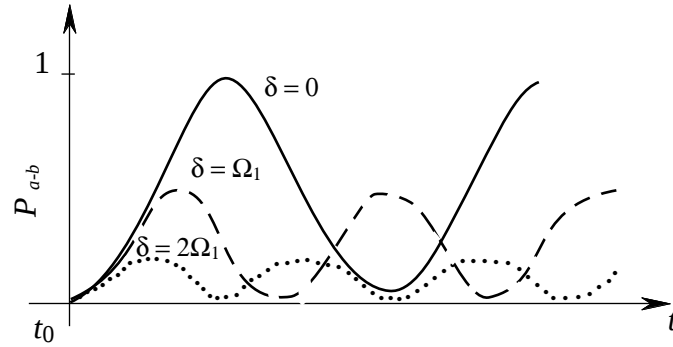


FIG. 2.7: **Oscillation de Rabi** dans le cas exactement résonnant (trait plein), et pour un désaccord $\delta = \Omega_1$ (tirets) et $\delta = 2\Omega_1$ (pointillés).

b. Oscillation de Rabi

La solution (C.27) nous permet de calculer la probabilité $P_{a \rightarrow b}(t_0, t)$ qu'un atome initialement dans l'état $|a\rangle$ soit passé à l'instant t dans l'état $|b\rangle$ sous l'effet de l'onde

électromagnétique quasi-résonnante :

$$\Rightarrow P_{a \rightarrow b}(t_0, t) = |\tilde{\gamma}_b(t)|^2 = \frac{\Omega_1^2}{\Omega_1^2 + \delta^2} \sin^2 \frac{\Omega}{2}(t - t_0) . \quad (2.81)$$

Aux notations près, cette expression est identique à celle de l'équation (1.59), établie dans le cas d'une interaction constante, et non pas oscillante comme ici.

Comme on le voit sur la figure 2.7, la probabilité de transition oscille à la pulsation de Rabi généralisée $\Omega = \sqrt{\Omega_1^2 + \delta^2}$ (Éq. 2.77b), entre la valeur 0 et une valeur maximale

$$P_{a \rightarrow b}^{\max} = \frac{\Omega_1^2}{\Omega_1^2 + \delta^2} . \quad (2.82)$$

La probabilité maximale de transition $P_{a \rightarrow b}^{\max}$ a un comportement résonnant lorsque la fréquence d'excitation ω varie autour de la fréquence de Bohr ω_0 . La loi de variation correspondante (Fig. 2.8) est une loi Lorentzienne, de largeur à mi-hauteur $2\Omega_1$. On constate que la largeur de la courbe de résonance est proportionnelle à l'amplitude du champ électrique.

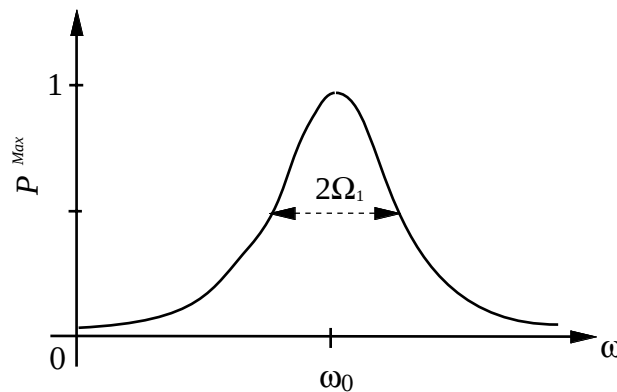


FIG. 2.8: **Probabilité maximale** $P_{a \rightarrow b}^{\max}$ de trouver l'atome dans l'état supérieur de la transition au cours de l'oscillation de Rabi, en fonction du désaccord à résonance. On a une variation résonnante de type Lorentzien, de largeur proportionnelle à la racine carrée de l'intensité de l'onde. À résonance, il est possible de transférer la totalité de la population du niveau $|a\rangle$ vers le niveau $|b\rangle$.

Notons que si on se place à résonance, la valeur maximale $P_{a \rightarrow b}^{\max}$ vaut 1. Il est donc possible de *transférer la totalité des atomes* du niveau $|a\rangle$ vers le niveau $|b\rangle$, en choisissant une durée d'interaction égale à π/Ω_1 (cf. figure 2.7). Une impulsion ayant cette durée s'appelle *impulsion* π .

c. Exemples

L'oscillation de Rabi est-elle observable dans le domaine des ondes lumineuses ? Généralisant la discussion du paragraphe C.1.c, nous admettrons que le calcul ci-dessus n'est valable que si la durée d'interaction considérée est plus courte que la durée de vie des

niveaux $|a\rangle$ et $|b\rangle$. Par ailleurs, pour observer l'oscillation de Rabi, il faut une durée d'interaction au moins égale à une période d'oscillation. Les deux conditions sont compatibles à condition d'avoir une pulsation de Rabi assez grande, de l'ordre de 10^9s^{-1} , ce qui s'obtient avec un faisceau laser de quelques dizaines de milliwatts concentré sur un millimètre carré. L'observation peut alors se faire en éclairant les atomes avec une impulsion lumineuse de durée $t - t_0$ contrôlée dans la gamme de quelques nanosecondes, et en observant la fluorescence qui suit, proportionnelle au nombre d'atomes qui ont été portés dans l'état excité. Une telle expérience est possible, mais elle présente de sérieuses difficultés techniques, et l'oscillation de Rabi n'est observable dans le domaine des ondes lumineuses que si on s'est placé dans des conditions particulièrement favorables.

Dans le domaine des radiofréquences, nous savons (§ 2.2.5) qu'il est possible de provoquer des transitions entre sous-niveaux atomiques fondamentaux, de durée de vie infinie, par couplage dipolaire magnétique. C'est par exemple le cas de la *transition d'horloge* à 9, 192 631 770 GHz entre les niveaux hyperfins $F = 3$ et $F = 4$ de l'état fondamental $6^2\text{S}_{1/2}$ de l'atome de *Césium*. Avec une intensité de quelques watts par centimètre carré, on peut produire des fréquences de Rabi $\Omega_1/2\pi$ de l'ordre d'une centaine de kilohertz. Comme on sait mesurer la population atomique dans chaque niveau hyperfin, on a pu vérifier ainsi toutes les caractéristiques de l'oscillation de Rabi présentées au paragraphe précédent, notamment le caractère résonnant. Notons que l'on utilise couramment des impulsions π pour transférer tous les atomes d'un niveau hyperfin à l'autre.

Un autre exemple très important est celui d'un spin $1/2$ plongé dans un champ magnétique statique. Il s'agit ici encore d'un système à deux niveaux, dont l'écart énergétique est proportionnel au champ magnétique appliqué. Ce système pourra interagir avec un champ électromagnétique par couplage dipolaire magnétique. Dans la *résonance magnétique nucléaire*, il s'agit le plus souvent des spins $1/2$ des noyaux d'hydrogène qui, plongés dans un champ magnétique intense de quelques teslas, résonnent lorsqu'ils sont soumis à une onde radiofréquence à quelques centaines de MHz. La valeur précise de la fréquence de résonance donne des renseignements sur *l'environnement* des noyaux d'hydrogène. Ce phénomène a des applications très importantes en *analyse chimique*, ou en *imagerie médicale*.

d. Impulsion $\pi/2$. Transitoires cohérents

Reprenons la situation étudiée au paragraphe 2.3.2.a, dans le cas d'une excitation exactement résonnante ($\delta = 0$), et supposons que l'on cesse d'appliquer le champ électromagnétique à l'instant

$$t_1 = t_0 + \frac{\pi}{2\Omega_1} . \quad (2.83)$$

Ce créneau de champ s'appelle une *impulsion* $\pi/2$.

Nous nous intéressons à l'évolution ultérieure de l'atome. Son état à l'instant t_1 est

alors donné par (C.27) :

$$\gamma_a(t_1) = \tilde{\gamma}_a(t_1) = \frac{1}{\sqrt{2}} , \quad (2.84a)$$

$$\gamma_b(t_1) = \tilde{\gamma}_b(t_1) = -\frac{i}{\sqrt{2}} e^{-i\varphi} . \quad (2.84b)$$

D'après (2.72), l'évolution libre ultérieure est décrite par

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}}|a\rangle - \frac{i}{\sqrt{2}}e^{-i(\omega_0 t + \varphi)}|b\rangle \quad (2.85)$$

ce qui implique que la probabilité de présence dans chacun des états $|a\rangle$ ou $|b\rangle$ n'évolue plus.

Il ne faut pourtant pas en conclure que le système est figé. Pour nous en convaincre, choisissons une observable ayant des éléments de matrice entre $|a\rangle$ et $|b\rangle$. Par exemple, dans le cas d'une transition dipolaire électrique où $\hat{\mathbf{D}}$ a des éléments de matrice non nuls, considérons la composante $\hat{D}_\varepsilon = \hat{\mathbf{D}} \cdot \boldsymbol{\varepsilon}$ de l'opérateur dipolaire électrique, représentée dans la base $\{|a\rangle, |b\rangle\}$ par

$$\hat{D}_\varepsilon = \begin{pmatrix} 0 & d \\ d & 0 \end{pmatrix} , \quad (2.86)$$

et calculons la valeur moyenne de $\hat{D}_\varepsilon$ dans l'état (2.85). On obtient

$$\langle \hat{D}_\varepsilon \rangle(t) = \langle \psi(t) | \hat{D}_\varepsilon | \psi(t) \rangle = -d \sin(\omega_0 t + \varphi) . \quad (2.87)$$

On constate donc que le dipôle oscille à la fréquence de Bohr ω_0 . Dans le cas d'une transition optique, ceci se traduit par une émission de lumière à cette fréquence.

Bien qu'émise à la fréquence ω_0 , cette lumière a des propriétés différentes de la lumière de fluorescence produite par émission spontanée. Ces propriétés particulières sont liées à la cohérence de l'émission. La formule (2.87) montre en effet que l'oscillation du dipôle se produit avec une phase reliée de façon univoque à celle de l'onde excitatrice. Si on a une assemblée d'atomes préparés par la même impulsion $\pi/2$, ils vont donc tous émettre *avec la même phase*, à la différence d'une assemblée d'atomes émettant par émission spontanée avec des phases aléatoires. Il est possible d'observer les conséquences expérimentales de ces propriétés de cohérence : directivité de l'émission ; apparition d'un battement entre la lumière de fluorescence et un faisceau cohérent avec le laser d'excitation. De tels comportements, appelés *transitoires cohérents*, s'observent sur une échelle de temps plus courte que l'émission spontanée⁸.

⁸Ces phénomènes ont donné lieu à de très belles expériences. Voir par exemple : R.G. Brewer, page 341 in « Aux frontières de la spectroscopie laser » (Les Houches, Session XXVII) édité par R. Balian, S. Haroche, et S. Liberman, North-Holland (1977).

Remarque

L'énergie emportée par le rayonnement est évidemment prélevée sur l'énergie atomique, de sorte que la probabilité de présence dans le niveau supérieur décroît au cours du temps, contrairement à ce que pourrait laisser penser la formule (2.85). La raison est que la formule (2.85) a été établie en négligeant toute interaction entre les atomes et le champ qu'ils rayonnent. C'est en fait la réaction du champ émis sur les atomes émetteurs qui entraîne une décroissance de $|\gamma_b(t)|$ avec le temps.

Dans le domaine des radiofréquences, et pour des transitions entre niveaux de très longues durées de vie, les impulsions $\pi/2$ sont couramment utilisées pour mettre un système en oscillation libre (à la pulsation ω_0), avec une phase contrôlée par l'onde excitatrice. Les méthodes de *résonance magnétique nucléaire en impulsion* en font un très large usage. C'est aussi une technique de base pour les horloges atomiques, et plus généralement pour de nombreuses méthodes utilisant des franges de Ramsey.

Remarque

Il ne faut pas croire qu'il soit essentiel dans ces expériences de réaliser une impulsion dont la durée soit très exactement celle d'une impulsion $\pi/2$. Le point important est de préparer l'atome dans un état superposition linéaire de $|a\rangle$ et $|b\rangle$ avec des poids du même ordre, de sorte que l'oscillation du dipôle (2.87) ait une grande amplitude.

2.3.3 Transitions multiphotoniques*a. Traitement perturbatif*

Considérons la situation représentée sur la figure 2.9, dans laquelle un atome, initialement dans un état $|i\rangle$ est soumis à une onde électromagnétique de fréquence ω voisine de la moitié de la fréquence de Bohr atomique associée à la transition entre $|i\rangle$ et $|k\rangle$: $\omega_{ki} = |E_k - E_i|/\hbar$.

Comme aucun niveau atomique n'a une énergie voisine de $E_i + \hbar\omega$ un calcul perturbatif au premier ordre donne une probabilité très faible que l'atome effectue une transition vers un autre niveau. En revanche, en poussant le calcul *perturbatif au deuxième ordre* (§ B.5 du chapitre 1 généralisé au cas d'une perturbation sinusoïdale), on trouve, à la limite quasi-résonnante, une probabilité de transition après un temps d'interaction T :

$$P_{i \rightarrow k}(T) = T \frac{2\pi}{\hbar} \left| \frac{1}{4} \sum_{j \neq i} \frac{W_{kj} W_{ji}}{E_i - E_j + \hbar\omega} \right|^2 \delta_T(E_k - E_i - 2\hbar\omega) . \quad (2.88)$$

Ici encore, on voit apparaître la fonction $\delta_T(E)$ piquée autour de 0, de demi-largeur $\pi\hbar/T$ et de hauteur $T/2\pi\hbar$ (de surface unité). Elle traduit la condition de résonance :

$$E_k - E_i = \hbar\omega_{ki} = 2\hbar\omega . \quad (2.89)$$

Comme le suggère la figure 2.9, un tel processus peut s'interpréter comme une *absorption à deux photons* portant l'atome de l'état $|i\rangle$ à l'état $|k\rangle$. Si l'atome avait été initialement dans l'état $|k\rangle$, il aurait évidemment pu passer dans l'état $|i\rangle$ d'énergie inférieure, sous l'effet de la même onde : ce processus inverse est l'*émission stimulée à deux photons*.

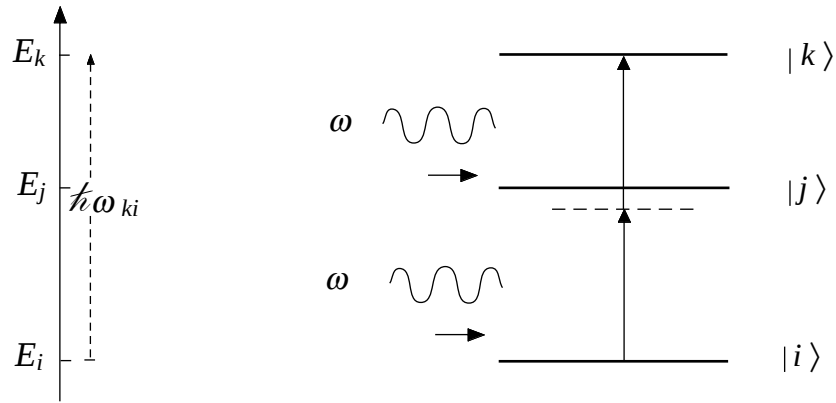


FIG. 2.9: **Transition à deux photons** de $|i\rangle$ vers $|k\rangle$. La fréquence de l'onde ω est voisine de $\omega_{ki}/2$, mais notablement différente de ω_{ji} . Le niveau $|j\rangle$ s'appelle un niveau « relais » pour le processus à deux photons.

Le processus décrit ici est un processus non-linéaire d'ordre 2. En effet, chacun des éléments de matrice W_{kj} et W_{ji} est proportionnel au champ électrique de l'onde lumineuse (Équation 2.59b), et la probabilité de transition (2.88) est proportionnelle au *carré de l'intensité lumineuse*.

On peut naturellement généraliser le résultat ci-dessus en considérant des ordres plus élevés de la théorie des perturbations. Par exemple, si $\omega = |E_k - E_i|/3\hbar$, on pourra avoir une transition à 3 photons (cf. Figure 2.5), de probabilité proportionnelle au cube de l'intensité lumineuse.

Remarque

Dans un atome, un élément de matrice dipolaire électrique n'est différent de zéro qu'entre deux niveaux de parités opposées (cf. Éq. 3 du complément II.1). Une transition dipolaire électrique à deux photons n'est donc possible qu'entre niveaux de même parité, pour lesquels une transition directe à un photon est impossible.

b. Ordre de grandeur

Il est facile de comparer la probabilité de transition à deux photons quasi-résonnante (2.88) à la probabilité de transition à un photon (2.58a). On constate (en prenant $|W_{ki}| \simeq |W_{kj}| \simeq |W_{ji}|$) qu'à résonance le rapport des deux quantités est de l'ordre de

$$\frac{W_{kj}W_{ji}}{(E_i - E_j + \hbar\omega)^2}, \quad (2.90)$$

c'est-à-dire le produit des pulsations de Rabi des transitions à un photon, divisé par le carré du désaccord dans l'état intermédiaire. À l'exception de cas très particuliers où le niveau relais $|j\rangle$ est au milieu des niveaux $|i\rangle$ et $|k\rangle$, le désaccord dans l'état intermédiaire est beaucoup plus grand que la pulsation de Rabi de chaque transition à un photon. Les transitions à plusieurs photons sont donc en général beaucoup moins probables que les transitions à un photon. Elles n'ont une probabilité notable que si les intensités lumineuses sont suffisamment grandes, ce qui explique qu'on n'observe facilement ces effets qu'en

utilisant des lasers très focalisés, ou des lasers en impulsion⁹.

Remarque

Le calcul ci-dessus n'a un sens que pour des durées d'interaction petites devant la durée de vie des niveaux $|i\rangle$ et $|k\rangle$. Il faut également que le désaccord dans l'état intermédiaire $|\omega - (E_i - E_j)/\hbar|$ soit grand devant l'inverse de la durée de vie de l'état intermédiaire $|j\rangle$, et devant la pulsation de Rabi $\hat{W}_{ij}/\hbar$, car sinon la transition à un photon de $|i\rangle$ vers $|j\rangle$ serait pratiquement résonnante, et le problème ne pourrait pas être traité aussi simplement.

c. Hamiltonien effectif d'une transition à deux photons

Lorsque l'intensité de l'onde est suffisamment grande, et que la transition à deux photons est quasi-résonnante, on peut donner un traitement plus exact de l'évolution du système. Pour rendre les notations plus claires, nous noterons $|a\rangle$ et $|c\rangle$ les deux niveaux connectés par la transition à deux photons ($\omega_0 = (E_c - E_a)/\hbar$ voisin de 2ω), et $|j\rangle$ les niveaux relais, dont nous supposons qu'aucun n'est résonnant. Nous développons $|\psi(t)\rangle$ sur $|a\rangle, |c\rangle$ et sur les états $|j\rangle$, de façon analogue à (2.72), en introduisant des coefficients $\gamma_a(t), \gamma_c(t), \gamma_j(t)$. L'équation de Schrödinger s'écrit alors sous une forme analogue à (C.20). On peut ainsi calculer les $\gamma_j(t)$ sous forme perturbative, aucun niveau j n'étant résonnant. En supposant que l'onde est branchée lentement, on se débarrasse des constantes d'intégration (cf. § 1.2.5 du chapitre 1), et on trouve

$$\gamma_j(t) = -\frac{W_{ja}}{2(E_j - E_a - \hbar\omega)}\gamma_a(t)e^{-i(E_j - E_a - \hbar\omega)t/\hbar} - \frac{W_{ja}}{2(E_j - E_a + \hbar\omega)}\gamma_a(t)e^{-i(E_j - E_a + \hbar\omega)t/\hbar} \quad (2.91)$$

+ {termes analogues avec $a \rightarrow c$ } .

Si on reporte ces expressions dans les équations d'évolution de $\gamma_a(t)$ et $\gamma_c(t)$, on trouve que tout se passe comme si on avait un système à deux niveaux $\{|a\rangle, |c\rangle\}$, soumis à un *hamiltonien effectif*¹⁰ d'éléments de matrice non-diagonaux :

$$W_{ca}^{\text{eff}} = -\sum_{j \neq a, c} \frac{W_{cj}W_{ja}}{4(E_c - E_j - \hbar\omega)}e^{-i(2\omega - \omega_0)t} = \frac{\hbar\Omega_1^{\text{eff}}}{2}e^{-i(2\omega - \omega_0)t}, \quad (2.92a)$$

et d'éléments de matrice diagonaux :

$$W_{cc}^{\text{eff}} = \frac{1}{4} \left\{ \sum_{j \neq c} \frac{|W_{cj}|^2}{E_c - E_j - \hbar\omega} + \sum_{j \neq c} \frac{|W_{cj}|^2}{E_c - E_j + \hbar\omega} \right\}, \quad (2.92b)$$

$$W_{aa}^{\text{eff}} = \frac{1}{4} \left\{ \sum_{j \neq a} \frac{|W_{aj}|^2}{E_a - E_j - \hbar\omega} + \sum_{j \neq a} \frac{|W_{aj}|^2}{E_a - E_j + \hbar\omega} \right\}. \quad (2.92c)$$

Il suffit alors de résoudre le système d'équations différentielles couplées du premier ordre, par exemple en suivant une démarche analogue à celle du § 2.3.2, pour obtenir

$$P_{a \rightarrow c}(t_0, t) = \left(\frac{\Omega_1^{\text{eff}}}{\Omega^{\text{eff}}} \right)^2 \sin^2 \left(\frac{\Omega^{\text{eff}}}{2}(t - t_0) \right), \quad (2.93a)$$

⁹Voir G. Grynberg, B. Cagnac et F. Biraben (1980) : Coherent Nonlinear Optics, p. 111 (édité par M.S. Feld et V.S. Letokhov), Springer Verlag (Berlin 1980).

¹⁰La dénomination « Hamiltonien effectif » s'utilise dans une très grande variété de situations, et pas seulement pour le problème traité ici. Sa signification peut donc différer suivant le contexte.

formule dans laquelle la pulsation de Rabi généralisée effective Ω^{eff} est définie par

$$(\Omega^{\text{eff}})^2 = (\Omega_1^{\text{eff}})^2 + \left(\frac{E_c + W_{cc}^{\text{eff}} - E_a - W_{aa}^{\text{eff}}}{\hbar} - 2\omega \right)^2. \quad (2.93b)$$

On trouve donc une oscillation de Rabi entre les niveaux $|a\rangle$ et $|c\rangle$, sous l'influence du couplage quasi-résonnant à deux photons. L'ensemble des commentaires faits au § 2.3.2 s'applique ici directement. Mais on peut faire quelques remarques supplémentaires.

Notons tout d'abord (Éq. (2.93a)) que l'atome peut passer totalement dans l'état $|c\rangle$, alors que les coefficients $\gamma_j(t)$ ont tous un module petit devant 1, et donc que la probabilité de trouver l'atome dans un *niveau relais* reste donc toujours négligeable. Ceci est naturellement dû au fait que la transition à deux photons entre $|a\rangle$ et $|c\rangle$ est résonnante, alors que les transitions à un photon vers les niveaux relais $|j\rangle$ sont non-résonnantes.

Une autre nouveauté du résultat (C.40) est le rôle des termes diagonaux du hamiltonien effectif, qui donnent lieu aux déplacements lumineux. Ce phénomène important est étudié un peu plus en détail dans le prochain paragraphe.

2.3.4 Déplacements lumineux

Les formules (C.40) montrent que la transition à deux photons est résonnante pour

$$2\omega = \frac{E_c + W_{cc}^{\text{eff}} - E_a - W_{aa}^{\text{eff}}}{\hbar}. \quad (2.94)$$

Tout se passe comme si les deux niveaux $|a\rangle$ et $|c\rangle$ avaient été déplacés respectivement de W_{aa}^{eff} et W_{cc}^{eff} . En revenant aux expressions (C.39), on constate que ces déplacements sont proportionnels au carré du couplage entre $|a\rangle$ (ou $|c\rangle$) et les niveaux relais $|j\rangle$, sous l'effet de l'onde. On les appelle *déplacements lumineux*¹¹ (*light-shift* en anglais, ou encore effet Stark dynamique). Ces déplacements sont proportionnels à l'intensité de l'onde, et inversement proportionnels au désaccord par rapport au niveau relais, le sens du déplacement étant indiqué sur la figure 2.10.

Les déplacements lumineux n'apparaissent pas seulement dans le cas des transitions à deux photons. Il s'agit d'un phénomène tout à fait général, que nous allons retrouver dans le cas d'une transition dipolaire électrique entre deux niveaux atomiques. Repartant de (C.20), on peut dans le cas non-résonnant effectuer une résolution perturbative analogue à celle présentée au paragraphe 2.3.3.c ci-dessus, et obtenir (on prend $\varphi = 0$)

$$\gamma_b(t) = \frac{\Omega_1}{2} \left\{ \frac{e^{-i(\omega-\omega_0)t}}{\omega - \omega_0} - \frac{e^{i(\omega+\omega_0)t}}{\omega + \omega_0} \right\} \gamma_a(t). \quad (2.95)$$

¹¹Les déplacements lumineux ont été prédits et observés pour la première fois par C. Cohen Tannoudji. En utilisant des lampes spectrales et les techniques du pompage optique, il a pu mettre en évidence un déplacement de 1Hz seulement (C. Cohen-Tannoudji et A. Kastler (1966) : Progress in Optics, Vol. V, p. 1 (édité par E. Wolf), North Holland (Amsterdam 1966)). De nos jours, avec des lasers, on observe des déplacements supérieurs au mégahertz, voire au gigahertz.

Lorsqu'on reporte dans (2.73a), on obtient quatre termes, dont deux non-oscillants (termes séculaires) qui donnent donc la contribution dominante. En ne gardant que le plus grand, on trouve

$$i \frac{d\gamma_a}{dt} = \frac{\Omega_1^2}{4} \frac{1}{\omega - \omega_0} \gamma_a(t) . \quad (2.96)$$

Tout se passe comme si l'énergie du niveau $|a\rangle$, qui valait initialement $E_a = 0$, avait été déplacée de la quantité

$$W_{aa}^{\text{eff}} = \hbar \frac{\Omega_1^2}{4(\omega - \omega_0)} = \frac{|W_{ab}|^2}{4(\hbar\omega - E_b + E_a)} , \quad (2.97a)$$

qui est analogue à (2.92c) en ne gardant que le terme dominant pour le cas $j = b$. Le déplacement lumineux du niveau $|b\rangle$ s'obtient de façon analogue. Il vaut

$$W_{bb}^{\text{eff}} = -\hbar \frac{\Omega_1^2}{4(\omega - \omega_0)} . \quad (2.97b)$$

On retrouve donc que les niveaux $|a\rangle$ et $|b\rangle$ s'éloignent l'un de l'autre pour un désaccord négatif, et se rapprochent pour un désaccord positif (Fig. 2.10).

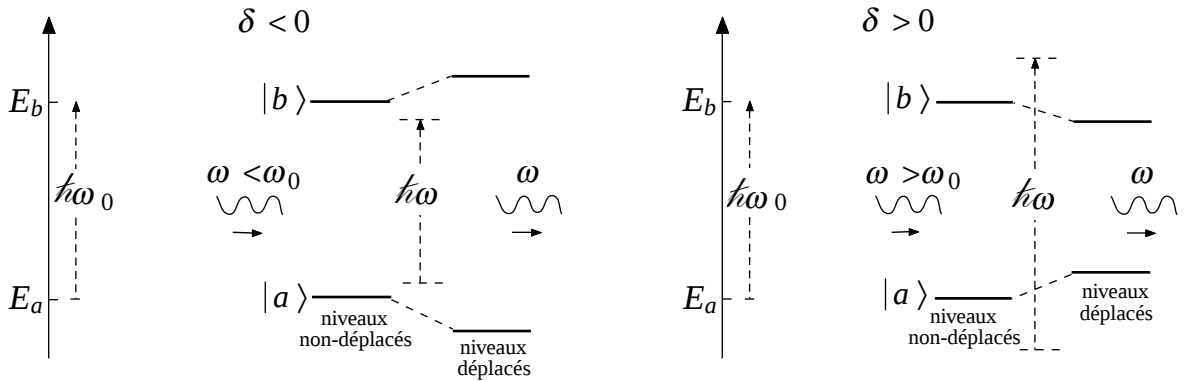


FIG. 2.10: **Déplacements lumineux** dans un atome à deux niveaux couplés par un laser de fréquence ω . Le signe des déplacements dépend du signe du désaccord

Depuis l'avènement des lasers, les déplacements lumineux jouent un rôle très important dans tout le domaine de l'interaction atome-laser¹¹. C'est ainsi qu'il est indispensable d'en tenir compte dans les expériences de spectroscopie de haute résolution (cf. Complément III.5). Ils sont également à la base de plusieurs mécanismes de refroidissement d'atomes par laser¹².

¹²A. Aspect et J. Dalibard, La Recherche (Janvier 1994). C. Cohen-Tannoudji in « Systèmes fondamentaux en optique quantique », édité par J. Dalibard, J.M. Raimond, et J. Zinn-Justin, North-Holland (1992). (Actes de l'école d'été Les Houches, 1990).

Remarques

(i) Les formules ci-dessus, établies dans le cas de niveaux de durée de vie infinie, restent valables tant que le désaccord δ est grand devant la largeur Γ des niveaux couplés. Le résultat (C.44) a donc en fait un domaine d'application très vaste.

(ii) Pour observer les déplacements lumineux des niveaux a ou b de la figure 2.10, il faut disposer d'un deuxième faisceau laser (laser sonde), de fréquence variable et d'intensité plus faible que celle du laser quasi-résonnant sur la transition $a \rightarrow b$, de sorte que l'on puisse négliger les déplacements lumineux dus à ce deuxième faisceau. On peut alors mesurer l'absorption de ce faisceau sonde, par exemple au voisinage d'une transition vers un troisième niveau c , soit depuis a , soit depuis b , et constater que la fréquence de résonance dépend de l'intensité et de la fréquence du premier laser suivant les relations (C.44).

2.4 Absorption entre deux niveaux de durée de vie finie

2.4.1 Présentation du modèle utilisé

Dans la partie 2.3, nous avons vu comment un atome (ou une molécule), soumis à un champ électromagnétique quasi-résonnant entre deux niveaux d'énergie, est susceptible de passer d'un niveau à l'autre par absorption ou émission induite. Nous avons indiqué que notre traitement n'était valable que dans la mesure où les niveaux considérés avaient une durée de vie très longue.

En fait, on a très souvent à considérer des transitions entre niveaux dont l'un au moins a une durée de vie courte. Par exemple, dans le domaine optique, le niveau supérieur de la transition se désexcite par émission spontanée. Il peut en être de même du niveau inférieur, s'il ne s'agit pas d'un niveau fondamental. Bien d'autres processus peuvent être à l'origine d'une durée de vie finie. C'est le cas des collisions avec d'autres atomes, ou avec les parois de la cellule contenant une vapeur atomique, ou avec des phonons pour un ion dans une matrice solide. Signalons aussi que lorsque le mouvement des atomes les fait sortir du volume d'interaction avec l'onde électromagnétique, tout se passe comme si les atomes eux-mêmes avaient une durée de vie finie et disparaissaient.

Nous souhaitons donc prendre en compte la durée de vie des niveaux atomiques. Malheureusement, le traitement général de l'interaction entre une onde électromagnétique et un atome ayant des niveaux de durée de vie finie nécessite des outils plus sophistiqués que ceux dont nous disposons dans ce chapitre. Il faut en effet faire appel au formalisme de la matrice densité, et utiliser les *équations de Bloch optiques* (cf. Complément II.2).

Le problème est que les processus qui rendent un niveau d'énergie instable (émission spontanée, ou collisions avec d'autres atomes) provoquent un couplage avec l'extérieur du système « atome-onde électromagnétique », ce qui se traduit par une dissipation (non-conservation de l'énergie du système). Ils ne peuvent donc être traités par l'équation de Schrödinger, qui n'est applicable qu'à des processus conservatifs, descriptibles par un hamiltonien, comme l'absorption ou l'émission induite.

Il se trouve pourtant qu'un certain nombre de résultats intéressants peuvent être obtenus à partir de l'équation de Schrödinger, en considérant un modèle particulier où les deux états atomiques $|a\rangle$ et $|b\rangle$, couplés par une onde électromagnétique quasi-résonnante, ont la même durée de vie Γ_D^{-1} (Figure 2.11). En appelant N_a et N_b les nombres d'atomes¹³ dans les états $|a\rangle$ ou $|b\rangle$ présents dans le volume d'interaction V à un instant donné, on admet donc que les vitesses de départ s'écrivent

$$\left[\frac{dN_a}{dt} \right]_{\text{depart}} = -\Gamma_D N_a , \quad (2.98a)$$

$$\left[\frac{dN_b}{dt} \right]_{\text{depart}} = -\Gamma_D N_b , \quad (2.98b)$$

ce qui entraîne que le nombre total d'atomes dans le volume V

$$N = N_a + N_b \quad (2.98c)$$

obéit à la même équation

$$\left[\frac{dN}{dt} \right]_{\text{depart}} = -\Gamma_D N . \quad (2.98d)$$

Une telle situation peut se rencontrer dans le cas de deux niveaux excités instables ayant le même taux d'émission spontanée, par exemple deux sous-niveaux d'un même niveau électronique. Il peut aussi s'agir de niveaux intrinsèquement stables, mais dans une situation où les atomes passent seulement un temps de l'ordre de Γ_D^{-1} dans la zone d'interaction avec l'onde électromagnétique (atomes traversant un faisceau laser par exemple).

Si nous voulons obtenir un régime stationnaire, il faut alimenter ces niveaux, ce qui peut par exemple se faire par excitation collisionnelle avec des particules chargées (décharge électrique) ou neutres, ou par pompage optique¹⁴. On peut aussi avoir des atomes qui pénètrent dans le volume d'interaction. Dans tous ces cas, nous introduisons des vitesses d'alimentation

$$\left[\frac{dN_a}{dt} \right]_{\text{alim}} = \Lambda_a , \quad (2.99a)$$

$$\left[\frac{dN_b}{dt} \right]_{\text{alim}} = \Lambda_b , \quad (2.99b)$$

qui sont a priori différentes.

¹³En fait, N_a et N_b doivent être compris comme des valeurs moyennes – au sens statistique – à l'instant t . Nous supposons que ces nombres sont suffisamment grands pour qu'il n'y ait pas lieu de se préoccuper des fluctuations statistiques.

¹⁴Voir Chapitre 3, Partie 3.2, et Complément II.1.

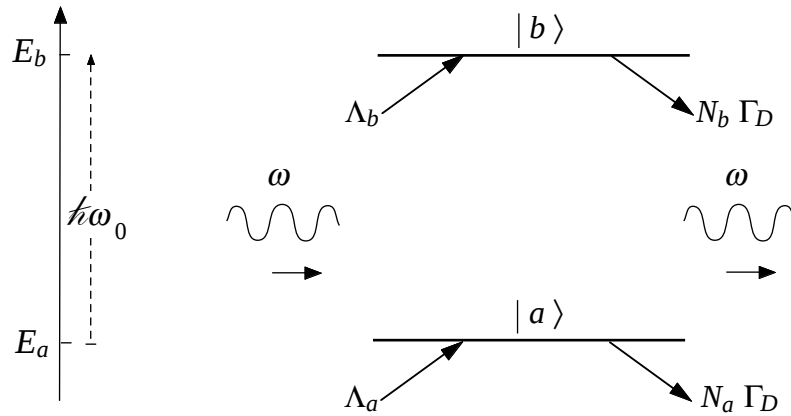


FIG. 2.11: **Modèle d'atome** où les deux états couplés par l'onde électromagnétique ont la même durée de vie Γ_D^{-1} . Les deux états sont alimentés avec des taux Λ_a et Λ_b .

En régime stationnaire, il y a compensation entre l'alimentation totale et le départ total, et le nombre total d'atome est donc

$$N = \frac{\Lambda_a + \Lambda_b}{\Gamma_D} . \quad (2.100)$$

Notons que cette équation est vraie quels que soient les transferts de population qui se produisent entre $|a\rangle$ et $|b\rangle$ sous l'effet de l'onde électromagnétique. La simplicité de cette équation, qui découle de (2.98d), est due à l'égalité des durées de vie des deux niveaux : elle est *spécifique* de notre modèle particulier¹⁵.

Nous considérons dans cette partie 2.4 le cas où Λ_b est *nul*, seul le niveau inférieur $|a\rangle$ étant alimenté. Sous l'effet de l'onde électromagnétique de pulsation ω , un certain nombre d'atomes vont passer dans le niveau $|b\rangle$: il s'agit du processus d'absorption. Nous nous proposons de décrire la situation obtenue en régime stationnaire. Nous calculerons d'une part le nombre N_b d'atomes portés dans le niveau excité, et d'autre part l'effet de ces transitions sur une onde électromagnétique, dont nous verrons qu'elle est atténuée lors de la propagation.

2.4.2 Population excitée

a. Principe du calcul

Lorsque seul le niveau $|a\rangle$ est alimenté, et en l'absence de champ électromagnétique,

¹⁵Nous donnons au paragraphe 2.4.5 quelques indications sur le cas très important d'une transition fermée entre un niveau inférieur stable et un niveau supérieur instable.

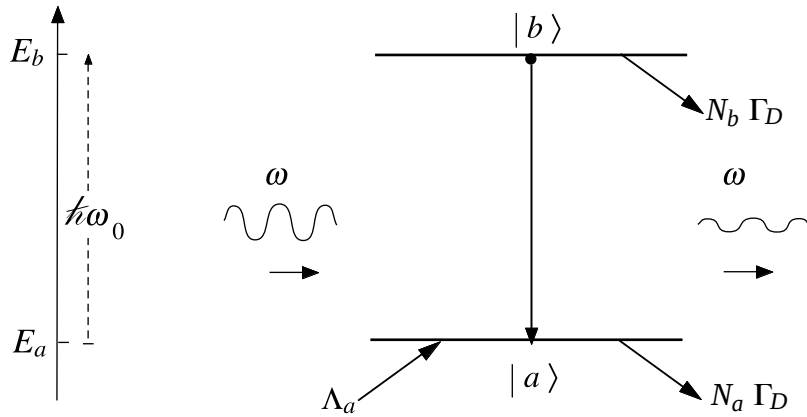


FIG. 2.12: **Absorption par des atomes à 2 niveaux de durée de vie finie**, où seul le niveau $|a\rangle$ est alimenté. Sous l'effet du rayonnement, un certain nombre d'atomes passent dans le niveau $|b\rangle$. Le rayonnement est donc absorbé.

les populations stationnaires sont

$$\begin{aligned} (N_a)_0 &= \frac{\Lambda_a}{\Gamma_D}, \\ (N_b)_0 &= 0. \end{aligned} \quad (2.101)$$

Si maintenant on applique un champ électromagnétique quasi résonnant, une fraction des atomes va passer dans l'état excité mais, en vertu de (2.100) on aura toujours

$$N_a + N_b = N = \frac{\Lambda_a}{\Gamma_D}, \quad (2.102)$$

En régime perturbatif, seule une petite fraction des atomes est excitée, de telle sorte que N_b est petit devant N_a , qui reste donc lui-même voisin de N .

Évaluons la probabilité qu'un atome qui est apparu à l'instant t_0 dans le niveau $|a\rangle$, soit dans le niveau $|b\rangle$ à l'instant t . Cette probabilité est le produit de deux termes : d'une part

$$P_{a \rightarrow b}(t_0, t) \quad (2.103a)$$

qui est la *probabilité d'excitation* d'un atome de $|a\rangle$ vers $|b\rangle$ sous l'effet du champ électromagnétique ; d'autre part

$$e^{-\Gamma_D(t-t_0)} \quad (2.103b)$$

qui est la *probabilité de survie* soit dans $|a\rangle$ soit dans $|b\rangle$.

Le nombre total d'atomes dans le niveau $|b\rangle$ à l'instant t s'obtient en sommant les contributions ci-dessus pour tous les instants t_0 antérieurs. Sachant que $\Lambda_a dt_0$ atomes sont apparus dans le niveau $|a\rangle$ entre les instants t_0 et $t_0 + dt_0$, nous trouvons

$$N_b(t) = \int_{-\infty}^t dt_0 \Lambda_a P_{a \rightarrow b}(t_0, t) e^{-\Gamma_D(t-t_0)} \quad (2.104)$$

b. Résultat perturbatif

La probabilité $P_{a \rightarrow b}(t_0, t)$ a été calculée plus haut, et elle est donnée en régime perturbatif par les formules (2.58a) et (2.58b). En substituant dans (2.104), et en posant $t - t_0 = T$, on obtient

$$N_b = \Lambda_a \frac{|W_{ba}|^2}{4\hbar^2} \int_0^\infty dT \left(\frac{\sin(\omega - \omega_0)T/2}{(\omega - \omega_0)/2} \right)^2 e^{-\Gamma_D T}. \quad (2.105)$$

L'intégration se fait en développant le sinus en exponentielles complexes, et elle donne

$$N_b = \frac{\Lambda_a}{\Gamma_D} \frac{|W_{ba}|^2}{2\hbar^2} \frac{1}{(\omega - \omega_0)^2 + \Gamma_D^2}. \quad (2.106a)$$

En remplaçant $\frac{|W_{ba}|}{\hbar}$ par la pulsation de Rabi Ω_1 , et en remplaçant Λ_a/Γ_D le nombre total d'atomes N (équation 2.102), on obtient :

$$\Rightarrow N_b = \frac{N}{2} \frac{\Omega_1^2}{(\omega - \omega_0)^2 + \Gamma_D^2}. \quad (2.106b)$$

Nous trouvons donc que la fraction d'atomes portés dans le niveau supérieur est proportionnelle à l'intensité de l'onde électromagnétique (terme Ω_1^2). De plus, si on fait varier la fréquence du laser, on a un comportement résonnant autour de la fréquence ω_0 , suivant une loi Lorentzienne de largeur $2\Gamma_D$ (Figure 2.13).

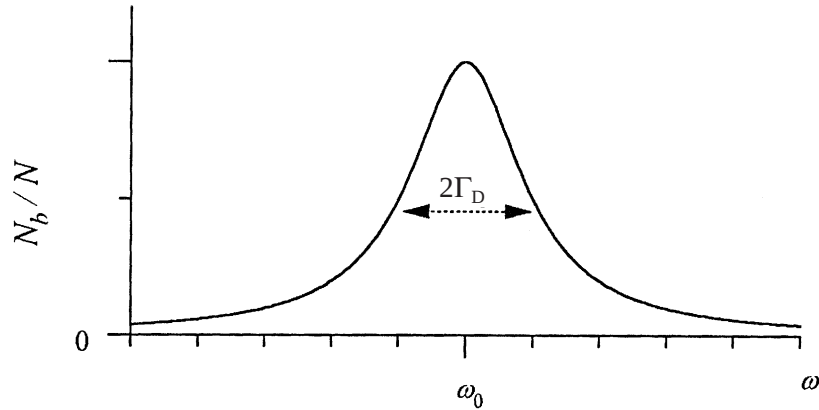


FIG. 2.13: **Variation résonnante** de la population du niveau $|b\rangle$, en fonction de la fréquence du champ incident. Dans le cas non perturbatif, cette courbe est élargie par saturation, et sa largeur devient $2\Gamma_D \sqrt{1 + \Omega_1^2/\Gamma_D^2}$.

Le calcul ci-dessus utilise l'expression perturbative de la probabilité de transition $P_{a \rightarrow b}(t_0, t)$. Il n'est légitime que si la fraction d'atomes N_b/N dans l'état $|b\rangle$ reste petite devant 1, c'est-à-dire si Ω_1 reste petit devant Γ_D ou devant le désaccord $\delta = \omega - \omega_0$. Nous allons maintenant donner l'expression non perturbative de N_b .

c. Résultat non perturbatif

Utilisons maintenant l'expression non perturbative (2.81) de la probabilité de transition $P_{a \rightarrow b}(t_0, t)$, et portons la dans (2.104), ce qui donne

$$N_b = \Lambda_a \frac{\Omega_1^2}{\Omega^2} \int_0^\infty dT \sin^2 \frac{\Omega T}{2} e^{-\Gamma_D T}, \quad (2.107)$$

avec (cf. Éq. (2.77b))

$$\Omega^2 = \Omega_1^2 + (\omega - \omega_0)^2$$

L'intégration se fait comme ci-dessus, et on obtient

$$\Rightarrow N_b = \frac{N}{2} \frac{\Omega_1^2}{(\omega - \omega_0)^2 + \Omega_1^2 + \Gamma_D^2}. \quad (2.108)$$

À la limite des basses intensités ($\Omega_1^2 \ll \Gamma_D^2 + (\omega - \omega_0)^2$) cette expression coïncide avec le résultat perturbatif (2.106b) : la population N_b croît linéairement avec Ω_1^2 , c'est-à-dire avec l'intensité de l'onde électromagnétique. En revanche, à haute intensité, la fraction d'atomes dans le niveau supérieur $|b\rangle$ croît de moins en moins vite avec Ω_1^2 , pour tendre vers la valeur asymptotique $1/2$. On dit qu'il y a *saturation de l'excitation*. Par ailleurs, la formule (2.108) montre que la largeur de la résonance augmente avec l'intensité, et devient $2\sqrt{\Gamma_D^2 + \Omega_1^2}$ (*élargissement par saturation*).

L'effet de saturation apparaît de façon particulièrement claire en introduisant le paramètre de saturation s qui, pour une transition entre deux niveaux de même durée de vie s'écrit

$$s = \frac{\Omega_1^2}{(\omega - \omega_0)^2 + \Gamma_D^2}. \quad (2.109)$$

La formule (2.108) peut alors se récrire

$$N_b = \frac{N}{2} \frac{s}{1 + s}. \quad (2.110)$$

La fonction $\frac{s}{1 + s}$, équivalente à s lorsque s est petit devant 1, et tendant asymptotiquement vers 1 lorsque s est grand devant 1, caractérise de façon générale le phénomène de saturation (Figure 2.14).

2.4.3 Susceptibilité diélectrique

Nous nous proposons maintenant de déterminer la réponse diélectrique de l'assemblée d'atomes du § 2.4.2, lorsqu'elle est soumise à un champ électromagnétique couplant le niveau alimenté $|a\rangle$ au niveau supérieur de la transition $|b\rangle$ (qui n'est pas alimenté de l'extérieur). Nous calculerons d'abord le dipôle électrique induit, puis la susceptibilité du milieu. Nous utiliserons pour ce calcul les résultats exacts (non perturbatifs) établis au paragraphe 2.3.2.

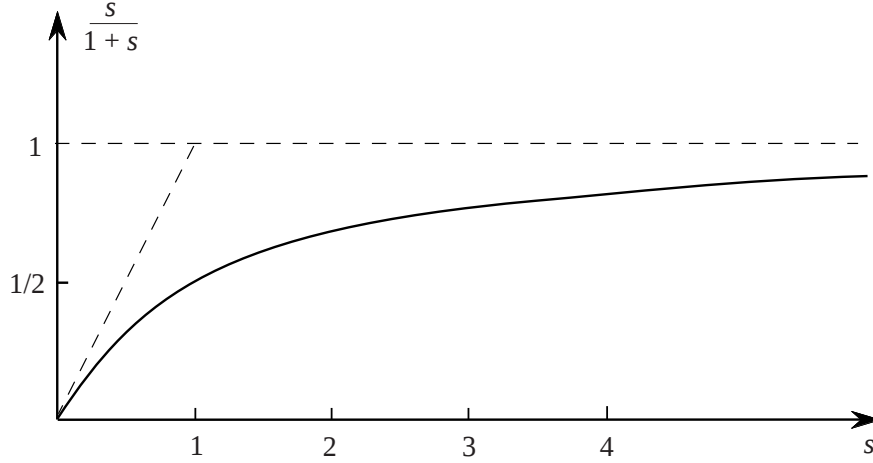


FIG. 2.14: **Phénomène de saturation.** La fonction $\frac{s}{1+s}$ est équivalente à s pour $s \ll 1$ (régime linéaire) puis elle sature pour $s \gtrsim 1$.

a. Valeur moyenne du dipôle électrique induit

Calculons d'abord la valeur moyenne (quantique) du dipôle électrique d'un atome. Pour un état atomique $|\psi\rangle$, on sait que

$$\langle \hat{\mathbf{D}} \rangle = \langle \psi | \hat{\mathbf{D}} | \psi \rangle . \quad (2.111)$$

Nous avons calculé au paragraphe 2.3.2 l'état d'un atome à deux niveaux soumis à un champ

$$\mathbf{E}(t) = \mathbf{E}_0 \cos(\omega t + \varphi) = \vec{\epsilon} E_0 \cos(\omega t + \varphi) \quad (2.112)$$

entre les instants t_0 et t , sachant qu'il était dans l'état $|a\rangle$ à l'instant t_0 . Nous avons trouvé (Éqs. C.27) :

$$|\psi(t)\rangle = \gamma_a(t)|a\rangle + \gamma_b(t)e^{-i\omega_0 t}|b\rangle , \quad (2.113a)$$

avec

$$\gamma_a(t) = \left[\cos \frac{\Omega}{2}(t - t_0) - i \frac{\delta}{\Omega} \sin \frac{\Omega}{2}(t - t_0) \right] \exp \left(i \frac{\delta}{2} t \right) , \quad (2.113b)$$

$$\gamma_b(t) = -i \left[\frac{\Omega_1}{\Omega} \sin \frac{\Omega}{2}(t - t_0) \right] \exp -i \left(\frac{\delta}{2} t + \varphi \right) , \quad (2.113c)$$

et

$$\Omega = \sqrt{\Omega_1^2 + (\omega - \omega_0)^2} . \quad (2.113d)$$

On peut alors calculer la valeur moyenne à l'instant t de la composante $\hat{D}_i$ ($i = x, y, z$) du dipôle d'un atome apparu dans l'état $|a\rangle$ à l'instant t_0 . On obtient

$$\langle \hat{D}_i \rangle(t_0, t) = d_i \gamma_a^* \gamma_b e^{-i\omega_0 t} + c.c. , \quad (2.114)$$

où d_i est l'élément de matrice $\langle a | \hat{D}_i | b \rangle$. Cet atome ayant une probabilité de survie $e^{-\Gamma_D(t-t_0)}$ (cf. Éq. 2.103b), on obtient le dipôle total de l'échantillon par une procédure analogue à celle du paragraphe 2.4.2 ci-dessus : on somme sur t_0 les contributions (2.114) pondérées par le terme de survie. En introduisant le volume V de l'échantillon, et en appelant $\mathbf{P}$ la *polarisation diélectrique* (dipôle par unité de volume) on écrit donc sa composante P_i

$$P_i = \frac{\Lambda_a}{V} \int_{-\infty}^t \langle D_i \rangle(t_0, t) e^{-\Gamma_D(t-t_0)} dt_0 . \quad (2.115)$$

Cette expression se calcule en utilisant les résultats ci-dessus, et on obtient :

$$P_i = -\frac{Nd_i}{V} \frac{\Omega_1}{2} \frac{(\omega - \omega_0) + i\Gamma_D}{\Gamma_D^2 + \Omega_1^2 + (\omega - \omega_0)^2} e^{-i(\omega t + \varphi)} + c.c. . \quad (2.116)$$

Si l'échantillon atomique est *totalelement isotrope*, la polarisation diélectrique $\mathbf{P}$ est nécessairement parallèle au champ électrique de l'onde incidente. En appelant d l'élément de matrice de la composante du dipôle suivant cet axe, et en se souvenant que $\hbar\Omega_1 = -dE_0$, on trouve finalement

$$\Rightarrow \quad \mathbf{P} = \frac{N}{V} \frac{d^2}{\hbar} \frac{\omega_0 - \omega + i\Gamma_D}{\Gamma_D^2 + \Omega_1^2 + (\omega - \omega_0)^2} \frac{\mathbf{E}_0}{2} e^{-i(\omega t + \varphi)} + c.c. . \quad (2.117)$$

Remarque

Lorsque l'hypothèse d'isotropie de la vapeur n'est pas vérifiée, la relation entre le champ électrique incident et la polarisation diélectrique résultante est tensorielle, et non pas scalaire comme en (2.117). La vapeur atomique est alors biréfringente. Cette situation peut se produire par exemple si l'état $|a\rangle$ est un sous-niveau Zeeman particulier d'un niveau atomique de moment cinétique différent de zéro, et si l'alimentation (anisotrope) peuple sélectivement ce sous-niveau. Cependant, l'hypothèse d'isotropie est souvent vérifiée par suite d'effets de moyenne, et nous nous limiterons dans la suite à des situations de ce type où le formalisme est beaucoup plus simple.

b. Susceptibilité linéaire

Nous venons de calculer la polarisation diélectrique $\mathbf{P}$ de l'échantillon soumis au champ électrique $\mathbf{E}_0 \cos(\omega t + \varphi)$. On en déduit la susceptibilité diélectrique complexe

$$\chi = \chi' + i\chi'' , \quad (2.118a)$$

définie par

$$\Rightarrow \quad \mathbf{P} = \varepsilon_0 \chi \frac{\mathbf{E}_0}{2} e^{-i(\omega t + \varphi)} + c.c. , \quad (2.118b)$$

soit encore

$$\Rightarrow \quad \mathbf{P} = \varepsilon_0 [\chi' \mathbf{E}_0 \cos(\omega t + \varphi) + \chi'' \mathbf{E}_0 \sin(\omega t + \varphi)] . \quad (2.118c)$$

À la limite des faibles intensités, nous pouvons négliger le terme Ω_1^2 au dénominateur de l'équation (2.117), et nous obtenons la *susceptibilité diélectrique linéaire*

$$\Rightarrow \chi_1 = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega + i\Gamma_D}{\Gamma_D^2 + (\omega - \omega_0)^2} . \quad (2.119a)$$

L'indice 1 rappelle qu'il s'agit de la réponse au premier ordre en champ.

On en tire les expressions explicites des parties réelle et imaginaire de la susceptibilité, dont nous verrons qu'elles caractérisent respectivement la dispersion et l'absorption de la vapeur atomique

$$\Rightarrow \chi'_1 = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega}{\Gamma_D^2 + (\omega - \omega_0)^2} , \quad (2.119b)$$

$$\Rightarrow \chi''_1 = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\Gamma_D}{\Gamma_D^2 + (\omega - \omega_0)^2} . \quad (2.119c)$$

Remarque

L'expression (2.119) de la susceptibilité linéaire obtenue en régime perturbatif à partir de notre modèle quantique simple, a la même forme que le résultat d'un calcul classique décrivant le milieu où se propage la lumière comme un ensemble d'électrons élastiquement lié (cf. le complément II.3). Il en est de même pour tout modèle quantique, tant que l'on reste en régime perturbatif (voir le paragraphe B.2 du complément II.2). Avant l'invention des lasers, l'intensité de la lumière était toujours trop faible pour que l'on sorte du régime perturbatif, ce qui explique le succès de la description des propriétés optiques des milieux matériels par le modèle de l'électron élastiquement lié.

c. Saturation

Utilisons maintenant l'expression exacte (2.117), pour calculer la valeur de la susceptibilité diélectrique définie par (D.17). On obtient

$$\chi = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega + i\Gamma_D}{\Gamma_D^2 + \Omega_1^2 + (\omega - \omega_0)^2} . \quad (2.120)$$

On constate que χ décroît avec Ω_1^2 , et donc avec l'intensité, lorsque celle-ci croît suffisamment. En utilisant le paramètre de saturation s (Équation (2.109)), on obtient la formule remarquable

$$\Rightarrow \chi = \frac{\chi_1}{1 + s} , \quad (2.121)$$

où χ_1 est la susceptibilité linéaire (Équation (2.119a)). La susceptibilité tend vers 0 à haute intensité, sous l'effet de la saturation de la transition.

2.4.4 Propagation d'une onde électromagnétique : absorption et dispersion

a. Propagation atténuée : absorption

On sait que dans un milieu de susceptibilité linéaire $\chi_1(\omega)$ peuvent se propager des ondes électromagnétiques de vecteur d'onde $\mathbf{k} = k \mathbf{e}_k$ avec

$$k = n \frac{\omega}{c}, \quad (2.122a)$$

l'indice de réfraction n étant relié à la susceptibilité par

$$\Rightarrow n^2 = 1 + \chi_1(\omega). \quad (2.122b)$$

Si la susceptibilité a une partie imaginaire différente de zéro, il en est de même de l'indice. Pour un milieu atomique dilué, χ est petit devant 1, et on peut écrire

$$n = 1 + \frac{\chi'_1}{2} + i \frac{\chi''_1}{2}. \quad (2.122c)$$

Le vecteur d'onde complexe s'écrit alors

$$k = k' + i k'', \quad (2.123a)$$

avec

$$k' \simeq \left(1 + \frac{\chi'_1}{2}\right) \frac{\omega}{c} \quad (2.123b)$$

et

$$k'' \simeq \frac{\chi''_1}{2k'} \frac{\omega}{c} \simeq \frac{k'}{|k'|} \frac{\chi''_1}{2} \frac{\omega}{c} \quad (2.123c)$$

(comme χ''_1 est positif, on note que k'' a le même signe que k').

Une onde de vecteur d'onde $\mathbf{k}$ dirigé suivant Oz a pour expression

$$\mathbf{E}(z, t) = \frac{\mathbf{E}_0}{2} e^{i(kz - \omega t)} + c.c. = \mathbf{E}_0 e^{-k''z} \cos(\omega t - k'z). \quad (2.124)$$

Comme k' et k'' ont le même signe, on constate que l'amplitude du champ décroît exponentiellement suivant la direction de propagation. On a une *propagation atténuée*.

La partie *imaginaire* χ''_1 de la susceptibilité diélectrique linéaire d'un milieu caractérise donc l'*absorption*. L'expression (2.119c) montre que l'absorption linéaire est proportionnelle à la densité atomique N/V (nombre d'atomes par unité de volume). Elle montre également un comportement résonnant Lorentzien avec la fréquence (voir Fig. 2.15.a). Ces caractéristiques de l'absorption peuvent être obtenues à partir du modèle d'électron élastiquement lié (cf. la remarque après les formules 2.119).

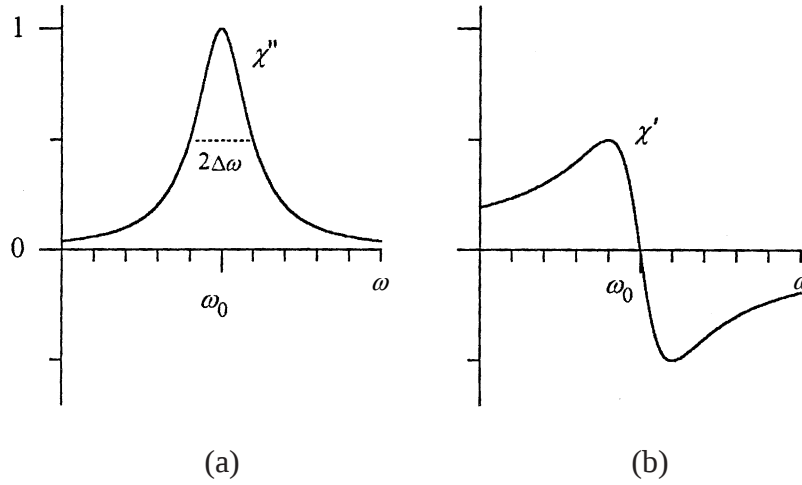


FIG. 2.15: **Absorption et dispersion autour d'une résonance de largeur $\Delta\omega$:**

(a) Profil Lorentzien : $\chi_1''(\omega) = \frac{\chi_1''(\omega_0)}{1 + \left(\frac{\omega - \omega_0}{\Delta\omega}\right)^2}$; (b) Profil de dispersion : $\chi_1'(\omega) = \chi_1''(\omega_0) \frac{-\left(\frac{\omega - \omega_0}{\Delta\omega}\right)}{1 + \left(\frac{\omega - \omega_0}{\Delta\omega}\right)^2}$

b. Dispersion

Comme le montre l'équation (2.122c), la partie réelle de la susceptibilité linéaire χ_1' est reliée à la partie réelle de l'indice de réfraction, qui caractérise la vitesse de phase d'une onde électromagnétique dans la vapeur. L'expression (2.119b), nous donne la variation de χ_1' avec la fréquence représentée sur la figure 2.15.b. Il s'agit d'un comportement dispersif typique : dans les « ailes » de la raie l'indice de réfraction croît avec la fréquence (dispersion normale), alors qu'à la traversée de la résonance le sens de variation s'inverse (dispersion anormale). Notons qu'à grand désaccord $|\delta| = |\omega - \omega_0| \gg \Delta\omega$, les effets de dispersion décroissent en $|\delta|^{-1}$, beaucoup moins vite que les effets d'absorption qui décroissent en $|\delta|^{-2}$.

Ici encore, on peut remarquer que dans le régime linéaire les effets dispersifs sont exactement ceux prévus par un modèle d'électron classique élastiquement lié.

c. Propagation en régime saturé

Lorsque le paramètre de saturation s (équation 2.109) n'est pas petit devant 1, la polarisation diélectrique n'est pas proportionnelle au champ électrique appliqué, et la question de la propagation d'une onde électromagnétique dans un tel milieu est en général très complexe. Le problème se simplifie beaucoup si on se place dans le cadre de l'« *approximation de l'enveloppe lentement variable* », où le module de l'amplitude complexe du champ électrique varie lentement à l'échelle de la longueur d'onde. Il en est alors de même de l'intensité lumineuse, et donc du paramètre de saturation s qui peut être considéré comme constant dans une tranche dz grande devant la longueur d'onde lumineuse. Les résultats ci-dessus se généralisent alors sans difficulté en remplaçant χ_1 par $\chi_1/(1 + s)$ (cf. équation 2.121). L'effet de la saturation est donc de réduire l'absorption, ainsi que l'écart entre la

vitesse de phase $c/(1 + \frac{\chi'}{2})$ et la vitesse de la lumière dans le vide c .

Remarque

Lorsque l'absorption et l'indice de réfraction dépendent de l'intensité lumineuse, par exemple à cause de la saturation, on voit apparaître des effets d'optique et de spectroscopie non linéaires (voir complément III.5).

2.4.5 Cas d'un système fermé à deux niveaux

De nombreux résultats obtenus ci-dessus à partir de notre modèle particulier se généralisent au cas où les deux niveaux $|a\rangle$ et $|b\rangle$ n'ont pas la même durée de vie. Un cas important est celui du *système fermé à deux niveaux* où le niveau du bas $|a\rangle$ est stable (durée de vie infinie) tandis que celui du haut ne peut se désexciter que vers celui du bas. La population totale est alors conservée en l'absence d'alimentation extérieure. Cette situation se rencontre chaque fois que l'on a un ensemble d'atomes dans leur état fondamental (vapeur à température modérée) en interaction avec un rayonnement quasi-résonnant sur une transition connectant ce niveau fondamental à un niveau excité qui peut se désexciter par émission spontanée (durée de vie radiative Γ_{sp}^{-1}). Un tel modèle rend également bien compte de nombreuses propriétés optiques des diélectriques dans le visible lorsque les premières transitions résonnantes électroniques sont à des fréquences ultraviolettes de fréquences supérieures à ω . Il suffit dans ce cas de considérer la limite $\omega \ll \omega_0$ des résultats obtenus.

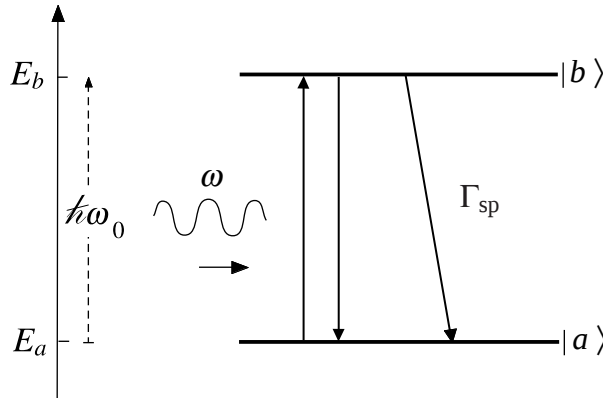


FIG. 2.16: **Système fermé à deux niveaux.** Le champ électromagnétique quasi résonnant (ω voisin de $\omega_0 = (E_b - E_a)/\hbar$) couple les deux niveaux $|a\rangle$ et $|b\rangle$, qui ne sont pas alimentés de l'extérieur. Le niveau $|a\rangle$ est stable, et le niveau $|b\rangle$ ne peut se désexciter spontanément que vers $|a\rangle$.

Pour traiter ce problème, il faut utiliser les équations de Bloch optiques (voir Complément II.2, § C.2). On trouve une susceptibilité ayant la même forme que celle obtenue en (2.120) au remplacement près de Γ_D par $\Gamma_{\text{sp}}/2$ et de Ω_1^2 par $\Omega_1^2/2$ (voir Complément

C.2, Équation C.13), ce qui donne

$$\chi = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega + i \frac{\Gamma_{sp}}{2}}{\frac{\Gamma_{sp}^2}{4} + \frac{\Omega_1^2}{2} + (\omega - \omega_0)^2} . \quad (2.125)$$

La similarité de cette formule avec la formule (2.120) donne beaucoup d'intérêt au modèle plus simple étudié dans cette partie 2.4, dont les résultats restent généralement valables, avec les substitutions indiquées ci-dessus.

La similitude est encore plus frappante si on introduit un paramètre de saturation s valant ici

$$s = \frac{\Omega_1^2/2}{(\omega - \omega_0)^2 + \Gamma_{sp}^2/4} . \quad (2.126)$$

Les expressions (2.110) et (2.121) restent alors formellement valables. Comme la susceptibilité linéaire χ_1 est identique à celle obtenue dans le modèle classique de l'électron élastiquement lié, on comprend les succès de ce modèle classique dans l'interprétation des phénomènes optiques en régime non saturant ($s \ll 1$) qui était le seul régime observable avant l'invention des lasers.

En fait, les expressions obtenues avec notre modèle simple sont souvent valables sans aucune modification. Par exemple, dans la limite très non-résonnante, la partie imaginaire de la susceptibilité devient négligeable, tandis qu'une bonne approximation de sa partie réelle est (cf. Équation (2.119b)) :

$$\chi'_1 = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{1}{\omega_0 - \omega} . \quad (2.127)$$

Cette expression, valable pour toute transition, peut être évaluée numériquement en introduisant la « force d'oscillateur » f_{ab} de la transition, définie par

$$f_{ab} = \frac{2m\omega_0 d^2}{\hbar q^2} . \quad (2.128a)$$

C'est un nombre sans dimension, souvent voisin de 1 pour les transitions « de résonance » (couplant le niveau fondamental d'un atome à un niveau excité), et que l'on peut trouver dans les tables de données atomiques. L'équation (2.127) s'écrit alors

$$\chi'_1 = \frac{N}{V} \frac{q^2 f_{ab}}{2m\varepsilon_0\omega_0(\omega_0 - \omega)} . \quad (2.128b)$$

Cette expression peut encore être transformée en introduisant l'inverse du temps d'amortissement classique d'un électron lié élastiquement

$$\Gamma_{cl} = \frac{q^2}{6\pi\varepsilon_0} \frac{\omega_0^2}{mc^3} , \quad (2.129)$$

ce qui conduit à

$$\chi'_1 = 3\pi \frac{N}{V} \frac{c^3}{\omega_0^3} \frac{\Gamma_{cl} f_{ab}}{\omega_0 - \omega} , \quad (2.130a)$$

ou encore, en fonction de la longueur d'onde $\lambda_0 = 2\pi c/\omega_0$

$$\chi'_1 = \frac{3}{8\pi^2} \frac{N\lambda_0^3}{V} \frac{\Gamma_{cl} f_{ab}}{\omega_0 - \omega} . \quad (2.130b)$$

On peut montrer que pour un atome à deux niveaux la quantité $f_{ab}\Gamma_{cl}$ est égale au taux Γ_{sp} d'amortissement par émission spontanée du niveau excité, quantité que l'on trouve également dans les tables de données. On obtient ainsi la formule utile

$$\chi'_1 = \frac{3}{8\pi^2} \frac{N\lambda_0^3}{V} \frac{\Gamma_{sp}}{\omega_0 - \omega} \quad (2.130c)$$

Cette expression permet de comprendre les ordres de grandeur de l'indice de réfraction des vapeurs ou des solides transparents. En effet, loin des zones d'absorption, le désaccord $|\omega - \omega_0|$ est très grand devant le taux d'émission spontanée Γ_{sp} , et il faut une densité atomique N/V beaucoup plus grande que $1/\lambda_0^3$ pour obtenir une susceptibilité qui ne soit pas très petite devant 1. On comprend ainsi que l'indice de réfraction d'une vapeur atomique ou moléculaire, qui est lié à χ'_1 (voir le paragraphe 4 ci-dessous), diffère très peu de 1 dans les zones de transparence : par exemple $n - 1$ vaut environ 3×10^{-4} dans le cas de l'air pour de la lumière visible. En revanche, pour des solides tels que le verre ou plus généralement les diélectriques, la densité atomique est beaucoup plus grande que $1/\lambda_0^3$ (typiquement 10^{28} m^{-3} comparés à 10^{19} m^{-3} pour des longueurs d'onde visibles), et on a des matériaux réfringents bien que transparents.

Remarque

Nous verrons dans la partie 2.5 que l'absorption peut s'exprimer en régime linéaire à l'aide d'une section efficace σ_{abs} , reliée à la partie imaginaire de la susceptibilité par la relation

$$\frac{\omega}{c} \chi''_1 = \frac{N}{V} \sigma_{abs} \quad (2.131)$$

(rappelons que nous sommes ici dans le cas où on n'alimente que le niveau du bas, et en régime linéaire $N_b \ll N_a$).

Pour un atome à deux niveaux fermés, et lorsque la relaxation du niveau $|b\rangle$ vers le niveau $|a\rangle$ est uniquement due à l'émission spontanée, la section efficace d'absorption prend la forme remarquable

$$\sigma_{abs} = \frac{3\lambda_0^2}{2\pi} \frac{1}{1 + \left(\frac{\omega - \omega_0}{\Gamma_{sp}/2}\right)^2} , \quad (2.132)$$

où λ_0 est la longueur d'onde à résonance ($\lambda_0 = 2\pi c/\omega_0$). On retiendra qu'à résonance exacte la section efficace d'absorption est de l'ordre du carré de la longueur d'onde.

2.5 Amplification laser

2.5.1 Alimentation du niveau supérieur : émission induite

Nous considérons à nouveau le modèle présenté au paragraphe 2.4.1 (Figure 2.11) mais nous supposons maintenant que c'est le niveau supérieur de la transition qui est alimenté (figure 2.17).

$$\Lambda_a = 0 \quad , \quad \Lambda_b \neq 0 . \quad (2.133)$$

Nous pouvons reprendre les calculs de la partie 2.4 en partant d'atomes initialement dans le niveau $|b\rangle$, et en calculant en régime stationnaire les populations N_b et N_a , et la susceptibilité χ .

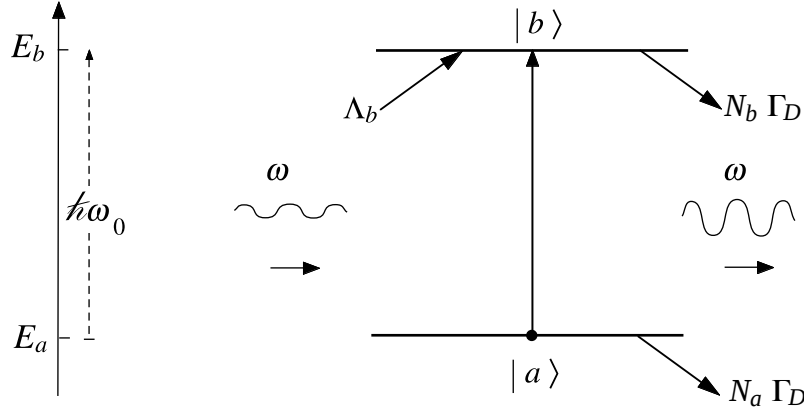


FIG. 2.17: **Emission induite**, pour des atomes à 2 niveaux de durée de vie finie, lorsque seul le niveau supérieur est alimenté. Sous l'effet du rayonnement, un certain nombre d'atomes passent de $|b\rangle$ à $|a\rangle$, et le rayonnement est amplifié.

a. Population transférée par émission induite

En l'absence de champ électromagnétique, tous les atomes sont dans le niveau supérieur :

$$(N_b)^{(0)} = N = \frac{\Lambda_b}{\Gamma} , \quad (2.134a)$$

$$(N_a)^{(0)} = 0 . \quad (2.134b)$$

Sous l'effet du champ, des atomes passent dans le niveau du bas : c'est le phénomène d'*émission induite*. En régime perturbatif, la population stationnaire dans $|a\rangle$ prend une forme analogue à (2.106b) :

$$N_a = \frac{N}{2} \frac{\Omega_1^2}{(\omega - \omega_0)^2 + \Gamma_D^2} = \frac{N}{2} s . \quad (2.135)$$

Dans le cas général (non perturbatif), la population stationnaire dans le niveau non alimenté prend la valeur analogue à (2.110)

$$N_a = \frac{N}{2} \frac{s}{1 + s} , \quad (2.136)$$

où

$$s = \frac{\Omega_1^2}{(\omega - \omega_0)^2 + \Gamma_D^2}$$

est le paramètre de saturation défini par (2.109). Notons que quelle que soit la valeur de s , la population N_a du niveau du bas reste toujours inférieure à celle du niveau du haut N_b .

b. Susceptibilité

Le calcul donne ici un résultat analogue à celui du paragraphe 2.4.3 (voir Équation 2.120), mais *avec un signe inversé*

$$\chi = -\frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega + i\Gamma_D}{(\omega - \omega_0)^2 + \Gamma_D^2} \frac{1}{1+s} . \quad (2.137)$$

En régime perturbatif (faible saturation) nous avons donc

$$\chi = -\chi_1 , \quad (2.138)$$

où χ_1 est la susceptibilité linéaire définie pour des atomes dans l'état fondamental (Équation 2.119a). Rappelons que cette quantité χ_1 est de la forme obtenue avec le modèle classique de l'électron élastiquement lié, et qu'elle est associée au comportement observé habituellement (absorption, dispersion normale). Nous allons voir que le *changement de signe entraîne des conséquences remarquables*.

2.5.2 Propagation amplifiée : effet laser

Reprenons le calcul du paragraphe 2.4.4 dans la situation où seul le niveau du haut est alimenté. Dans un milieu de susceptibilité linéaire complexe $\chi_1 = \chi'_1 + i\chi''_1$, une onde électromagnétique monochromatique a un vecteur d'onde complexe

$$\mathbf{k} = \mathbf{e}_k(k' + i k'') . \quad (2.139)$$

Le vecteur unitaire $\mathbf{e}_k$ caractérise la direction de propagation. Dans le cas usuel où $|\chi_1| \ll 1$, on a

$$k'' \simeq \frac{k'}{|k'|} \frac{\chi''_1}{2} \frac{\omega}{c} . \quad (2.140)$$

L'expression ci-dessus met en évidence la relation entre les signes de k' et k'' . Si nous prenons la propagation suivant Oz , nous trouvons un champ de la forme (2.124)

$$\begin{aligned} \mathbf{E}(z, t) &= \frac{\mathbf{E}_0}{z} \exp\{i(kz - \omega t)\} + \text{c.c.} \\ &= \mathbf{E}_0 e^{-k''z} \cos(\omega t - kz) . \end{aligned} \quad (2.141)$$

Mais comme χ''_1 est maintenant négatif (Équation 2.137), l'équation 2.140 montre que k'' et k' ont des signes opposés, c'est-à-dire que l'amplitude du champ *croît* suivant la direction de propagation. L'onde électromagnétique est donc *amplifiée*. Le phénomène à la base de cette amplification est évidemment l'émission induite, qui est le phénomène majoritaire dans la situation où N_b est très grand devant N_a .

L'amplification d'une onde lumineuse grâce à l'émission induite est l'effet LASER (Light Amplification by Stimulated Emission of Radiation). On le caractérise par un

gain par unité de longueur g , caractérisant la croissance de l'intensité lumineuse $I(z)$ (proportionnelle au carré de l'amplitude du champ)

$$\frac{dI(z)}{dz} = g I(z) . \quad (2.142)$$

Dans le régime perturbatif considéré ici, les équations (2.138), (2.140) et (2.141) nous donnent le *gain par unité de longueur*

$$g = -2k'' = \frac{\omega}{c} \chi_1'' . \quad (2.143)$$

Remarque

L'inversion du signe de la partie réelle χ_1' de la susceptibilité se traduit également par un comportement inhabituel de l'indice de réfraction $n' = 1 + \frac{\chi_1'}{2}$: il décroît avec la fréquence dans les zones de transparence, et croît dans les raies d'absorption. L'observation d'une telle dispersion inversée (parfois appelée « anormale ») est une indication expérimentale de la possibilité d'obtenir une amplification.

2.5.3 Généralisation : alimentation des deux niveaux et saturation

Supposons maintenant que les deux niveaux $|a\rangle$ et $|b\rangle$ sont simultanément alimentés (taux Λ_a et Λ_b). Comment les résultats précédents sont-ils modifiés ?

En fait, le traitement reste simple. Comme les entrées d'atomes dans les niveaux $|a\rangle$ ou $|b\rangle$ sont des événements indépendants, on pourra généraliser les calculs des paragraphes (2.2.4.3) et (2.2.5.1), en ajoutant indépendamment les contributions des atomes injectés en $|a\rangle$ et celles des atomes injectés en $|b\rangle$. On trouve ainsi, pour les populations

$$N_a = \frac{\Lambda_a}{\Gamma_D} + \frac{1}{2} \frac{s}{1+s} \frac{\Lambda_b - \Lambda_a}{\Gamma_D} , \quad (2.144a)$$

$$N_b = \frac{\Lambda_b}{\Gamma_D} + \frac{1}{2} \frac{s}{1+s} \frac{\Lambda_a - \Lambda_b}{\Gamma_D} . \quad (2.144b)$$

En notant comme précédemment

$$(N_a)^{(0)} = \frac{\Lambda_a}{\Gamma_D} , \quad (2.145a)$$

$$(N_b)^{(0)} = \frac{\Lambda_b}{\Gamma_D} , \quad (2.145b)$$

les populations stationnaires obtenues en l'absence de rayonnement ($s = 0$), ces équations se récrivent sous la forme plus générale :

$$N_a = (N_a)^{(0)} + \frac{1}{2} \frac{s}{1+s} (N_b - N_a)_0 \quad (2.146a)$$

$$N_b = (N_b)^{(0)} + \frac{1}{2} \frac{s}{1+s} (N_a - N_b)_0 . \quad (2.146b)$$

On notera alors la relation importante pour la suite

$$\Rightarrow N_b - N_a = \frac{(N_b - N_a)^{(0)}}{1 + s}. \quad (2.147)$$

La différence des populations diminue, et tend vers 0, par saturation.

Le calcul de la susceptibilité s'effectue de la même façon, en généralisant l'équation (2.117). En notant χ_1 la susceptibilité linéaire (équations (D.22), avec $N = (N_a)_0 + (N_b)_0$), on obtient

$$\Rightarrow \chi = \frac{(N_a - N_b)^{(0)}}{(N_a + N_b)^{(0)}} \frac{\chi}{1 + s} = \frac{N_a - N_b}{(N_a + N_b)^{(0)}} \chi_1 \quad (2.148)$$

(on a utilisé la relation (2.113)).

2.5.4 Gain laser et inversion de population

Le raisonnement du paragraphe (2.5.2) se généralise à la situation du paragraphe (2.5.3), dans le cadre de l'approximation de l'enveloppe lentement variable si le paramètre de saturation n'est pas petit devant 1. On constate sur l'expression (2.148) que la partie imaginaire χ'' de la susceptibilité prend des valeurs négatives si $N_b > N_a$. On aura alors une amplification de la lumière, par effet laser. La situation où un niveau d'énergie élevée est plus peuplé qu'un niveau d'énergie moins élevée est manifestement *hors d'équilibre thermodynamique*. Pour cette raison, on dit qu'il y a *inversion de population*.

En raisonnant comme au paragraphe (2.5.2), on peut définir le gain par unité de longueur

$$g = \frac{1}{I} \frac{dI}{dz} = \frac{\omega}{c} \chi_1'' \frac{N_b - N_a}{N} = \frac{g_0}{1 + s}, \quad (2.149)$$

où on a introduit le gain non saturé¹⁶ ($s \ll 1$)

$$g^{(0)} = \frac{\omega}{c} \chi_1'' \frac{(N_b - N_a)^{(0)}}{N} \quad (2.150)$$

Ces formules montrent clairement que le gain laser est lié à l'inversion de population. De plus, elles font apparaître le phénomène de *saturation du gain* : le gain décroît vers 0 lorsque l'intensité du rayonnement décroît.

Ces résultats peuvent manifestement s'interpréter en considérant la compétition entre l'émission induite, qui permet l'amplification, et l'absorption, qui s'y oppose. Le premier phénomène est lié à la population du niveau $|b\rangle$, tandis que le deuxième est lié à celle du niveau $|a\rangle$. Nous allons maintenant développer cette interprétation d'abord en faisant un bilan énergétique lors de la propagation, puis en écrivant des équations cinétiques, tant pour le rayonnement que pour les atomes.

¹⁶Il ne faut pas confondre l'indice 0 qui, dans le cas du gain $g^{(0)}$, ou des populations $(N_a)^{(0)}$ et $(N_b)^{(0)}$, fait référence à une situation non saturée ($s \ll 1$), avec la notation ω_0 qui désigne la fréquence de résonance. Pour éviter les confusions, nous noterons parfois $g^{(0)}(\omega_M)$ la valeur du gain non saturé à la fréquence de résonance $\omega_M = \omega_0$.

2.5.5 Bilan énergétique au cours de la propagation : absorption et émission induite

a. Absorption

Considérons d'abord une onde plane se propageant suivant les z positifs dans un milieu où les atomes sont injectés dans l'état $|a\rangle$ (cf. § 2.4), et où l'intensité est suffisamment faible pour se contenter du traitement perturbatif. Il est instructif de calculer la puissance dissipée par absorption dans une tranche d'épaisseur dz , et de section A , perpendiculaire à Oz . Le vecteur de Poynting moyen dans le plan d'abscisse z est dirigé suivant Oz , et il a pour norme

$$\Pi = \frac{\varepsilon_0 c E_0^2 e^{-2k''z}}{2} . \quad (2.151)$$

La puissance absorbée dans la tranche dz est donc

$$[-d\Phi]_{\text{abs}} = A[\Pi(z) - \Pi(z + dz)] = A dz \varepsilon_0 c [E(z)]^2 k'' , \quad (2.152)$$

avec

$$E(z) = E_0 e^{-k''z} . \quad (2.153)$$

En utilisant (2.123c) et en remplaçant χ_1'' par sa valeur (2.119c), on trouve

$$[-d\Phi]_{\text{abs}} = A dz \frac{N}{V} \frac{\omega d^2 [E(z)]^2}{2\hbar} \frac{\Gamma_D}{\Gamma_D^2 + (\omega - \omega_0)^2} . \quad (2.154)$$

On peut récrire cette expression en utilisant l'expression (2.106b) de la population atomique dans le niveau $|b\rangle$. On obtient alors

$$[-d\Phi]_{\text{abs}} = A dz \frac{N_b}{V} \hbar \omega \Gamma_D . \quad (2.155)$$

Cette expression s'interprète très simplement, si on remarque que $\Gamma_D A dz N_b/V$ est le nombre d'atomes dans le niveau $|b\rangle$ qui disparaissent du volume $A dz$ par unité de temps. Comme les atomes ont été injectés dans l'état $|a\rangle$, ils ont *absorbé une énergie* $\hbar\omega$ pour être portés dans l'état supérieur, et l'équation ci-dessus exprime simplement le *bilan énergétique* de ce transfert de l'onde électromagnétique vers les atomes.

Remarque

Il peut sembler paradoxal que le transfert des atomes du niveau $|a\rangle$ au niveau $|b\rangle$ soit associé à une énergie absorbée $\hbar\omega$ et non $\hbar\omega_0 = E_b - E_a$. Une analyse détaillée montre que pour chaque processus élémentaire c'est bien un quantum d'énergie $\hbar\omega$ qui est prélevé sur le rayonnement, ce qui s'interprète naturellement dans un modèle totalement quantique où le rayonnement est décrit comme un ensemble de photons d'énergie $\hbar\omega$. Mais on peut alors s'interroger sur la conservation de l'énergie lors du passage de l'atome de $|a\rangle$ à $|b\rangle$. En fait l'atome ne va rester qu'un temps fini Δt dans l'état atteint après l'absorption d'un photon, ce temps étant de l'ordre de la durée de vie Γ_b^{-1} de l'état $|b\rangle$ dans le cas d'une diffusion exactement résonnante, et de $|\omega - \omega_0|^{-1}$ hors de résonance. Il n'y a aucun paradoxe

à considérer que pendant ce temps Δt l'état atomique a une énergie différente de E_b par $\Delta E \simeq \hbar/\Delta t$ (relation temps-énergie de Heisenberg). Notons par contre que si on considère l'état final du système atome rayonnement lorsque l'atome est passé du niveau $|b\rangle$ à un autre niveau supposé stable, on doit avoir conservation exacte de l'énergie. Par exemple, dans le cas d'un cycle de fluorescence où un atome à deux niveaux fermés (§ 2.2.4.5) prélève un photon du rayonnement incident et le diffuse dans une direction différente en se retrouvant dans l'état initial $|a\rangle$, le rayonnement diffusé a une fréquence strictement égale à celle du rayonnement incident.

b. Emission induite

Considérons maintenant le cas où les atomes sont injectés dans le niveau $|b\rangle$, et plaçons-nous encore en régime perturbatif (cf. § 2.5.2). On sait qu'il y a amplification du rayonnement par effet laser. Un calcul analogue au calcul ci-dessus montre alors que la puissance gagnée par le rayonnement dans la tranche dz peut s'écrire

$$[d\Phi]_{\text{sti}} = A dz \frac{N_a}{V} \hbar \omega \Gamma_D . \quad (2.156)$$

Ce résultat s'interprète évidemment de façon analogue. La quantité $\Gamma_D N_a A dz/V$ est le nombre d'atomes dans le niveau $|a\rangle$ qui disparaissent de la tranche dz par unité de temps. Comme ils ont été injectés dans le niveau $|b\rangle$, chacun a cédé une énergie $\hbar \omega$ au rayonnement. L'équation (2.156) exprime ici encore un bilan énergétique.

c. Cas général

On peut généraliser les calculs ci-dessus dans le cas non perturbatif, et lorsque les deux niveaux sont alimentés. En utilisant l'expression (2.148) de la susceptibilité, on trouve que la puissance gagnée par le rayonnement dans la tranche dz , pour un faisceau de section A , vaut

$$d\phi = A dz \frac{\Lambda_b - \Lambda_a}{V} \frac{1}{2} \frac{1}{1+s} \hbar \omega . \quad (2.157)$$

En utilisant l'expression (2.146b) de la population N_b , on peut écrire

$$\Lambda_b - \Gamma_D N_b = \frac{\Lambda_b - \Lambda_a}{2} \frac{1}{1+s} . \quad (2.158)$$

Or le premier membre ci-dessus n'est autre que le nombre d'atomes qui, chaque seconde, sont passé de $|b\rangle$ à $|a\rangle$, puisqu'il s'agit de la différence entre le taux entrant Λ_b et le taux de disparition $\Gamma_D N_b$ dans l'état $|b\rangle$. En multipliant ce nombre par $\hbar \omega$, on obtient la puissance cédée par les atomes au rayonnement, et l'équation (2.157) représente bien le bilan d'énergie échangée entre atomes et rayonnement dans le volume $A dz$.

2.5.6 Équations cinétiques pour les atomes

Nous avons vu au paragraphe précédent que l'absorption ou l'amplification de l'onde peuvent se comprendre en faisant le bilan des échanges d'énergie entre les atomes et

le rayonnement : la puissance gagnée par le rayonnement résulte de la différence entre l'énergie apportée par les processus d'émission induite, et celle prélevée par les processus d'absorption.

Nous allons voir ici, de façon analogue, que les valeurs N_a et N_b obtenues pour les populations des niveaux $|a\rangle$ et $|b\rangle$ peuvent s'interpréter en considérant la compétition entre les divers processus à l'œuvre : absorption, émission induite, relaxation. On peut rendre compte de cette compétition en introduisant des *équations cinétiques*, encore appelées *équations de pompage* ou « équations de taux » (*rate equations* en anglais), qui sont en fait de simples équations de bilan entre taux de départ et d'arrivée, que nous écrirons

$$\frac{dN_b}{dt} = \Gamma_D \frac{s}{2} N_a - \Gamma_D \frac{s}{2} N_b - \Gamma_D N_b + \Lambda_b, \quad (2.159a)$$

$$\frac{dN_a}{dt} = -\Gamma_D \frac{s}{2} N_a + \Gamma_D \frac{s}{2} N_b - \Gamma_D N_a + \Lambda_a. \quad (2.159b)$$

Des équations de ce type ont été écrites pour la première fois par Einstein, qui a introduit la notion d'émission induite, et qui a montré que les coefficients associés à l'absorption (terme $\Gamma_D \frac{s}{2} N_a$) et à l'émission induite (terme $\Gamma_D \frac{s}{2} N_b$) sont nécessairement égaux. La validité de telles équations n'est pas facile à établir à partir des premiers principes. Ici, nous nous contenterons de vérifier qu'elles conduisent à des résultats en accord avec ceux obtenus plus haut dans un cadre plus rigoureux. Leur intérêt est qu'elles permettent une interprétation très simple (Fig. 2.18), et qu'elles se généralisent aisément.

La structure des équations (2.159a) et (2.159b) est claire. Le taux de variation de chaque population (N_a et N_b) est la somme des taux associés à chaque processus physique identifié sur la figure 2.18. On reconnaît ainsi les termes d'alimentation (Λ_a et Λ_b) et les termes de relaxation ($-\Gamma_D N_a$ et $-\Gamma_D N_b$). De plus, les processus d'absorption et d'émission induite sont décrits par les taux

$$\left[\frac{dN_a}{dt} \right]_{\text{abs}} = - \left[\frac{dN_b}{dt} \right]_{\text{abs}} = -\Gamma_D \frac{s}{2} N_a \quad (2.160a)$$

et

$$\left[\frac{dN_b}{dt} \right]_{\text{sti}} = - \left[\frac{dN_a}{dt} \right]_{\text{sti}} = -\Gamma_D \frac{s}{2} N_b \quad (2.160b)$$

Le régime stationnaire s'obtient immédiatement en annulant le premier membre des équations (2.159a) et (2.159b). On en déduit par addition que le nombre total d'atomes N vaut

$$N = N_a + N_b = \frac{\Lambda_a}{\Gamma_D} + \frac{\Lambda_b}{\Gamma_D}, \quad (2.161)$$

ce que l'on avait obtenu en (2.100). Par ailleurs, l'inversion de population stationnaire s'obtient par différence et elle vaut

$$N_b - N_a = \frac{\Lambda_b - \Lambda_a}{\Gamma_D} \frac{1}{1+s} = N \frac{\Lambda_b - \Lambda_a}{\Lambda_a + \Lambda_b} \frac{1}{1+s} = \frac{(N_b - N_a)^{(0)}}{1+s}. \quad (2.162)$$

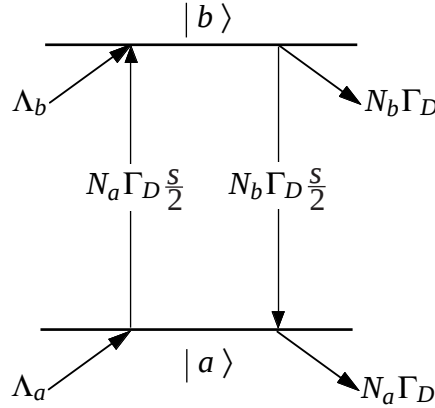


FIG. 2.18: **Taux intervenant dans les équations cinétiques**, décrivant l'alimentation (Λ_a et Λ_b), la relaxation vers l'extérieur ($N_a \Gamma_D$ et $N_b \Gamma_D$), l'absorption ($N_a \Gamma_D \frac{s}{2}$), l'émission induite ($N_b \Gamma_D \frac{s}{2}$).

On constate que l'on retrouve bien le résultat (2.147), qui avait été obtenu à partir d'une méthode plus rigoureuse. Ce calcul donne aussi les populations stationnaires N_a et N_b , que l'on trouve identiques au résultat (2.144a) et (2.144b) du calcul plus rigoureux.

L'intérêt des équations cinétiques pour les populations atomiques est d'abord de mettre en relief la compétition entre les termes d'absorption et d'émission induite, compétition qui est responsable du phénomène de saturation de l'inversion de population.

Un autre intérêt des équations cinétiques est de se généraliser sans difficulté à des cas plus complexes : temps de relaxation inégaux pour les niveaux $|a\rangle$ et $|b\rangle$; prise en compte d'une relaxation de $|b\rangle$ vers $|a\rangle$. Il reste généralement aisé d'écrire et de résoudre les équations cinétiques, alors que les équations de Bloch optiques deviennent rapidement inextricables.

Il faut néanmoins ne pas perdre de vue que les équations cinétiques ne constituent qu'un traitement approché, qui n'est en général justifié que lorsque le temps de relaxation des dipôles atomiques est court par rapport aux temps caractéristiques de variation des populations atomiques. De plus, les équations cinétiques sont incapables de fournir la valeur des dipôles atomiques, qui ne peuvent être décrits que par un formalisme prenant en compte les relations de phase entre les coefficients γ_a et γ_b du développement (2.72) de l'état atomique. Ainsi, dans l'expression (2.148) de la susceptibilité, que nous pouvons récrire

$$\chi = \frac{N_a - N_b}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega + i\Gamma_D}{\Gamma_D^2 + (\omega - \omega_0)^2}, \quad (2.163)$$

les équations cinétiques peuvent fournir le facteur $N_a - N_b$, mais les autres termes ne peuvent a priori être obtenus que par un traitement calculant les dipôles atomiques et prenant en compte leur interaction avec le rayonnement. En fait nous allons voir que les

équations cinétiques sont plus puissantes, à condition de leur associer un modèle microscopique de l'interaction entre atomes et photons qui va permettre d'écrire des équations cinétiques pour les photons.

2.5.7 Interaction atomes photons. Section efficace laser

La description par les équations cinétiques (2.160) des processus d'absorption et d'émission stimulée, schématisée sur la figure 2.18, suggère l'existence de processus élémentaires d'absorption ou d'émission stimulée, dans lesquels l'atome change de niveau tandis qu'un photon est absorbé ou au contraire est ajouté au rayonnement. Si cette description ne peut être justifiée que dans le cadre d'un traitement où le rayonnement est quantifié, elle est pourtant tellement pratique et fructueuse que nous la présentons ici. Nous admettons qu'une onde lumineuse progressive monochromatique (pulsation ω), transportant un flux de puissance par unité de surface Π (vecteur de Poynting, cf. Équation (2.61c)), correspond à un flux de photons par unité de surface, c'est-à-dire un nombre de photons passant par unité de temps dans une surface unité :

$$\Pi_{\text{phot}} = \frac{\Pi}{\hbar\omega} . \quad (2.164)$$

Nous admettons alors que la probabilité par unité de temps qu'un atome interagisse avec le rayonnement est proportionnelle à Π_{phot} : le coefficient de proportionnalité, qui a les dimensions d'une surface, s'appelle *section efficace* d'interaction, encore appelée dans ce contexte section efficace laser. On la note σ_L , et en suivant Einstein on admet qu'elle a la même valeur pour l'absorption et l'émission stimulée. Si N_a atomes sont dans l'état $|a\rangle$, le nombre de processus d'absorption par seconde vaut $N_a\sigma_L \Pi_{\text{phot}}$ et les taux de variation de populations correspondantes s'écrivent

$$\left[\frac{dN_a}{dt} \right]_{\text{abs}} = - \left[\frac{dN_b}{dt} \right]_{\text{abs}} = -N_a\sigma_L \frac{\Pi}{\hbar\omega} . \quad (2.165a)$$

De même, si on a N_b atomes dans l'état $|b\rangle$, les processus d'émission stimulée sont décrits par des taux de variation

$$\left[\frac{dN_b}{dt} \right]_{\text{sti}} = - \left[\frac{dN_a}{dt} \right]_{\text{sti}} = -N_b\sigma_L \frac{\Pi}{\hbar\omega} . \quad (2.165b)$$

Ces taux de variation ont la même forme que les expressions (2.160). En remplaçant s par son expression en fonction du vecteur de Poynting, on peut trouver la valeur de la section efficace σ_L . Ainsi pour le modèle simple étudié plus haut, on trouve

$$\sigma_L = \frac{d^2}{\hbar\epsilon_0 c} \frac{\Gamma_D \omega}{(\omega - \omega_0)^2 + \Gamma_D^2} . \quad (2.166)$$

Dans le cas d'une transition laser quelconque, on peut le plus souvent écrire des équations de taux dans lesquels les termes d'absorption et d'émission stimulée sont encore

décrits par les équations (2.165), la section efficace laser étant une donnée empirique qui présente une variation résonnante autour de ω_0 , et dont on mesure expérimentalement la valeur à résonance $\sigma_L(\omega_0)$ et la largeur de raie Γ_2 . Elle s'écrit alors

$$\sigma_L(\omega) = \sigma_L(\omega_0) \frac{1}{1 + \left(\frac{\omega - \omega_0}{\Gamma_2/2} \right)^2} . \quad (2.167)$$

Connaissant la structure des niveaux, et la façon dont ils sont alimentés, on peut écrire un système d'équations cinétiques pour les populations N_a et N_b , et chercher sa solution stationnaire. On trouve alors en général une inversion de population stationnaire présentant le phénomène de saturation, c'est-à-dire se mettant sous la forme

$$N_b - N_a = \frac{(N_b - N_a)_0}{1 + s} \quad (2.168)$$

avec

$$s = \frac{\Pi}{\Pi_{\text{sat}}} \frac{1}{1 + \left(\frac{\omega - \omega_0}{\Gamma_2/2} \right)^2} . \quad (2.169)$$

Pour notre modèle, on vérifie aisément que l'intensité de saturation s'exprime en fonction de la section efficace à résonance $\sigma_L(\omega_0)$ et de la constante de départ Γ_D par la relation

$$\Pi_{\text{sat}} = \frac{\Gamma_D \hbar \omega}{2 \sigma_L(\omega_0)} . \quad (2.170)$$

Dans les systèmes réels, l'intensité de saturation est une grandeur empirique déterminée expérimentalement.

2.5.8 Équations cinétiques pour les photons. Gain laser

Si on considère maintenant les processus d'absorption du point de vue des photons, chaque processus d'absorption supprime un photon du faisceau. Un faisceau de section S se propageant dans une tranche d'épaisseur dz (figure 2.19) interagit avec $dz S N_a/V$ atomes dans l'état $|a\rangle$, et le nombre de photons absorbés par unité de temps vaut

$$\left[\frac{d\mathcal{N}}{dt} \right]_{\text{abs}} = -dz S \frac{N_a}{V} \sigma_L \Pi_{\text{phot}} . \quad (2.171)$$

Cette quantité représente la variation du flux de photons $S \Pi_{\text{phot}}$ lors de la traversée de la tranche dz , et on peut écrire

$$\left[d(S \Pi_{\text{phot}}) \right]_{\text{abs}} = \left[\frac{d\mathcal{N}}{dt} \right]_{\text{abs}} = -dz S \frac{N_a}{V} \sigma_L \Pi_{\text{phot}} , \quad (2.172)$$

d'où la relation, remarquable de simplicité

$$\left[\frac{d\Pi_{\text{phot}}}{dz} \right]_{\text{abs}} = -\sigma_L \left(\frac{N_a}{V} \right) \Pi_{\text{phot}} , \quad (2.173a)$$

que l'on écrit habituellement en introduisant la densité $n_a = N_a/V$ des atomes absorbants

$$\Rightarrow \left[\frac{d\Pi_{\text{phot}}}{dz} \right]_{\text{abs}} = -\sigma_L n_a \Pi_{\text{phot}} . \quad (2.173b)$$

Si on a uniquement des absorbeurs, avec une densité n_a uniforme, on retrouve la décroissance exponentielle habituelle (loi de « Beer-Lambert »), caractérisée par la longueur de décroissance (à $1/e$)

$$\ell = \frac{1}{\sigma_L n_a} . \quad (2.174)$$

Dans le cas où on a uniquement des atomes dans l'état supérieur $|b\rangle$, un raisonnement analogue conduit à

$$\Rightarrow \left[\frac{d\Pi_{\text{phot}}}{dz} \right]_{\text{sti}} = \sigma_L \frac{N_b}{V} \Pi_{\text{phot}} = \sigma_L n_b \Pi_{\text{phot}} , \quad (2.175)$$

avec $n_b = N_b/V$ la densité volumique des atomes dans l'état $|b\rangle$.

Malgré sa similarité avec le raisonnement relatif à l'absorption, il faut se rendre compte que le raisonnement conduisant à (2.175) implique une hypothèse forte, à savoir que les photons émis par émission stimulée viennent se rajouter au faisceau incident avec la même direction de propagation, la même fréquence, la même polarisation, et la même phase que les photons incidents. Cette propriété des photons obtenus par émission stimulée est un résultat fondamental du traitement complètement quantique de l'interaction matière rayonnement. Ici nous considérerons qu'elle est justifiée par le fait que nous obtenons un résultat identique à celui du traitement semi-classique présenté plus haut.

Dans le cas général d'un milieu laser où les densités volumiques des atomes dans l'état inférieur $|a\rangle$ et l'état supérieur $|b\rangle$ sont respectivement n_a et n_b , on a l'équation cinétique globale

$$\Rightarrow \frac{d\Pi_{\text{phot}}}{dz} = (n_b - n_a) \sigma_L \Pi_{\text{phot}} , \quad (2.176a)$$

ou encore, en multipliant par $\hbar\omega$

$$\Rightarrow \frac{d\Pi}{dz} = (n_b - n_a) \sigma_L \Pi . \quad (2.176b)$$

Le gain par unité de longueur est donc

$$\Rightarrow g = \frac{1}{\Pi} \frac{d\Pi}{dz} = (n_b - n_a) \sigma_L . \quad (2.177)$$

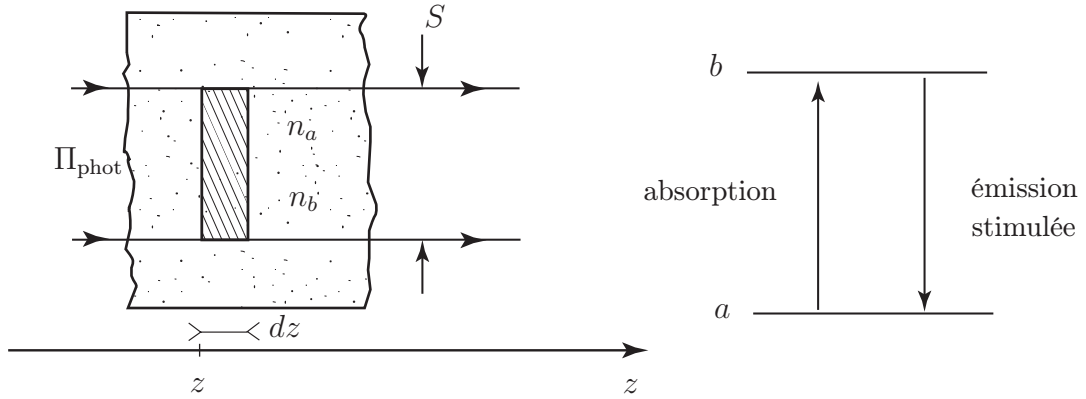


FIG. 2.19: *Description par équations cinétiques des processus d'absorption et d'émission stimulée.* Dans la tranche d'épaisseur dz , le faisceau de section S interagit avec $n_a S dz$ atomes dans l'état $|a\rangle$, et $n_b S dz$ dans l'état $|b\rangle$. Pour chaque atome dans l'état $|a\rangle$ (respectivement $|b\rangle$) le taux d'absorption (respectivement d'émission stimulée) est $\sigma_L \Pi_{phot}$, où Π_{phot} est le nombre de photons par unité de temps et de surface. Les processus d'absorption éliminent les photons de faisceau, tandis que les processus d'émission stimulée rajoutent des photons au faisceau incident.

Cette formule généralise la formule (2.149) que nous avons obtenu dans le cadre de notre modèle simple. C'est une formule fondamentale dans la modélisation des lasers.

En effet, pour chaque transition laser, on trouve des données sur la section efficace à résonance $\sigma_L(\omega_0)$ ainsi que sur la largeur de raie laser Γ_2 , et on peut en déduire $\sigma_L(\omega)$ à toute fréquence par la formule (2.167). Par ailleurs, il est généralement assez facile d'écrire des équations cinétiques permettant de calculer l'inversion de population $n_b - n_a$. En régime stationnaire, on obtient alors une expression de la forme (2.168), où le paramètre de saturation peut se mettre sous la forme (2.169), dans laquelle l'intensité de saturation Π_{sat} (qui a la dimension d'une puissance par unité de surface) est une autre donnée empirique caractérisant la transition laser.

2.6 Conclusion : modèle semi-classique ou équations cinétiques ?

Dans ce chapitre nous avons d'abord présenté le modèle semi-classique de l'interaction matière-rayonnement où l'atome est quantifié et le champ classique. Grâce à un modèle à deux niveaux dont les durées de vie sont égales, il nous a été possible de prendre en compte les phénomènes de relaxation sans avoir recours au formalisme élaboré de la matrice densité (complément II.2). En nous intéressant à la propagation d'une onde électromagnétique dans un milieu comportant un grand nombre de ces systèmes à deux niveaux, nous avons vu apparaître le phénomène d'amplification laser, à condition que le système présente une inversion de population.

Les calculs aboutissant à ce résultat sont loin d'être immédiats, et nous avons vu que l'on pouvait arriver de façon beaucoup plus simple au résultat en décrivant l'interaction atome-rayonnement par des équations cinétiques. Il s'agit d'équations de bilan des processus microscopiques d'absorption et d'émission stimulée, dans lesquels l'atome change de niveau en enlevant ou en rajoutant un photon à l'onde considérée comme un courant de photons. Ce point de vue est particulièrement utile car il permet d'écrire des équations très simples utilisant des paramètres connus des transitions laser : leur section efficace à résonance $\sigma_L(\omega_0)$, leur largeur de raie Γ_2 , leur intensité de saturation Π_{sat} . Cette simplicité ne doit pas nous faire oublier que ce résultat n'est valable que moyennant un certain nombre d'hypothèses, un fait souvent vérifiées. De plus, ce modèle basé sur des équations cinétiques avec des photons fait perdre de vue que nous avons une onde électromagnétique avec une fréquence, et une phase bien déterminées, caractéristiques essentielles des faisceaux laser dont la cohérence est une propriété cruciale. L'avantage du modèle semi-classique, qui décrit le rayonnement par une onde, est de mettre l'accent sur ces caractéristiques. Si l'on veut décrire de façon synthétique les aspects ondulatoires et corpusculaires de la lumière, il faudra recourir à un modèle quantique du rayonnement. Mais ce formalisme est nettement plus difficile à mettre en œuvre, et la modélisation des lasers est le plus souvent faite avec les outils développés dans ce chapitre.

La pratique de ces outils montre qu'en fait les phénomènes de propagation, interférence, diffraction se décrivent simplement dans le point de vue de l'électromagnétisme classique. En revanche l'interaction matière-rayonnement est plus facilement décrite par les équations cinétiques, qui mettent l'accent sur l'aspect quantique des phénomènes, et qui permettent aisément d'obtenir les expressions du gain des ondes électromagnétiques en fonction de l'inversion de population. Cette inversion de population atomique, et sa saturation quand l'intensité de l'onde augmente, peut elle aussi s'obtenir facilement à partir des équations cinétiques. Cela ne doit pas nous surprendre dans la mesure où ces deux phénomènes sont étroitement liés à la nature quantique de la matière, et sans équivalent classique simple.

À l'issue de ce survol, nous comprenons donc tout l'intérêt conceptuel du modèle semi-classique, qui met l'accent sur la cohérence des ondes sans occulter l'aspect quantique de la matière. Mais d'un point de vue pragmatique, l'introduction de la notion de photon entrant en collision avec les atomes, permet d'écrire des équations cinétiques remarquables de simplicité, et à l'interprétation limpide.

Complément II.1

Polarisation du rayonnement et transitions dipolaires électriques. Application à la résonance optique et au pompage optique

1 Règles de sélection et polarisation

1.1 Transition dipolaire électrique interdite

Comme on l'a vu dans le chapitre II, le rayonnement ne peut provoquer une transition entre deux états atomiques $|a\rangle$ et $|b\rangle$, au premier ordre de la théorie des perturbations, que dans la mesure où l'hamiltonien d'interaction $\hat{H}_I$ entre l'atome et le rayonnement a un élément de matrice non-nul entre les états $|a\rangle$ et $|b\rangle$

$$\langle b|\hat{H}_I|a\rangle \neq 0 \quad (1)$$

Lorsque cet élément de matrice est nul, c'est en général pour des raisons de symétrie fondamentales. Par exemple, deux états atomiques de même parité ne sont pas couplés par l'hamiltonien dipolaire électrique qui est impair. On le vérifie aisément dans le cas d'un atome à un seul électron de position $\mathbf{r}$, pour lequel l'hamiltonien dipolaire électrique vaut

$$\hat{H}_I = -\hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{0}, t) = -q\hat{\mathbf{r}} \cdot \mathbf{E}(\mathbf{0}, t) \quad (2)$$

En effet, en remarquant que la fonction $\psi_b^*(\mathbf{r})\psi_a(\mathbf{r})$ est une fonction paire si ψ_a et ψ_b ont la même parité, on trouve

$$\langle b|\hat{H}_I|a\rangle = -\langle b|\hat{\mathbf{D}} \cdot \mathbf{E}(\mathbf{0}, t)|a\rangle = -q\mathbf{E}(\mathbf{0}, t) \cdot \int d^3r \mathbf{r} \psi_b^*(\mathbf{r}) \psi_a(\mathbf{r}) = 0 \quad (3)$$

puisqu'on intègre une fonction impaire sur tout l'espace. On dit que *la transition entre deux niveaux de même parité est interdite* vis à vis du couplage dipolaire électrique, qui

ne pourra provoquer ni absorption ni émission de lumière. Ce résultat s'appelle une *règle de sélection*.

Remarques

- (i) L'opérateur $\hat{\mathbf{p}}$ est également un opérateur impair car chacune des composantes cartésiennes de l'opérateur gradient vérifie

$$\frac{\partial}{\partial x} = -\frac{\partial}{\partial(-x)} \quad (4)$$

Comme on pouvait s'y attendre, une transition interdite par couplage dipolaire électrique est également interdite par couplage $\mathbf{A} \cdot \mathbf{p}$.

- (ii) Rappelons que lorsque le couplage dipolaire entre deux états est nul, il ne faut pas en conclure que toute absorption ou émission de lumière par l'atome est strictement impossible. On sait que le terme d'interaction dipolaire électrique (ou le terme équivalent $\mathbf{A} \cdot \mathbf{p}$) n'est que le terme prépondérant d'un développement en série de l'hamiltonien d'interaction (§ II.B.4). Il arrive fréquemment qu'une transition interdite en couplage dipolaire électrique soit permise en couplage dipolaire magnétique (§ II.B.5) ou quadrupolaire électrique, ou par transition multiphotonique (cf. § II.A.5 et II.C.3). Mais les éléments de matrice correspondants sont plus petits par plusieurs ordres de grandeur, et nous négligeons ces phénomènes chaque fois que la transition dipolaire électrique est autorisée.

- (iii) Si le noyau est à la position $\mathbf{r}_0$ et non à l'origine des coordonnées, les raisonnements ci-dessus restent bien entendu valables, en remplaçant dans l'équation (3) $\mathbf{r}$ par $\mathbf{r} - \mathbf{r}_0$. L'équation (4) devient alors

$$\frac{\partial}{\partial x} = \frac{\partial}{\partial(x - x_0)} = -\frac{\partial}{\partial(x_0 - x)} \quad (5)$$

Nous allons voir dans ce complément d'autres règles de sélection, qui restreignent les transitions possibles entre sous-niveaux de moment cinétique déterminé, en fonction de la *polarisation de la lumière*. Les techniques d'analyse et de production de lumière polarisée étant très élaborées, ces règles de sélection ont permis de développer des méthodes expérimentales à la fois subtiles et puissantes, dans lesquelles se sont en particulier illustrés les physiciens français Jean Brossel et Alfred Kastler. La *résonance optique* donne des informations sur les atomes, par analyse de la lumière de fluorescence qu'ils réémettent à la suite d'une excitation par de la lumière quasi-résonnante. Le *pompage optique* permet de porter une assemblée d'atomes hors d'équilibre thermodynamique : on peut ainsi obtenir une inversion de population entre deux niveaux d'énergies voisines. Développée initialement pour le niveau fondamental des atomes, la méthode du pompage optique a pu être généralisée au cas de niveaux d'énergies très différentes, permettant d'obtenir une amplification laser (cf. Chapitre III).

Nous nous limiterons dans ce complément au cas d'un atome à un seul électron, où les règles de sélection sont faciles à établir. Mais il faut savoir que les résultats obtenus sont d'une portée beaucoup plus générale, car on peut les établir par des arguments de symétrie indépendants du modèle particulier utilisé ici¹.

¹A. Messiah, *Mécanique Quantique*, Dunod, Paris (1964) ; nouvelle édition : 1995. Traduction anglaise : *Quantum Mechanics*, North Holland, Amsterdam (1962) ; Wiley, New York.

1.2 Lumière polarisée linéairement

Considérons un atome à un seul électron (atome d'hydrogène), dont nous prendrons le noyau à l'origine des coordonnées. Nous supposons que cet atome interagit avec de la lumière polarisée linéairement suivant Oz , dont le champ électrique en $\mathbf{r} = \mathbf{0}$ vaut

$$\mathbf{E}(\mathbf{0}, t) = \mathbf{e}_z E_0 \cos(\omega t + \varphi) \quad (6)$$

$\mathbf{e}_z$ étant le vecteur unitaire de Oz . L'hamiltonien d'interaction se met donc sous la forme

$$\hat{H}_I(t) = \hat{W} \cos(\omega t + \varphi) \quad (7a)$$

avec

$$\hat{W} = -qE_0 \hat{\mathbf{r}} \cdot \mathbf{e}_z = -qE_0 \hat{z} = -\hat{\mathbf{D}} \cdot \mathbf{e}_z E_0 \quad (7b)$$

Étudions les transitions possibles à partir du niveau fondamental de l'atome :

$$|a\rangle = |n = 1, \ell = 0, m = 0\rangle \quad (8)$$

n étant le nombre quantique principal, l le nombre quantique orbital, et m le nombre quantique magnétique. Pour un état excité quelconque repéré par les nombres quantiques n, l, m

$$|b\rangle = |n, l, m\rangle \quad (9)$$

l'élément de matrice de transition s'écrit

$$W_{ba} = \langle b | \hat{W} | a \rangle = -qE_0 \langle b | \hat{z} | a \rangle. \quad (10)$$

Rappelons l'expression des fonctions d'onde d'un atome à un seul électron, en coordonnées sphériques² :

$$\psi_{nlm}(r, \theta, \varphi) = R_{nl}(r) Y_l^m(\theta, \varphi) \quad (11)$$

Cette expression fait apparaître les harmoniques sphériques $Y_l^m(\theta, \varphi)$. L'élément de matrice (10) prend la forme

$$W_{ba} = -qE_0 \int r^2 \sin \theta dr d\theta d\varphi R_{nl}^*(r) Y_l^{m*}(\theta, \varphi) z R_{10}(r) Y_0^0(\theta, \varphi) \quad (12)$$

En coordonnées sphériques, z s'écrit

$$z = r \cos \theta = r \sqrt{\frac{4\pi}{3}} Y_1^0(\theta, \varphi) \quad (13a)$$

et par ailleurs

$$Y_0^0(\theta, \varphi) = \frac{1}{\sqrt{4\pi}} \quad (13b)$$

²Voir par exemple BD, Chapitre XI; ou CDL 1, Chapitres VI et VII.

ce qui permet d'isoler, dans l'intégrale (12), une partie angulaire

$$I_{ang} = \iint \sin \theta d\theta d\varphi Y_l^{m*}(\theta, \varphi) Y_1^0(\theta, \varphi) \quad (14)$$

qui est tout simplement le produit hilbertien des deux harmoniques sphériques. Compte tenu des relations d'orthonormalisation des harmoniques sphériques², on a donc

$$I_{ang} = \delta_{l1} \delta_{m0} \quad (15a)$$

et

$$W_{ba} = -q \frac{E_0}{\sqrt{3}} \delta_{l1} \delta_{m0} \int_0^\infty dr r^3 R_{nl}^*(r) R_{10}(r) \quad (15b)$$

Nous trouvons donc que seuls les états

$$|b\rangle = |n, l = 1, m = 0\rangle \quad (16)$$

sont susceptibles d'être excités depuis l'état fondamental $|1, 0, 0\rangle$ par de la lumière polarisée linéairement. De même, les seuls états excités susceptibles de se désexciter vers l'état fondamental $|1, 0, 0\rangle$ sous l'effet d'une onde polarisée linéairement suivant Oz , sont de la forme (16). Les mêmes règles de sélection s'appliquent donc à l'absorption et à l'émission induite.

Si au lieu de partir de l'état fondamental on était parti d'un état $|i\rangle$ quelconque, on aurait pu démontrer, en utilisant les propriétés des harmoniques sphériques, les règles de sélection suivantes pour une transition de $|i\rangle$ vers $|k\rangle$

$$\left\{ \begin{array}{l} \text{polarisation linéaire} \\ \text{suivant l'axe de quantification} \end{array} \right\} \Rightarrow \left\{ \begin{array}{l} l_k - l_i = \pm 1 \\ m_k = m_i \end{array} \right. \quad (17)$$

Une telle transition, provoquée par la lumière polarisée linéairement, s'appelle une *transition* π . On la représente souvent par une ligne verticale dans un diagramme où les sous-niveaux Zeeman³ de même nombre quantique m sont dessinés sur une même verticale (Figure 1).

Remarques

(i) Les nombres quantiques magnétiques m_i et m_k caractérisent la composante $\hat{L}_z$ du moment cinétique atomique par rapport à l'axe Oz de quantification. La règle de sélection (17) ne s'applique que si l'axe de quantification est pris suivant la polarisation de la lumière. Si la polarisation était suivant un axe $O\xi$, la règle de sélection s'appliquerait entre sous-niveaux Zeeman relatifs à l'axe $O\xi$, c'est-à-dire entre états propres de l'opérateur

$$\hat{L}_\xi = \hat{\mathbf{L}} \cdot \mathbf{e}_\xi \quad (18)$$

où $\mathbf{e}_\xi$ est le vecteur unitaire de $O\xi$.

³On appelle niveau un ensemble d'états de même énergie. Si n et l sont fixés, les $2l + 1$ états $|n, l, m\rangle$ avec $-l \leq m \leq l$ ont la même énergie, et on les appelle « sous-niveaux Zeeman ». Ce nom rappelle que leur dégénérescence en énergie est levée sous l'action d'un champ magnétique (effet Zeeman).

(ii) Une onde polarisée linéairement suivant Oz se propage nécessairement suivant un axe perpendiculaire à Oz .

(iii) Pour un atome à plusieurs électrons, les états atomiques sont caractérisés par des nombres quantiques J et m_J associés aux valeurs propres des opérateurs de moment cinétique total, $\hat{\mathbf{J}}$ et $\hat{J}_z$. Il est possible de montrer que les règles de sélection pour de la lumière polarisée linéairement suivant Oz sont

$$\left\{ \begin{array}{c} \text{polarisation linéaire} \\ \text{suivant l'axe de quantification} \end{array} \right\} \Rightarrow \left\{ \begin{array}{c} J_k - J_i = \pm 1 \text{ ou } 0 \\ m_k = m_i \end{array} \right. \quad (19)$$

Ces règles peuvent s'interpréter en termes de conservation du moment cinétique total lors de l'absorption ou de l'émission du rayonnement. On est ainsi amené à attribuer la valeur

$$m_z = 0 \quad (20)$$

à la projection sur Oz du moment cinétique du photon d'une onde polarisée linéairement suivant Oz (une telle onde se propage perpendiculairement à Oz).

Si l'atome possède un spin nucléaire I , les règles de sélection (19) se généralisent en introduisant le moment cinétique total $\hat{\mathbf{F}} = \hat{\mathbf{J}} + \hat{\mathbf{I}}$ et sa composante $\hat{F}_z$.

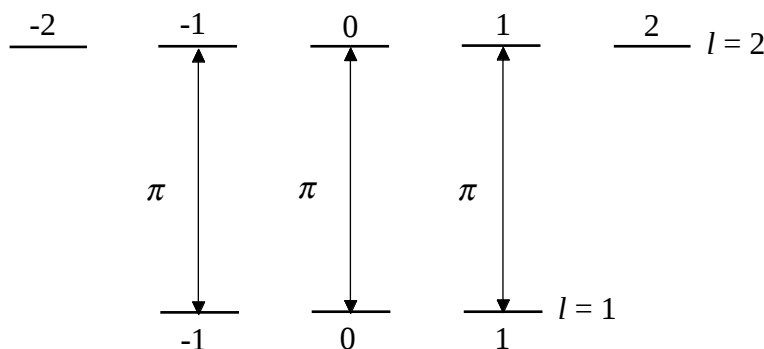


FIG. 1: **Transitions π entre un niveau $l = 1$ et un niveau $l = 2$.** Sous l'effet d'une lumière polarisée linéairement suivant l'axe de quantification, seules les transitions indiquées sont possibles. Pour un atome à plusieurs électrons, on remplace l et m_l par les nombres quantiques J et m_J du moment cinétique total $\hat{\mathbf{J}}$ et de sa composante $\hat{J}_z$ suivant la direction de polarisation de la lumière. Si l'atome possède un spin nucléaire I , c'est le moment cinétique total $\hat{\mathbf{F}} = \hat{\mathbf{I}} + \hat{\mathbf{J}}$ qui remplace $\hat{\mathbf{L}}$, et $\hat{F}_z$ qui remplace $\hat{L}_z$.

1.3 Lumière polarisée circulairement

a. Définition

Considérons maintenant une onde lumineuse dont le champ électrique en $\mathbf{r} = 0$ vaut

$$\mathbf{E} = -\frac{E_0}{\sqrt{2}}[\mathbf{e}_x \cos(\omega t + \varphi) + \mathbf{e}_y \sin(\omega t + \varphi)] \quad (21)$$

ce qui s'écrit encore

$$\mathbf{E} = \text{Re} \{ \vec{\epsilon}_+ E_0 e^{-i(\omega t + \varphi)} \} \quad (22a)$$

avec

$$\vec{\epsilon}_+ = -\frac{\mathbf{e}_x + i\mathbf{e}_y}{\sqrt{2}} \quad (22b)$$

Ce champ électrique tourne à la vitesse angulaire ω autour de Oz , dans le sens direct. Il s'agit de lumière polarisée circulairement dans le plan (xOy) , appelée *polarisation* σ_+ relativement à Oz .

La polarisation circulaire dans (xOy) tournant dans le sens rétrograde autour de Oz s'appelle *lumière polarisée* σ_- . Elle peut se définir en remplaçant dans (22a) le vecteur polarisation complexe $\vec{\epsilon}_+$ par le vecteur

$$\vec{\epsilon}_- = \frac{\mathbf{e}_x - i\mathbf{e}_y}{\sqrt{2}} \quad (22c)$$

Remarques

(i) Le champ (22a) correspond à une onde se propageant le long de Oz . Dans le cas où la propagation s'effectue vers les z positifs, on a une hélicité droite, alors que pour une propagation vers les z négatifs l'hélicité est gauche. Mais quel que soit le sens de propagation, la polarisation est circulaire σ_+ vis à vis de Oz orienté. *Il ne faut pas confondre circularité et hélicité.*

(ii) Les vecteurs unitaires $\vec{\epsilon}_+$ et $\vec{\epsilon}_-$ qui caractérisent les polarisations circulaires σ_+ et σ_- sont définis à une constante multiplicative de module 1 près. Les choix faits en (22b) et (22c) correspondent à la définition des composantes standard d'un opérateur vectoriel. Ce choix facilite l'utilisation de théorèmes généraux comme le théorème de Wigner-Eckart⁴.

b. Règles de sélection

Pour l'onde polarisée circulairement σ_+ , donnée par l'expression (22a), l'hamiltonien d'interaction dipolaire électrique (2) s'écrit

$$\hat{H}_I = \frac{qE_0}{\sqrt{2}} [\hat{x} \cos(\omega t + \varphi) + \hat{y} \sin(\omega t + \varphi)] \quad (23)$$

On peut identifier dans (23) la partie résonnante, en $e^{-i\omega t}$ et la partie antirésonnante en $e^{i\omega t}$, en écrivant

$$\hat{H}_I = \frac{1}{2} \hat{W} e^{-i(\omega t + \varphi)} + \frac{1}{2} \hat{W}^+ e^{i(\omega t + \varphi)} \quad (24a)$$

avec

$$\hat{W} = \frac{qE_0}{\sqrt{2}} (\hat{x} + i\hat{y}) = -qE_0 \hat{\mathbf{r}} \cdot \vec{\epsilon}_+ = -\hat{\mathbf{D}} \cdot \vec{\epsilon}_+ E_0 \quad (24b)$$

Notons le parallélisme étroit entre les équations (A.24) et les équations (A.7), la différence étant entièrement contenue dans le vecteur polarisation $\vec{\epsilon}_+$ qui est complexe, alors que $\mathbf{e}_z$ ne l'était pas. Ici, $\hat{W}$ n'est plus un opérateur autoadjoint, et on doit donc faire apparaître son conjugué hermitien $\hat{W}^+$.

⁴A. Messiah, cf. Ref. 1. Voir aussi par exemple CDL 2, Exercice X.8.

Les règles de sélection pour la polarisation σ_+ vont s'obtenir en étudiant l'élément de matrice de transition associé au terme résonnant de (A.24) :

$$W_{ba} = \langle b | \hat{W} | a \rangle = \frac{qE_0}{\sqrt{2}} \langle b | \hat{x} + i\hat{y} | a \rangle \quad (25)$$

Comme au paragraphe 2, nous chercherons les transitions de l'atome d'hydrogène possibles à partir du niveau fondamental $|a\rangle = |1, 0, 0\rangle$, vers un état excité quelconque $|b\rangle = |n, l, m\rangle$. La démarche est analogue : les fonctions d'onde ψ_a et ψ_b sont exprimées en fonction des harmoniques sphériques, et on utilise la relation

$$\frac{x + iy}{\sqrt{2}} = \frac{r}{\sqrt{2}} \sin \theta e^{i\varphi} = -r \sqrt{\frac{4\pi}{3}} Y_1^1(\theta, \varphi) \quad (26)$$

Le calcul explicite de W_{ba} fait alors intervenir l'intégrale angulaire

$$I_{\text{ang}} = \iint \sin \theta d\theta d\varphi Y_l^{m*}(\theta, \varphi) Y_1^1(\theta, \varphi) = \delta_{l1} \delta_{m1} \quad (27)$$

et on obtient finalement

$$W_{ba} = -\frac{aE_0}{\sqrt{3}} \delta_{l1} \delta_{m1} \int_0^\infty dr r^3 R_{nl}^*(r) R_{10}(r) \quad (28)$$

qui ressemble étroitement à (15b). La différence, essentielle, est que les seuls états excités couplés à $|1, 0, 0\rangle$ par polarisation σ_+ ont un nombre quantique magnétique $m = 1$

$$|b\rangle = |n, l = 1, m = 1\rangle \quad (29)$$

alors que ce nombre valait $m = 0$ pour une polarisation π .

On peut généraliser les résultats ci-dessus en cherchant les règles de sélection pour les transitions par absorption depuis un état $|i\rangle$ vers un état $|k\rangle$ d'énergie supérieure. On trouve

$$\left\{ \begin{array}{c} E_k > E_i \\ \text{polarisation } \sigma_+ \end{array} \right\} \Rightarrow \left\{ \begin{array}{c} l_k - l_i = \pm 1 \\ m_k = m_i + 1 \end{array} \right. \quad (30)$$

Si réciproquement on étudie la possibilité d'une transition par émission induite d'un état $|k\rangle$ vers un état $|j\rangle$ d'énergie inférieure sous l'effet d'un rayonnement polarisé circulairement σ_+ , on trouve que le nombre quantique m doit diminuer d'une unité, l variant d'une unité. En définitive, avec les notations choisies, *les règles de sélection relatives à l'émission induite prennent la même forme que celles relatives à l'absorption*, et elles restent données par les relations (30).

Habituellement, une transition σ_+ se représente par une ligne oblique montante vers la droite dans un diagramme où apparaissent les sous-niveaux Zeeman (Figure 2), l'absorption se produisant de bas en haut et l'émission induite de haut en bas.

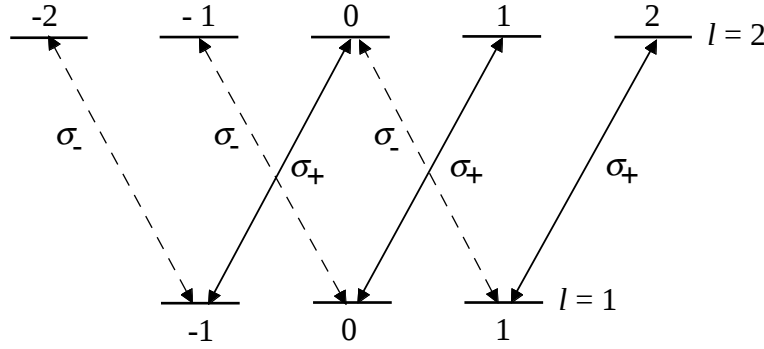


FIG. 2: **Transitions σ_+ (lignes continues) sous l'effet d'une lumière polarisée circulaire droite** (relativement à l'axe de quantification ayant permis de définir le nombre quantique m). Ces règles de sélection s'appliquent à l'absorption et à l'émission induite. De façon analogue, les lignes pointillées correspondent aux transitions σ_- permises sous l'effet d'une polarisation circulaire gauche.

Remarques

- (i) L'ensemble des considérations relatives à la polarisation σ_+ se transpose directement au cas de la lumière polarisée circulairement σ_- dont le vecteur polarisation $\vec{\epsilon}_-$ est défini en (22c). Le champ électrique correspondant s'écrit

$$\mathbf{E} = \frac{E_0}{\sqrt{2}} [\mathbf{e}_x \cos(\omega t + \varphi) - \mathbf{e}_y \sin(\omega t + \varphi)] \quad (31a)$$

$$\mathbf{E} = \text{Re} \left\{ \vec{\epsilon}_- E_0 e^{-i(\omega t + \varphi)} \right\} \quad (31b)$$

On obtient alors les règles de sélection pour une transition induite par σ_- entre des états $|i\rangle$ et $|k\rangle$, tels que l'énergie de $|k\rangle$ soit supérieure à celle de $|i\rangle$ (absorption)

$$\left\{ \begin{array}{l} E_k > E_i \\ \text{polarisation } \sigma_- \end{array} \right\} \Rightarrow \left\{ \begin{array}{l} l_k - l_i = \pm 1 \\ m_k = m_i - 1 \end{array} \right. \quad (32)$$

Une telle transition se représente par une ligne oblique montant vers la gauche (voir Figure 2).

- (ii) Pour un atome à plusieurs électrons, où les nombres J et m_J associés au moment cinétique total $\hat{\mathbf{J}}$ sont des bons nombres quantiques, les règles de sélection pour la lumière polarisée circulairement σ_+ ou σ_- autour de l'axe de quantification Oz sont

$$\left\{ \begin{array}{l} E_k > E_i \\ \text{polarisation } \sigma_+ \end{array} \right\} \Rightarrow \left\{ \begin{array}{l} J_k - J_i = \pm 1 \text{ ou } 0 \\ m_k = m_i + 1 \end{array} \right. \quad (33a)$$

(comparer à (30)), ainsi que

$$\left\{ \begin{array}{l} E_k > E_i \\ \text{polarisation } \sigma_- \end{array} \right\} \Rightarrow \left\{ \begin{array}{l} J_k - J_i = \pm 1 \text{ ou } 0 \\ m_k = m_i - 1 \end{array} \right. \quad (33b)$$

Pour un atome de moment cinétique total $\mathbf{F} = \mathbf{I} + \mathbf{J}$, les règles de sélection s'écrivent de façon analogue avec les nombres quantiques F et m_F .

Ces règles peuvent s'interpréter comme la conservation du moment cinétique lors de l'absorption d'un photon polarisé circulairement $\sigma_{\pm}$, dont la projection du moment cinétique le long de l'axe de quantification vaut $\pm \hbar$ suivant le sens de la polarisation circulaire.

- (iii) Une onde polarisée linéairement suivant une direction perpendiculaire à Oz , par exemple suivant Ox , est dite polarisée « σ linéaire ». On peut la considérer comme la superposition

d'une onde polarisée circulairement σ_+ et d'une onde polarisée circulairement σ_- , ce que traduit l'équation

$$\mathbf{e}_x = \frac{\vec{\epsilon}_- - \vec{\epsilon}_+}{\sqrt{2}} \quad (34)$$

(cf. Equations (22b) et (22c)). Par absorption depuis $|a\rangle = |1, 0, 0\rangle$ une telle polarisation excite l'atome dans une superposition d'états reproduisant la superposition (34) des polarisations σ_+ et σ_- :

$$|b\rangle = \frac{1}{\sqrt{2}}(|n, l = 1, m = -1\rangle - |n, l = 1, m = +1\rangle) \quad (35a)$$

On peut vérifier qu'il s'agit en fait de l'état

$$|b\rangle = |n, l = 1, m_x = 0\rangle \quad (35b)$$

c'est-à-dire de l'état propre associé à la valeur propre $m_x = 0$ de la composante $\hat{L}_x$ du moment cinétique $\hat{\mathbf{L}}$. Ce résultat aurait pu s'obtenir directement à partir des résultats du paragraphe (2) en choisissant l'axe Ox comme axe de quantification.

Il est important de noter que cette décomposition d'une onde polarisée linéairement en deux ondes polarisées circulairement n'est pas qu'un exercice académique. Il est souvent intéressant de prendre comme axe de quantification la direction de propagation de l'onde, et une polarisation linéaire est alors nécessairement σ et non π .

1.4 Emission spontanée

L'émission spontanée ne peut être traitée correctement que dans le cadre d'une théorie quantique du rayonnement. Nous nous contentons ici de donner les règles de sélection et les caractéristiques des diagrammes d'émission correspondants. Nous considérons d'emblée le cas d'un atome à plusieurs électrons, où les nombres J et m_J associés au moment cinétique total $\hat{\mathbf{J}}$ sont de bons nombres quantiques.

Un atome dans un état excité $|k\rangle$, de moment cinétique (J_k, m_k) (m_k est le nombre quantique associé à la composante $\hat{J}_z$ de $\hat{\mathbf{J}}$) peut se désexciter spontanément vers un état $|i\rangle$ d'énergie inférieure, caractérisé par un moment cinétique (J_i, m_i) , tel que

$$J_i - J_k = \pm 1 \text{ ou } 0 \quad (36a)$$

et

$$\begin{aligned} m_i &= m_k & (\text{transition } \pi) \\ m_i &= m_k - 1 & (\text{transition } \sigma_+) \\ m_i &= m_k + 1 & (\text{transition } \sigma_-) \end{aligned} \quad (36b)$$

Ces règles se représentent graphiquement par la réunion des figures 1 et 2.

La distribution de la lumière émise en fonction de la direction d'émission est caractérisée par un diagramme de rayonnement identique à celui du dipôle oscillant classique correspondant (complément II.3). Cette règle de correspondance illustrée ci-dessous donne aussi la polarisation de la lumière émise dans la direction considérée.

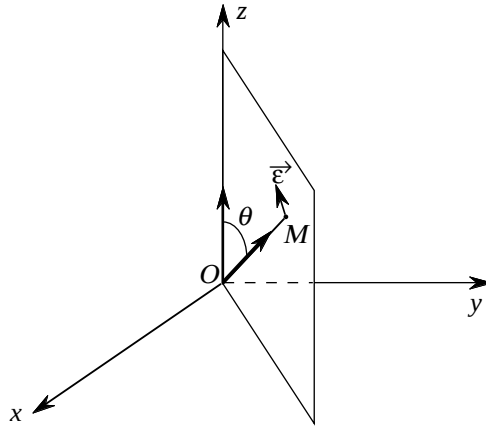


FIG. 3: **Émission spontanée pour une transition π** . Les caractéristiques sont analogues à celles du rayonnement émis par un dipôle classique oscillant suivant Oz . L'émission a lieu préférentiellement dans les directions perpendiculaires à Oz , avec une polarisation linéaire $\vec{e}$ dans le plan contenant Oz et la direction d'émission.

Appliquons les règles ci-dessus au cas d'une *transition π* (Figure 3). Le dipôle classique correspondant oscille suivant Oz . Le diagramme de rayonnement varie donc en $\sin^2 \theta$: il présente un minimum nul suivant Oz et un maximum dans toute direction perpendiculaire à Oz . La lumière émise est polarisée linéairement suivant la projection de Oz sur un plan perpendiculaire à la direction de propagation. On retrouve ce résultat dans un traitement quantique de l'émission spontanée.

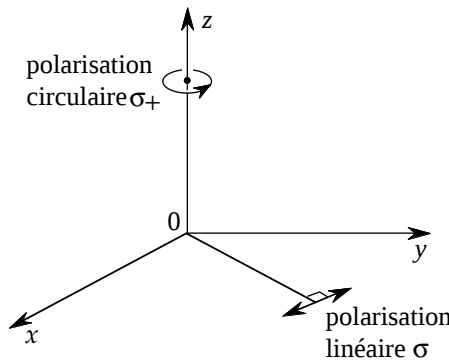


FIG. 4: **Émission spontanée $\sigma_{\pm}$** . Les caractéristiques sont analogues à celles d'un rayonnement émis par un dipôle classique tournant dans le plan (xOy) . Pour une observation suivant Oz , la polarisation est circulaire. Pour une observation perpendiculaire à Oz , la polarisation est linéaire dans le plan (xOy) perpendiculaire à Oz (polarisation appelée « linéaire σ »). Pour une direction d'observation quelconque, la polarisation est elliptique.

Dans le cas d'une transition σ_+ ou σ_- , le dipôle classique correspondant est un dipôle tournant dans le plan (xOy) , le sens de rotation étant *direct par rapport à Oz pour une transition σ_+* , et *rétrograde* pour σ_- . Pour une observation dans la direction Oz , la

lumière est polarisée circulairement dans le sens correspondant à la rotation du dipôle. Pour une observation perpendiculaire à Oz , la polarisation est linéaire perpendiculaire à Oz et à la direction d'observation. Pour une direction d'observation quelconque, on a une polarisation elliptique, dont les caractéristiques sont données par la projection du dipôle tournant sur le plan perpendiculaire à la direction d'observation.

Ces diverses caractéristiques de l'émission spontanée, et en particulier la façon dont la polarisation de la lumière émise est liée à la variation Δm_J du nombre quantique magnétique, donnent des informations très riches sur le comportement des atomes. Les parties B et C de ce complément en donnent des exemples.

2 Résonance optique

2.1 Principe de l'expérience

On dispose d'une cellule contenant une vapeur atomique, éclairée par un faisceau lumineux dont la fréquence est quasi-résonnante avec une transition connectant le niveau fondamental – que nous supposons de moment cinétique nul ($J_a = 0$) – à un niveau excité de moment cinétique $J_b = 1$. Sous l'effet de l'onde incidente, les atomes sont portés dans le niveau excité, d'où ils retombent par émission spontanée : à résonance exacte, on peut observer une forte illumination de la cellule, phénomène appelé « fluorescence de résonance ». Il est dû au processus de diffusion résonnante présenté au paragraphe II.A.4, auquel nous pouvons appliquer les règles de sélection ci-dessus en le considérant comme un processus comportant une étape d'absorption et une étape d'émission spontanée (voir la figure II.A.3, et la remarque (i) du paragraphe II.A.4).

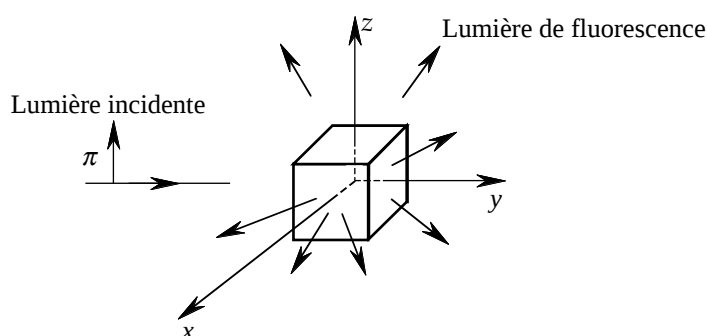


FIG. 5: **Résonance optique.** La cellule en verre, contenant une vapeur atomique, est éclairée par de la lumière résonnante avec une transition atomique partant du niveau fondamental. On observe une réémission de lumière de fluorescence dans toutes les directions.

Supposons que la lumière incidente, se propageant suivant Oy , soit polarisée suivant Oz (Fig. 5). Il s'agit d'une excitation π par rapport à l'axe Oz , qui porte les atomes

dans l'état $m_b = 0$ vis à vis de cet axe (Figure 6). La transition spontanée qui suit est également de type π , et si on observe la lumière de fluorescence suivant la direction Ox on la trouve polarisée linéairement suivant Oz , ce qui est facile à vérifier avec précision en constatant l'annulation de l'intensité transmise à travers un polariseur linéaire « croisé » avec Oz , c'est-à-dire éteignant la lumière polarisée suivant Oz . Toute modification de la polarisation de la lumière de fluorescence pourra donc être observée avec une grande sensibilité puisqu'elle se traduira par la réapparition de lumière transmise par le polariseur croisé. Nous allons donner quelques exemples d'effets que l'on peut étudier ainsi.

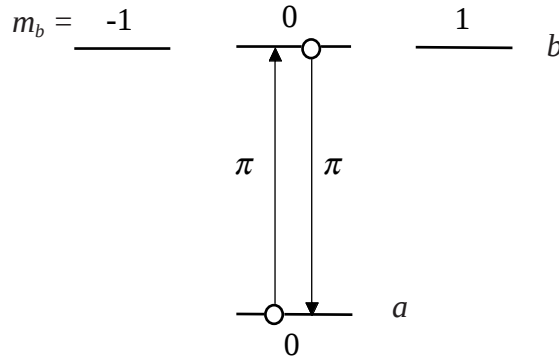


FIG. 6: **Absorption et émission spontanée** pour un atome isolé, excité par de la lumière polarisée π .

2.2 Transfert de population dans l'état excité : collisions et effet Hanle

Supposons qu'un processus physique fasse passer certains atomes excités du sous-niveau Zeeman $m_b = 0$ à un autre sous-niveau de l'état excité $m_b = \pm 1$ (Figure 7). On peut alors avoir une désexcitation spontanée σ_+ ou σ_- , ce qui s'observe par le fait que la lumière de fluorescence émise dans la direction Ox possède une composante de polarisation suivant Oy (« σ linéaire ») observable à travers le polariseur croisé avec Oz . On peut également noter l'apparition de lumière de fluorescence émise dans la direction Oz dans laquelle une transition π n'émet rien.

Les collisions sont susceptibles d'induire de tels transferts de population dans l'état excité. On observe effectivement que le rapport entre les intensités I_y et I_z des composantes de polarisation de la lumière de fluorescence suivant Ox ou Oy croît avec la pression dans la cellule de résonance. Sans rentrer dans les détails du traitement quantitatif, indiquons que le paramètre important est le nombre de collisions qu'un atome excité subit pendant une durée de vie radiative Γ_{sp}^{-1} de l'état excité. Le taux de dépolarisation I_y/I_z augmente avec ce paramètre. On obtient ainsi des informations importantes sur les potentiels d'interaction interatomiques mis en jeu lors des collisions.

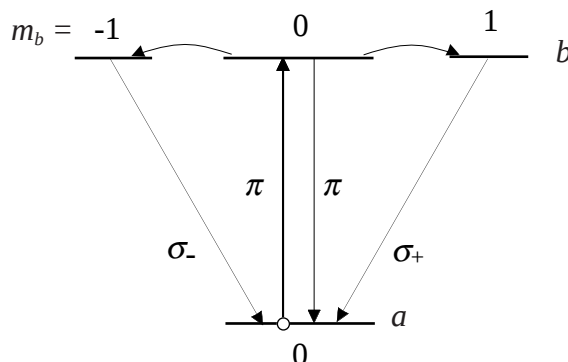


FIG. 7: **Transfert de population dans l'état excité.** Bien que l'excitation π ne peuple que $m_b = 0$, un processus (collision, champ magnétique) susceptible de provoquer un transfert de population vers $m_b = \pm 1$ conduit à une émission spontanée σ_+ ou σ_- ce qui modifie la polarisation de la lumière émise.

Un champ magnétique perpendiculaire à Oz peut également provoquer des transitions entre les sous-niveaux de l'état excité. Par exemple, un champ magnétique statique $\mathbf{B}_0$ parallèle à Ox induit des transitions entre $m_b = 0$ et $m_b = \pm 1$. Elles sont dues au terme supplémentaire qui apparaît dans l'hamiltonien atomique (terme Zeeman)

$$\hat{H}_x^{\text{Ze}} = -g \frac{q}{2m} B_0 \hat{J}_x \quad (37)$$

(g est le facteur de Landé, qui vaut 1 pour un moment cinétique purement orbital, et dont la valeur dépend de l'état considéré lorsque le moment cinétique total est une combinaison des moments orbital et de spin⁵). Cet hamiltonien Zeeman possède des éléments de matrice non-diagonaux, par exemple

$$\langle J_b = 1, m_b = 1 | H_x^{\text{Ze}} | J_b = 1, m_b = 0 \rangle = \frac{\hbar \omega_B}{\sqrt{2}} \quad (38a)$$

expression dans laquelle apparaît la *pulsation de Larmor*

$$\omega_B = -g \frac{q}{2m} B_0 \quad (38b)$$

(Il est utile de savoir que $\frac{\omega_B}{2\pi} \frac{1}{gB_0}$ vaut environ 14 MHz/mT).

Tant que l'émission spontanée ne se produit pas, la population des niveaux excités oscille périodiquement entre $m_b = 0$ et $m_b = \pm 1$. On peut alors voir apparaître dans la lumière de fluorescence une polarisation suivant Oy . Mais si l'émission spontanée se produit avant que le transfert entre sous-niveaux ait eu lieu, la lumière de fluorescence reste polarisée suivant Oz . Le paramètre pertinent est donc le rapport $\omega_B/\Gamma_{\text{sp}}$. Il peut être aisément contrôlé en choisissant la valeur B_0 du champ magnétique appliqué, et on peut ainsi observer l'apparition d'une composante I_y dans la lumière de fluorescence lorsque B_0 croît.

⁵Voir CDL 2, Complément D_x , § 3, ou Ex. G_X

Ce phénomène est l'effet Hanle, observé pour la première fois en 1923, avant l'élaboration d'une théorie quantique cohérente. Les expériences de Hanle semblent avoir eu une influence considérable sur la pensée de Heisenberg, qui était en train de développer la « mécanique des matrices ». L'effet Hanle est toujours très largement utilisé sous de multiples variantes. Il permet notamment de déterminer la durée de vie des états excités puisque la mesure du taux de polarisation de la lumière de fluorescence est reliée à $\omega_B/\Gamma_{\text{sp}}$, et que la pulsation de Larmor ω_B est connue en fonction du champ magnétique appliqué.

Remarque

Au-delà de sa simplicité technique (elle permet de déterminer des durées de vie de quelques dizaines de nanosecondes sans aucun appareil de détection rapide), cette méthode de mesure de Γ_{sp} a l'avantage d'être insensible à l'effet Doppler lié au mouvement des atomes, et de donner ainsi accès à la largeur naturelle de la raie. Il s'agit d'un avantage décisif sur les méthodes de spectroscopie ordinaires dont la résolution ultime est limitée par l'élargissement Doppler des raies atomiques.

2.3 Double résonance

Dans cette méthode sophistiquée, inventée par J. Brossel et F. Bitter en 1949, on reprend le schéma de la Figure 5 et on applique en plus un champ magnétique statique $\mathbf{B}_0$ suivant Oz . La dégénérescence entre les sous-niveaux excités est levée par effet Zeeman, puisqu'on doit rajouter à l'hamiltonien un terme

$$\hat{H}_z^{\text{Ze}} = -g \frac{q}{2m} B_0 \hat{J}_z \quad (39a)$$

Ce terme est diagonal dans la base $|J_b = 1, m_b = 1, 0, -1\rangle$, et il provoque donc un déplacement des sous-niveaux Zeeman excités

$$\Delta E_{m_b} = m_b \hbar \omega_B \quad (39b)$$

ω_B étant la pulsation de Larmor (38b). La situation est représentée sur la figure 8.

Si on éclaire les atomes par de la lumière polarisée linéairement suivant Oz , on excite sélectivement le sous-niveau $m_b = 0$, et la lumière de fluorescence réémise dans la direction Ox est polarisée linéairement suivant Oz . Il est cependant possible d'induire des transitions entre les sous-niveaux excités en appliquant une onde radiofréquence, qui interagit par couplage dipolaire magnétique (cf. § II.B.5), décrit par un hamiltonien d'interaction

$$\hat{H}_I'' = -g \frac{q}{2m} B_1 \hat{J}_x \cos \Omega t \quad (40)$$

B_1 étant l'amplitude du champ magnétique de l'onde radiofréquence supposé parallèle à Ox , et Ω étant sa fréquence. Ces transferts se produiront pour la condition de résonance

$$\Omega = \omega_B \quad (41)$$

Ils se manifesteront par l'apparition d'une polarisation I_y dans la lumière de fluorescence.

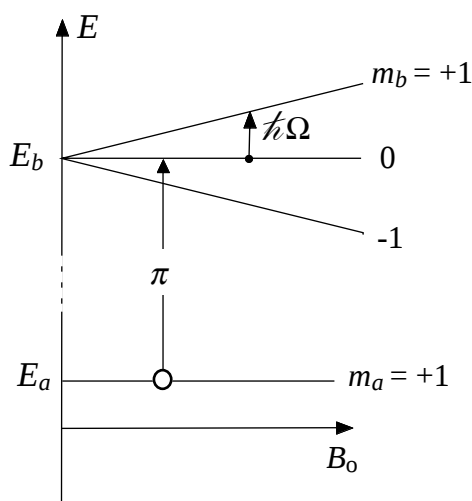


FIG. 8: **Niveaux d'énergie dans une expérience de double résonance.** Le champ magnétique $\mathbf{B}_0$ parallèle à Oz provoque une levée de dégénérescence des sous-niveaux Zeeman de l'état excité. La lumière excitatrice, polarisée suivant Oz , peuple exclusivement le sous-niveau $m_b = 0$. Par application d'un champ radiofréquence résonnant, on peut induire des transferts vers $m_b = \pm 1$, détectables par observation de la polarisation de la lumière de fluorescence.

Si on balaie le champ magnétique B_0 (ou la fréquence Ω de la radiofréquence), l'analyse de la polarisation de la lumière de fluorescence fait donc apparaître des comportements résonnants qui donnent des informations précises sur la structure de l'état excité, en particulier la valeur du facteur de Landé.

Remarque

Comme l'effet Hanle, la méthode de la double résonance présente l'avantage d'être quasiment insensible à l'effet Doppler dû au mouvement des atomes. En effet, pour une onde radiofréquence dont la fréquence $\omega/2\pi$ est typiquement de 100 MHz, et un atome se déplaçant à $v = 300$ m/s, le décalage Doppler vaut

$$\frac{\Delta\omega_D}{2\pi} = \frac{\omega}{2\pi} \frac{v}{c} = 100 \text{ Hz}$$

ce qui est négligeable devant les structures à étudier (de quelques MHz).

Le même calcul effectué pour une onde lumineuse donnerait un décalage Doppler typique de 500 MHz, grand devant les écarts entre sous-niveaux Zeeman provoqués par un champ de quelques milliteslas. On comprend donc que le peuplement sélectif du niveau $m_b = 0$ (figure 7) repose bien sur les règles de sélection liées à la polarisation de la lumière excitatrice, et non sur une résonance en fréquence qui ne saurait discriminer entre les divers sous-niveaux excités dont l'écartement est très inférieur à l'élargissement Doppler.

3 Pompage optique

L'idée de base du pompage optique, imaginé par Kastler en 1949, est d'étendre aux niveaux fondamentaux l'utilisation des règles de sélection associées à la polarisation, afin

de contrôler la distribution des population des divers sous-niveaux fondamentaux. Le nombre d'applications de cette idée est considérable, et nous nous contentons ici de présenter quelques exemples simples.

3.1 Transition $J = 1 \rightarrow J = 0$ excitée en polarisation linéaire

Considérons une transition atomique entre un niveau fondamental de moment cinétique total $J_a = 1$ et un état excité de moment cinétique total $J_b = 0$ (Figure 9). En l'absence d'onde lumineuse, les trois sous-niveaux Zeeman fondamentaux qui ont la même énergie sont équipés : $1/3$ des atomes sont dans chaque état $m_a = -1, 0, +1$:

$$\pi_1 = \pi_0 = \pi_{-1} = \frac{1}{3} \quad (42)$$

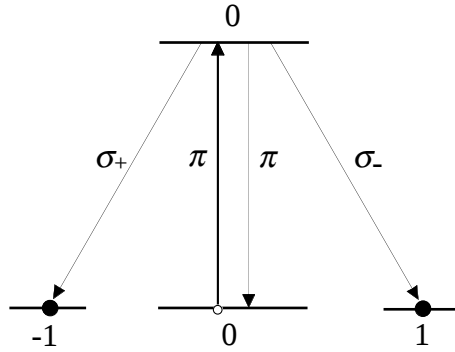


FIG. 9: **Pompage optique pour une transition $J_a = 1 \leftrightarrow J_b = 0$ excitée en polarisation linéaire.** Le sous-niveau fondamental $m_a = 0$ est vidé après quelques cycles de fluorescence.

A l'instant $t = 0$, on envoie sur la vapeur atomique un faisceau de lumière polarisée linéairement suivant Oz , dont on va mesurer l'intensité transmise. La lumière polarisée linéairement interagit avec les atomes dans le sous-niveau fondamental $m_a = 0$ (en prenant comme axe de quantification la direction Oz de polarisation), et on observe donc une absorption. Les atomes $m_a = 0$ sont portés dans l'état excité, d'où ils se désexcitent par émission spontanée. Pour une transition $1 \leftrightarrow 0$, les probabilités d'émission spontanée vers les trois sous-niveaux fondamentaux sont égales, et seulement $1/3$ des atomes excités retombent dans le sous-niveau $m_a = 0$. Au bout d'un cycle de fluorescence, la fraction des atomes dans le sous-niveau fondamental $m_a = 0$ est donc passée de $\pi_0 = \frac{1}{3}$ à $\pi_0 = \frac{1}{9}$. Le processus continuant, on tend rapidement vers la situation stationnaire

$$\pi_0(\infty) = 0 \quad (43a)$$

$$\pi_1(\infty) = \pi_{-1}(\infty) = \frac{1}{2} \quad (43b)$$

Dans cette situation, la lumière polarisée linéairement n'interagit pas avec les atomes, et elle cesse donc d'être absorbée par la vapeur atomique. En résumé, à partir de l'instant

d'application du faisceau lumineux polarisé sur la vapeur atomique, l'intensité lumineuse transmise montre d'abord une absorption, puis elle augmente et tend vers sa valeur incidente avec une constante de temps τ_p (Fig. 10). On définit le taux de pompage Γ_p comme l'inverse de τ_p .

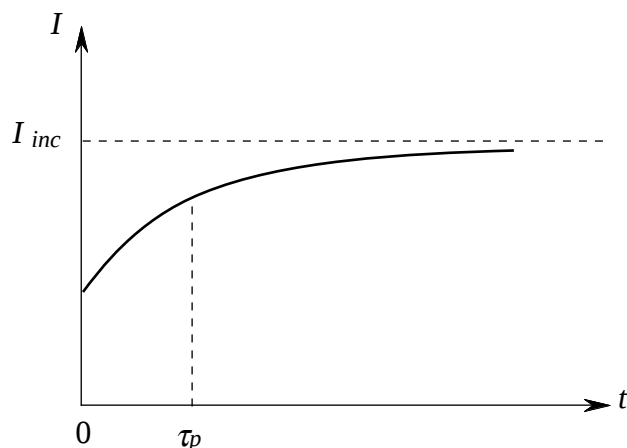


FIG. 10: **Évolution de l'intensité lumineuse transmise**, pour une onde lumineuse polarisée π envoyée à $t = 0$ sur une vapeur atomique possédant une transition $J_a = 1 \leftrightarrow J_b = 0$ résonnante. La vapeur, initialement absorbante, devient transparente au bout de quelques constantes de temps de pompage τ_p .

La situation (C.2) obtenue comme résultat du pompage optique est une situation hors d'équilibre thermodynamique. Tout processus de relaxation va provoquer un retour à l'équilibre, qui se traduit par une augmentation de l'absorption. Comme la résonance optique, les méthodes du pompage optique permettent donc d'étudier les processus de relaxation (collisions, champs magnétiques). Mais il s'agit ici de relaxation dans l'état fondamental et non dans l'état excité. La durée de vie d'un sous-niveau Zeeman de l'état fondamental étant très longue, on conçoit que l'on peut accéder à des effets très petits, ce qui donne aux méthodes de pompage optique une très grande sensibilité. Citons par exemple les magnétomètres à pompage optique, capables de mesurer des champs magnétiques de 10^{-11} Tesla (soit 10^{-6} fois le champ terrestre).

Remarques

- (i) Le raisonnement ci-dessus se transpose directement au cas d'une polarisation circulaire. Par exemple, si la lumière incidente est polarisée σ_+ , le sous-niveau fondamental $m_a = -1$ (compté suivant un axe de quantification Oz parallèle à la direction de propagation de la lumière) est vidé par pompage optique (Figure 11), et la vapeur devient transparente à la lumière polarisée σ_+ au bout de quelques temps de pompage τ_p .
- (ii) Les populations des sous-niveaux Zeeman résultent en pratique d'un équilibre entre les effets de pompage optique, qui tendent à écarter le système de l'équilibre thermodynamique, et les effets de relaxation qui l'y ramènent (par exemple les collisions). Il s'ensuit que ces populations, et par conséquent la transmission du faisceau de pompage, dépendent de l'intensité de l'onde incidente (en présence de relaxation, il faudrait une intensité infinie pour atteindre les valeurs asymptotiques (C.2)). L'intensité transmise dépend donc de façon non-

linéaire de l'intensité incidente. Cette propriété est exploitée dans de nombreuses expériences d'optique non-linéaire.

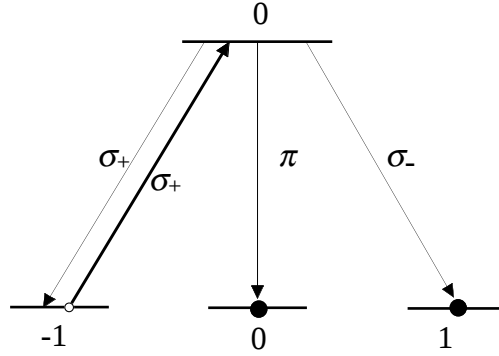


FIG. 11: **Pompage optique** pour une transition $J_a = 1 \leftrightarrow J_b = 0$ excitée en polarisation circulaire σ_+ . Le sous-niveau fondamental $m_a = -1$ est vidé au bout de quelques cycles de pompage.

3.2 Équations de pompage

Il est possible de trouver les populations stationnaires des divers sous-niveaux en présence de rayonnement et d'émission spontanée. Le principe du calcul est d'écrire des équations de pompage – encore appelées « équations cinétiques » ou « rate equations » en anglais – qui ont une structure simple, et d'en chercher les solutions. Si la justification rigoureuse des équations de pompage est généralement délicate⁶, on peut en revanche souvent les écrire facilement car ce sont des équations cinétiques qui ont une interprétation physique simple. Nous allons illustrer cette procédure sur un exemple concret.

Considérons une transition entre deux niveaux $|a\rangle$ et $|b\rangle$ ayant chacun des sous-niveaux. Pour fixer les idées, nous choisirons une transition $J_a = 1 \leftrightarrow J_b = 2$. Lorsque cette transition est soumise à une excitation par de la lumière quasi résonnante de polarisation déterminée, les populations des divers sous-niveaux évoluent vers une nouvelle situation, hors d'équilibre thermique, que nous nous proposons de déterminer.

Si l'intensité lumineuse est assez faible, on peut donner de ce problème un traitement basé sur des équations de pompage analogues aux équations présentées dans le paragraphe D.6 du chapitre II ou dans la partie B du chapitre III. Le principe en est le suivant.

Il est possible de définir pour la transition $a \leftrightarrow b$ une matrice (Γ_{ij}) qui caractérise la « force » des couplages entre divers sous-niveaux. Par exemple, pour la transition

⁶Cette étape est importante dans la mesure où toutes les situations ne relèvent pas d'équations de pompage optique. Un traitement exact doit être basé sur des Equations de Bloch Optiques (cf. Complément II.2). On peut en fait montrer que les équations de pompage sont valables dans le cas d'une excitation en raie large, ou plus généralement lorsque le taux de relaxation des cohérences est beaucoup plus grand que celui des populations.

$J_a = 1 \leftrightarrow J_b = 2$ considérée, on écrit une matrice dont les colonnes correspondant aux cinq sous-niveaux $m_b = -2, \dots, 2$, et les lignes sont relatives à $m_a = -1, 0, 1$

$$(\Gamma_{ij}) = \Gamma_{\text{sp}}(C_{ij}) = \Gamma_{\text{sp}} \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{6} & 0 & 0 \\ 0 & \frac{1}{2} & \frac{2}{3} & \frac{1}{2} & 0 \\ 0 & 0 & \frac{1}{6} & \frac{1}{2} & 1 \end{pmatrix} \quad (44)$$

Dans cette expression, Γ_{sp} est l'inverse de la durée de vie radiative de l'état excité, et les coefficients C_{ij} peuvent s'obtenir à partir du théorème de Wigner-Eckart, qui généralise les règles de sélection (les C_{ij} sont en fait des carrés de coefficients de Clebsch-Gordan). On les a représentés sur la figure 12.

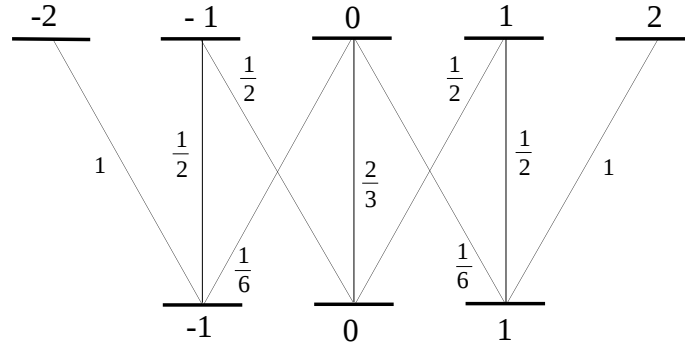


FIG. 12: **Représentation des forces relatives des diverses transitions** entre sous-niveaux Zeeman pour deux niveau $J_a = 1 \leftrightarrow J_b = 2$ (voir (44)).

Chaque coefficient Γ_{ij} donne le taux de désexcitation spontanée du sous-niveau supérieur j vers le sous-niveau inférieur i . Dans le cas particulier d'une transition $J_a \leftrightarrow J_b = J_a + 1$ considéré ici, les sous-niveaux supérieurs extrêmes ($m_b = \pm 2$) ne peuvent se désexciter que vers un seul sous-niveau inférieur à cause des règles de sélection (A.36). En revanche, un sous-niveau tel que $m_b = 0$ peut se désexciter soit vers $m_a = 0$ (transition π) avec une probabilité $2/3$, soit vers $m_a = \pm 1$ (transitions $\sigma_{\pm}$) avec une probabilité $1/6$. Dans ce cas, le taux de désexcitation du sous-niveau j est la somme de tous les taux de désexcitation Γ_{ij} . On peut alors vérifier sur (44) que

$$\sum_{i=1}^3 \Gamma_{ij} = \Gamma_{\text{sp}} \quad (45)$$

ce qui montre que tous les sous-niveaux d'un même niveau excité ont la même durée de vie radiative.

Les coefficients Γ_{ij} permettent également de calculer les taux de départ de chaque sous-niveau fondamental sous l'effet d'une onde quasi-résonnante d'intensité I . On définit

d'abord un taux de pompage global

$$\Gamma_p = \frac{\Gamma_{sp}}{2} \frac{I}{I_{sat}} \frac{1}{1 + 4 \frac{\delta^2}{\Gamma_{sp}^2}} \quad (46)$$

qui fait intervenir l'écart à résonance δ ainsi que l'intensité lumineuse incidente normalisée par l'intensité de saturation. Cette formule n'est valable que si le paramètre de saturation de la transition

$$s = \frac{I}{I_{sat}} \frac{1}{1 + 4 \frac{\delta^2}{\Gamma_{sp}^2}} \quad (47a)$$

est petit devant 1, soit $\Gamma_p \ll \Gamma_{sp}$. En pratique, on utilise souvent les éclairagements (puissance lumineuse par unité de surface, c'est-à-dire vecteur de Poynting Π) et on écrit

$$\frac{I}{I_{sat}} = \frac{\Pi}{\Pi_{sat}} \quad (47b)$$

L'éclairement de saturation Π_{sat}

$$\Pi_{sat} = \pi \frac{hc}{3} \frac{\Gamma_{sp}}{\lambda_0} \quad (47c)$$

(parfois appelé « intensité » de saturation) est typiquement de quelques milliwatts par centimètres carrés.

Remarques

(i) Pour justifier les formules ci-dessus, considérons la transition σ_+ connectant le sous-niveau $m_a = 1$ au sous niveau $m_b = 2$. Comme le montre la figure 12, la force du couplage vaut alors 1. Or un tel système constitue un système à deux niveaux pour lequel les résultats du chapitre II sont applicables. L'éclairement de saturation correspond à une situation où l'élargissement par saturation est égal à la largeur naturelle ou encore un paramètre de saturation à résonance égal à 1. D'après l'équation (2.126), cette situation est réalisée pour une pulsation de Rabi

$$(\Omega_{sat})^2 = \frac{\Gamma_{sp}^2}{2} \quad (48a)$$

L'éclairement de saturation vaut donc

$$\Pi_{sat} = \varepsilon_0 c \frac{E^2}{2} = \varepsilon_0 c \frac{(\hbar \Omega_{sat})^2}{2d^2} = \varepsilon_0 c \frac{(\hbar \Gamma_{sp})^2}{4d^2} \quad (48b)$$

Il est possible de relier d^2 et Γ_{sp} à l'aide de la formule (VI.C.21.c) qui donne

$$d^2 = \hbar \Gamma_{sp} 3\pi \varepsilon_0 \frac{c^3}{\omega_0^3} \quad (48c)$$

En reportant dans (48b), et en introduisant la longueur d'onde $\lambda_0 = 2\pi c/\omega_0$, on obtient (47c).

(ii) La formule (46) peut s'exprimer en fonction de la pulsation de Rabi Ω_1 d'une transition de force égale à 1 sous la forme

$$\Gamma_p = \frac{\Omega_1^2}{4} \frac{\Gamma_{sp}}{\delta^2 + \frac{\Gamma_{sp}^2}{4}} \quad (49)$$

Supposons que la polarisation de l'onde incidente couple les sous-niveaux i et j . Le taux de départ de i vers j vaut alors

$$(\Gamma_p)_{ij} = \Gamma_p C_{ij} \quad (50)$$

Comme les coefficients C_{ij} donnent les taux de désexcitation à partir des sous-niveaux excités (44), on peut, en combinant (44) et (50), obtenir des taux de transfert entre sous-niveaux fondamentaux, sous l'effet de l'irradiation résonnante. Par exemple, le taux de transfert entre $|a_i\rangle$ et $|a_{i'}\rangle$ sous l'effet d'une onde couplant $|a_i\rangle$ à $|b_j\rangle$ s'écrit simplement :

$$(\Gamma_p)_{ii'} = \Gamma_p C_{ij} C_{i'j} \quad (51)$$

On peut écrire des taux de transfert de ce type entre tous les sous-niveaux fondamentaux. En admettant qu'il est possible de faire un simple *bilan* additif entre les divers termes de départ et d'alimentation pour chaque sous-niveau, on obtient des *équations de pompage*.

A titre d'exemple, nous écrirons ces équations lorsque la transition $J_a = 1 \leftrightarrow J_b = 2$ est soumise à une lumière polarisée linéairement, de sorte que seules les excitations correspondant à des lignes verticales de la figure 12 peuvent se produire. En utilisant les coefficients (44), on obtient les équations d'évolution des populations des sous-niveaux fondamentaux

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_{-1}) = -\pi_{-1} C_{-1-1} C_{0-1} + \pi_0 C_{00} C_{-10} \quad (52a)$$

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_0) = -\pi_0 C_{00} (C_{-10} + C_{10}) + \pi_{-1} C_{-1-1} C_{0-1} + \pi_1 C_{11} C_{01} \quad (52b)$$

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_1) = -\pi_1 C_{11} C_{01} + \pi_0 C_{00} C_{10} \quad (52c)$$

Compte-tenu des valeurs numériques des coefficients (44), ces équations s'écrivent

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_{-1}) = -\frac{1}{4}\pi_{-1} + \frac{1}{9}\pi_0 \quad (53a)$$

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_0) = -\frac{2}{9}\pi_0 + \frac{1}{4}\pi_{-1} + \frac{1}{4}\pi_1 \quad (53b)$$

$$\frac{1}{\Gamma_p} \frac{d}{dt}(\pi_1) = -\frac{1}{4}\pi_1 + \frac{1}{9}\pi_0 \quad (53c)$$

Il est facile de résoudre ces équations à partir de conditions initiales quelconques, en tenant compte de la condition

$$\pi_{-1} + \pi_0 + \pi_1 = 1 \quad (54)$$

La solution stationnaire s'obtient immédiatement à partir des équations (C.12) et (54) dans lesquelles le premier membre est pris égal à zéro. On trouve :

$$\pi_{-1}(\infty) = \pi_1(\infty) = \frac{4}{17} \quad (55a)$$

$$\pi_0(\infty) = \frac{9}{17} \quad (55b)$$

Le régime transitoire n'est guère plus difficile à obtenir. C'est une combinaison d'exponentielles amorties dont les inverses des constantes de temps s'obtiennent en diagonalisant la matrice associée aux deuxièmes membres de (C.12)

$$\begin{pmatrix} -\frac{1}{4} & \frac{1}{9} & 0 \\ \frac{1}{4} & -\frac{2}{9} & \frac{1}{4} \\ 0 & \frac{1}{9} & -\frac{1}{4} \end{pmatrix} \quad (56)$$

ce qui donne :

$$\begin{aligned} \Gamma' &= 0 \\ \Gamma'' &= 0,10 \Gamma_p \\ \Gamma''' &= 0,63 \Gamma_p \end{aligned} \quad (57)$$

L'existence d'une valeur propre nulle reflète simplement l'existence d'une « constante du mouvement », à savoir la somme des populations des sous-niveaux fondamentaux (Equation (54)).

L'exemple ci-dessus nous a montré comment les équations de pompage permettent d'obtenir simplement la valeur des populations au cours du temps. Il se généralise aisément à de nombreuses situations.

Remarques

(i) Les équations de pompage ne sont justifiées que lorsque la polarisation de l'onde est telle qu'elle couple un sous-niveau fondamental à un seul sous-niveau excité. Par ailleurs, rappelons qu'il est nécessaire d'être dans un cas où l'intensité lumineuse est non saturante.

(ii) Le traitement ci-dessus met en évidence le fait que l'on peut avoir des constantes de temps d'évolution très longues devant Γ_{sp}^{-1} si I est très petit devant I_{sat} .

(iii) Si on écrit les équations de pompage pour la situation de la figure 9, on trouve comme solution stationnaire $\pi_0(\infty) = 0$, ce qu'on avait deviné intuitivement.

De même, l'excitation de la transition $J_a = 1 \leftrightarrow J_b = 2$ de la figure 12 par une lumière polarisée circulaire σ_+ aboutit à l'accumulation de toute la population dans le sous-niveau $m_a = +1$, de sorte que l'on doit trouver $\pi_1(\infty) = 1$.

Il est fréquent de pouvoir deviner ainsi, par des arguments simples, l'état stationnaire d'une situation de pompage optique.

(iv) Dans de nombreuses situations, les quantités importantes pour l'expérience ne sont pas les populations des différents sous-niveaux Zeeman mais seulement quelques grandeurs moyennes comme

$$\frac{\langle \hat{J}_z \rangle}{\hbar} = \sum_{m_a} m_a \pi_{m_a} \quad (58)$$

qui est appelée *orientation*, ou encore

$$\frac{\langle \hat{J}_z^2 \rangle}{\hbar^2} - \frac{1}{3} \frac{\langle \hat{\mathbf{J}}^2 \rangle}{\hbar^2} = \sum_{m_a} m_a^2 \pi_{m_a} - \frac{J_a(J_a + 1)}{3} \quad (59)$$

qui est appelée *alignement*.

En l'absence de pompage optique, quand tous les sous-niveaux Zeeman sont également peuplés, l'orientation aussi bien que l'alignement sont nuls ; le système a une symétrie sphérique. L'effet du pompage optique est de rendre ce type de grandeur non nulle : on brise alors la symétrie sphérique.

Complément II.2

Matrice densité et équations de Bloch optiques

Les développements du chapitre II, ainsi que des chapitres ultérieurs, reposent sur le formalisme du vecteur d'état et l'équation de Schrödinger. Une telle approche est mal adaptée aux situations où le couplage entre l'atome et son environnement (collisions avec d'autres atomes, émission spontanée de photons dans des modes initialement vides...) ne peut pas être négligé. Si on ne s'intéresse pas aux corrélations apparaissant entre l'atome et son environnement par suite de ces interactions, mais seulement à l'évolution de l'atome, il existe un formalisme quantique bien adapté, celui de la *matrice densité*. Ce formalisme permet de décrire à chaque instant l'état de l'atome, dans une situation où il n'existe pas de vecteur d'état pour l'atome seul.

Dans la partie 1, nous présentons succinctement le formalisme de la matrice densité, et nous montrons comment on rend compte des effets de l'environnement sur l'atome en introduisant des termes de *relaxation* dans les équations d'évolution. Dans la partie B, nous montrons comment la théorie des *perturbations* introduite dans le chapitre I pour l'équation de Schrödinger peut être généralisée au cas de la matrice densité. Nous traitons dans ce cadre l'interaction d'un atome avec un champ électromagnétique d'intensité modérée, ce qui permet d'obtenir les susceptibilités linéaires pour une vapeur atomique. Il existe un autre cas où les équations d'évolution de la matrice densité sont suffisamment simples pour donner des résultats analytiques utilisables, c'est celui où deux niveaux atomiques seulement jouent un rôle dans l'interaction avec le rayonnement (Partie C). On obtient alors un ensemble d'équations non-perturbatives, appelées équations de Bloch optiques, qui sont exactement solubles même lorsque le rayonnement est très intense, et qui permettent de décrire une grande variété de phénomènes physiques.

1 Fonction d'onde et matrice densité

1.1 Système isolé et système en interaction

L'évolution d'un système isolé, dont l'état à l'instant $t = 0$ est $|\psi_0\rangle$, est déterminée par l'équation de Schrödinger

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = \hat{H} |\psi(t)\rangle \quad (1a)$$

avec comme condition initiale

$$|\psi(0)\rangle = |\psi_0\rangle \quad (1b)$$

Si le système est composé de deux sous-systèmes A et B , dont les états à l'instant $t = 0$ sont respectivement $|\psi_A\rangle$ et $|\psi_B\rangle$, la résolution de l'équation de Schrödinger avec la condition initiale $|\psi(0)\rangle = |\psi_A(0)\rangle \otimes |\psi_B(0)\rangle$ donne généralement un état $|\psi(t)\rangle$ qu'il n'est pas possible de factoriser en un produit de deux états correspondant aux sous-systèmes A et B

$$|\psi(t)\rangle \neq |\psi_A(t)\rangle \otimes |\psi_B(t)\rangle$$

L'interaction entre les systèmes A et B a créé *des corrélations quantiques* qui vont subsister même en l'absence d'interaction ultérieure entre A et B . Si l'on doit ensuite étudier l'interaction entre le système A et un autre système C , il n'est plus possible de postuler que l'état initial du système est de la forme $|\psi'_A\rangle \otimes |\psi_C\rangle$ et il est, en toute rigueur, nécessaire de travailler dans l'espace produit tensoriel des états relatifs aux systèmes A , B et C . Lorsque le système A interagit ainsi successivement avec plusieurs systèmes B , C , D ... X , l'espace à considérer croît à chaque étape par suite des corrélations quantiques accumulées lors de chacune des interactions. Ainsi, pour un atome A subissant des collisions avec d'autres atomes au rythme d'une collision toutes les microsecondes, la fonction d'onde décrivant l'atome A au bout d'une seconde doit également prendre en compte l'état du million d'atomes qu'il a rencontrés. Dans cette situation, l'ambition de décrire l'état d'un système A au moyen d'un vecteur d'état va se trouver très rapidement limitée par des problèmes évidents d'encombrement !

1.2 Description en terme de matrice densité

En pratique, pour la plupart des problèmes, l'information contenue dans la fonction d'onde globale décrivant le système et son environnement est surdimensionnée pour les prévisions physiques relatives au seul système. Plutôt que de décrire chaque interaction, on se contente alors de décrire leur effet moyen sur le système A . Une telle *approche statistique* nécessite l'emploi de *l'opérateur densité* $\hat{\sigma}$ plutôt qu'un vecteur d'état pour décrire le système A . Nous présentons brièvement quelques propriétés de l'opérateur densité ¹.

¹Pour plus de détails, voir le cours de Physique de Roger Balian, Chapitre 2 ; ou CDL I Chapitre III, Complément E.

Dans le cas – appelé « cas pur » – où l'état du système A peut être décrit au moyen d'un vecteur d'état $|\psi_A\rangle$, son opérateur densité est le projecteur $\hat{\sigma} = |\psi_A\rangle\langle\psi_A|$. L'objet mathématique $\hat{\sigma}$ est un opérateur hermitien agissant dans l'espace des états du système A . Cette propriété reste vraie dans le cas général où l'opérateur densité n'est plus un projecteur. Dans chaque base, l'opérateur densité s'exprime sous forme de matrice, la matrice densité : par abus de langage, on emploie souvent indifféremment les noms « matrice densité » et « opérateur densité ».

La probabilité de trouver le système dans un état $|i\rangle$ est égale à $\sigma_{ii} = \langle i|\hat{\sigma}|i\rangle$. (Cette propriété générale est manifestement vraie pour un cas pur). La probabilité de trouver le système dans un état quelconque de l'espace étant normalisée à 1, ceci entraîne que

$$\text{Tr}\hat{\sigma} = \sum_i \sigma_{ii} = 1 \quad (2)$$

L'élément diagonal σ_{ii} s'appelle *population* de l'état $|i\rangle$.

Pour un cas pur, on vérifie aisément qu'une observable $\hat{O}$ du système a pour valeur moyenne

$$\langle\hat{O}\rangle = \text{Tr}(\hat{\sigma}\hat{O}) \quad (3)$$

Cette propriété reste vraie dans le cas général.

L'évolution de la matrice densité d'un système A isolé, initialement dans un cas pur, peut se déduire de l'équation de Schrödinger (1a). Elle est décrite par l'équation :

$$\frac{d\hat{\sigma}}{dt} = \frac{1}{i\hbar}[\hat{H}, \hat{\sigma}] \quad (4)$$

qui fait intervenir le commutateur entre l'hamiltonien du système et la matrice densité. Cette équation d'évolution est également valable pour un système décrit initialement par une matrice densité quelconque, et qui est soumis à des interactions descriptibles par un hamiltonien (ou qui reste isolé).

Supposons que le système A subisse de surcroît, à des instants aléatoires mais fréquents, de nombreuses collisions brèves et peu intenses² avec des systèmes $B, C \dots X \dots$. L'effet *moyen* des collisions peut être décrit au moyen d'un opérateur de relaxation que l'on rajoutera à l'équation (4). Les termes de relaxation relatifs aux populations σ_{ii} s'écrivent :

$$\left\{ \frac{d}{dt} \sigma_{ii} \right\}_{\text{relax}} = - \left(\sum_{j \neq i} \Gamma_{i \rightarrow j} \right) \sigma_{ii} + \sum_{j \neq i} \Gamma_{j \rightarrow i} \sigma_{jj} \quad (5a)$$

Cette équation exprime que la population σ_{ii} de l'état i diminue par suite des transferts vers les autres états j , et croît par suite des transitions effectuées des autres états j vers l'état i . On notera que cette équation assure la conservation de la population totale $\sum_i \sigma_{ii}$.

²Pour approfondir la théorie de la relaxation, voir par exemple CDG 2, Chapitre IV. Des équations de la forme (A.5) peuvent, en particulier, être obtenues si le système évolue peu pendant la durée d'une collision.

Pour les éléments non-diagonaux $\sigma_{ij} = \langle i|\hat{\sigma}|j\rangle$ (que l'on appelle *cohérences* car ils prennent en compte la phase relative des composantes relatives aux deux états $|i\rangle$ et $|j\rangle$), on écrit les termes de relaxation suivants :

$$\left\{ \frac{d}{dt} \sigma_{ij} \right\}_{\text{relax}} = -\gamma_{ij} \sigma_{ij} \quad (5b)$$

Notons que si les coefficients $\Gamma_{i \rightarrow j}$ sont réels (et positifs), rien n'empêche les coefficients γ_{ij} d'être complexes (ils vérifient néanmoins la relation $\gamma_{ij} = \gamma_{ji}^*$). Les coefficients $\Gamma_{i \rightarrow j}$ et γ_{ij} peuvent être calculés lorsque l'hamiltonien d'interaction entre le système A et son environnement est connu. Ici nous les considérerons comme des coefficients phénoménologiques qu'il est possible de déterminer expérimentalement.

Remarque

L'équation de relaxation des cohérences (5b) suppose implicitement que toutes les fréquences de Bohr du système sont différentes. Dans le cas où deux ou plusieurs fréquences de Bohr coïncideraient, il peut être nécessaire de rajouter à (5b) des termes couplant les cohérences associées à ces fréquences. Ces termes de transfert de cohérence sont importants chaque fois qu'une symétrie de l'hamiltonien entraîne l'égalité exacte de plusieurs fréquences de Bohr. Un exemple frappant est celui de l'oscillateur harmonique où les transferts de cohérence sous l'effet de la relaxation jouent un rôle essentiel dans la dynamique³.

1.3 Systèmes à deux niveaux

Beaucoup de situations physiques peuvent être modélisées par un système quantique à deux niveaux⁴. Il est donc essentiel de savoir écrire l'effet de la relaxation dans un tel système. Nous décrirons successivement le cas d'un *système fermé* qui correspond à une situation où le niveau d'énergie le plus bas est le niveau fondamental stable, et le cas d'un *système ouvert* qui correspond à une situation où les deux niveaux sont des niveaux excités instables, et qui décrit la plupart des transitions laser.

a. Système fermé

Considérons un atome à deux niveaux, le niveau fondamental étant appelé a et le niveau excité b . Dans le cas où la seule cause de relaxation est l'émission spontanée (voir le chapitre II pour une présentation phénoménologique), les équations (5a) s'écrivent

$$\left\{ \frac{d}{dt} \sigma_{bb} \right\}_{\text{relax}} = -\Gamma_{\text{sp}} \sigma_{bb} \quad (6a)$$

$$\left\{ \frac{d}{dt} \sigma_{aa} \right\}_{\text{relax}} = \Gamma_{\text{sp}} \sigma_{bb} \quad (6b)$$

où Γ_{sp}^{-1} est la durée de vie radiative du niveau excité. L'équation (6a) décrit bien la décroissance exponentielle de la population du niveau excité avec une durée de vie Γ_{sp}^{-1} .

³Voir CDG 2, Complément B IV, § 3.

⁴R.P. Feynmann, R. Sands, and R. Leighton, « Lectures on Physics », Vol. III, Chap. 11. Voir aussi CDL 1, Chapitre IV.

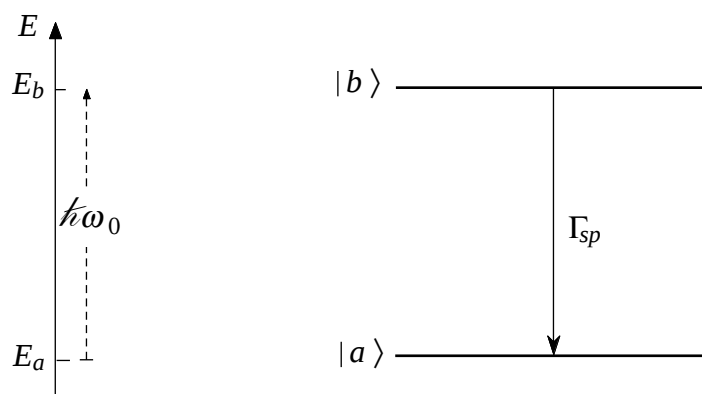


FIG. 1: **Système à deux niveaux fermé.** La relaxation du niveau b aboutit nécessairement au niveau inférieur a . La somme des populations des deux niveaux est constante.

L'équation (6b) montre que tout atome quittant le niveau b aboutit nécessairement dans le niveau a . La somme des populations est donc constante : vis-à-vis de la relaxation, ce système à deux niveaux est fermé.

En ce qui concerne la cohérence σ_{ba} la relaxation s'écrit

$$\left\{ \frac{d\sigma_{ba}}{dt} \right\}_{\text{relax}} = -\gamma\sigma_{ba} \quad (6c)$$

avec $\gamma = \Gamma_{\text{sp}}/2$ lorsque la relaxation est due exclusivement à l'émission spontanée.

Remarques

(i) Cet exemple illustre la démarche conduisant à l'utilisation du formalisme de la matrice densité. L'espace des états où évolue le système global formé de l'atome et du photon spontané émis fait intervenir les degrés de liberté du rayonnement, mettant en jeu un très grand espace. En fait, les équations (6a) résultent d'une *trace partielle* sur les variables du rayonnement⁵, ce qui permet de se ramener à un espace de dimension deux (atome). L'inconvénient de cette démarche est de faire disparaître du formalisme les corrélations pouvant exister entre l'atome et les photons spontanés émis (il existe en fait des méthodes permettant de remonter à certaines corrélations).

(ii) Il arrive souvent que l'émission spontanée ne soit pas la cause unique de relaxation. Nous avons déjà cité un autre phénomène rencontré dans les vapeurs atomiques, la *relaxation collisionnelle*. Considérons des atomes à deux niveaux interagissant avec des atomes perturbateurs. Il peut arriver que les collisions induisent des transferts de population b vers a (« quenching ») ou de a vers b (excitation collisionnelle) mais ces processus sont fréquemment négligeables devant l'émission spontanée de sorte que la relaxation des populations reste correctement décrite par les équations (6a) et (6b). En revanche, l'évolution des cohérences peut être fortement perturbée parce que durant chaque collision la fréquence de Bohr ω_{ba} change, par suite de l'effet du potentiel d'interaction entre l'atome et le perturbateur. Ces collisions appelées *collisions déphasantes* conduisent à un terme supplémentaire dans le taux de relaxation de la cohérence, qui s'écrit :

$$\gamma = \frac{\Gamma_{\text{sp}}}{2} + \gamma_{\text{coll}} \quad (7a)$$

⁵Voir CDG 2, Chapitre IV Parties A et E.

Le premier terme $\Gamma_{\text{sp}}/2$ est relatif à l'émission spontanée, tandis que γ_{coll} est associé à la relaxation collisionnelle. Ce dernier terme est proportionnel au nombre d'atomes perturbateur par unité de volume $\frac{N}{V}$, et à leur vitesse v

$$\gamma_{\text{coll}} = \frac{N}{V} \sigma_{\text{coll}} v \quad (7b)$$

Pour un potentiel d'interaction donné, la section efficace de collisions déphasantes σ_{coll} (ne pas faire de confusion avec un élément de la matrice densité!) ne dépend que de la température, et le taux de relaxation collisionnelle varie donc proportionnellement à la pression, à température fixée.

b. Système ouvert

Supposons à présent que les deux niveaux b et a sont deux niveaux excités, le niveau a ayant une énergie inférieure au niveau b . La relaxation des populations de ces niveaux s'écrit :

$$\left\{ \frac{d}{dt} \sigma_{bb} \right\} = -\Gamma_b \sigma_{bb} - \Gamma_{b \rightarrow a} \sigma_{bb} \quad (8a)$$

$$\left\{ \frac{d}{dt} \sigma_{aa} \right\}_{\text{relax}} = -\Gamma_a \sigma_{aa} + \Gamma_{b \rightarrow a} \sigma_{bb} \quad (8b)$$

Dans ces équations, Γ_a et Γ_b décrivent la relaxation vers l'extérieur, et $\Gamma_{b \rightarrow a}$ décrit la retombée du niveau b au niveau a (qui peut être due à l'émission spontanée). L'évolution de la cohérence σ_{ba} s'écrit

$$\left\{ \frac{d}{dt} \sigma_{ba} \right\}_{\text{relax}} = -\gamma \sigma_{ba} \quad (9a)$$

avec, dans le cas d'une relaxation due exclusivement à l'émission spontanée

$$\gamma = \gamma_{\text{sp}} = \left(\frac{\Gamma_b^{\text{sp}}}{2} + \frac{\Gamma_a^{\text{sp}}}{2} \right) \quad (9b)$$

Contrairement aux équations (A.6) qui conservent la population globale (et qui sont donc compatibles avec la condition $\text{Tr} \hat{\sigma} = 1$), les équations (8a) conduisent à une variation de la population totale des niveaux a et b avec le temps. Un tel système physique apparaît donc comme un sous-système d'un ensemble plus vaste (d'où le nom de système ouvert). Ainsi, nous pouvons considérer que les niveaux excités a et b se vident vers un niveau fondamental f , à partir duquel ils peuvent aussi être alimentés par un mécanisme de pompage. Nous supposons que ce pompage, qui peut être produit par un bombardement électronique, par une source lumineuse annexe, ou tout autre mécanisme (voir Chapitre III, partie B), est décrit par les équations d'alimentation suivantes :

$$\left\{ \frac{d}{dt} \sigma_{bb} \right\}_{\text{alim}} = \tilde{\Lambda}_b \sigma_{ff} \quad (10a)$$

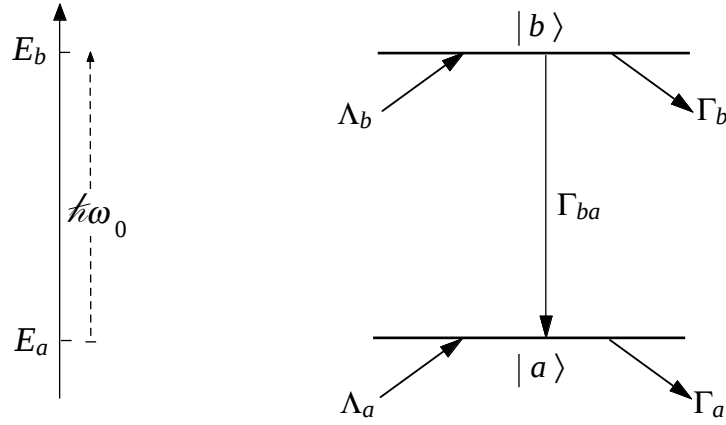


FIG. 2: **Système à deux niveaux ouvert.** Les deux niveaux a et b peuvent se désexciter vers un ou plusieurs niveaux d'énergie inférieure. Ils peuvent également être alimentés à partir de l'extérieur. La somme des populations des deux niveaux a et b n'est pas constante.

$$\left\{ \frac{d}{dt} \sigma_{aa} \right\}_{\text{alim}} = \tilde{\Lambda} \sigma_{ff} \quad (10b)$$

Dans la pratique, les taux d'alimentation⁶ $\tilde{\Lambda}_a$ et $\tilde{\Lambda}_b$ sont généralement petits devant Γ_a et Γ_b de sorte que les populations dans les niveaux excités restent très petites devant la population du niveau fondamental. On a donc $\sigma_{ff} \approx 1$, ce qui permet de remplacer les équations (10a) par les expressions approchées

$$\left\{ \frac{d}{dt} \sigma_{bb} \right\}_{\text{alim}} \approx \tilde{\Lambda}_b \quad (11a)$$

$$\left\{ \frac{d}{dt} \sigma_{aa} \right\}_{\text{alim}} \approx \tilde{\Lambda}_a \quad (11b)$$

En revanche, le pompage ne crée généralement pas de cohérence entre deux niveaux électroniques différents, de sorte qu'il n'y a pas pour σ_{ba} de terme source dû au pompage.

Remarques

(i) Malgré sa ressemblance avec le modèle utilisé dans les parties 2.4 et 2.5 du chapitre 2, le système considéré ici en diffère par plusieurs points. D'une part les taux de relaxation des niveaux a et b ne sont pas égaux. D'autre part on considère ici un terme de relaxation interne de b vers a . Dans cette situation beaucoup plus générale, le recours au formalisme des équations de Bloch optiques est indispensable.

(ii) Le système à deux niveaux fermé apparaît comme un cas particulier du système à deux niveaux ouvert, où l'on fait $\Gamma_a = \Gamma_b = 0$ et $\tilde{\Lambda}_a = \tilde{\Lambda}_b = 0$.

⁶Les taux d'alimentation $\tilde{\Lambda}_a$ et $\tilde{\Lambda}_b$ introduits ici sont relatifs à un seul atome. Ils sont reliés aux taux d'alimentation Λ_i relatifs à une assemblée de N atomes, utilisés au chapitre II (Éq. II.D.2), par : $\Lambda_i = N\tilde{\Lambda}_i$.

2 Traitement perturbatif

2.1 Résolution par itération de l'équation d'évolution de la matrice densité

a. Présentation du problème

Le but de cette partie est de généraliser au formalisme de la matrice densité la méthode des perturbations dépendant du temps établie dans le chapitre I pour l'évolution des vecteurs d'état. Comme au chapitre I, nous supposons que l'hamiltonien se décompose en un terme principal $\hat{H}_0$ et un hamiltonien de perturbation $\hat{H}_1(t)$. L'hamiltonien $\hat{H}_0$ indépendant du temps a pour états propres les états $|n\rangle$ d'énergie E_n . Pour faire ressortir les différents ordres du développement en série de perturbations, le terme de perturbation $\hat{H}_1(t)$ sera écrit, comme dans la formule (1.10) du chapitre I, sous la forme $\lambda\hat{H}'_1(t)$, le paramètre réel λ étant supposé très petit alors que $\hat{H}'_1(t)$ est de l'ordre de $\hat{H}_0$.

L'équation d'évolution de la matrice densité qui se substitue à l'équation de Schrödinger (Equation (1.11) du chapitre I), est de la forme :

$$\frac{d\hat{\sigma}}{dt} = \frac{1}{i\hbar}[\hat{H}_0 + \hat{H}_1(t), \hat{\sigma}] + \left\{ \frac{d\hat{\sigma}}{dt} \right\} \quad (12)$$

Le premier terme est le commutateur (4) traduisant l'évolution hamiltonienne. Le terme $\left\{ \frac{d\hat{\sigma}}{dt} \right\}$ est la somme des termes de relaxation et de pompage relatifs aux populations et aux cohérences. Utilisons la notation $\hat{H}_1(t) = \lambda\hat{H}'_1(t)$ dans l'équation (12)

$$\frac{d\hat{\sigma}}{dt} = \frac{1}{i\hbar}[\hat{H}_0, \hat{\sigma}] + \left\{ \frac{d\hat{\sigma}}{dt} \right\} + \frac{\lambda}{i\hbar}[\hat{H}'_1(t), \hat{\sigma}] \quad (13)$$

et introduisons le développement suivant de $\hat{\sigma}$:

$$\hat{\sigma} = \hat{\sigma}^{(0)} + \lambda\hat{\sigma}^{(1)} + \lambda^2\hat{\sigma}^{(2)} + \dots \quad (14)$$

qui généralise le développement (1.16) du chapitre I. En reportant (14) dans (13) et en identifiant les termes de même ordre en λ , nous trouvons

– à l'ordre 0

$$\frac{d\hat{\sigma}^{(0)}}{dt} - \frac{1}{i\hbar}[\hat{H}_0, \hat{\sigma}^{(0)}] - \left\{ \frac{d\hat{\sigma}^{(0)}}{dt} \right\} = 0 \quad (15a)$$

– à l'ordre 1

$$\frac{d\hat{\sigma}^{(1)}}{dt} - \frac{1}{i\hbar}[\hat{H}_0, \hat{\sigma}^{(1)}] - \left\{ \frac{d\hat{\sigma}^{(1)}}{dt} \right\} = \frac{1}{i\hbar}[\hat{H}'_1(t), \hat{\sigma}^{(0)}] \quad (15b)$$

– à l'ordre r

$$\frac{d\hat{\sigma}^{(r)}}{dt} - \frac{1}{i\hbar}[\hat{H}_0, \hat{\sigma}^{(r)}] - \left\{ \frac{d\hat{\sigma}^{(r)}}{dt} \right\} = \frac{1}{i\hbar}[\hat{H}'_1(t), \hat{\sigma}^{(r-1)}] \quad (15c)$$

b. Solution à l'ordre 0

A l'ordre 0, et dans la base des états propres de $\hat{H}_0$, les équations (15a) pour les populations $\sigma_{jj}^{(0)}$ se ramènent aux équations de relaxation (5a). Les solutions stationnaires de ces équations donnent les populations des divers états en l'absence de perturbation. (En particulier, pour un système à l'équilibre thermodynamique, les coefficients $\Gamma_{i \rightarrow j}$ sont tels que les populations $\sigma_{jj}^{(0)}$ correspondent à une distribution de Boltzmann.)

Quant à la cohérence $\sigma_{jk}^{(0)}$, elle est, d'après (15a) et (5b), solution de l'équation

$$\frac{d}{dt}\sigma_{jk}^{(0)} + i\omega_{jk}\sigma_{jk}^{(0)} + \gamma_{jk}\sigma_{jk}^{(0)} = 0 \quad (16)$$

où on a introduit la fréquence de Bohr pour la transition $j \rightarrow k$:

$$\omega_{jk} = (E_j - E_k)/\hbar$$

Si toutes les cohérences $\sigma_{jk}^{(0)}(t)$ sont nulles à l'instant initial, elles le restent au cours du temps. Si ce n'est pas le cas, elles présentent des oscillations amorties avec des constantes de temps égales à $(\text{Re } \{\gamma_{jk}\})^{(-1)}$ (où $\text{Re } \{\}$ signifie « partie réelle de »).

Nous supposons dans la suite que $\hat{\sigma}^{(0)}$ a atteint le régime stationnaire avant l'application de $\hat{H}_1(t)$ de sorte que les seuls termes non-nuls de $\hat{\sigma}^{(0)}$ sont les populations $\sigma_{jj}^{(0)}$. L'état initial du système est déterminé par la donnée de ces populations que nous supposons connues.

c. Solution à l'ordre 1

Appliquons d'abord l'équation (15b) au cas des populations. Il vient

$$\frac{d}{dt}\sigma_{jj}^{(1)} - \left\{ \frac{d\sigma_{jj}^{(1)}}{dt} \right\} = \frac{1}{i\hbar} \sum_{\ell} [\langle j | \hat{H}'_1(t) | \ell \rangle \sigma_{\ell j}^{(0)} - \sigma_{j \ell}^{(0)} \langle \ell | \hat{H}'_1(t) | j \rangle] \quad (17)$$

Le terme relatif à $\ell = j$ dans le second membre de (17) est nul, les deux termes se compensant. Les autres termes sont également nuls puisque toutes les cohérences d'ordre 0 $\sigma_{\ell j}^{(0)}$ sont nulles. Le second membre de l'équation (17) est donc nul. Il s'ensuit qu'en régime permanent, la solution de l'équation (17) est simplement

$$\sigma_{jj}^{(1)}(t) = 0 \quad (18)$$

A l'ordre 1, les populations n'évoluent pas.

Considérons maintenant les cohérences. En utilisant les équations (15b) et (5b), et compte-tenu de la nullité des cohérences à l'ordre 0, nous trouvons :

$$\frac{d}{dt}\sigma_{jk}^{(1)} + i\omega_{jk}\sigma_{jk}^{(1)} + \gamma_{jk}\sigma_{jk}^{(1)} = \frac{1}{i\hbar} \langle j | \hat{H}'_1(t) | k \rangle (\sigma_{kk}^{(0)} - \sigma_{jj}^{(0)}) \quad (19)$$

La solution de cette équation vérifiant la condition initiale $\sigma_{jk}^{(1)}(t_0) = 0$ est

$$\sigma_{jk}^{(1)}(t) = \frac{\sigma_{kk}^{(0)} - \sigma_{jj}^{(0)}}{i\hbar} \int_{t_0}^t e^{-(i\omega_{jk} + \gamma_{jk})(t-t')} \langle j | \hat{H}_1'(t') | k \rangle dt' \quad (20)$$

Rappelons que le terme d'ordre 1 de la série de perturbations (14) est $\lambda \hat{\sigma}^{(1)}$. Les corrections au premier ordre à la matrice densité $\hat{\sigma}^{(0)}$ sont donc

$$\lambda \sigma_{jj}^{(1)}(t) = 0 \quad (21a)$$

$$\lambda \sigma_{jk}^{(1)}(t) = \frac{\sigma_{kk}^{(0)} - \sigma_{jj}^{(0)}}{i\hbar} \int_{t_0}^t e^{-(i\omega_{jk} + \gamma_{jk})(t-t')} \langle j | \hat{H}_1(t') | k \rangle dt' \quad (21b)$$

d. Perturbation aux ordres supérieurs

Pour obtenir les termes suivants du développement (14), il suffit de résoudre successivement les équations (15c) pour des r croissants. Ainsi partant des valeurs au premier ordre (18) pour les populations et (20) pour les cohérences, nous obtenons les populations et les cohérences au second ordre. Partant de ces valeurs à l'ordre 2, il est possible de trouver les termes d'ordre 3 etc ...

Remarque

Aux ordres supérieurs à 1, les termes relatifs aux populations et aux cohérences sont en général simultanément non-nuls.

2.2 Atome interagissant avec un champ oscillant : réponse linéaire

a. Perturbation sinusoïdale

Nous reprenons dans cette partie l'étude d'un atome décrit par un hamiltonien $\hat{H}_0$ (indépendant du temps) et interagissant avec le champ électrique oscillant d'une onde électromagnétique. L'hamiltonien d'interaction $\hat{H}_1(t)$ est de la forme $\hat{W} \cos \omega t$ avec $\hat{W} = -\hat{\mathbf{D}} \cdot \mathbf{E}_0$ (voir formules (1.7) et (1.8) du chapitre I). En portant cette expression dans la formule (21b), nous trouvons en régime permanent (c'est-à-dire pour $t - t_0 \gg 1/\gamma_{jk}$) :

$$\lambda \sigma_{jk}^{(1)}(t) = \frac{\sigma_{kk}^{(0)} - \sigma_{jj}^{(0)}}{2i\hbar} W_{jk} \left[\frac{e^{-i\omega t}}{i(\omega_{jk} - \omega) + \gamma_{jk}} + \frac{e^{i\omega t}}{i(\omega_{jk} + \omega) + \gamma_{jk}} \right] \quad (22)$$

Dans le cas d'une excitation quasi-résonnante (voir Chapitre I, § B.4.d), un des dénominateurs d'énergie est beaucoup plus petit que l'autre. Plaçons nous par exemple dans le cas où ω_{jk} est positif (niveau j au dessus de k), et supposons que

$$|\omega_{jk} - \omega|, |\gamma_{jk}| \ll |\omega_{jk} + \omega| \quad (23a)$$

Il est alors possible de négliger le terme anti-résonnant de la formule (22) qui prend la forme simplifiée

$$\lambda\sigma_{jk}^{(1)}(t) \approx \frac{\sigma_{kk}^{(1)} - \sigma_{jj}^{(0)}}{2i\hbar} W_{jk} \frac{e^{-i\omega t}}{\gamma_{jk} + i(\omega_{jk} - \omega)} \quad (23b)$$

Les expressions (22) et (23b), à l'ordre 1 en W_{jk} (qui est proportionnel au champ électrique), donnent la réponse linéaire de l'atome.

b. Valeur moyenne du dipôle électrique. Susceptibilité linéaire

La connaissance de la matrice densité permet de calculer la valeur moyenne du dipôle électrique

$$\langle \hat{\mathbf{D}} \rangle = \text{Tr}(\hat{\sigma} \hat{\mathbf{D}}) \quad (24)$$

Pour un atome isolé, les éléments de matrice diagonaux $\mathbf{D}_{ii}$ sont nuls (à cause de l'invariance de l'hamiltonien $\hat{H}_0$ par réflexion d'espace, voir le complément II.1) de sorte que seules les cohérences σ_{jk} apparaissent dans le développement (24) :

$$\langle \hat{\mathbf{D}} \rangle = \sum_{j,k} \sigma_{jk} \mathbf{D}_{kj} \quad (25)$$

Comme le premier terme non-nul de la série de perturbations pour les cohérences est d'ordre 1, nous en déduisons que le dipôle électrique à l'ordre le plus bas est $\langle \hat{\mathbf{D}}^{(1)} \rangle$, et qu'il est linéaire en champ $\mathbf{E}$. Il est égal à

$$\langle \hat{\mathbf{D}}^{(1)} \rangle = \sum_{j,k} \sigma_{jk}^{(1)} \mathbf{D}_{kj} \quad (26)$$

Nous supposons à présent pour simplifier que le champ $\mathbf{E}$ est parallèle à Oz , et que l'état initial – et donc $\sigma^{(0)}$ – est invariant par rotation. Pour des raisons de symétrie le dipôle atomique induit est parallèle à Oz . Nous omettrons désormais la notation vectorielle sachant que les résultats sont relatifs à la composante de $\mathbf{D}$ le long de l'axe Oz . En regroupant les termes relatifs à $\sigma_{jk}^{(1)}$, nous trouvons, à partir de (22) et (26) (et en supposant pour simplifier les taux de relaxation γ_{jk} réels) :

$$\begin{aligned} \langle \hat{D}^{(1)} \rangle = \sum_{j>k} \frac{\sigma_{kk}^{(0)} - \sigma_{jj}^{(0)}}{\hbar} (-D_{jk} E) D_{kj} \\ \left\{ \left(\frac{\omega - \omega_{jk}}{(\omega - \omega_{jk})^2 + \gamma_{jk}^2} - \frac{\omega + \omega_{jk}}{(\omega + \omega_{jk})^2 + \gamma_{jk}^2} \right) \cos \omega t + \left(\frac{\gamma_{jk}}{(\omega + \omega_{jk})^2 + \gamma_{jk}^2} - \frac{\gamma_{jk}}{(\omega - \omega_{jk})^2 + \gamma_{jk}^2} \right) \sin \omega t \right\} \end{aligned} \quad (27)$$

Remarques

(i) Nous avons supposé que γ_{jk} est réel pour écrire la formule (27). Si γ_{jk} a une composante réelle γ'_{jk} et une composante imaginaire $i\gamma''_{jk}$, il faut dans la formule (27) remplacer γ_{jk} par γ'_{jk} et ω_{jk} par $\omega_{jk} + \gamma''_{jk}$. Le centre de la résonance est dans ce cas déplacé par les interactions provoquant la relaxation.

(ii) L'écriture de la formule (27) suppose implicitement que la relaxation peut être décrite par les mêmes coefficients γ_{jk} quel que soit l'écart à résonance. Ceci implique que les courbes

relatives à la composante en quadrature du dipôle électrique (terme en $\sin \omega t$ dans (27)), sont des lorentziennes pour toute valeur de $|\omega_{jk} - \omega|$ ce qui n'est pas vrai en toute généralité. La forme lorentzienne est exacte dans un certain domaine au voisinage de la résonance, généralement très grand comparé à γ_{jk} , mais néanmoins limité à un intervalle plus ou moins vaste selon le mécanisme de relaxation⁷.

En utilisant la définition (2.118a, 2.118b, 2.118c) du chapitre II, il est possible d'exprimer le résultat de la formule (27) en faisant apparaître les parties réelle χ' et imaginaire χ'' de la susceptibilité. Nous noterons N/V le nombre total d'atomes par unité de volume et $N_k^{(0)} = N\sigma_{kk}^{(0)}$ le nombre d'atomes par unité de volume dans l'état $|k\rangle$ en absence de champ électromagnétique. Dans le cas d'une excitation quasi-résonnante (ω voisin de ω_{jk} , cf. Éq. (23a)), on trouve

$$\chi' = \frac{N_k^{(0)} - N_j^{(0)}}{V} \frac{|D_{jk}|^2}{\varepsilon_0 \hbar} \frac{(\omega_{jk} - \omega)}{(\omega_{jk} - \omega)^2 + \gamma_{jk}^2} \quad (28)$$

$$\chi'' = \frac{N_k^{(0)} - N_j^{(0)}}{V} \frac{|D_{jk}|^2}{\varepsilon_0 \hbar} \frac{\gamma_{jk}}{(\omega_{jk} - \omega)^2 + \gamma_{jk}^2} \quad (29)$$

Il apparaît clairement que les formules (28) et (29) généralisent les formules (2.119b, 2.119c) établies dans le chapitre II. En fait, l'étude du chapitre II apparaît comme un cas particulier de l'étude précédente, dans une situation où les taux de relaxation sont égaux ($\Gamma_j = \Gamma_k = \gamma_{jk} = \Gamma$) et où la retombée est négligeable ($\Gamma_{j \rightarrow k} = 0$). Le formalisme développé ici a donc une *portée beaucoup plus vaste* que le modèle simple du chapitre II. Par exemple, comme nous le verrons au chapitre III, l'obtention d'une *inversion de population* dans de nombreux milieux amplificateurs repose souvent sur une *différence entre taux de relaxation*, et le traitement théorique de ces situations relève du formalisme présenté ici.

c. Lien avec la théorie classique

Les résultats ci-dessus ont un équivalent en physique classique : autour d'une résonance ω_{jk} , les formules (28) et (29) sont l'analogue des formules relatives à un électron élastiquement lié (cf. complément II.3) de fréquence de résonance égale à ω_{jk} .

Loin de résonance, il faut revenir au résultat (27) mettant en jeu toutes les transitions atomiques. Mais il est alors possible de négliger γ_{jk} devant $|\omega - \omega_{jk}|$ et $|\omega + \omega_{jk}|$ de sorte que la partie imaginaire de la susceptibilité χ'' est négligeable, et que sa partie réelle χ' vaut

$$\chi' = \sum_{j>k} \frac{N_k^{(0)} - N_j^{(0)}}{V} \frac{|D_{jk}|}{\varepsilon_0 \hbar} \frac{2\omega_{jk}}{(\omega_{jk}^2 - \omega^2)} \quad (30)$$

Ce résultat peut être réexprimé à l'aide des forces d'oscillateur f_{kj} (qui sont des quantités sans dimension)

$$f_{kj} = \frac{2m\omega_{jk}}{\hbar q^2} |D_{jk}|^2 \quad (31)$$

⁷Pour plus de détails, voir CDG 2, Complément B.VI.

(m et q étant la masse et la charge de l'électron). Nous obtenons ainsi :

$$\chi' = \sum_{j>k} \frac{N_k^{(0)} - N_j^{(0)}}{V} \frac{q^2}{m\varepsilon_0} \frac{f_{kj}}{(\omega_{jk}^2 - \omega^2)} \quad (32)$$

En particulier dans la situation habituelle où le seul niveau peuplé est le niveau fondamental, que nous appelons ici a , nous pouvons remplacer tous les $N_j^{(0)}$ par 0, sauf $N_a^{(0)}$ qui est égal à N . La susceptibilité est alors égale à :

$$\chi' = \sum_j \frac{N}{V} \frac{q^2}{m\varepsilon_0} \frac{f_{aj}}{(\omega_{ja}^2 - \omega^2)} \quad (33)$$

La formule (33) apparaît comme l'équivalent quantique du résultat obtenu classiquement en utilisant le modèle de l'électron élastiquement lié :

$$\chi'_{\text{cl}} = \frac{N}{V} \frac{q^2}{m\varepsilon_0} \frac{1}{(\omega_0^2 - \omega^2)} \quad (34)$$

Dans cette expression, ω_0 est la pulsation de l'oscillation libre de l'électron.

L'analogie entre les formules (33) et (34) est d'autant plus frappante que les forces d'oscillateur vérifient la relation suivante (appelée relation de Reiche-Thomas-Kuhn) :

$$\sum_j f_{aj} = 1 \quad (35)$$

Le résultat (33) correspond, en quelque sorte, à une situation où l'oscillateur classique aurait plusieurs fréquences de résonance, chacune ayant un poids f_{aj} . En fait, les formules (33) et (35) ont été introduites empiriquement, à partir des résultats expérimentaux, avant l'avènement de la mécanique quantique. Ce fut d'ailleurs un des premiers succès de cette théorie d'expliquer rigoureusement l'origine de ces formules⁸.

3 Atome à deux niveaux : Traitement non-perturbatif par équations de Bloch optiques

3.1 Introduction

Considérons un atome à deux niveaux a et b , soumis à un champ électromagnétique quasi-résonnant. Nous allons voir qu'il est possible de donner un traitement *non-perturbatif* de ce problème, comme on a su le faire en l'absence de relaxation (cf. § II.C.2). En fait, l'interaction d'un atome à deux niveaux avec un champ électrique oscillant a la même forme mathématique que l'interaction entre un champ magnétique oscillant et un spin 1/2

⁸A. Sommerfeld, Optics, Academic Press, New-York (1954)

plongé dans un champ magnétique statique⁹ qui lève la dégénérescence entre les niveaux $m = \pm 1/2$. Ce dernier problème, qui est celui de la résonance magnétique nucléaire¹⁰, a été étudié notamment par Bloch. Sa transposition au domaine optique conduit aux équations de Bloch optiques, outil de base de l'optique quantique. Nous en présentons brièvement les grandes lignes¹¹.

Considérons d'abord les termes hamiltoniens provenant de l'équation (4). L'hamiltonien $\hat{H}$ est la somme de l'hamiltonien atomique $\hat{H}_0 = \hbar\omega_0|b\rangle\langle b|$ (l'origine des énergies est choisie au niveau a) et de l'hamiltonien dipolaire électrique $\hat{H}_1(t) = -\hat{D}E \cos \omega t$. Nous trouvons ainsi :

$$\frac{d\sigma_{bb}}{dt} = i\Omega_1 \cos \omega t (\sigma_{ba} - \sigma_{ab}) \quad (36a)$$

$$\frac{d\sigma_{aa}}{dt} = -i\Omega_1 \cos \omega t (\sigma_{ba} - \sigma_{ab}) \quad (36b)$$

$$\frac{d\sigma_{ab}}{dt} = i\omega_0 \sigma_{ab} - i\Omega_1 \cos \omega t (\sigma_{bb} - \sigma_{aa}) \quad (37a)$$

$$\frac{d\sigma_{ba}}{dt} = -i\omega_0 \sigma_{ba} + i\Omega_1 \cos \omega t (\sigma_{bb} - \sigma_{aa}) \quad (37b)$$

Nous avons introduit ici la pulsation de Rabi à résonance égale à $\Omega_1 = -D_{ba}E/\hbar$ (voir Chapitre II, Éq. (2.70)). On remarquera que ces équations ne sont pas totalement indépendantes. D'une part $\frac{d}{dt}(\sigma_{aa} + \sigma_{bb}) = 0$, ce qui correspond au fait que $\text{Tr} \hat{\sigma} = \sigma_{aa} + \sigma_{bb} = 1$ pour un système fermé. D'autre part les équations (37a) et (37b) sont complexes conjuguées, car la matrice densité est hermitienne et $\sigma_{ab}^* = \sigma_{ba}$.

A l'approximation quasi-résonnante, seule une des composantes exponentielles du $\cos \omega t$ contribue de façon significative à l'évolution. Par exemple, si Ω_1 tend vers 0, σ_{ab} varie en $\exp(i\omega_0 t)$ comme le montre (37a). On transforme donc les équations (C.1) et (C.2) en ne gardant que les termes oscillant le plus lentement, qui donneront les termes prédominants après intégration. Après cette simplification, les équations précédentes deviennent :

$$\frac{d\sigma_{bb}}{dt} = \frac{i\Omega_1}{2} (e^{i\omega t} \sigma_{ba} - e^{-i\omega t} \sigma_{ab}) \quad (38a)$$

$$\frac{d\sigma_{aa}}{dt} = -\frac{i\Omega_1}{2} (e^{i\omega t} \sigma_{ba} - e^{-i\omega t} \sigma_{ab}) \quad (38b)$$

⁹L'équivalence entre un système à deux niveaux quelconque et un spin 1/2 a été soulignée par Feynmann, « Lectures on Physics », Vol. III, Chap. 11. Voir aussi CDL 1, Complément C.IV.

¹⁰Voir par exemple A. Abragam « Les Principes du Magnétisme Nucléaire », Presses universitaires de France, Paris (1961).

¹¹Pour plus de détails, voir par exemple CDG 2, Chapitre V ; ou L. Allen et J.H. Eberly, « Optical Resonance and Two-Level Atoms », Wiley Interscience, New-York (1975).

$$\frac{d\sigma_{ab}}{dt} = i\omega_0\sigma_{ab} - i\frac{\Omega_1}{2}e^{i\omega t}(\sigma_{bb} - \sigma_{aa}) \quad (39a)$$

$$\frac{d\sigma_{ba}}{dt} = -i\omega_0\sigma_{ba} + i\frac{\Omega_1}{2}e^{-i\omega t}(\sigma_{bb} - \sigma_{aa}) \quad (39b)$$

En ce qui concerne les termes de relaxation, nous les introduirons séparément pour le système fermé et le système ouvert.

3.2 Système fermé

Dans le cas du système fermé (voir § A.3.a), les équations (A.6) jointes aux équations (C.3) et (C.4) conduisent aux équations

$$\frac{d\sigma_{bb}}{dt} = i\frac{\Omega_1}{2}(e^{i\omega t}\sigma_{ba} - e^{-i\omega t}\sigma_{ab}) - \Gamma_{\text{sp}}\sigma_{bb} \quad (40a)$$

$$\frac{d\sigma_{aa}}{dt} = -i\frac{\Omega_1}{2}(e^{i\omega t}\sigma_{ba} - e^{-i\omega t}\sigma_{ab}) + \Gamma_{\text{sp}}\sigma_{bb} \quad (40b)$$

$$\frac{d\sigma_{ab}}{dt} = i\omega_0\sigma_{ab} - i\frac{\Omega_1}{2}e^{i\omega t}(\sigma_{bb} - \sigma_{aa}) - \gamma\sigma_{ab} \quad (41a)$$

$$\sigma_{ba}^* = \sigma_{ab} \quad (41b)$$

Nous supposons γ réel.

Pour résoudre ce système, il est possible d'éliminer la dépendance temporelle rapide en faisant le changement de variables :

$$\sigma'_{ba} = e^{-i\omega t}\sigma_{ba} \quad (42a)$$

$$\sigma'_{ab} = e^{i\omega t}\sigma_{ab} \quad (42b)$$

$$\sigma'_{bb} = \sigma_{bb} \quad (42c)$$

$$\sigma'_{aa} = \sigma_{aa} \quad (42d)$$

Avec ces nouvelles variables, les équations (C.5) et (C.6) deviennent :

$$\frac{d\sigma'_{bb}}{dt} = i\frac{\Omega_1}{2}(\sigma'_{ba} - \sigma'_{ab}) - \Gamma_{\text{sp}}\sigma'_{bb} \quad (43a)$$

$$\frac{d\sigma'_{aa}}{dt} = -i\frac{\Omega_1}{2}(\sigma'_{ba} - \sigma'_{ab}) + \Gamma_{\text{sp}}\sigma'_{bb} \quad (43b)$$

$$\frac{d\sigma'_{ab}}{dt} = i(\omega_0 - \omega)\sigma'_{ab} - i\frac{\Omega_1}{2}(\sigma'_{bb} - \sigma'_{aa}) - \gamma\sigma'_{ab} \quad (44)$$

La solution stationnaire est obtenue en annulant les dérivées par rapport au temps dans les équations (C.8) et (44) et en utilisant $\sigma'_{aa} + \sigma'_{bb} = 1$ (conservation de la population totale). En régime permanent, nous trouvons donc :

$$\sigma'_{bb} = \frac{1}{2} \frac{\Omega_1^2 \frac{\gamma}{\Gamma_{sp}}}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\Gamma_{sp}}} \quad (45a)$$

$$\sigma'_{aa} = 1 - \sigma'_{bb} \quad (45b)$$

$$\sigma'_{ab} = i \frac{\Omega_1}{2} \frac{\gamma - i(\omega - \omega_0)}{\gamma^2 + (\omega_0 - \omega)^2 + \Omega_1^2 \frac{\gamma}{\Gamma_{sp}}} \quad (45c)$$

$$\sigma'_{ba} = \sigma'^{*}_{ab} \quad (45d)$$

On peut en déduire la cohérence $\sigma_{ba}^{(1)}$ à l'ordre 1 de l'interaction (linéaire en Ω_1). En utilisant l'ordre 1 de (45c), et en reportant dans (42a), on trouve

$$\sigma_{ba}^{(1)} = -\frac{i}{2} \frac{\Omega_1}{\gamma + i(\omega_0 - \omega)} e^{-i\omega t} \quad (46)$$

Ce résultat coïncide bien avec (23a) puisque $\sigma_{aa}^{(0)} - \sigma_{bb}^{(0)} = 1$ et $W_{ba}/\hbar = \Omega_1$.

Mais l'intérêt des formules (45a) à (45d) est qu'elles sont valables pour de grandes valeurs du champ incident. Elles permettent donc de décrire correctement les effets de saturation qui apparaissent par la présence du terme en Ω_1^2 (proportionnel à l'intensité) au dénominateur. Ainsi, (45a) montre que la population du niveau supérieur de la transition passe de 0 à 1/2 quand l'intensité (et donc Ω_1^2) croît. Elle reste donc toujours inférieure ou au plus égale à la population du niveau fondamental. Pour un système à deux niveaux fermé, il n'est donc pas possible d'obtenir, *en régime permanent*, une inversion de population entre le niveau fondamental et le niveau excité, sous l'action d'un champ électromagnétique. Cette remarque est essentielle pour la recherche de milieux amplificateurs laser (cf. Chapitre III).

De même, en utilisant (45c) et (42a), et en reportant dans l'expression (25), on obtient la valeur moyenne du dipôle atomique :

$$\langle \hat{D} \rangle = i \frac{|D_{ba}|^2}{2\hbar} \frac{\gamma + i(\omega - \omega_0)}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\Gamma_{sp}}} E e^{-i\omega t} + c.c \quad (47)$$

Pour un milieu contenant N/V atomes par unité de volume, le dipôle moyen par unité de volume vaut $P = \frac{N}{V} \langle D \rangle$. En utilisant les définitions (2.118a) et (2.118b) du chapitre II, et en posant $D_{ba} = d$, nous trouvons la partie réelle (dispersion) et imaginaire (absorption) de la susceptibilité :

$$\chi' = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\omega_0 - \omega}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\Gamma_{sp}}} \quad (48a)$$

$$\chi'' = \frac{N}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\gamma}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\Gamma_{sp}}} \quad (48b)$$

Ces expressions non-perturbatives de la susceptibilité d'une assemblée d'atomes sont très utiles dans de nombreux problèmes d'optique non-linéaire. Elles permettent de décrire les phénomènes à haute intensité tout en prenant correctement en compte la relaxation dans un système fermé. Dans le cas où la seule cause de relaxation est l'émission spontanée, on a $\gamma = \Gamma_{\text{sp}}/2$, et on trouve ainsi la formule (II.2.125) donnée sans justification au chapitre II.

Rappelons que les formules (48a) et (48b) sont associées au régime stationnaire. Il est clair qu'il existe d'autres solutions intéressantes des équations de Bloch, dans le régime transitoire. Par exemple, lorsque le champ électromagnétique est appliqué brusquement, tous les atomes sont à l'instant initial dans le niveau fondamental. La solution des équations de Bloch permet alors de déterminer le régime transitoire¹². Cette solution transitoire est en fait une oscillation de Rabi amortie par la relaxation. On l'appelle un « transitoire cohérent » (cf. Chapitre II, § C.2.d).

3.3 Système ouvert

Pour finir, nous considérons le cas d'un système ouvert, ce qui permet de faire le lien avec le modèle simple du chapitre II, qui avait la particularité de pouvoir être traité sans le formalisme de la matrice densité. Ce cas est particulièrement important car il constitue une description réaliste de nombreuses transitions utilisées dans les amplificateurs laser. Nous écrirons les termes de pompage et de relaxation en supposant que la retombée du niveau b au niveau a est négligeable ($\Gamma_{b \rightarrow a} \ll \Gamma_b$) de sorte que nous prendrons $\Gamma_{b \rightarrow a}$ nul dans l'équation (8b). En utilisant les équations (C.3) et (C.4), nous trouvons les équations de Bloch pour un système ouvert :

$$\begin{aligned} \frac{d\sigma_{bb}}{dt} &= i\frac{\Omega_1}{2}(e^{i\omega t}\sigma_{ba} - e^{-i\omega t}\sigma_{ab}) - \Gamma_b\sigma_{bb} + \tilde{\Lambda}_b \\ \frac{d\sigma_{aa}}{dt} &= -i\frac{\Omega_1}{2}(e^{i\omega t}\sigma_{ba} - e^{-i\omega t}\sigma_{ab}) - \Gamma_a\sigma_{aa} + \tilde{\Lambda}_a \end{aligned} \quad (49a)$$

$$\frac{d\sigma_{ab}}{dt} = i\omega_0\sigma_{ab} - i\frac{\Omega_1}{2}e^{i\omega t}(\sigma_{bb} - \sigma_{aa}) - \gamma\sigma_{ab} \quad (50)$$

Tout comme dans le cas précédent, ce système se résout en faisant le changement de variable (C.7) de façon à rendre le deuxième membre des équations indépendant du temps. Il est commode d'introduire les quantités p_b et p_a qui, dans le cas du pompage faible ($\tilde{\Lambda}_b \ll \Gamma_b, \tilde{\Lambda}_a \ll \Gamma_a$), s'interprètent comme les probabilités stationnaires de trouver l'atome dans les niveaux b et a en absence de champ incident :

$$p_b = \frac{\tilde{\Lambda}_b}{\Gamma_b} \quad \text{et} \quad p_a = \frac{\tilde{\Lambda}_a}{\Gamma_a} \quad (51)$$

¹²L. Allen et J.H. Eberly, « Optical Resonance and Two-Level Atoms », Wiley Interscience, New-York (1975).

Définissons le taux de relaxation moyen $\bar{\Gamma}$ des populations :

$$\frac{2}{\bar{\Gamma}} = \frac{1}{\Gamma_a} + \frac{1}{\Gamma_b} \quad (52)$$

Avec ces notations, la solution stationnaire des équations de Bloch pour les populations s'écrit :

$$\sigma_{bb} - \sigma_{aa} = (p_b - p_a) \left(1 - \frac{\Omega_1^2 \frac{\gamma}{\bar{\Gamma}}}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\bar{\Gamma}}} \right) \quad (53)$$

$$\sigma_{bb} + \sigma_{aa} = p_b + p_a + \frac{\Omega_1^2}{2} \frac{\left(\frac{\gamma}{\Gamma_a} - \frac{\gamma}{\Gamma_b} \right) (p_b - p_a)}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\bar{\Gamma}}} \quad (54)$$

Remarquons d'abord que dans le cas particulier $\gamma = \Gamma_a = \Gamma_b$ les formules (53) et (54) coïncident avec celles obtenues dans le modèle simple du chapitre II. Notons ensuite la non conservation de $(\sigma_{bb} + \sigma_{aa})$, liée au caractère ouvert du système.

De la valeur stationnaire de la cohérence

$$\sigma_{ba} = i \frac{\Omega_1}{2} \frac{(p_b - p_a)[\gamma + i(\omega - \omega_0)]}{(\omega - \omega_0)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\bar{\Gamma}}} e^{-i\omega t} \quad (55)$$

nous déduisons le dipôle atomique, et donc les susceptibilités χ' et χ'' pour une vapeur comportant N/V atomes par unité de volume :

$$\chi' = \frac{N_a^{(0)} - N_b^{(0)}}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{(\omega_0)}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\bar{\Gamma}}} \quad (56a)$$

$$\chi'' = \frac{N_a^{(0)} - N_b^{(0)}}{V} \frac{d^2}{\varepsilon_0 \hbar} \frac{\gamma}{(\omega_0 - \omega)^2 + \gamma^2 + \Omega_1^2 \frac{\gamma}{\bar{\Gamma}}} \quad (56b)$$

où $N_b^{(0)} = p_b N$ et $N_a^{(0)} = p_a N$ sont les nombres d'atomes dans les niveaux b et a en l'absence du champ oscillant de pulsation ω . Les formules (56) généralisent les résultats (2.120) et (2.148) du chapitre II obtenus à partir du modèle simple où les taux de relaxation des deux niveaux sont égaux. Elles montrent qu'un milieu est amplificateur si la population stationnaire au seuil $N_b^{(0)}$ du niveau supérieur de la transition est plus grande que la population stationnaire au seuil $N_a^{(0)}$ du niveau inférieur. D'après (51), une telle inversion de population peut être obtenue soit en pompant le niveau supérieur plus énergiquement que le niveau inférieur, soit en utilisant une transition pour la quelle le niveau inférieur a une durée de vie très courte.

4 Conclusion

Les équations de Bloch optiques constituent un outil fondamental de l'interaction atome-rayonnement. Dans le cadre de l'approximation résonnante qui correspond à un

grand nombre de situations expérimentales importantes, ce sont des équations exactes, qui décrivent quantiquement l'évolution atomique, tout en permettant de prendre en compte les phénomènes de relaxation (émission spontanée, collisions ...).

Par rapport à des traitements approchés plus simples, elles ont l'avantage de décrire correctement deux types de phénomènes essentiels dans l'interaction atome-laser. Il s'agit d'une part des effets d'ordre supérieur apparaissant à forte intensité, comme la saturation, ou plus généralement les effets non-linéaires. D'autre part, elles permettent de traiter correctement les cohérences entre niveaux atomiques, qui sont responsables de très nombreux effets. Ce sont notamment ces cohérences qui déterminent la fréquence et la phase du dipôle électrique atomique induit par le rayonnement incident sur l'atome (cf. Équation (25)). Comme le champ réémis est directement relié à ce dipôle, on peut ainsi prévoir les caractéristiques les plus fines du rayonnement diffusé par un atome soumis à un champ incident quelconque¹³.

Les équations de Bloch optiques se généralisent naturellement aux systèmes comportant plus de deux niveaux. Mais elles conduisent souvent à des systèmes d'équations inextricables, dès que l'on doit considérer plus de trois niveaux atomiques¹⁴. Il faut alors revenir à des méthodes approchées plus simples : à faible intensité, la méthode des perturbations donne la réponse linéaire ; à forte intensité, on peut avoir recours à des équations de pompage (Chapitre 2, § 2.5.6 et complément II.1), qui décrivent correctement les phénomènes non-linéaires tels que la saturation, mais qui prennent très mal en compte les phénomènes liés aux cohérences. En fait, l'emploi de telle ou telle approximation se fait souvent de façon heuristique, en s'appuyant sur l'expérience acquise en appliquant les équations de Bloch optiques aux systèmes à deux et trois niveaux.

¹³CDG 2, Chapitre V.

¹⁴Notons que dans ce cas la détermination des termes de relaxation est loin d'être évidente. Il existe par exemple des transferts de cohérence par émission spontanée, entre sous-niveaux excités et sous-niveaux fondamentaux. Les règles qui gouvernent ces transferts ont une très grande importance dans de nombreux phénomènes. Elles ont été établies et explicitées dans C. Cohen-Tannoudji, Thèse, Annales de Physique (Paris) 7, 423 (1962). (Reproduit dans C. Cohen-Tannoudji, « Atoms in Electromagnetic fields », World Scientific (1994)).

Complément II.3

Modèle de l'électron élastiquement lié

Nous nous proposons dans ce complément de calculer, dans le cadre de l'électrodynamique classique, le champ électromagnétique rayonné par une charge mobile rappelée vers une position d'équilibre par une force proportionnelle à la distance. Dans le cas où la charge n'est soumise à aucune autre force (oscillations libres), nous verrons que son *mouvement d'oscillation est amorti*, puisque l'énergie émise sous forme de rayonnement électromagnétique est prélevée sur l'énergie mécanique. Dans le cas où la charge est soumise aussi à l'action d'un champ électromagnétique extérieur oscillant à une fréquence ω , nous calculerons les caractéristiques du champ qu'elle rayonne dans tout l'espace en régime d'oscillations forcées. Le problème traité est alors celui de la *diffusion d'une onde électromagnétique*. On distinguera les régimes de diffusion Rayleigh, Thomson, ou résonnante, selon que la pulsation ω du rayonnement incident est inférieure, supérieure, ou voisine de la pulsation propre ω_0 du mouvement de la charge élastiquement liée.

Le problème abordé ici est entièrement classique, et constitue donc une sorte d'étape 0 dans l'étude de l'interaction matière-rayonnement, l'étape 1 étant le traitement semi-classique présenté dans le chapitre 2 de ce cours, tandis que l'étape 2 – le traitement complètement quantique de l'atome et du rayonnement, ne sera vue que dans le cadre du cours d'« Optique Quantique 2 ». Ce traitement classique de la matière semble a priori peu adapté pour décrire le cas où la charge mobile est un électron à l'intérieur d'un atome, système lié microscopique dans lequel l'aspect quantique est déterminant. Nous montrerons cependant dans ce complément que si l'on prend comme fréquence propre de l'électron élastiquement lié la fréquence de Bohr atomique $\omega_0 = (E_b - E_a)/\hbar$, les résultats du modèle de l'électron élastiquement lié sont en accord avec les résultats quantiques que l'on obtient pour le système à deux niveaux $|a\rangle$ et $|b\rangle$ en interaction avec un champ électromagnétique classique de fréquence ω (modèle développé dans ce chapitre), dans le cas où le champ électromagnétique est assez faible pour que le phénomène de saturation soit négligeable (régime linéaire). On justifie ainsi que de très nombreuses propriétés de la propagation des ondes électromagnétiques dans les diélectriques se décrivent très bien en représentant le diélectrique par un ensemble d'électrons élastiquement liés (le modèle de

Sellmeier de l'indice de réfraction repose sur cette approche). De plus, un certain nombre de propriétés du rayonnement émis par la matière, par exemple sa polarisation, sont très faciles à visualiser dans le cadre de l'émission par un électron élastiquement lié. Or elles sont strictement identiques à celles que l'on obtient dans le cadre semi-classique ou totalement quantique, pour le cas d'une transition entre niveaux atomiques de moments cinétiques $J = 0$ et $J = 1$ (voir le complément III.1, partie A). Avoir bien assimilé la correspondance entre les divers modèles dans ce cas simple peut constituer un moyen utile d'aborder les cas plus complexes où le formalisme quantique est indispensable.

1. Notations et hypothèses simplificatrices

Considérons un électron de charge q ($q = -1,6 \times 10^{-19}\text{C}$) oscillant autour du point O à la pulsation ω . On décrira le mouvement de l'électron de position $\mathbf{r}_e$ par l'amplitude complexe $\mathcal{S}$:

$$\mathbf{r}_e = \text{Re} \{ \mathcal{S} \} \quad (1a)$$

$$\mathcal{S} = \mathcal{S}_0 e^{-i\omega t} \quad (1b)$$

Cette notation permet de traiter aussi bien le cas d'un mouvement linéaire que celui d'un mouvement circulaire ou elliptique. Par exemple, un électron oscillant linéairement le long de Oz correspond à $\mathcal{S}_0 = a\mathbf{e}_z$, et un électron ayant une trajectoire circulaire dans le plan xOy correspond à $\mathcal{S}_0 = a(\mathbf{e}_x + i\mathbf{e}_y)/\sqrt{2}$ (a est un nombre complexe quelconque, et $\mathbf{e}_x, \mathbf{e}_y$ et $\mathbf{e}_z$ sont les vecteurs unitaires des axes Ox, Oy et Oz).

On notera $\mathcal{D}$ la quantité

$$\mathcal{D} = q\mathbf{s}_0 e^{-i\omega t} = \mathcal{D}_0 e^{-i\omega t} \quad (1c)$$

qui a les dimensions d'un dipôle électrique, et qui peut être interprétée comme telle dans le cas où l'électron considéré fait partie d'un ensemble neutre, dont le barycentre de toutes les charges autres que l'électron se trouve en O. Rappelons que le dipôle réel est

$$\mathbf{D}(t) = \text{Re} \{ \mathcal{D}_0 e^{-i\omega t} \} \quad (1d)$$

Pour calculer le rayonnement de ce dipôle, nous nous placerons dans le cadre d'hypothèses simplificatrices :

- (i) *approximation dipolaire électrique* : on se limite au cas où les charges restent localisées dans un volume nettement plus petit que la longueur d'onde λ du rayonnement émis :

$$|\mathcal{S}_0| \ll \lambda \quad (2)$$

Notons que cette hypothèse implique que la vitesse de l'électron, qui est de l'ordre de $\omega|\mathbf{s}_0|$, reste petite devant la vitesse de la lumière c .

- (ii) *champ à grande distance* : on se contentera de calculer les champs rayonnés à des distances grandes devant la longueur d'onde :

$$|\mathbf{r}| \gg \lambda \quad (3)$$

Dans le cas général, le champ rayonné se met sous la forme d'un développement en puissances de λ/r ; l'approximation faite ici consiste à ne garder que le premier terme.

Si le mouvement de l'électron n'est pas purement sinusoïdal, les deux hypothèses que nous venons d'énoncer devront être vérifiées pour toutes les composantes de Fourier de $\mathbf{D}(t)$.

2. Rayonnement dipolaire électrique

a. Potentiels retardés

La formulation de l'électromagnétisme la mieux adaptée à ce problème (c'est-à-dire celle qui permet de faire intervenir le plus rapidement les hypothèses simplificatrices) est celle des potentiels retardés en jauge de Lorentz¹. Nous utiliserons donc l'expression du potentiel vecteur retardé :

$$\mathbf{A}_L(\mathbf{r}, t) = \frac{1}{4\pi\epsilon_0 c^2} \int d^3r' \frac{\mathbf{j}(\mathbf{r}', t - |\mathbf{r} - \mathbf{r}'|)c}{|\mathbf{r} - \mathbf{r}'|} \quad (4)$$

Pour calculer le rayonnement du dipôle, nous ferons la simplification

$$|\mathbf{r} - \mathbf{r}'| \approx |\mathbf{r}| = r \quad (5)$$

en vertu des hypothèses (i) et (ii) ci-dessus.

Il nous reste à déterminer l'expression de la densité de courant associée à l'électron en mouvement. En toute rigueur, on a :

$$\mathbf{j}(\mathbf{r}', t) = q\delta(\mathbf{r}' - \mathbf{r}_e(t))\dot{\mathbf{r}}_e(t) \quad (6a)$$

puisque l'on a une charge ponctuelle de vitesse $\dot{\mathbf{r}}_e$ située au point $\mathbf{r}_e$. Utilisant à nouveau les hypothèses (i) et (ii), nous prendrons un courant localisé en O, c'est-à-dire :

$$\mathbf{j}(\mathbf{r}', t) \cong q\delta(\mathbf{r}')\dot{\mathbf{r}}_e(t) \quad (6b)$$

En reportant cette expression dans (4), nous trouvons l'expression suivante du potentiel vecteur retardé :

$$\Rightarrow \quad \mathbf{A}_L(\mathbf{r}, t) = \frac{1}{4\pi\epsilon_0 c^2} \frac{\dot{\mathbf{D}}(t - r/c)}{r} \quad (7)$$

¹Voir par exemple le paragraphe I.C.6 du cours PHY557A (Optique quantique 2 : photons).

La condition de Lorentz¹ nous permet d'obtenir le potentiel scalaire associé, $U_L(\mathbf{r}, t)$:

$$\begin{aligned}\frac{\partial U_L(\mathbf{r}, t)}{\partial t} &= -c^2 \nabla \cdot \mathbf{A}_L(r, t) \\ &= \frac{r}{4\pi\epsilon_0 cr^3} \cdot [c\dot{\mathbf{D}}(t - r/c) + r\ddot{\mathbf{D}}(t - r/c)]\end{aligned}\quad (8)$$

(Pour obtenir ce résultat, on a utilisé

$$\frac{\partial r}{\partial x} = \frac{x}{r} \quad ; \quad \frac{\partial r}{\partial y} = \frac{y}{r} \quad ; \quad \frac{\partial r}{\partial z} = \frac{z}{r} \quad (9)$$

formules très utiles dans tout ce calcul). En intégrant par rapport au temps et en ignorant le terme électrostatique qui ne joue aucun rôle dans la suite, on obtient

$$\Rightarrow U_L(\mathbf{r}, t) = \frac{\mathbf{r}}{4\pi\epsilon_0} \cdot \left[\frac{\mathbf{D}(t - r/c)}{r^3} + \frac{\dot{\mathbf{D}}(t - r/c)}{cr^2} \right] \quad (10)$$

expression dans laquelle le dernier terme est prédominant, compte tenu de l'hypothèse (ii).

Remarque

Le passage de (6a) à (6b) n'est valable que parce que la vitesse de l'électron reste petite devant celle de la lumière, compte tenu de l'hypothèse (ii). Dans le cas relativiste, le champ d'une charge en mouvement serait donné par les potentiels de Lienard et Wiechert (voir par exemple Feynman, Leighton, Sands, Lectures in Physics, Vol. 2, Chapitre 21, § 21.5).

b. Champ rayonné à grande distance par une charge accélérée

Le champ magnétique s'obtient aisément en prenant le rotationnel de $\mathbf{A}_L$ (équation (7)), et en utilisant les résultats intermédiaires

$$\nabla \left(\frac{1}{r} \right) = -\frac{\mathbf{r}}{r^3} \quad (11a)$$

$$\nabla \times \dot{\mathbf{D}}(t - r/c) = -\frac{1}{rc} \mathbf{r} \times \ddot{\mathbf{D}}(t - r/c) \quad (11b)$$

On trouve

$$\mathbf{B}(\mathbf{r}, t) = \frac{-\mathbf{r}}{4\pi\epsilon_0 c^2} \times \left[\frac{\dot{\mathbf{D}}(t - r/c)}{r^3} + \frac{\ddot{\mathbf{D}}(t - r/c)}{cr^2} \right] \quad (12a)$$

expression dont on ne gardera que le dernier terme, en vertu de l'hypothèse (ii) :

$$\mathbf{B}(\mathbf{r}, t) = \frac{1}{4\pi\epsilon_0 c^3} \frac{\ddot{\mathbf{D}}(t - r/c)}{r^2} \times \mathbf{r} \quad (12b)$$

Le champ électrique s'obtient à partir de

$$\mathbf{E} = -\frac{\partial \mathbf{A}_L}{\partial t} - \nabla U_L \quad (13)$$

Le calcul général est fastidieux. A grande distance, on se limite aux termes de degré le plus bas en $1/r$, ce qui donne

$$\begin{aligned} \Rightarrow \quad \mathbf{E}(\mathbf{r}, t) &= \frac{1}{4\pi\epsilon_0 c^2} \left[-\frac{\ddot{\mathbf{D}}(t - r/c)}{r} + \frac{\mathbf{r}[\mathbf{r} \cdot \ddot{\mathbf{D}}(t - r/c)]}{r^3} \right] \\ &= \frac{1}{4\pi\epsilon_0 c^2 r^3} \times [\mathbf{r} \times \ddot{\mathbf{D}}(t - r/c)] \end{aligned} \quad (14)$$

En comparant cette expression du champ électrique avec l'équation (12b), on voit que :

$$\mathbf{B}(\mathbf{r}, t) = \frac{\mathbf{r}}{rc} \times \mathbf{E}(\mathbf{r}, t) \quad (15)$$

Les champs $\mathbf{E}$ et $\mathbf{B}$ sont donc tous deux orthogonaux à $\mathbf{r}$. Ils sont de surcroît en phase. On constate par ailleurs que le champ électromagnétique rayonné à grande distance est proportionnel à $\ddot{\mathbf{D}}$, c'est-à-dire à l'accélération de la charge qui rayonne, mais avec un retard r/c correspondant au temps de propagation. Dans une direction $\mathbf{r}/r$ donnée, le champ électromagnétique ou magnétique décroît en $1/r$, et non pas en $1/r^2$ ou en $1/r^3$ comme en électrostatique et en magnétostatique. Il s'agit d'une onde sphérique, non uniforme, qui se propage à partir de l'origine. Le champ électromagnétique au point M a donc localement la même structure qu'une onde plane progressive se propageant suivant $\mathbf{r}$.

Le vecteur de Poynting $\mathbf{\Pi}$ au point M s'écrit

$$\Rightarrow \quad \mathbf{\Pi}(\mathbf{r}, t) = \epsilon_0 c^2 \mathbf{E} \times \mathbf{B} = \frac{1}{16\pi^2 \epsilon_0 c^3} \frac{\mathbf{r}[\mathbf{r} \times \ddot{\mathbf{D}}(t - r/c)]^2}{r^5} \quad (16)$$

La puissance rayonnée décroît suivant une loi en $1/r^2$. Dans une direction donnée, elle est proportionnelle au carré de l'accélération de la charge.

3. Dipôle oscillant sinusoïdalement. Polarisation

a. Champ rayonné

Revenons maintenant au cas d'un dipôle oscillant sinusoïdalement (Équation 1b). En notation complexe, le champ électromagnétique rayonné à grande distance, de pulsation ω s'obtient à partir de (12b) et (14) :

$$\Rightarrow \quad \mathcal{E}(\mathbf{r}, t) = -\frac{\omega^2}{4\pi\epsilon_0 c^2} \frac{\mathbf{r} \times (\mathbf{r} \times \mathcal{D}_0)}{r^3} e^{-i\omega(t-r/c)} \quad (17a)$$

$$\Rightarrow \quad \mathcal{B}(\mathbf{r}, t) = \frac{1}{c} \frac{\mathbf{r}}{r} \times \mathcal{E}(\mathbf{r}, t) = \frac{\omega^2}{4\pi\epsilon_0 c^3} \frac{(\mathbf{r} \times \mathcal{D}_0)}{r^2} e^{-i\omega(t-r/c)} \quad (17b)$$

et le vecteur de Poynting *moyenné sur une période optique* vaut donc

$$\Rightarrow \quad \overline{\mathbf{\Pi}(\mathbf{r}, t)} = \frac{\omega^4}{32\pi^2 \epsilon_0 c^3} \frac{|\mathbf{r} \times \mathcal{D}_0|^2}{r^5} \mathbf{r} \quad (18)$$

b. Polarisation

Nous allons maintenant étudier successivement deux cas particuliers importants : tout d'abord celui d'un dipôle oscillant suivant Oz (cas appelé π), puis celui d'un dipôle polarisé circulairement dont l'extrémité décrit, à la vitesse angulaire ω , un cercle dans le plan xOy (cas appelé σ_+ ou σ_- suivant le sens de rotation).

(i) Dans le cas de la polarisation π , on a

$$\mathcal{D}_0 = d_0 \mathbf{e}_z \quad (19a)$$

En introduisant les coordonnées sphériques (r, θ, φ) et en posant $k = \omega/c$, on obtient (voir figure 1) :

$$\mathcal{E}_\pi(\mathbf{r}, t) = -\frac{\omega^2 d_0}{4\pi\epsilon_0 c^2} \frac{\sin \theta}{r} e^{-i(\omega t - kr)} \mathbf{e}_\theta \quad (19b)$$

$$\mathcal{B}_\pi(\mathbf{r}, t) = -\frac{\omega^2 d_0}{4\pi\epsilon_0 c^3} \frac{\sin \theta}{r} e^{-i(\omega t - kr)} \mathbf{e}_\varphi \quad (19c)$$

$$\overline{\Pi} = \frac{\omega^4 d_0^2}{32\pi^2 \epsilon_0 c^3} \frac{\sin^2 \theta}{r^2} \mathbf{e}_r \quad (19d)$$

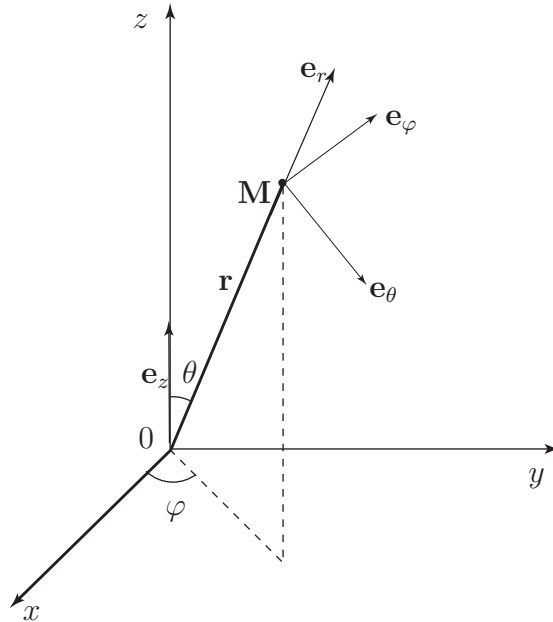


FIG. 1 : Coordonnées sphériques.

Un dipôle linéaire émet donc préférentiellement dans un plan perpendiculaire à son axe de vibration ($\theta = \pi/2$). Le champ électrique est polarisé linéairement dans le plan contenant le dipôle et la direction de propagation (figure 2). La puissance totale Φ rayonnée

à travers une surface fermée entourant le dipôle s'obtient en intégrant (19d). On trouve un résultat indépendant de r :

$$\Rightarrow \Phi = \frac{\omega^4 d_0^2}{12\pi\epsilon_0 c^3} \quad (20)$$

(ii) Considérons maintenant le cas d'un dipôle tournant dans le sens direct (cas σ_+)

$$\mathcal{D}_0 = d_0 \vec{\epsilon}_+ \quad \text{avec} \quad \vec{\epsilon}_+ = -\frac{1}{\sqrt{2}}(\mathbf{e}_x + i\mathbf{e}_y) . \quad (21)$$

La polarisation du rayonnement émis suivant une direction OM quelconque est en général elliptique. Il existe cependant des directions particulières où elle est simple.

Considérons tout d'abord le champ en un point M situé sur l'axe Oz :

$$\mathcal{E}(r\mathbf{e}_z, t) = \frac{\omega^2 d_0}{4\pi\epsilon_0 c^2 r} \vec{\epsilon}_+ e^{i(kr - \omega t)} \quad (22)$$

Le champ électrique est polarisé circulairement dans le même sens que le dipôle ; il en est naturellement de même du champ magnétique $\mathcal{B}$ qui reste perpendiculaire à $\mathcal{E}$ à tout instant. C'est ce qui caractérise un champ polarisé circulairement dans le sens positif.

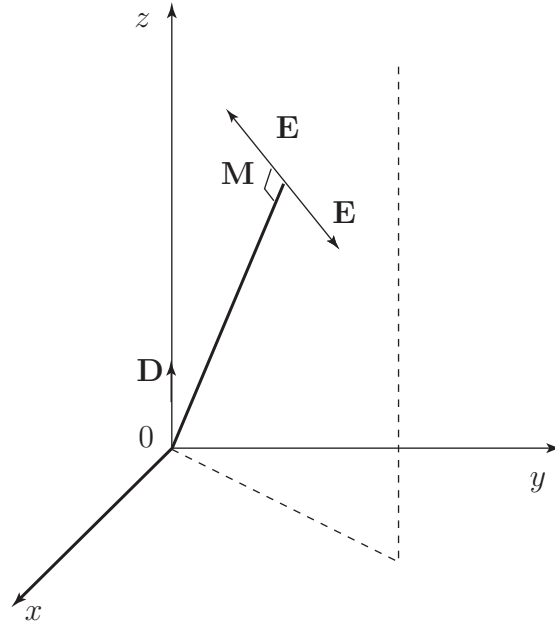


FIG. 2: Rayonnement émis par un dipôle oscillant suivant Oz . Le champ électrique en M est polarisé linéairement perpendiculairement à OM , dans le plan (Oz, OM) . L'émission est nulle suivant Oz , et maximale dans le plan (xOy) .

Considérons maintenant, toujours pour le dipôle σ_+ , le champ en un point M situé

dans le plan xOy , dans une direction $\mathbf{e}_r$ caractérisée par l'angle φ . Alors

$$\mathbf{e}_r \times (\mathbf{e}_r \times \mathcal{D}_0) = -\frac{i}{\sqrt{2}} e^{i\varphi} d_0 \mathbf{e}_\varphi \quad (23a)$$

$$\mathcal{E}(\mathbf{r}, t) = \frac{i}{\sqrt{2}} e^{i\varphi} \frac{\omega^2 d_0}{4\pi\epsilon_0 c^2} \frac{1}{r} e^{i(kr-\omega t)} \mathbf{e}_\varphi \quad (23b)$$

Le champ électrique est polarisé linéairement suivant le vecteur $\mathbf{e}_\varphi$, perpendiculaire à la fois à la direction de propagation OM et à l'axe de rotation du dipôle, et il oscille dans le plan xOy . Ceci est assez intuitif, le dipôle tournant étant vu « par la tranche ». En un point M qui n'est ni sur l'axe Oz , ni dans le plan xOy , le rayonnement émis est polarisé elliptiquement, et le diagramme de rayonnement a une symétrie cylindrique.

Bien que ce diagramme de rayonnement soit différent de celui du dipôle linéaire, la puissance totale émise dans tout l'espace est la même (Eq. 20). Ceci peut se démontrer en intégrant l'équation (18) sur toutes les directions d'émission. Plus simplement, on peut remarquer que les champs émis par les composantes du dipôle suivant Ox et Oy sont en tout point en quadrature : les puissances moyennes rayonnées par chacune d'elles s'ajoutent sans terme croisé.

Le cas d'un dipôle tournant dans le sens rétrograde (cas σ_-), caractérisé par :

$$\mathcal{D}_0 = d_0 \vec{\varepsilon}_- \quad \text{avec} \quad \vec{\varepsilon}_- = \frac{1}{\sqrt{2}} (\mathbf{e}_x - i\mathbf{e}_y) \quad (24)$$

conduit de même à un rayonnement polarisé circulairement σ_- pour une émission dans la direction Oz , et à un rayonnement polarisé linéairement pour une émission dans le plan xOy .

Remarque

Un dipôle linéaire oscillant suivant Ox peut être considéré comme la superposition d'un dipôle tournant σ_+ et d'un dipôle tournant σ_- . Il est facile de vérifier que le champ rayonné, obtenu par addition des deux champs σ_+ et σ_- , est identique au rayonnement π par rapport à l'axe Ox . C'est donc un rayonnement polarisé linéairement, parfois appelé « σ linéaire ».

c. Cas de l'atome quantifié

Considérons un atome (ou une assemblée d'atomes) émettant de la lumière en passant d'un niveau $|b\rangle$ à un niveau $|a\rangle$, séparés par l'écart en énergie $E_b - E_a = \hbar\omega_0$. Le champ rayonné par ce système s'obtient en faisant un calcul quantique de la valeur moyenne du dipôle atomique $\mathbf{D}(t) = \langle \psi(t) | \hat{\mathbf{D}} | \psi(t) \rangle$, et en portant cette valeur dans les formules (17) et (18). On obtient les mêmes diagrammes de rayonnement que ceux décrits dans ce complément : plus précisément, si la transition a lieu entre niveaux ayant le même nombre quantique magnétique suivant Oz (transition $\Delta m = 0$ selon Oz), on a un diagramme de rayonnement du type π (voir le Complément II.1) ; si la transition fait varier m d'une unité ($\Delta m = \pm 1$), on a un diagramme de rayonnement de type σ_+ ou σ_- relativement à Oz . Notons enfin que l'élément de matrice dipolaire électrique est nul si $|\Delta m| > 1$ (règles de sélection des transitions dipolaires électriques, voir le Complément II.1), et qu'il n'y a alors pas de rayonnement dipolaire électrique émis dans ce cas.

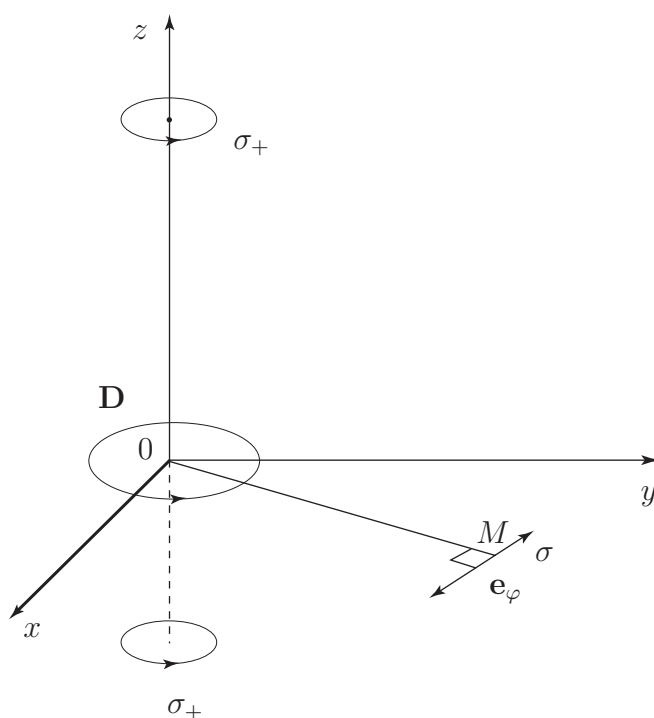


FIG. 3: *Emission dipolaire σ_+ . Pour un dipôle tournant en O , la lumière émise suivant Oz est polarisée circulairement ; elle est polarisée linéairement pour une émission dans le plan xOy du dipôle tournant, et elliptiquement dans une direction quelconque. Notons que le sens de rotation σ_+ est défini par rapport à l'orientation de l'axe Oz , et non par rapport à la direction de propagation ($\pm Oz$).*

4. Réaction de rayonnement. Amortissement radiatif. Moment cinétique de la lumière

a. Réaction de rayonnement

Par suite de la conservation de l'énergie, l'énergie électromagnétique rayonnée par l'électron doit être prélevée sur son énergie mécanique. Il doit donc exister une interaction entre l'électron et le champ qu'il émet en son propre emplacement. On peut effectivement montrer qu'il existe une force de ce type², appelée réaction de rayonnement. Cette force comporte un terme³ proportionnel à $\ddot{\mathbf{r}}$, que l'on peut interpréter comme une correction de la masse de l'électron, et un terme proportionnel $\dot{\mathbf{r}}$, qui s'écrit :

$$\mathbf{F}_{\text{RR}} = \frac{q^2}{6\pi\epsilon_0 c^3} \ddot{\mathbf{r}} = \frac{2mr_0}{3c} \ddot{\mathbf{r}} \quad (25a)$$

On a introduit la quantité r_0 , homogène à une longueur, et appelée « rayon classique de

²Voir CDG 1 p. 70, ou J.D. Jackson « Classical Electrodynamics », chapitre 17 (Wiley 1975).

³Comme dans ce paragraphe on ne s'intéresse qu'au mouvement de l'électron, on notera pour simplifier $\mathbf{r}$ et non $\mathbf{r}_e$ la position de l'électron.

l'électron » :

$$r_0 = \frac{q^2}{4\pi\epsilon_0 mc^2} \approx 2,8 \times 10^{-15} \text{ m} \quad (25b)$$

Nous ne chercherons pas à établir l'expression (25a), mais nous allons par contre vérifier qu'elle est compatible avec les résultats établis au § 2. Calculons pour cela la puissance moyenne fournie par cette force à un électron oscillant sinusoïdalement avec une amplitude a :

$$\frac{dW}{dt} = \overline{\mathbf{F}_{RR} \cdot \dot{\mathbf{r}}} = -\frac{q^2 \omega^4 a^2}{12\pi\epsilon_0 c^3} \quad (26)$$

On trouve une puissance reçue par l'électron négative, exactement opposée à la puissance moyenne rayonnée par l'électron (Eq. (20)), ce qui assure que le choix de $\mathbf{F}_{RR}$ présenté en (25a) est cohérent avec la conservation de l'énergie : l'énergie emportée sous forme de rayonnement électromagnétique est prélevée à l'énergie mécanique de l'électron qui est freiné. Cette force de freinage $\mathbf{F}_{RR}$ n'est autre que la contrepartie, du point de vue du mouvement de l'électron, de son rayonnement.

b. Amortissement radiatif d'un électron élastiquement lié

Pour évaluer l'effet de ce freinage, revenons à notre modèle d'électron élastiquement lié, rappelé vers l'origine par une force $-m\omega_0^2 \mathbf{r}$. Il obéit à l'équation du mouvement

$$m\ddot{\mathbf{r}} + m\omega_0^2 \mathbf{r} - \frac{2}{3} \frac{mr_0}{c} (\ddot{\mathbf{r}}) = 0 \quad (27)$$

Cherchons une solution s'écrivant (en notation complexe) $\mathcal{S}_0 e^{-i\Omega t}$. La pulsation Ω est alors donnée par :

$$\Omega^2 = \omega_0^2 - i \frac{2}{3} \frac{r_0}{c} \Omega^3 \quad (28)$$

On n'envisage ici que des fréquences Ω très inférieures à $c/r_0 \approx 10^{23} \text{ s}^{-1}$, et le terme imaginaire de (28) est petit devant le terme réel⁴. On peut alors chercher les solutions de (28) sous forme d'un développement perturbatif.

$$\Omega = \omega_0 + \Delta\Omega^{(1)} + \Delta\Omega^{(2)} + \dots \quad (29a)$$

Le terme du premier ordre donne

$$2\omega_0 \Delta\Omega^{(1)} = -i \frac{2}{3} \frac{r_0}{c} \omega_0^3 \quad (29b)$$

On obtient ainsi la solution au premier ordre en r_0/c que l'on écrit sous la forme

$$\Delta\Omega^{(1)} = -i \frac{\Gamma_{cl}}{2} = -i \frac{r_0}{3c} \omega_0^2 \quad (29c)$$

⁴Le domaine d'application de la théorie présentée dans ce complément est en fait limité à des fréquences bien inférieures à c/r_0 . Par exemple, les effets quantiques relativistes imposent une limite $\Omega \ll mc^2/\hbar = \alpha c/r_0$, où $\alpha \approx 1/137$ est la constante de structure fine.

Le mouvement de l'électron est donc finalement un mouvement oscillant amorti :

$$\mathbf{r} = \mathcal{S}_0 \exp(-\Gamma_{\text{cl}} t/2) \cos \omega_0 t \quad (30a)$$

avec :

$$\Gamma_{\text{cl}} = \frac{2}{3} \frac{r_0}{c} \omega_0^2 = \frac{4\pi}{3} \frac{r_0}{\lambda_0} \omega_0 \quad (30b)$$

formule dans laquelle nous avons introduit la longueur d'onde $\lambda_0 = 2\pi c/\omega_0$. La puissance rayonnée est donc amortie avec la constante de temps Γ_{cl}^{-1} que l'on appelle *durée de vie radiative classique*.

Il est intéressant de calculer la valeur de Γ_{cl}^{-1} pour des fréquences $\omega_0/2\pi$ correspondant à de la lumière visible. Pour le milieu du domaine visible (5×10^{14} Hz, soit une longueur d'onde $\lambda_0 = 0,6 \mu\text{m}$), on trouve une durée de vie radiative :

$$\Gamma_{\text{cl}}^{-1} \approx 16 \text{ ns} \quad (\lambda_0 = 0,6 \mu\text{m})$$

Dans le domaine des rayons X, on trouve des durées de vie beaucoup plus courts, puisque Γ_{cl} varie comme λ_0^{-2} . Notons néanmoins que, même dans ce domaine, le *facteur de qualité* $\omega_0/\Gamma_{\text{cl}}$ de la résonance, qui caractérise le nombre de périodes émises par l'électron avant que son mouvement ne soit notablement amorti, reste extrêmement grand devant 1.

Le résultat obtenu à partir du modèle classique peut être utilisé dans le cas d'un atome excité dans un état $|b\rangle$ et se désexcitant par émission spontanée vers un état $|a\rangle$, en émettant un rayonnement de pulsation ω_0 . La probabilité de trouver l'atome dans l'état $|b\rangle$ décroît exponentiellement au cours du temps, la durée de vie radiative Γ_{sp}^{-1} caractérisant cette décroissance étant de l'ordre de la valeur donnée par la formule (30b). Les valeurs calculées ci-dessus donnent donc généralement un bon ordre de grandeur pour les durées de vie radiatives atomiques. Mais un calcul exact requiert un traitement complètement quantique⁵.

c. Moment cinétique de la lumière polarisée circulairement

Considérons un électron tournant dans le plan xOy dans le sens direct (cas σ_+) à la fréquence ω . Un tel électron perd de l'énergie comme le montre la formule (26) (qui est valable également pour un mouvement circulaire de rayon $a/\sqrt{2}$), mais également du moment cinétique. Calculons la variation de moment cinétique par unité de temps $d\mathbf{L}/dt$. De la relation fondamentale de la dynamique, on déduit

$$\frac{d\mathbf{L}}{dt} = \mathbf{r} \times \mathbf{F}_{\text{RR}} \quad (31)$$

Les coordonnées de l'électron à l'instant t sont :

$$x = \frac{a}{\sqrt{2}} \cos \omega t \quad (32a)$$

$$y = \frac{a}{\sqrt{2}} \sin \omega t \quad (32b)$$

⁵Voir le cours de majeure 2 PHY557.

et les projections F_x et F_y de $\mathbf{F}_{\text{RR}}$ sont d'après (25a) et (32) égales à

$$F_x = \frac{2mr_0}{3c} \frac{a}{\sqrt{2}} \omega^3 \sin \omega t \quad (33a)$$

$$F_y = -\frac{2mr_0}{3c} \frac{a}{\sqrt{2}} \omega^3 \sin \omega t \quad (33b)$$

La variation du moment cinétique par unité de temps est donc

$$\frac{d}{dt} L_z = -\frac{1}{3} \frac{mr_0}{c} a^2 \Omega^3 \quad (34)$$

En remplaçant r_0 par sa définition (Eq. (25b)), on trouve

$$\frac{d}{dt} L_z = -\frac{q^2 \omega^3 a^2}{12\pi \varepsilon_0 c^3} \quad (35)$$

La comparaison de cette formule et de l'équation (26) montre que les pertes par rayonnement de moment cinétique et d'énergie sont dans un rapport constant :

$$\frac{dL_z}{dt} = \frac{1}{\omega} \frac{dW}{dt} \quad (36a)$$

Lorsque le mouvement de l'électron est totalement amorti, l'énergie W et le moment cinétique L_z du champ rayonné au cours du mouvement sont liés par la relation déduite de (36a) par intégration temporelle :

$$L_z = \frac{W}{\omega} \quad (36b)$$

Cette relation, obtenue par un calcul purement classique, est interprétable dans un cadre quantique : si l'état initial du système est un atome excité unique, l'énergie totale rayonnée W est égale à $\hbar\omega$, et donc *le moment cinétique du photon σ_+ ainsi émis a une composante L_z le long de Oz égale à $\hbar$.*

Le même raisonnement fait avec une rotation de l'électron dans le sens opposé (σ_-) conduirait à une égalité analogue à (36b), mais avec un signe opposé, et donc à un moment cinétique le long de Oz égal à $-\hbar$ pour un photon σ_- .

5. Diffusion Rayleigh, diffusion Thomson, diffusion résonnante

Le modèle de l'électron élastiquement lié se prête très bien à l'étude de la diffusion d'une onde électromagnétique. Il suffit en effet de déterminer le mouvement forcé de l'électron sous l'effet d'une onde incidente, et d'en déduire le rayonnement diffusé.

Supposons donc que l'on a un rayonnement incident

$$\mathbf{E}(t) = \mathbf{E}_0 \cos \omega t \quad (37a)$$

Si on néglige l'effet du champ magnétique, dont l'effet est en v/c par rapport à l'effet du champ électrique, l'amplitude du mouvement forcé s'obtient en résolvant une équation différentielle dont le membre de gauche coïncide avec celui de l'équation (27) et le membre de droite est $q\mathbf{E}(t)$. On obtient, en notation complexe (et en utilisant la définition (30b))

$$\mathcal{S}_0 = \frac{q\mathbf{E}_0}{m} \frac{1}{\omega_0^2 - \omega^2 - i\Gamma_{cl}\omega^3/\omega_0^2} \quad (37b)$$

d'où une puissance diffusée Φ_d (équation 20) qui vaut :

$$\Phi_d = \frac{q^4 \mathbf{E}_0^2}{12\pi\epsilon_0 c^3 m^2} \frac{\omega^4}{(\omega_0^2 - \omega^2)^2 + \Gamma_{cl}^2 \omega^6 / \omega_0^4} \quad (37c)$$

On rappelle qu'il est possible de définir la section efficace totale de diffusion $\sigma(\omega)$ par le rapport de la puissance totale diffusée au flux incident par unité de surface :

$$\Phi_d = \frac{d\Phi_i}{dS} \sigma(\omega) \quad (38a)$$

formule dans laquelle $\frac{d\Phi_i}{dS}$ est plus précisément la puissance moyenne incidente par unité de surface, c'est-à-dire

$$\frac{d\Phi_i}{dS} = \Pi_i = \frac{1}{2} \epsilon_0 c \mathbf{E}_0^2 \quad (38b)$$

On en déduit

$$\Rightarrow \sigma(\omega) = \frac{8\pi}{3} r_0^2 \frac{\omega^4}{(\omega_0^2 - \omega^2)^2 + \Gamma_{cl}^2 \omega^6 / \omega_0^4} \quad (39)$$

(i) Considérons tout d'abord le régime $\omega \ll \omega_0$ (*diffusion Rayleigh*). La section efficace vaut alors :

$$\Rightarrow \sigma(\omega) \approx \frac{8\pi}{3} r_0^2 \left(\frac{\omega}{\omega_0} \right)^4 \quad (40)$$

Elle croît très rapidement avec la fréquence du rayonnement incident. C'est le cas de la diffusion de la lumière solaire visible par les molécules de l'atmosphère, dont les transitions électroniques sont dans l'ultraviolet. La partie haute fréquence du spectre visible de la lumière solaire est donc la plus diffusée. Le ciel apparaît cependant bleu et non violet à l'œil humain d'une part parce que l'œil est plus sensible dans le bleu que dans le violet, et d'autre part à cause de la décroissance du spectre de la lumière solaire du bleu au violet.

(ii) À l'opposé, le régime de la *diffusion Thomson* est celui des fréquences très supérieures à la fréquence de résonance, $\omega \gg \omega_0$. La formule (39) montre qu'alors la section efficace de diffusion tend vers une constante qui vaut⁶ :

$$\sigma(\omega) \approx \frac{8\pi}{3} r_0^2 \approx 6,5 \times 10^{-29} \text{ m}^2 \quad (41)$$

⁶En vertu de l'hypothèse (ii), le modèle est limité à $\Gamma_{cl}\omega/\omega_0^2 \ll 1$, ce qui permet de négliger le second terme du dénominateur ($\omega_0^2/\Gamma_{cl} \approx 10^{23} \text{ s}^{-1}$ dans le domaine optique).

Ce régime est celui de la diffusion des rayons X par la matière. Cette formule a joué un rôle important dans la détermination, au début du XX^{ème} siècle, du nombre d'électrons de certains atomes à partir de mesures d'absorption de rayons X.

(iii) Considérons enfin le cas de fréquences voisines de la fréquence d'oscillation propre de l'électron, $\omega \approx \omega_0$ (*diffusion résonnante*). En ne gardant que les termes d'ordre le plus bas en $\omega - \omega_0$, et en se souvenant que Γ_{cl} est petit devant ω_0 , on obtient :

$$\sigma(\omega) = \frac{8\pi}{3} r_0^2 \frac{\omega_0^2}{4(\omega_0 - \omega)^2 + \Gamma_{cl}^2} \quad (42)$$

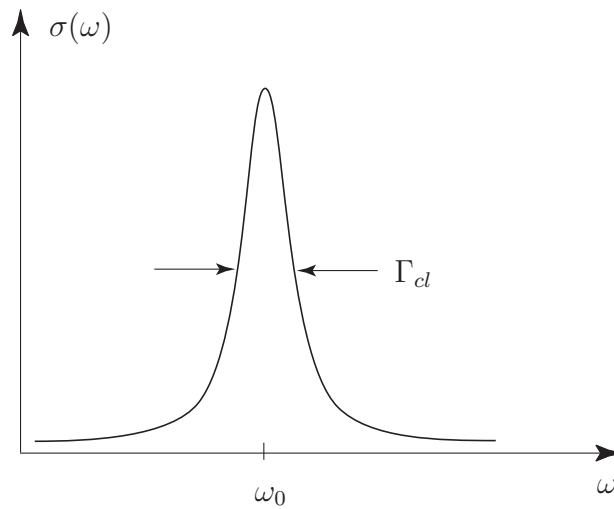


FIG. 4: Section efficace de diffusion résonnante. On a une variation Lorentzienne, de largeur à mi-hauteur Γ_{cl} . Pour de la lumière visible, le facteur de qualité de la résonance ω_0/Γ_{cl} est de l'ordre de 10^8 . Ces résultats restent qualitativement corrects en théorie quantique.

En reportant la définition (30b) de r_0 , on trouve la forme remarquable :

$$\Rightarrow \quad \sigma(\omega) = \frac{3\lambda_0^2}{2\pi} \frac{\Gamma_{cl}^2/4}{(\omega - \omega_0)^2 + \Gamma_{cl}^2/4} \quad (43)$$

La section efficace a un comportement résonnant de forme Lorentzienne autour de ω_0 . La figure 4 montre l'allure de la résonance, de largeur à mi-hauteur Γ_{cl} . Notons que la section efficace à résonance, $3\lambda_0^2/2\pi$ a pour ordre de grandeur le carré de la longueur d'onde du rayonnement à résonance. Elle est considérablement plus élevée que la section efficace Thomson (par 18 ordres de grandeur dans le domaine visible !) : il s'agit d'une *résonance exceptionnellement aiguë*.

Le calcul quantique donne des résultats analogues⁷ : on trouve en particulier que dans les deux cas extrêmes des diffusions Rayleigh et Thomson (Equations (40) et (41)) les expressions des sections efficaces sont pratiquement inchangées. Quant à la section

⁷Cours PHY557 de majeure 2 (Optique quantique 2).

efficace résonnante, elle garde la forme (42) à condition de remplacer Γ_{cl} par la largeur naturelle Γ_{sp} de la transition considérée. Notons que la valeur à résonance $3\lambda_0^2/2\pi$ est exactement la même.

6. Susceptibilité

Nous nous proposons d'établir l'expression de la susceptibilité d'un milieu constitué d'« atomes classiques » comprenant chacun un électron élastiquement lié, avec une densité atomique égale à N/V . D'après (37b), la polarisation du milieu est égale à :

$$\mathbf{P}(t) = \text{Re} \left\{ \frac{N}{V} \frac{q^2 \mathbf{E}_0 e^{-i\omega t}}{m} \frac{1}{\omega_0^2 - \omega^2 - i\Gamma_{cl}\omega^3/\omega_0^2} \right\} \quad (44)$$

En utilisant la définition (II.D.21b) de la susceptibilité en notation complexe, nous trouvons donc :

$$\chi = \frac{N}{V} \frac{q^2}{m\varepsilon_0} \frac{1}{\omega_0^2 - \omega^2 - i\Gamma_{cl}\omega^3/\omega_0^2} \quad (45)$$

soit, pour les parties réelle et imaginaire de la susceptibilité :

$$\chi' = \frac{N}{V} \frac{q^2}{m\varepsilon_0} \frac{\omega_0^2 - \omega^2}{(\omega_0^2 - \omega^2)^2 + (\Gamma_{cl}\omega^3/\omega_0^2)^2} \quad (46)$$

$$\chi'' = \frac{N}{V} \frac{q^2}{m\varepsilon_0} \frac{\Gamma_{cl}\omega^3/\omega_0^2}{(\omega_0^2 - \omega^2)^2 + (\Gamma_{cl}\omega^3/\omega_0^2)^2} \quad (47)$$

Au voisinage de la résonance, c'est-à-dire lorsque $\Gamma_{cl} \ll |\omega - \omega_0| \ll \omega_0$, la partie réelle de la susceptibilité prend la forme simple :

$$\chi' = \frac{N}{V} \frac{q^2}{2m\varepsilon_0} \frac{1}{\omega_0(\omega_0 - \omega)} = \frac{\chi'_Q}{f_{ab}} \quad (48a)$$

où χ'_Q désigne la partie réelle de la susceptibilité linéaire obtenue dans le cadre de la théorie quantique à faible intensité et l'approximation résonnante (équation II.D.22b). Le facteur sans dimension f_{ab} est appelé *force d'oscillateur* de la transition. Il vaut :

$$f_{ab} = \frac{2m\omega_0 z_{ab}^2}{\hbar} \quad (48b)$$

Par ailleurs, en utilisant (25b) et (30b), on peut transformer l'expression donnant la partie imaginaire de la susceptibilité pour trouver :

$$\chi'' = \frac{N}{V} \sigma(\omega) \frac{c}{\omega} \quad (49)$$

où la section efficace $\sigma(\omega)$ est donnée par l'expression (39). L'équation (49) montre la relation étroite existant entre le coefficient d'absorption de l'onde incidente (proportionnel à χ'') et la section efficace de diffusion $\sigma(\omega)$. On retrouve ainsi dans ce cas particulier un résultat très général⁸, connu sous le nom de *théorème optique*.

⁸Voir par exemple J. Jackson *Classical Electrodynamics* § 9.14, Wiley (New York 1975), ou A. Messiah, *Mécanique Quantique* § XIX-31, Dunod (Paris 1964).

7. Lien entre le modèle classique de l'électron élastiquement lié et le modèle quantique de l'atome à deux niveaux

Nous allons montrer que les résultats que nous venons d'obtenir dans le cadre du modèle de l'électron classique élastiquement lié, de fréquence caractéristique ω_0 , se transposent au cas d'un atome (ou d'une assemblée d'atomes), traité de façon quantique à l'approximation du *système à deux niveaux* $|a\rangle$ et $|b\rangle$, séparés par l'écart en énergie

$$E_b - E_a = \hbar\omega_0, \quad (50)$$

à condition que le système reste peu excité (loin de la saturation).

Utilisons pour cela l'approche semi-classique développée dans le chapitre 2, à l'approximation dipolaire électrique. L'hamiltonien du système vaut alors :

$$\hat{H} = \hat{H}_0 - q\hat{\mathbf{r}} \cdot \mathbf{E}_0 \cos \omega t \quad (51)$$

où $\hat{H}_0$ est l'hamiltonien de l'atome. La restriction de cet hamiltonien à l'espace sous-tendu par les deux niveaux $|a\rangle$ et $|b\rangle$ (où $|a\rangle$ est le niveau fondamental pris comme zéro d'énergie) s'écrit alors, si $\mathbf{E}_0$ est pris parallèle à Oz :

$$\hat{H} = \hbar\omega_0 \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} - qz_{ab}E_0 \cos \omega t \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad (52)$$

avec $z_{ab} = \langle a|\hat{z}|b\rangle$, supposé réel. Récrivons $\hat{H}$ en fonction de la matrice identité $\hat{I}$ et des matrices de Pauli⁹ $\hat{\sigma}_x, \hat{\sigma}_y$ et $\hat{\sigma}_z$:

$$\hat{\sigma}_x = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad \hat{\sigma}_y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \quad \hat{\sigma}_z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \quad (53)$$

On obtient :

$$\hat{H} = \frac{\hbar\omega_0}{2}(\hat{I} - \hat{\sigma}_z) - qz_{ab}\hat{\sigma}_xE_0 \cos \omega t. \quad (54)$$

Rappelons les relations de commutation des matrices de Pauli, qui s'obtiennent par permutation circulaire à partir de :

$$[\hat{\sigma}_x, \hat{\sigma}_y] = 2i\hat{\sigma}_z. \quad (55)$$

Pour comparer avec le calcul classique, il nous faut déterminer l'équation qui régit l'évolution de la *valeur moyenne de la position* $\hat{\mathbf{r}}$ de l'électron atomique, et plus précisément, de la coordonnée $\hat{z}$. Nous allons pour cela utiliser le théorème d'Ehrenfest :

$$\begin{aligned} \frac{d}{dt}\langle \hat{z} \rangle &= \frac{1}{i\hbar}\langle [\hat{z}, \hat{H}] \rangle = \frac{z_{ab}}{i\hbar}\langle [\hat{\sigma}_x, \hat{H}] \rangle \\ &= \frac{z_{ab}\omega_0}{2i}\langle [\hat{\sigma}_x, \hat{\sigma}_z] \rangle = \omega_0 z_{ab}\langle \hat{\sigma}_y \rangle \end{aligned} \quad (56)$$

⁹Voir par exemple JLB Chapitre VI.1.

On en déduit la dérivée seconde de $\langle \hat{z} \rangle$ par rapport au temps :

$$m \frac{d^2}{dt^2} \langle \hat{z} \rangle = m \omega_0 z_{ab} \frac{d}{dt} \langle \hat{\sigma}_y \rangle = \frac{m \omega_0 z_{ab}}{i \hbar} \langle [\hat{\sigma}_y, \hat{H}] \rangle \quad (57)$$

Compte tenu de l'expression (54) de $\hat{H}$, le commutateur présente deux termes non nuls :

$$\begin{aligned} m \frac{d^2}{dt^2} \langle \hat{z} \rangle &= \frac{m \omega_0 z_{ab}}{i \hbar} \left(-\frac{\hbar \omega_0}{2} \langle [\hat{\sigma}_y, \hat{\sigma}_z] \rangle - q E_0 z_{ab} \cos \omega t \langle [\hat{\sigma}_y, \hat{\sigma}_x] \rangle \right) \\ &= m \omega_0^2 z_{ab} \langle \hat{\sigma}_x \rangle + \frac{2 m \Omega_0 z_{ab}^2}{\hbar} q E_0 \cos \omega t \langle \hat{\sigma}_z \rangle \end{aligned} \quad (58)$$

Cette équation s'exprime finalement en fonction de la *force d'oscillateur* f_{ab} de la transition $|a\rangle \rightarrow |b\rangle$, donnée par l'expression (48b) de ce complément, et de la force classique $F_{cl} = q E_0 \cos \omega t$ agissant sur l'électron :

$$m \frac{d^2}{dt^2} \langle \hat{z} \rangle = -m \omega_0^2 \langle \hat{z} \rangle + f_{ab} \langle \hat{\sigma}_z \rangle F_{cl} \quad (59)$$

Cette équation coïncide avec l'équation fondamentale de la dynamique classique pour la position moyenne de l'électron, avec la force de rappel harmonique $-m \omega_0^2 \langle \hat{z} \rangle$, à deux différences près :

- la présence de la « force d'oscillateur » f_{ab} , qui corrige la force classique d'un facteur sans dimension dépendant de la transition considérée. Pour la raie de résonance d'atomes alcalins comme le sodium par exemple, cette force d'oscillateur est peu différente de un.
- la présence du facteur $\langle \hat{\sigma}_z \rangle$, qui est proportionnel à $P_a - P_b$, P_a et P_b étant les probabilités pour l'atome d'être respectivement dans l'état a ou b .

Si le rayonnement incident sur l'atome est peu intense, on a $P_a \simeq 1$ et $P_b \simeq 0$ et on retrouve donc exactement l'équation de la dynamique classique. Comme nous l'avions annoncé, le modèle de l'électron élastiquement lié est dans ce cas équivalent au traitement semi-classique dans lequel l'atome est quantifié, tant qu'on ne s'intéresse qu'à des quantités qui dépendent de la valeur moyenne du dipôle électrique atomique. On comprend aussi les succès de ce modèle pour décrire l'interaction matière lumière, dominée par l'interaction dipolaire électrique.

Lorsque le rayonnement incident est suffisamment intense pour que la population P_b de l'état excité ne soit pas négligeable, on voit que tout se passe comme si la force appliquée étant réduite par le facteur $P_b - P_a = 1 - 2P_a$. Il s'agit en fait du terme de saturation $\frac{1}{1+s}$ que l'on a vu apparaître dans le chapitre 2, par exemple 2.2.121 ou 2.2.148, et qui vient simplement réduire la réponse linéaire caractérisée par la susceptibilité linéaire de χ_1 . Nous avons donc justifié l'affirmation suivant laquelle la réponse linéaire de la matière est correctement décrite par un modèle entièrement classique, mais le phénomène de saturation est caractéristique du modèle quantique de système à 2 niveaux.

Remarque

On peut construire des modèles classiques de la matière faisant apparaître des non-linéarités, par exemple en prenant un électron dans un potentiel enharmonique. Mais il n'existe pas de modèle classique simple pour le phénomène de saturation.

Complément II.4

Effet photoélectrique

L'effet photoélectrique est le phénomène d'éjection d'électrons hors de la matière sous l'effet d'un rayonnement incident. D'abord observé par Hertz dans le cas d'un métal (les électrodes de son « éclateur » se déchargeaient lorsqu'elles étaient éclairées par un rayonnement ultraviolet), l'effet photoélectrique apparaît aussi dans le cas d'atomes (ou de molécules) isolés qui peuvent être ionisés par un rayonnement ultraviolet ou X.

L'existence d'une *fréquence seuil* ω_s , en dessous de laquelle le rayonnement est incapable de provoquer l'éjection d'un électron, est une caractéristique de l'effet photoélectrique que la physique classique (plus précisément l'électrodynamique classique) ne pouvait expliquer. C'est ce qui conduisit Einstein à postuler, en 1905, qu'un rayonnement monochromatique de fréquence ω contient des « grains de lumière », que l'on appella plus tard des photons, ayant chacun une énergie¹

$$E_{\text{photon}} = \hbar\omega \quad (1)$$

Einstein interprète alors l'effet photoélectrique comme une collision entre un électron lié et un photon qui disparaît, cédant toute son énergie à l'électron. Si on appelle E_I l'énergie d'ionisation, c'est-à-dire l'énergie qu'il faut communiquer à l'électron pour l'arracher à son atome (ou le travail de sortie hors d'un métal), l'équation d'Einstein relie l'énergie cinétique de l'électron éjecté à la fréquence du rayonnement (Fig. 1) :

$$E_e = \hbar\omega - E_I \quad (2a)$$

Cette énergie cinétique étant nécessairement positive, la fréquence du rayonnement doit être supérieure à la fréquence seuil

$$\omega_s = \frac{E_I}{\hbar} \quad (2b)$$

¹A. Einstein, *Annalen der Physik* 17, 132 (1905). Traduction française dans A. Einstein, *Œuvres choisies*, Tome I : *Quanta* (Seuil-CNRS, Paris (1993)). Notons que l'hypothèse d'Einstein est beaucoup plus radicale que celle faite par Planck en 1900 pour interpréter le spectre du corps noir. Ce dernier supposait en effet seulement que les *échanges* d'énergie entre matière et rayonnement sont quantifiés.

Après la confirmation expérimentale indubitable par Millikan², en 1915, de la validité de l'équation (2), Einstein allait recevoir en 1921 le prix Nobel « pour ses services rendus à la physique théorique, et particulièrement pour sa découverte de la loi de l'effet photoélectrique ».

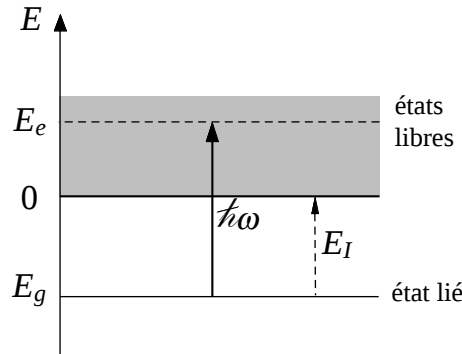


FIG. 1: **Interprétation d'Einstein de l'effet photoélectrique.** L'électron lié, d'énergie négative $E_g = -E_I$ peut être porté dans un état libre par un photon d'énergie $\hbar\omega$ supérieure à E_I .

Dans ce complément, nous allons pourtant étudier l'effet photoélectrique à partir d'un modèle semi-classique dans lequel la lumière n'est pas quantifiée. On sait en effet aujourd'hui que l'effet photoélectrique s'interprète très bien « sans photons », à condition de traiter la matière (l'atome) comme un objet quantique. Cela n'enlève naturellement rien au génie d'Einstein, qui avait reconnu la nécessité d'une quantification (en 1905, on ne savait pas plus quantifier la matière que la lumière, puisque l'atome de Bohr date de 1913). Mais il est intéressant de noter que les résultats que nous établirons ici sont parfaitement exacts, même s'il est plus élégant (et plus cohérent) de traiter l'effet photoélectrique dans le cadre de la théorie complètement quantique de l'interaction matière-rayonnement.

Ce complément permet de rendre concrète la question du couplage entre un niveau discret et un quasi-continuum, sous l'effet d'une perturbation dépendant sinusoïdalement du temps (cf. § C.4 du chapitre I). En ce sens, il peut être considéré comme un exercice corrigé, permettant d'explicitier les calculs, et d'obtenir des résultats quantitatifs.

²De nombreuses citations attestent que Milikan s'attendait à infirmer cette équation. Comme la majorité des physiciens de son époque, il jugeait l'hypothèse d'Einstein incompatible avec toutes les propriétés ondulatoires connues de la lumière (interférences, diffraction). Voir par exemple : A. Pais, *Albert Einstein, La vie et l'œuvre*, Chap. 18 (Interditions, Paris 1993).

1 Modèle utilisé

1.1 Etat atomique lié

On considère un électron atomique dans un niveau lié $|g\rangle$. C'est par exemple l'électron de l'atome d'hydrogène, ou l'électron de valence d'un alcalin, dans son état fondamental ou dans un état excité lié. Cela peut être aussi un électron d'une couche profonde d'un atome à plusieurs électrons. En prenant comme origine des énergies l'énergie potentielle de l'électron loin du noyau, l'énergie de l'état E_g est négative :

$$E_g = -E_I \quad (3)$$

E_I étant l'énergie d'ionisation.

Pour fixer les idées, donnons quelques ordres de grandeur. L'état fondamental de l'atome d'hydrogène a une énergie d'ionisation de 13,6 eV. La fréquence seuil associée est

$$\frac{\omega_s}{2\pi} = \frac{E_I}{h} = 3,3 \times 10^{15} \text{ Hz} \quad (4)$$

ce qui correspond à une longueur d'onde dans l'ultraviolet lointain ($\lambda_S = 91 \text{ nm}$).

Les états excités de l'atome d'hydrogène ont des énergies d'ionisation plus faibles. En revanche, les électrons internes des atomes « lourds » (de numéro atomique élevé) ont des énergies d'ionisation beaucoup plus grandes, pouvant dépasser le keV, ce qui correspond à des longueurs d'onde seuil inférieures au nanomètre (rayons X).

Lorsque nous aurons besoin d'expressions explicites pour la fonction d'onde de l'électron lié (ou sa transformée de Fourier), nous utiliserons les fonctions d'onde de l'atome d'hydrogène qui sont connues³.

1.2 Etats ionisés - Densité d'états

Les états finaux du processus de photoionisation sont des états d'énergie positive, correspondant à un état non-lié de l'électron, qui peut s'éloigner librement du noyau : ce sont les états ionisés de l'atome.

Dans le cas de l'atome d'hydrogène, ces états ont des expressions explicites³. Néanmoins, dans le but de discuter plus simplement les questions attachées au quasi-continuum, nous nous intéresserons uniquement aux états non-liés de grande énergie, qui sont très peu affectés par le potentiel attractif coulombien du noyau, de sorte qu'ils peuvent être décrits par des ondes planes de de Broglie :

$$\psi_e(\mathbf{r}) = \alpha \exp i\mathbf{k}_e \cdot \mathbf{r} \quad (5a)$$

³Voir par exemple : H.A. Bethe et E.E. Salpeter, *Quantum Mechanics of one - and two - electron atoms*, Plenum/Rosetta(New-York, 1977).

Ce sont donc des états propres de l'opérateur « impulsion » :

$$\hat{\mathbf{P}}|\psi_e\rangle = \hbar\mathbf{k}_e|\psi_e\rangle \quad (5b)$$

Leur énergie vaut :

$$E_e = \frac{\hbar^2 k_e^2}{2m} \quad (5c)$$

(m étant la masse de l'électron).

Pour faire des calculs explicites, il faut que les fonctions d'ondes soient normalisées, ce qui n'est a priori pas le cas des ondes planes (5a). Nous adopterons alors une démarche analogue à celle du paragraphe C.2.a du chapitre I, consistant à admettre que l'électron est confiné dans un cube d'arête L très grande devant les dimensions atomiques, et dont les faces sont perpendiculaires aux vecteurs unitaires $\mathbf{e}_x$, $\mathbf{e}_y$, $\mathbf{e}_z$ des trois axes. La fonction d'onde $\psi_e(\mathbf{r})$ étant nulle hors du cube, on en déduit la condition de normalisation

$$|\alpha| = \frac{1}{L^{3/2}} \quad (5d)$$

Par simplicité, on prendra α réel positif.

On sait que les fonctions d'ondes doivent respecter des conditions aux limites sur les faces du cube, ce qui conduit à la discrétisation des énergies et des impulsions. Au lieu de la condition « physique » associée à un véritable puits de potentiel ($\psi_e(\mathbf{r})$ nulle sur les faces du cube) on adoptera la *condition aux limites périodiques*⁴ qui est équivalente mais qui conduit à des calculs plus commodes pour la suite. On écrit donc

$$\psi_e(\mathbf{r} + L\mathbf{e}_x) = \psi_e(\mathbf{r}) \quad (6)$$

et de même suivant les autres axes. Les valeurs de $\mathbf{k}_e$ possibles sont alors

$$\mathbf{k}_e = \frac{2\pi}{L}(n_x\mathbf{e}_x + n_y\mathbf{e}_y + n_z\mathbf{e}_z) \quad (7a)$$

avec

$$(n_x, n_y, n_z) \text{ entiers relatifs} \quad (7b)$$

On peut représenter, dans l'espace des vecteurs d'onde (espace réciproque), les vecteurs d'onde $\mathbf{k}_e$ donnés par la condition (A.5). Leurs extrémités forment un *réseau cubique* de pas $2\pi/L$ (Figure 2).

La procédure de discrétisation ci-dessus aboutit donc à transformer le continuum des états libres en quasi-continuum : les écarts d'énergie entre les états tendent vers 0 quand la longueur L tend vers l'infini. En fait, la grandeur L est arbitraire, et les résultats physiques ne doivent pas en dépendre.

⁴Voir Complément I.1, § 4.

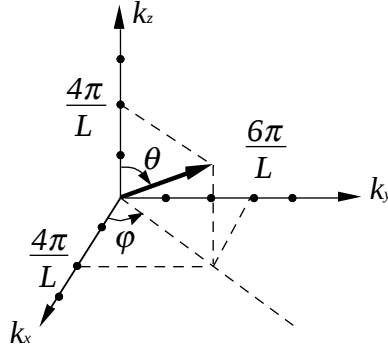


FIG. 2: **Représentation des vecteurs d'onde des électrons libres dans l'espace réciproque.** Compte tenu des conditions aux limites périodiques (période L) leurs extrémités forment un réseau cubique d'arête $2\pi/L$. A titre d'exemple, on a représenté le vecteur d'onde $\mathbf{k}_e$ associé à $(n_x = 2, n_y = 3, n_z = 2)$. Ce vecteur d'onde peut aussi être repéré par son module k_e et par sa direction (θ, φ) .

Pour pouvoir appliquer le formalisme du paragraphe C du chapitre I, et en particulier la règle d'or de Fermi, il faut connaître la densité d'états du quasi-continuum. Nous avons ici des états dégénérés en énergie, et il faut faire apparaître explicitement la densité d'états sous la forme $\rho(\beta, E)$ introduite au paragraphe C.3.c du chapitre I. Notons d'abord que la densité d'états dans l'espace réciproque est uniforme, et vaut

$$\rho(\mathbf{k}_e) = \left(\frac{L}{2\pi}\right)^2 \quad (8)$$

puisque chaque état est associé à une maille cubique de volume $2\pi/L^3$ (voir Figure 2). Pour faire apparaître l'énergie, nous repèrerons un vecteur d'onde électronique $\mathbf{k}_e$ par son module, qui est relié à l'énergie à l'aide de l'équation (5c) et par sa direction (θ, φ) . Le volume élémentaire de l'espace réciproque est alors l'intersection d'un élément d'angle solide $d\Omega$ et d'une coquille située entre les sphères de rayon k_e et $k_e + dk_e$. Il a pour mesure :

$$d^3k_e = k_e^2 dk_e d\Omega \quad (9a)$$

Le nombre d'états contenus dans ce volume est

$$d^3N = \left(\frac{L}{2\pi}\right)^3 k_e^2 dk_e d\Omega = \left(\frac{L}{2\pi\hbar}\right)^3 \sqrt{2E_e m^3} dE_e d\Omega \quad (9b)$$

La densité d'états ionisés associée à (E_e, Ω) s'écrit donc finalement

$$\rho(E_e, \theta, \varphi) = \frac{d^3N}{dE_e d\Omega} = \left(\frac{L}{2\pi\hbar}\right)^3 \sqrt{2E_e m^3} = \left(\frac{L}{2\pi\hbar}\right)^3 m\hbar k_e \quad (10)$$

C'est cette expression qui sera utilisée pour évaluer la probabilité d'éjecter un photoélectron d'énergie donnée dans un élément d'angle solide $d\Omega$ autour d'une direction déterminée.

Remarque

Si au lieu de caractériser l'élément d'angle solide par sa mesure $d\Omega$ on utilisait les éléments d'angle $d\theta$ et $d\varphi$, la densité d'états correspondante $\frac{d^3N}{dE_e d\theta d\varphi}$ serait obtenue en multipliant l'expression (10) de la densité d'états par $\sin \theta$.

1.3 Hamiltonien d'interaction

L'atome est irradié par une onde électromagnétique décrite, en jauge de Coulomb, par les potentiels

$$\begin{aligned}\mathbf{A}(\mathbf{r}, t) &= \vec{\epsilon} A_0 \cos(\omega t - \mathbf{k} \cdot \mathbf{r}) \\ U(\mathbf{r}, t) &= 0\end{aligned}\tag{11}$$

L'interaction entre l'atome et le champ sera décrite à l'ordre le plus bas par l'hamiltonien (cf. Équation (2.40d) du chapitre II) :

$$\hat{H}_{I1} = -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}(\hat{\mathbf{r}}, t)\tag{12}$$

q étant la charge de l'électron. Notons que nous ne faisons pas l'approximation des grandes longueurs d'onde. Notre traitement sera donc valable pour des rayons X, dont la longueur d'onde n'est pas grande devant les dimensions atomiques. Il s'appliquera également à la photoionisation à partir d'états atomiques très excités, dont les dimensions sont beaucoup plus grandes que les dimensions atomiques habituelles.

Nous allons chercher à calculer l'élément de matrice $\langle e | \hat{H}_{I1} | g \rangle$ responsable des transitions entre l'état lié $|g\rangle$ et un état ionisé $|e\rangle$:

$$\begin{aligned}\langle e | \hat{H}_{I1} | g \rangle &= \langle e | -\frac{q}{m} \hat{\mathbf{p}} \cdot \mathbf{A}(\hat{\mathbf{r}}, t) | g \rangle \\ &= -\frac{q}{2m} A_0 e^{-i\omega t} \langle e | \hat{\mathbf{p}} \cdot \vec{\epsilon} e^{i\mathbf{k} \cdot \hat{\mathbf{r}}} | g \rangle - \frac{q}{2m} A_0 e^{i\omega t} \langle e | \hat{\mathbf{p}} \cdot \vec{\epsilon} e^{-i\mathbf{k} \cdot \hat{\mathbf{r}}} | g \rangle\end{aligned}\tag{13}$$

Notons d'abord que l'énergie de l'état final étant supérieure à celle de l'état initial seul le terme en $\exp(-i\omega t)$ de (13) est susceptible de provoquer une transition résonnante, en donnant une amplitude de transition non-négligeable pour un état final d'énergie E_e :

$$E_e = \hbar\omega + E_g = \hbar\omega - E_I\tag{14}$$

(cf § B.4.d du chapitre I).

L'approximation résonnante consiste à ne garder que le premier terme du second membre de (13). Elle est valable dès que le temps d'interaction est très supérieur à $1/\omega_s$ (soit à $10^{-16}s$ pour une photoionisation à partir de l'état fondamental de l'atome d'hydrogène). On écrit donc

$$\langle e | \hat{H}_{I1} | g \rangle \approx \frac{1}{2} e^{-i\omega t} W_{eg}\tag{15a}$$

avec

$$W_{eg} = -\frac{q}{m} A_0 \langle e | \hat{\mathbf{p}} \cdot \vec{\varepsilon} e^{i\mathbf{k} \cdot \hat{\mathbf{r}}} | g \rangle \quad (15b)$$

L'approximation où l'état excité est une onde plane électronique décrite par la fonction d'onde (5a) permet de transformer simplement (15b) :

$$\begin{aligned} W_{eg} &\approx -\frac{q}{m} A_0 \hbar (\mathbf{k}_e \cdot \vec{\varepsilon}) L^{-\frac{3}{2}} \int d^3r e^{i(\mathbf{k}-\mathbf{k}_e) \cdot \mathbf{r}} \psi_g(\mathbf{r}) \\ &= -\left(\frac{2\pi}{L}\right)^{\frac{3}{2}} \frac{q A_0}{m} \hbar k_e (\mathbf{u}_e \cdot \vec{\varepsilon}) \tilde{\psi}_g(\mathbf{k}_e - \mathbf{k}) \end{aligned} \quad (16a)$$

On a introduit dans l'expression ci-dessus la transformée de Fourier $\tilde{\psi}_g(\mathbf{q})$ de la fonction d'onde $\psi_g(\mathbf{r})$

$$\tilde{\psi}_g(\mathbf{q}) = (2\pi)^{-\frac{3}{2}} \int d^3r e^{-i\mathbf{q} \cdot \mathbf{r}} \quad (16b)$$

(On prendra garde à ne pas faire de confusion entre le vecteur $\mathbf{q}$ de l'espace réciproque, intervenant dans les transformées de Fourier, et la charge q de l'électron qui apparaît dans le terme de couplage de l'hamiltonien d'interaction, ou dans le rayon de Bohr). C'est, à un facteur $\hbar^{-3/2}$ près, la fonction d'onde en représentation d'impulsion. Par ailleurs, on a noté

$$\mathbf{u}_e = \frac{\mathbf{k}_e}{k_e} \quad (16c)$$

le vecteur unitaire colinéaire à la direction d'éjection de l'électron.

Dans le cas de l'atome d'hydrogène, on connaît les transformées de Fourier $\tilde{\psi}_g(\mathbf{q})$ des états liés³. Pour l'état fondamental, on a

$$\tilde{\psi}_{n=1,l=0}(\mathbf{q}) = \frac{2\sqrt{2}}{\pi} \frac{1}{(\mathbf{q}^2 a_0^2 + 1)^2} a_0^{\frac{3}{2}} \quad (17a)$$

a_0 étant le rayon de Bohr :

$$a_0 = \frac{4\pi\varepsilon_0 \hbar^2}{q^2 m} = 0,53 \times 10^{-10} \text{ m} \quad (17b)$$

Les transformées de Fourier des états liés excités³ s'expriment à l'aide de fonctions spéciales ayant les valeurs asymptotiques suivantes pour les états de symétrie sphérique ($l=0$) :

$$\tilde{\psi}_{n,l=0}(\mathbf{q}=0) = (-1)^{n-1} \frac{2\sqrt{2}}{\pi} a_0^{\frac{3}{2}} n^{\frac{5}{2}} \quad (18a)$$

et, pour $n|\mathbf{q}|a_0 \gg 1$

$$\tilde{\psi}_{n,l=0}(\mathbf{q}) \approx \frac{2\sqrt{2}}{\pi} \left(\frac{a_0}{n}\right)^{\frac{3}{2}} \left(\frac{1}{|\mathbf{q}|a_0}\right)^4 \quad (18b)$$

Ces expressions nous permettront de calculer explicitement l'élément de matrice (15b), pourvu que l'amplitude A_0 de l'onde soit connue. Pour une onde incidente progressive, transportant une puissance par unité de surface (éclairage) Π , on sait que :

$$\Pi = \varepsilon_0 c \frac{\omega^2 A_0^2}{2} \quad (19)$$

puisque l'amplitude du champ électrique est égale à ωA_0 (cf. Éq. (2.61c)). Dans le cadre de l'interprétation d'Einstein de l'effet photoélectrique, il est utile d'introduire le flux de photons par unité de surface et de temps

$$\Pi_{\text{phot}} = \frac{\Pi}{\hbar\omega} = \varepsilon_0 c \frac{\omega A_0^2}{2\hbar} \quad (20)$$

Nous connaissons maintenant toutes les grandeurs nécessaires au calcul de la probabilité de photoionisation.

2 Taux de photoionisation, section efficace de photoionisation

2.1 Taux de photoionisation

Comme on l'a vu au paragraphe C du chapitre I, le traitement perturbatif du couplage entre un niveau discret et un quasi-continuum conduit, après une intégration sur les états finaux accessibles, à une probabilité de transition croissant linéairement avec le temps. On peut alors définir un *taux de transition* par unité de temps, proportionnel au carré du module de l'élément de matrice de l'hamiltonien de couplage, pris entre l'état discret initial et les états du quasi-continuum pour lesquels il y a résonance. Il s'agit ici d'états d'énergie donnée par la relation (14).

Pour calculer la valeur précise du taux de photoionisation, il suffit d'utiliser la règle d'or de Fermi. Plus précisément, nous nous référons à l'expression (1.93) du chapitre I, relative à un continuum dégénéré, et étendue au cas d'une perturbation sinusoïdale ce qui conduit à introduire un facteur 1/4 (cf. Éq. (1.98)). Le taux de photoionisation par élément d'angle solide s'écrit alors

$$\frac{d\Gamma}{d\Omega} = \frac{\pi}{2\hbar} |W_{eg}|^2 \rho(E_e, \theta, \varphi) \quad (21)$$

formule dans laquelle l'élément de matrice W_{eg} (16a), et la densité d'états finaux (10) sont pris pour une énergie finale (ou pour la valeur correspondante de k_e) donnée par la relation (14) :

$$E_e = \frac{\hbar^2 k_e^2}{2m} = \hbar\omega + E_S = \hbar\omega - E_I = \hbar(\omega - \omega_s) \quad (22)$$

Comme attendu, L s'élimine de l'expression finale, et on obtient le taux de photoéjection d'un électron dans une direction $\mathbf{u}_e$

$$\frac{d\Gamma}{d\Omega}(\mathbf{u}_e) = \frac{\pi q^2 A_0^2 k_e^3}{2m\hbar} (\mathbf{u}_e \cdot \vec{\epsilon})^2 |\tilde{\psi}_g(\mathbf{k}_e - \mathbf{k})|^2 \quad (23)$$

En intégrant sur toutes les directions d'émission possible, on obtient le taux global de photoionisation à partir de l'état $|g\rangle$:

$$\Gamma = \frac{\pi q^2 A_0^2}{2m\hbar} k_e^3 \int d\Omega (\mathbf{u}_e \cdot \vec{\epsilon})^2 |\tilde{\psi}_g(\mathbf{k}_e - \mathbf{k})|^2 \quad (24)$$

Les formules ci-dessus rendent parfaitement compte de toutes les observations relatives à l'effet photoélectrique. Tout d'abord, l'équation (22), qui donne la valeur de l'énergie cinétique de l'électron émis, n'est autre que l'équation d'Einstein, en parfait accord avec les mesures de Millikan. Mais notons que cette valeur de l'énergie finale n'est pas obtenue par un argument de conservation de l'énergie *a priori*, comme dans le raisonnement d'Einstein. Ici, la conservation de l'énergie apparaît comme une conséquence de la résonance de l'amplitude de transition lorsque les énergies initiale et finale sont liées par l'équation (14).

Par ailleurs, on trouve bien un taux de transition proportionnel au carré de l'amplitude de l'onde électromagnétique, c'est à dire à l'éclairement (19). On remarque également que le taux de photoionisation dans une direction $\mathbf{u}_e$ donnée varie comme $(\mathbf{u}_e \cdot \vec{\epsilon})^2$, c'est-à-dire comme le carré du cosinus de l'angle entre le champ électrique de l'onde et la direction d'éjection : or on observe effectivement que l'électron est éjecté préférentiellement suivant la direction du champ électrique.

Ainsi, sans introduire la notion de photon, nous pouvons rendre compte des caractéristiques essentielles de l'effet photoélectrique. Le traitement présenté ici permet en fait bien d'autres prédictions, dont nous donnons quelques exemples dans la suite.

Remarque

On dit quelque fois que l'observation de photoélectrons immédiatement après le début de l'illumination, même dans le cas d'un éclairage très faible, est en désaccord avec le traitement présenté ici. L'argument serait que le flux du vecteur de Poynting à travers une surface de l'ordre de la taille atomique est trop faible pour fournir l'énergie de l'électron en un temps aussi court. Cette observation imposerait donc d'introduire la notion du photon.

En fait, les taux de transition calculés ici ne sont autres que des probabilités quantiques, et il faut se souvenir qu'une probabilité faible ne veut pas dire qu'un événement ne se produit pas : cela veut dire que l'événement ne se produit que pour une faible fraction des cas. Si on répète un grand nombre de fois l'expérience d'illumination soudaine de l'échantillon, et que l'on construit l'histogramme des instants de détection des électrons en prenant comme origine l'instant d'illumination, on trouve un histogramme plat, en accord avec l'équation (23) ou (24). En fait, on peut interpréter la photoionisation à temps court comme une conséquence de la nature quantique de l'atome, et ici encore la quantification du rayonnement ne s'impose pas.

2.2 Section efficace de photoionisation

Même si elle n'est pas strictement imposée sur le plan logique à ce niveau, l'interprétation de la photoionisation comme une collision entre un photon incident et l'atome est élégante et commode. Il est donc judicieux d'introduire la quantité obtenue en divisant le taux de photoionisation (23) ou (24) par le flux de photons par unité de surface et de temps (Equation (20)). Cette quantité a les dimensions d'une surface : c'est la *section efficace de photoionisation* σ_g . Le taux de photoionisation prend alors la forme suggestive

$$\Gamma = \sigma_g \frac{\Pi}{\hbar\omega} = \sigma_g \Pi_{\text{phot}} \quad (25)$$

où le nombre de photoionisations par unité de temps apparaît comme le produit de la *section efficace totale de photoionisation* par le flux de photons par unité de surface et de temps.

En utilisant (24), on obtient

$$\sigma_g = \frac{\pi q^2}{m\varepsilon_0 c \omega} k_e^3 \int d\Omega (\mathbf{u}_e \cdot \boldsymbol{\varepsilon})^2 |\tilde{\psi}_g(\mathbf{k}_e - \mathbf{k})|^2 \quad (26a)$$

qui s'écrit encore

$$\sigma_g = 2\pi r_0 \lambda k_e^3 \int d\Omega (\mathbf{u}_e \cdot \boldsymbol{\varepsilon})^2 |\tilde{\psi}_g(\mathbf{k}_e - \mathbf{k})|^2 \quad (26b)$$

Dans l'équation ci-dessus, on a introduit la longueur d'onde λ du rayonnement incident, et le « rayon classique de l'électron » :

$$r_0 = \frac{q^2}{4\pi\varepsilon_0 m c^2} = \frac{1}{a_0} \frac{\hbar^2}{m^2 c^2} = 2,8 \times 10^{-15} m \quad (26c)$$

(on a utilisé l'expression (17b) du rayon de Bohr a_0).

Revenant au taux de photoionisation dans une direction donnée (23) on peut définir la *section efficace différentielle* de photoionisation dans la direction $\mathbf{u}_e$

$$\frac{d\sigma_g}{d\Omega}(\mathbf{u}_e) = 2\pi r_0 \lambda k_e^3 (\mathbf{u}_e \cdot \boldsymbol{\varepsilon})^2 |\tilde{\psi}_g(\mathbf{k}_e - \mathbf{k})|^2 \quad (27)$$

2.3 Comportement aux temps longs

A l'issue de ce calcul, il convient naturellement de vérifier que les approximations qui ont été faites sont valables. En particulier, il faut s'assurer qu'il existe une échelle de temps assez longs pour que l'approximation résonnante soit valide (cf. § A.3) et assez courts pour que la probabilité Γt d'ionisation par atome reste petite devant 1. Ceci impose une limite supérieure à l'intensité de l'onde incidente. Les valeurs très élevées dépassant cette limite ne pourraient en pratique être obtenues qu'avec des lasers extrêmement intenses.

Cette condition étant vérifiée, on peut néanmoins se demander ce qui se passe lorsque le temps d'interaction augmente au-delà du domaine de validité du traitement perturbatif. Par un traitement analogue à celui du paragraphe C.2.d du chapitre I (méthode de Weisskopf et Wigner), on montre que la probabilité de trouver l'atome dans l'état initial $|g\rangle$ décroît exponentiellement, suivant la loi

$$P_g(t) = e^{-\Gamma t} \quad (28)$$

3 Application : Photoionisation de l'atome d'hydrogène

En utilisant les expressions connues des fonctions d'onde hydrogénoïdes, il est possible de calculer complètement les valeurs des sections efficaces de photoionisation, qui peuvent ainsi être comparées aux résultats expérimentaux. Comme nous avons fait l'approximation de l'onde plane électronique pour les états ionisés (5a), nous nous restreindrons au cas où l'énergie de l'électron éjecté est très grande, c'est-à-dire que la fréquence du rayonnement incident est très supérieure à la fréquence seuil. C'est le cas pour l'hydrogène soumis à un rayonnement X. On a alors

$$k_e \approx \sqrt{\frac{2m\omega}{\hbar}} \gg \sqrt{\frac{2m\omega_s}{\hbar}} \quad (29)$$

Pour l'état fondamental de l'atome d'hydrogène, la fréquence seuil est

$$\omega_s = \frac{E_I}{\hbar} = \frac{q^2}{4\pi\epsilon_0} \frac{1}{2a_0\hbar} = \frac{\hbar}{2ma_0^2} \quad (30)$$

(on a utilisé l'expression (17b) du rayon de Bohr a_0). La relation (29) nous montre alors que

$$k_e a_0 \gg 1 \quad (31)$$

Par ailleurs, tant que

$$\hbar\omega \ll mc^2 \quad (32)$$

c'est-à-dire que l'on reste dans le domaine où il ne peut y avoir de création de paires électron-positron ($\hbar\omega \ll 0,5MeV$), on a d'après (29) et (32)

$$\frac{k_e}{k} \approx \sqrt{\frac{2mc^2}{\hbar\omega}} \gg 1 \quad (33)$$

Dans ces conditions, en utilisant la forme asymptotique (18b) de la transformée de Fourier de la fonction d'onde de l'état fondamental, et en remplaçant $\mathbf{k}_e - \mathbf{k}$ par $\mathbf{k}_e$, l'expression (26a,26b,26c) conduit à

$$\sigma_{n=1,l=0} = 64\alpha a_0^2 \left(\frac{\omega_s}{\omega}\right)^{7/2} \int (\mathbf{u}_e \cdot \boldsymbol{\varepsilon})^2 d\Omega = \frac{256\pi}{3} \alpha a_0^2 \left(\frac{\omega_s}{\omega}\right)^{7/2} \quad (34)$$

α étant la constante de structure fine

$$\alpha = \frac{q^2}{4\pi\epsilon_0\hbar c} = \frac{1}{a_0} \frac{\hbar}{mc} = \frac{1}{137.04} \quad (35)$$

On constate que la section efficace de photoionisation est nettement plus petite que la surface πa_0^2 d'un disque de rayon égal au rayon de Bohr.

Remarques

(i) Si l'atome est initialement dans un état lié excité, on peut utiliser l'expression (18b) de la transformée de Fourier de cet état, pour obtenir

$$\sigma_{n,l=0} = \frac{1}{n^3} \sigma_{n=1,l=0} \quad (36)$$

On constate donc que lorsque le niveau de départ est plus élevé – c'est-à-dire que l'électron est moins lié – la section efficace de photoionisation décroît. Elle tend vers 0 à la limite $n \rightarrow \infty$ c'est-à-dire lorsque l'électron est de moins en moins lié.

Ce comportement paradoxal peut être interprété en remarquant que la collision inélastique

$$\text{photon} + \text{électron libre} \rightarrow \text{électron libre}$$

est impossible, car une collision entre une particule massive et une particule de masse nulle ne peut conserver à la fois l'énergie et l'impulsion relativistes. En revanche, si elle met en jeu un troisième corps susceptible d'emporter l'impulsion nécessaire (le noyau dans le cas d'un électron lié), la collision

$$\text{photon} + \text{électron} + \text{noyau} \rightarrow \text{électron} + \text{noyau}$$

est possible.

(ii) Ayant obtenu des valeurs numériques, nous pouvons revenir sur la question de la validité du traitement du § B. Nous savons que l'approximation résonnante est valable si

$$t \gg \omega^{-1} \quad (37)$$

Par ailleurs, la condition de probabilité de photoionisation petite devant 1 s'écrit

$$\frac{\Pi}{\hbar\omega} \sigma t \ll 1 \quad (38)$$

On peut majorer σ en utilisant l'expression (34)

$$\sigma \ll \pi a_0^2 \approx 10^{-20} \text{ m}^2 \quad (39)$$

et la condition (38) donne, compte tenu de $\omega \gg \omega_s$

$$t \ll \frac{\hbar\omega}{10^{-20}\Pi} \quad (40)$$

Cette relation est compatible avec (37) si

$$\Pi \ll 10^{20} \hbar\omega^2 \quad (41)$$

soit, en prenant $\omega \approx 10 \omega_s$

$$\Pi \ll 10^{20} W / \text{m}^2 \quad (42)$$

Pour la très grande majorité des sources X disponibles, cette condition est toujours remplie. Les sources de rayonnement synchrotron peuvent fournir des intensités crête de l'ordre de $10^{13} W/m^2$, à des longueurs d'onde de l'ordre de 10 nm. Les plasmas obtenus par irradiation laser intense émettent jusqu'à $10^{15} W/m^2$. Par génération d'harmonique très élevées (d'ordre

supérieur à 100) à partir d'un laser infrarouge, on obtient des faisceaux X mous (longueur d'onde de l'ordre de 10 nm) cohérents, et pouvant donc être focalisés. On obtient ainsi des intensités encore plus grandes, mais encore inférieures à la limite (42).

Les premiers lasers à rayons X ont fonctionné en 1993, dans le domaine des X mous. Lorsque ces sources deviendront utilisables, on peut s'attendre à ce que les propriétés de cohérence des faisceaux permettent d'aller au delà de la limite (42). On abordera alors un régime nouveau de l'interaction matière rayonnement, où le temps de photoionisation sera inférieur à une période de l'onde incidente, ce qui nécessitera une nouvelle approche.

Complément II.5

Détecteurs infrarouge à multipuits quantiques

Les détecteurs infrarouges sont utilisés pour réaliser des images dans les gammes de longueurs d'ondes comprises entre 3 et 20 μm . Les gammes 3-5 μm et 8-12 μm sont particulièrement intéressantes car elles correspondent à deux fenêtres de transparence atmosphérique. Les domaines d'applications en sont la vision nocturne à longue portée, la surveillance de points chauds, la vision à travers des milieux très diffusants, ... La procédure d'acquisition des images est obtenue en projetant l'image à détecter sur une matrice (ou mosaïque) de détecteurs située au point focal d'un système de vision.

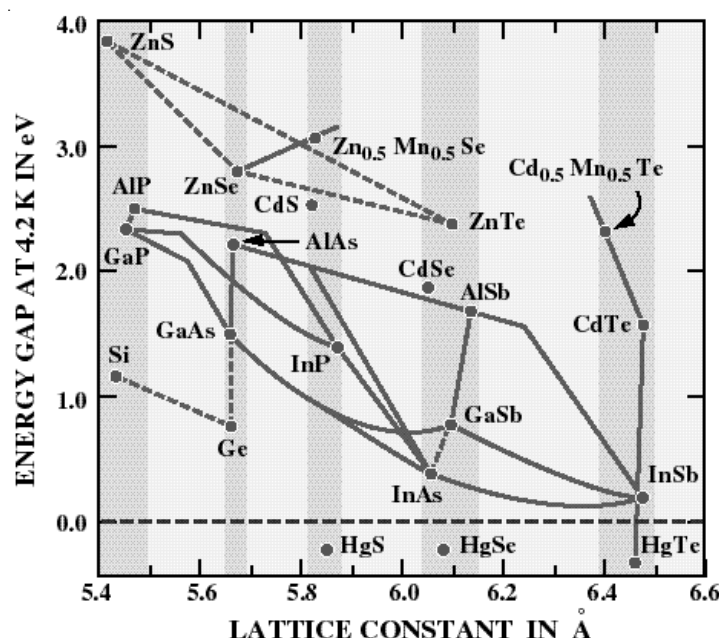


FIG. 1: Largeur de la bande interdite en fonction de la dimension de la maille cristalline dans les principaux semiconducteurs. Les colonnes définissent les familles de semiconducteurs qui présentent des accords épitactiques possibles.

L'existence des puits quantiques est liée à la variété des semiconducteurs composés III-V, ainsi désignés parce qu'ils sont obtenus à partir des éléments de la colonne III (Al, Ga, In,...) et de la colonne V (As, P, N,...) du tableau de Mendeleev. Ces composés ont pour la plupart la même structure cristalline dite de type « *blende de zinc* » qui est caractérisée par une seule grandeur : la dimension de la maille élémentaire du réseau cristallin variant entre 5,5 et 6,5 angströms. Comme il apparaît sur la figure 1, des composés différents peuvent présenter le même paramètre de maille : c'est le cas pour GaAs/AlAs, mais aussi InP/InGaAs... Il est alors possible d'effectuer, à partir d'un substrat semi-conducteur (GaAs, InP,...), la croissance de couche cristalline de composition différente mais de structure parfaitement compatible : cette opération est appelée *épitaxie*. L'accord épitactique est primordial, les propriétés semiconductrices de ces matériaux étant en effet détruites par la présence de densité infime de défauts cristallins. Les progrès des techniques de croissance épitactique, en ultra-vide notamment, permettent de contrôler les épaisseurs des couches déposées à une couche atomique près : ceci est un élément-clef de l'ingénierie quantique des matériaux.

Comme l'indique la figure 1, les composés III-V ont des largeurs de bande interdite (ou gap) différentes suivant leur composition. En intercalant des couches de faible gap (par exemple GaAs) entre des couches de gap plus grand (par exemple $\text{Al}_x\text{Ga}_{1-x}\text{As}$), on peut créer des puits de potentiel artificiels dans la matière pour les électrons (et les trous) (voir figure 2). La forme de ce puits de potentiel est caractérisée par deux paramètres : la largeur d du puits quantique et la hauteur de barrière ΔE entre les bandes de conduction des deux semiconducteurs et qui dépend plus ou moins linéairement de la concentration x en Al. Si la taille du puits de potentiel est suffisamment petite, le mouvement des électrons perpendiculairement aux interfaces est quantifié c'est à dire que l'écart entre les deux premiers niveaux quantifiés est supérieur à kT soit 25 meV à l'ambiante : on parle alors de *puits quantiques*. Cette composante se calcule en utilisant l'équation de Schrödinger :

$$-\frac{\hbar^2}{2m^*} \frac{d^2}{dz^2} \Phi(z) + (V(z) - E) \Phi(z) = 0 \quad (1)$$

où $\hbar$ est la constante de Planck, $V(z)$ est la forme du puits de potentiel et m^* est la *masse effective* de l'électron dans le composé qui décrit les interactions complexes de cet électron avec le réseau cristallin-hôte. $\Phi(z)$ est la fonction d'onde de l'électron. La résolution de l'équation d'onde (1) est très simple et montre que l'énergie ne peut prendre que des valeurs discrètes E_n pour des énergies inférieures à la discontinuité de bande ΔE . Les transitions optiques entre ces différents niveaux quantiques sont nommées *transitions inter-sousbande*.

La figure 3 montre une abaque où sont représentées les différences d'énergie $E_2 - E_1$ entre les deux premiers niveaux quantifiés pour le système GaAs/AlGaAs. On se limite aux compositions en aluminium variant de 0 à 40 % pour lesquels ΔE varie de 0 à 320 meV. Une ligne continue sépare deux parties du graphe. Au dessus de cette ligne (partie incolore), le puits quantique est suffisamment large pour accueillir deux niveaux quantiques. Les deux niveaux sont localisés c'est-à-dire que leurs fonctions d'onde (leur probabilité de présence) sont non nulles seulement près du puits. Les électrons sur ces niveaux sont piégés et ne peuvent conduire l'électricité. Une transition optique entre ces

deux niveaux est dite de type *lié-lié* et ne peut donner lieu à de la photoconduction. En dessous de la ligne continue (partie verte), le puits quantique est trop étroit pour accueillir deux niveaux quantiques : les seuls états excités accessibles à une excitation optique sont ceux du continuum d'énergie (états délocalisés). Les électrons dans ce continuum sont libres de se déplacer dans la structure. Cette transition optique est dite de type *lié-libre*. La ligne continue entre ces deux régions représente les couples largeur de puits – composition x pour lesquels le niveau E_2 est en résonance avec le bas de la bande de conduction de la barrière : il s'agit d'un état *quasi-lié*.

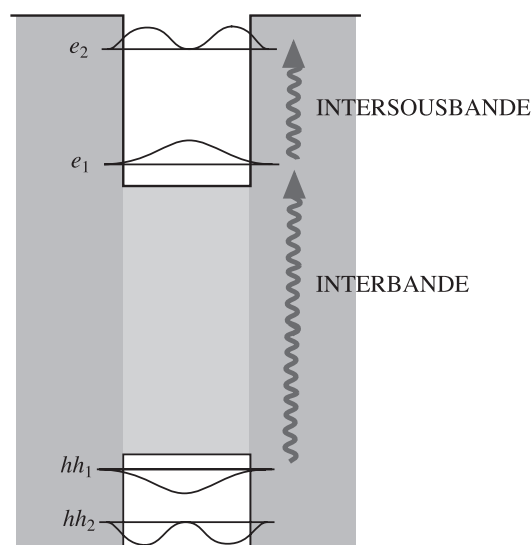


FIG. 2 : Transition interbande et transitions intersousbande dans les puits quantiques.

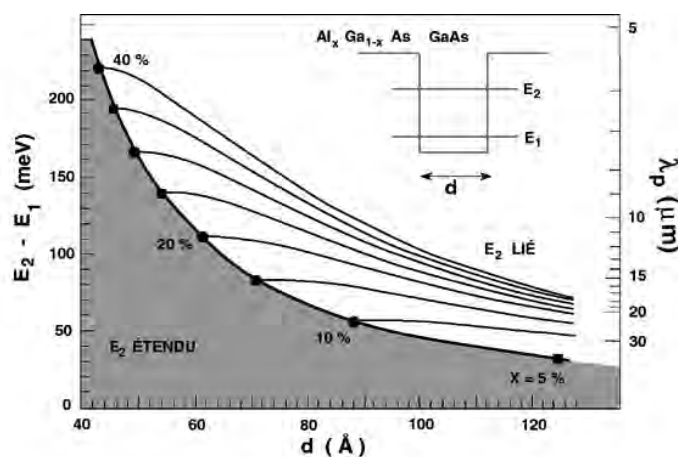


FIG. 3: Différence d'énergie entre les deux premiers niveaux quantifiés dans un puits quantique en fonction de la largeur du puits et pour différentes concentrations d'aluminium dans la barrière. L'échelle de droite indique les longueurs d'onde de coupure correspondantes. La partie claire indique les transitions lié-lié et la partie verte les transitions lié-libre.

On voit qu'en faisant varier les paramètres épaisseur du puits d et composition x dans des limites métallurgiques raisonnables, on peut synthétiser un matériau artificiel présentant des énergies et des types (*lié-lié*, *lié-libre*) de transitions optiques quasiment arbitraires entre 5 et 50 μm .

La meilleure situation pour un détecteur infrarouge est celle de la transition *lié-quasiliée* : on profite de la forte force d'oscillateur due au caractère quasiment lié du niveau excité (d'où une forte absorption) et en même temps de la délocalisation de ce niveau d'où la possibilité de photoconduction. La figure 4 représente le schéma de principe d'un détecteur à multipuits quantique (MPQ). La structure est faite d'un empilement de couches de GaAs prises en sandwich entre des couches barrière d'AlGaAs, typiquement entre 20 et 50 couches. Le couple $d-x$ est choisi de telle sorte que la transition optique soit du type *lié-quasilié* pour la longueur d'onde à détecter. Des couches de GaAs très fortement dopées assurent, de chaque côté de la structure, des contacts ohmiques qui permettront d'injecter des électrons et donc de mesurer le photocourant. Les puits quantiques sont dopés lors de la croissance : à température suffisamment basse, ces électrons issus du dopage sont piégés dans les puits quantiques et la structure est idéalement isolante. Sous éclairage, les électrons sont excités dans la bande de conduction de la barrière AlGaAs et balayés par le champ électrique appliqué à la structure. Ils vont induire un photocourant qui sera détecté grâce aux contacts ohmiques. L'énergie minimale des photons détectés, c'est-à-dire l'énergie de photoionisation $\Delta E - E_1$, est la *longueur d'onde de coupure* du photodétecteur. Nous montrons maintenant comment la théorie de la photoémission, appliquée à ce cas unidimensionnel, permet de décrire simplement l'absorption de la lumière infrarouge par les transitions *lié-libre* dans ces puits quantiques.

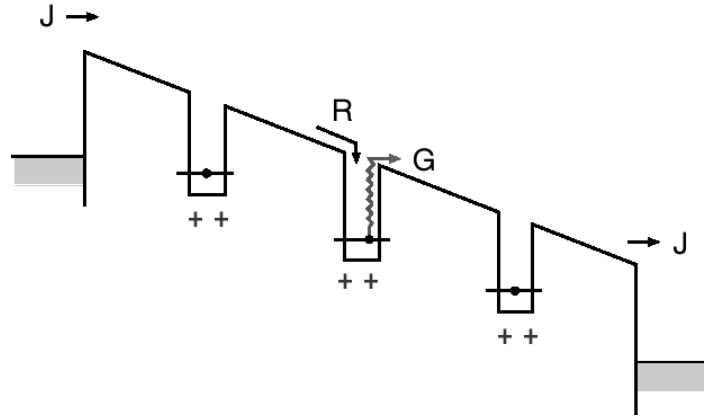


FIG. 4 : Une vision simplifiée d'un photodétecteur à multipuits quantique

On considère le puits quantique simplifié de largeur d représenté en figure 5. Ce puits quantique admet un niveau quantifié $|i\rangle$ décrit par une fonction d'onde $\phi_i(z)$ de carré sommable et d'énergie quantifiée $+E_I$, l'indice I désignant le terme d'ionisation. On fait l'hypothèse que ce puits est suffisamment profond pour que $\phi_i(z)$ soit considérée comme

la fonction d'onde du niveau fondamental du puits quantique infini, soit :

$$\phi_i(z) = \sqrt{\frac{2}{d}} \cos \frac{\pi}{d} z \quad (2)$$

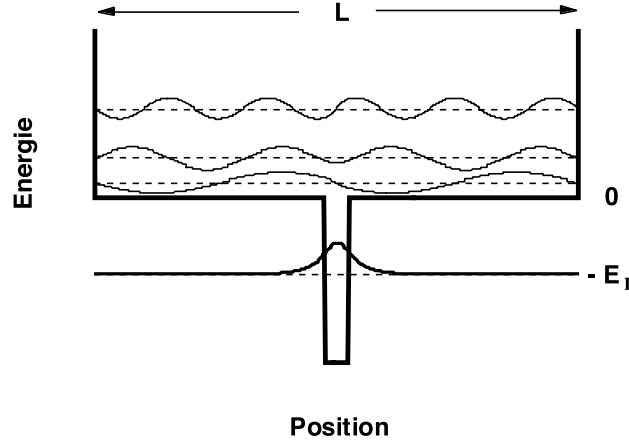


FIG. 5 : Puits quantique utilisé pour le calcul de photoémission.

Ce puits admet aussi des états délocalisés, les électrons pouvant prendre toute énergie positive dans la bande de conduction de la barrière d'AlGaAs. Nous négligerons l'influence du puits sur les électrons libres, c'est-à-dire que nous considérons que ces électrons libres sont soumis à un potentiel nul une fois dans le continuum. Néanmoins, afin de s'affranchir des problèmes de normalisation de ces fonctions d'onde, on introduit une boîte de potentiel fictive de longueur L dans laquelle ces électrons sont piégés. Les énergies et fonctions d'onde des états libres sont alors :

$$\begin{aligned} e_n &= n^2 e_0 \\ \phi_n(z) &= \sqrt{\frac{2}{L}} \sin nk_L z \quad \text{si } n \text{ pair} \\ \phi_n(z) &= \sqrt{\frac{2}{L}} \cos nk_L z \quad \text{si } n \text{ impair} \end{aligned} \quad (3)$$

e_0 est l'énergie de confinement du puits donnée par :

$$e_0 = \frac{\hbar^2}{2m^*} k_L^2 \quad (4)$$

avec le vecteur d'onde k_L :

$$k_L = \frac{\pi}{L} \quad (5)$$

Il est clair que si L est centimétrique, e_0 est de l'ordre de 10^{-11} eV et la boîte est bien fictive dans le sens où l'écart énergétique entre les niveaux de la boîte est infinitésimal en comparaison avec les énergies quantiques ou thermiques typiques (quelques dizaines de meV au moins). Les énergies données en équation (3) sont tellement rapprochées qu'il

est inutile de chercher à les distinguer. On va plutôt chercher à les compter *à la pelle*, c'est-à-dire à chercher des densités d'états à une dimension.

On se fixe un intervalle de vecteurs d'onde dk . Les états sont séparés, en vecteur d'onde, de π/L . *Sans tenir compte des problèmes de spin*, le nombre d'états dans cet intervalle est donc clairement :

$$dn = \frac{L}{\pi} dk \quad (6)$$

La densité d'état dn/dk est alors :

$$\rho(k) = \frac{dn}{dk} = \frac{L}{\pi} \quad (7)$$

Dans ce même intervalle, comme l'indique la différentiation de (3), l'énergie a varié de :

$$dk = \frac{1}{2} \frac{\sqrt{2m^*}}{h} \frac{dE}{E} \quad (8)$$

La densité énergétique d'états dn/dE est donc finalement :

$$\rho(E) = \frac{L}{2\pi} \frac{\sqrt{2m^*}}{h} \frac{1}{\sqrt{E}} \quad (9)$$

On constate que, lorsque L tend vers l'infini, $\rho(E)$ tend vers l'infini. Ceci est normal puisqu'on entasse de plus en plus d'états dans le même espace énergétique.

Nous calculons maintenant la probabilité de transition entre l'état quantifié initial $|i\rangle$ et le continuum sous l'effet d'une perturbation dipolaire sinusoïdale :

$$W(z, t) = -qFz \cos(\omega t) \quad (10)$$

D'après la règle d'or de Fermi, la probabilité de transition s'écrit :

$$G_{ic}(\hbar\omega) = \frac{\pi q^2 F^2}{2\hbar} |z_{if}|^2 \rho(\hbar\omega = E_f - E_i) \quad (11)$$

L'élément de transition z_{if} n'est non nul que pour les fonctions d'ondes finales impaires et est donné par :

$$|z_{if}|^2 = |\langle \Phi_i | z | \Phi_f \rangle|^2 = \left| \frac{2}{\sqrt{Ld}} \int_{-d/2}^{d/2} \cos\left(\frac{\pi}{d}z\right) \sin(k_f z) z dz \right|^2 \quad (12)$$

soit

$$|z_{if}|^2 = 4\sqrt{d^3} L f^2(E_f) \quad (13)$$

où $f(E)$ est la partie sans dimension de l'intégrale dans l'équation (11) soit :

$$f(E_f) = \frac{\pi}{\pi^2 - d^2 k_f^2} \sin \frac{k_f d}{2} - \frac{4k_f d}{\pi^2 - d^2 k_f^2} \cos \frac{k_f d}{2} \quad (14)$$

et k_f est le vecteur d'onde correspondant à l'énergie E_f :

$$E_f = h\omega - E_I = \frac{h^2 k_f^2}{2m^*} . \quad (15)$$

On reporte ensuite l'expression de $|z_{if}|$ dans (11). On constate ainsi que la largeur de la boîte fictive L qui apparaît au dénominateur de l'élément de transition et au numérateur de la densité d'état, se simplifie, comme annoncé. En termes plus physiques, plus la longueur du pseudo-puits quantique L est grande, plus la densité d'état final est élevée mais, corrélativement, plus la probabilité de présence de l'électron au dessus du puits quantique d'épaisseur d est dissoute dans L .

Prenant en compte le fait que la densité d'états finaux n'est que la moitié de ce que l'on a trouvé en (8) puisque les fonctions d'onde paires ne participent pas aux transitions, on trouve alors la probabilité de transition de l'état initial dans le continuum :

$$G_{if}(h\omega) = q^2 F^2 d^3 \frac{\sqrt{m^*}}{h^2} \frac{f^2(h\omega - E_I)}{\sqrt{2(h\omega - E_I)}} \quad (16)$$

Le comportement du système est donc bien indépendant de la taille de la boîte fictive d'où le nom de *pseudo-quantification* : il ne s'agit que d'un intermédiaire de calcul, mais néanmoins très puissant. La figure 6 montre la variation du taux de transition en fonction de la fréquence de l'excitation ω .

On constate ainsi l'existence d'un seuil d'ionisation pour la probabilité de transition : l'énergie de coupure des photons détectés est bien l'énergie d'ionisation E_I . De plus, l'absorption proche du seuil, c'est-à-dire pour $h\omega \approx E_I$, est donnée par :

$$G_{if} \propto \sqrt{h\omega - E_I} \quad \text{quand } h\omega \approx E_I$$

Un deuxième caractère est le caractère quasi-résonnant de la transition proche du seuil de photo-ionisation, quasi-résonance due à la diminution conjointe de la densité d'état (en k_f^{-1}) et du moment dipolaire z_{if} en (en k_f^{-2}) ce qui conduit à :

$$G_{if} \propto \frac{1}{(h\omega - E_I)^{5/2}} \quad \text{quand } h\omega \gg E_I$$

Toutes ces expressions rendent raisonnablement bien compte du comportement spectral des détecteurs à puits quantiques.

Références :

1- E. Rosencher et B. Vinter, Optoelectronique, Masson (1998). 2- B. F. Levine, J. Appl. Phys. **74**, R1-R81 (1993).

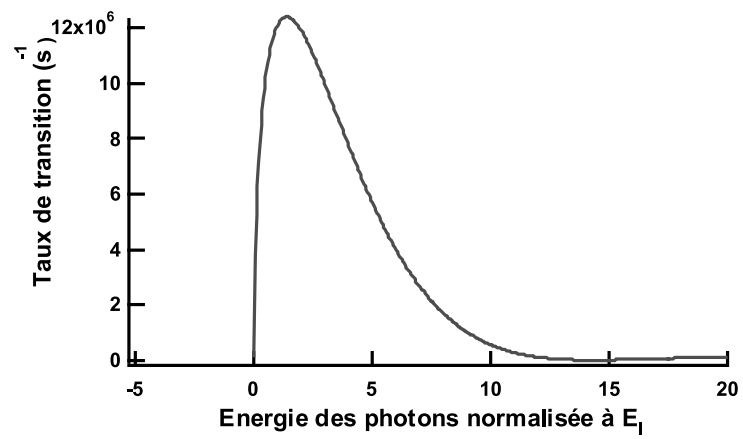


FIG. 6: Taux d'ionisation par seconde d'un puits quantique en fonction de l'énergie normalisée à l'énergie d'ionisation E_I .

Chapitre 3

Les lasers

Nous exposons dans ce chapitre une vue d'ensemble de la physique des lasers, permettant de comprendre leurs principales caractéristiques et nous donnons un bref aperçu sur les propriétés de la lumière émise par ces sources. Notre but n'est pas de fournir un catalogue exhaustif des lasers existant aujourd'hui. Nous nous servirons plutôt d'exemples concrets choisis parmi les lasers les plus courants actuellement pour présenter les principes généraux à la base du fonctionnement des lasers, et montrer comment ils ont été mis en pratique. Notre expérience est que ces principes nous ont toujours permis de comprendre le fonctionnement des nouveaux lasers que nous avons vu apparaître au fil des années.

Les idées physiques permettant de comprendre le fonctionnement des lasers apparaîtront souvent extrêmement simples. Cette impression tient au fait que les concepts essentiels sont désormais bien dégagés et que les détails accessoires et les errements des raisonnements incorrects ou incomplets, sont passés sous silence¹. Il est cependant intéressant de constater combien l'émergence de ces idées a été laborieuse. On considère généralement que la préhistoire du laser commence en 1917 lorsque Albert Einstein introduit la notion d'**émission induite** (encore appelée **émission stimulée**). En effet, en étudiant l'équilibre thermodynamique entre le rayonnement électromagnétique et des atomes à la température T , Einstein montre que l'on ne peut retrouver la *formule de Planck* pour le rayonnement du *corps noir* qu'à la condition d'introduire, en plus de l'*absorption* et de l'*émission spontanée*², un nouveau processus l'*émission induite*. Après les articles d'Einstein sur ce sujet, de nombreux physiciens s'intéressèrent à ce problème tant du point de vue théorique que du point de vue expérimental. Cet intérêt fut concrétisé par la mise en évidence expérimentale de l'émission induite dans une décharge de

¹Pour en savoir plus sur la chronologie des idées, on pourra consulter M. Bertolotti, *Masers and Lasers, An Historical Approach*, (Hilger, 1983), C.H. Townes dans les *Centennial Papers*, I.E.E.E. Journal of Quantum Electronics, **20**, pp. 545–615 (1984) ou J. Hecht, *Laser Pioneers*, (Academic Press, 1992). On lira aussi avec profit les conférences Nobel (1964) de C.H. Townes, N.G. Basov et A.M. Prokhorov, in « *Nobel Lectures, Physics 1963-1970* » (Elsevier) ; disponibles sur www.nobel.se/physics.

²Ces processus s'imposaient *empiriquement* sans toutefois s'inscrire dans un cadre théorique global bien défini.

néon par Ladenburg et Kopferman³ en 1928. Assez curieusement, l'intérêt pour l'émission induite décrut ensuite. En fait, les physiciens ne voyaient aucune application possible à ce phénomène parce qu'ils ne croyaient pas qu'un système puisse s'éloigner suffisamment de l'*équilibre thermodynamique* pour qu'un gain notable puisse être obtenu. En outre, peu d'entre eux étaient familiers avec l'électronique et en particulier avec les mécanismes de **réaction** conduisant à l'oscillation d'un système qui présente du gain. En d'autres termes, on n'imaginait des applications à l'émission induite que dans des conditions de très grande amplification, c'est-à-dire lorsque le gain de l'amplificateur est très grand devant 1, conditions qui apparaissaient impossible à remplir.

Il fallut attendre 1954 et la réalisation du *maser* à ammoniac par Townes et ses collaborateurs Gordon et Zeiger pour que la situation change. L'idée novatrice de Townes est que si on place le *milieu amplificateur* dans une **cavité résonnante**, une oscillation peut se produire même pour un gain faible, le point fondamental étant que le gain par émission induite soit suffisant pour compenser les pertes (très petites) de la cavité⁴.

Le passage du *maser* (« microwave amplification by stimulated emission of radiation ») au *laser* (« light amplification by stimulated emission of radiation »), c'est-à-dire le passage du domaine microonde au domaine visible, n'a pas été immédiat et a également donné lieu à des controverses. Ce n'est qu'en 1958 que Schawlow et Townes publièrent l'idée fondamentale que dans le domaine optique un système de *deux miroirs se faisant face* pouvait jouer le rôle de la *cavité résonnante*. Leur article déclencha un intense effort expérimental qui fut concrétisé dans la réalisation du premier laser par Maiman en 1960. Assez curieusement, ce laser utilisait le *rubis* comme milieu amplificateur, alors que certains avaient cru démontrer qu'un tel matériau ne pourrait pas conduire à une oscillation laser. Nous verrons en effet que le mécanisme de l'amplification dans un système « à trois niveaux » comme le rubis est relativement particulier et peu intuitif. Le laser à rubis fut suivi de près par le laser à mélange gazeux *hélium-néon* mis au point par Javan, puis par le laser à *dioxyde de carbone* réalisé par Patel. Parmi les progrès marquants qui ont suivi, il faut mentionner la mise au point des lasers de grande puissance comme le laser à *néodyme*, le développement des lasers *accordables* à colorant ou à amplificateur solide (saphir dopé au titane). La révolution des lasers à *semi-conducteur*, de caractéristiques remarquables au niveau prix de revient, encombrement, et rendement global, a permis d'introduire les lasers dans de nombreux systèmes, y compris dans des produits de grande diffusion (lecteurs de disques, de codes-barres, imprimantes, etc ...).

Nous avons vu dans le chapitre précédent qu'il est possible d'obtenir des *amplificateurs de lumière* en utilisant des milieux dans lesquels une **inversion de population** a été réalisée. Nous montrons dans la partie A comment transformer un amplificateur de lumière en *générateur de lumière* en plaçant le milieu amplificateur dans une cavité servant à réinjecter le faisceau dans cet amplificateur et nous étudions l'influence de cette cavité sur les caractéristiques de la lumière émise. Dans la partie B, nous présentons quelques méca-

³En fait, ces auteurs n'ont pas observé l'amplification du champ incident, mais la modification de l'indice en fonction de la différence de population entre les niveaux extrêmes de la transition, ce qui est une preuve indirecte de l'existence de l'émission induite (cf. Chapitre II, § 2.5.2, remarque).

⁴Le maser fut découvert indépendamment par les physiciens russes Basov et Prokhorov.

nismes généraux pour obtenir une inversion de population dans un milieu et obtenir ainsi une amplification de la lumière. Ces mécanismes sont établis sur des systèmes modèles (*systèmes à quatre niveaux*, *systèmes à trois niveaux*) et illustrés par quelques exemples portant sur des lasers ayant (ou ayant eu) une grande importance. Parmi les lasers ainsi brièvement décrits, citons les lasers à néodyme, le laser hélium-néon, le laser à colorant, le laser à semiconducteur, le laser à erbium et enfin pour son importance historique le laser à rubis. Dans la partie C, nous étudions les *propriétés spectrales* de la lumière émise par un laser. Nous montrons que selon les caractéristiques de la cavité et de la largeur de la courbe de gain, la lumière émise par un laser peut avoir une fréquence bien définie (fonctionnement *monomode*) ou être la superposition de plusieurs ondes ayant des fréquences différentes (fonctionnement *multimode*). Dans ce dernier cas, nous expliquons comment sélectionner une seule fréquence (et obtenir ainsi un fonctionnement monomode) et nous donnons quelques indications sur la *largeur en fréquence* ou pureté spectrale de la lumière ainsi émise. Dans la partie D, nous présentons quelques propriétés des *lasers en impulsion*. Ces derniers n'émettent pas la lumière de manière continue mais par impulsions de faible durée : l'échelle sera la *nanoseconde* (10^9 s), la *picoseconde* (10^{-12} s), voire la *femtoseconde* (10^{-15} s) et même moins. Ces impulsions ont des *puissances instantanées* très élevées : l'unité sera ici le *terawatt* (10^{12} W), voire le *petawatt* (10^{15} W). Dans la conclusion, nous insistons sur les propriétés importantes de la lumière laser et nous indiquons en quoi elle diffère fondamentalement de la lumière émise par une lampe classique incohérente. Ceci permet de mieux comprendre les applications des lasers, et en particulier celles qui sont présentées dans les compléments de ce chapitre.

Le complément III.1 rappelle les propriétés essentielles des *cavités de Fabry-Perot* qui ont été d'abord un outil important en spectroscopie à haute résolution avant d'être utilisées pour assurer la réinjection du champ dans le milieu amplificateur d'un laser.

Les développements mathématiques dans le corps de ce chapitre et du complément III.1 sont, pour l'essentiel, faits en supposant que les champs électromagnétiques peuvent être décrits par des ondes planes. En fait, les faisceaux lasers ont une *distribution transversale non uniforme d'intensité*. Nous introduisons ce sujet dans le complément III.2 et nous insistons plus particulièrement sur le mode transverse fondamental qui correspond à une distribution *gaussienne* d'intensité. Ce complément est particulièrement utile pour comprendre les propriétés de *cohérence spatiale* de la lumière émise par un laser.

Les compléments III.3 à III.5 sont consacrés à diverses *applications des lasers*. Les deux premiers décrivent des applications dans les domaines *industriel*, *médical* et *militaire*. Le complément III.5 décrit les avancées spectaculaires que le laser a permis en *spectroscopie*.

Nous revenons dans le complément III.6 sur le problème de la *largeur de raie* d'un laser et nous montrons qu'il existe une *limite inférieure fondamentale*, appelée largeur de Schawlow-Townes, liée à la diffusion aléatoire de la phase lors du processus d'émission spontanée.

Enfin, le complément III.7 donne quelques éléments sur la photométrie des sources incohérentes, et sur la notion de mode élémentaire du champ électromagnétique. Ces notions sont essentielles pour comprendre en profondeur la différence entre faisceau de lumière incohérente et faisceau laser.

3.1 Conditions d'oscillation

3.1.1 Seuil d'oscillation

Le laser, comme les oscillateurs développés en électronique, repose sur l'application d'une *boucle de réaction* à un système amplificateur. Pour le laser, le système amplificateur de lumière est un milieu dans lequel on a réalisé une inversion de population (voir Chapitre II, partie E). Le coefficient d'amplification en intensité G pour un faisceau lumineux traversant un milieu amplificateur de longueur L_A , vaut

$$G = \frac{I_{\text{sortie}}}{I_{\text{entrée}}} = \exp \left\{ \int_{z_\ell}^{z_\ell + L_A} g(z) dz \right\} \quad (3.1)$$

où $g(z)$ est le gain par unité de longueur dans la tranche d'abscisses z (cf. équation 2.2.177). Si l'intensité lumineuse circulant dans l'amplificateur laser est petite devant l'intensité de saturation I_{sat} (cf. § 2.2.5.7), le gain par unité de longueur est maximal en tout point. Il en est évidemment de même du *coefficient d'amplification*, qui prend alors sa *valeur non saturée* $G^{(0)}$. Pour un milieu amplificateur homogène, et un gain $G^{(0)} - 1$ petit devant 1, le gain non saturé s'écrit de façon approchée

$$G^{(0)} - 1 \simeq g^{(0)} L_A \quad (3.2)$$

où $g^{(0)}$ est le gain non saturé par unité de longueur.

L'amplificateur laser est placé dans une boucle de réaction généralement assurée par un ensemble de miroirs réalisant une *cavité laser*. La cavité laser représentée sur la figure 3.1 est constituée par trois miroirs permettant à un pinceau de rayons parallèles de suivre une *trajectoire fermée* par réflexion sur ces miroirs. Nous supposons que deux de ces miroirs sont totalement réfléchissants et que le troisième (appelé miroir de sortie) a des coefficients de réflexion et de transmission en intensité respectivement égaux à R et T (avec $R + T = 1$ pour un miroir sans pertes).

Considérons un pinceau lumineux circulant dans la cavité en anneau à partir du point A situé juste après le miroir M_S (Fig. 3.2). Si l'intensité du champ est I_A au point A, elle est égale à GI_A après un tour et juste avant la réflexion sur M_S . Après réflexion sur M_S , l'intensité du faisceau est égale à RGI_A .

Pour qu'une oscillation s'établisse dans la cavité, il faut que le champ après un tour soit plus grand que le champ initial. Pour une onde de faible intensité la *condition de démarrage de l'oscillation* est :

$$RG^{(0)} > 1 \quad (3.3)$$

soit encore

$$G^{(0)}(1 - T) > 1 \quad (3.4)$$

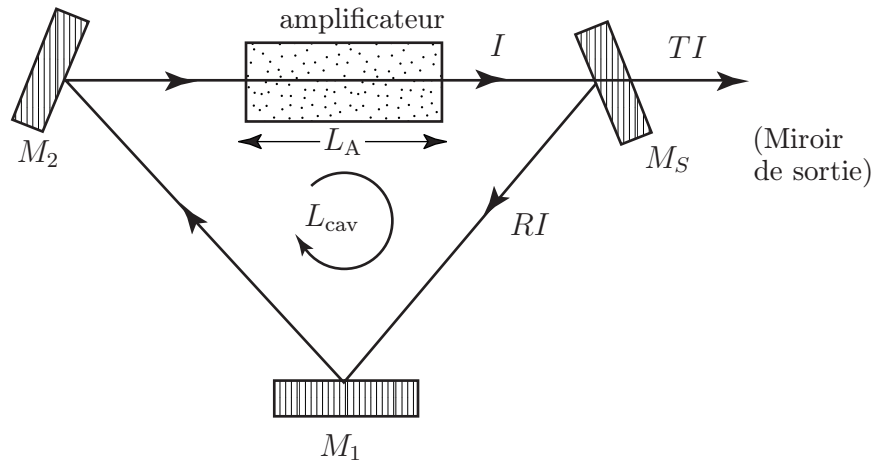


FIG. 3.1: **Laser en anneau.** Le miroir de sortie M_S semi-transparent renvoie une partie de la lumière vers l'entrée de l'amplificateur, via les miroirs totalement réfléchissants M_1 et M_2 . Si le gain de l'amplificateur est suffisant pour compenser les pertes, un faisceau lumineux s'établit dans la cavité. On désigne par I l'intensité dans la cavité avant le miroir de sortie (coefficients de réflexion et de transmission R et T). Le laser émet vers l'extérieur un faisceau d'intensité TI . La longueur optique pour un tour de cavité est L_{cav} . La longueur du milieu amplificateur est L_A .

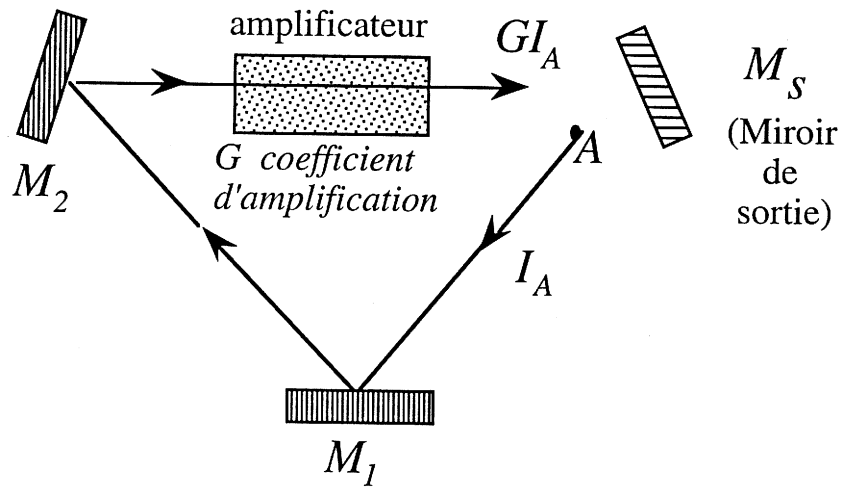


FIG. 3.2: **Valeurs successives de l'intensité au cours de la propagation, dans une cavité en anneau.**

En pratique, il faut ajouter aux pertes par transmission (qui sont les pertes utiles puisqu'elles correspondent au faisceau disponible par l'utilisateur) des pertes supplémentaires dues à l'absorption ou à la diffusion à l'intérieur de la cavité et sur les miroirs. Nous représenterons globalement ces pertes supplémentaires par un coefficient d'absorption A . La condition (3.4) est alors remplacée par

$$G_0(1 - T)(1 - A) > 1$$

ou encore :

$$G^{(0)} > \frac{1}{(1-T)(1-A)} \quad (3.5)$$

Cette équation constitue la condition de démarrage au seuil. Elle s'écrit encore dans la limite $A, T \ll 1$ (limite rencontrée fréquemment en pratique) :

$$\Rightarrow G^{(0)} - 1 > T + A \quad (3.6)$$

Cette inégalité signifie que le *gain non saturé de l'amplificateur* ($G^{(0)} - 1$) *doit être supérieur aux pertes totale* pour que l'oscillation laser démarre. Il est important de noter que dans un tel système, *il n'y a pas de champ incident sur la cavité*. L'oscillation s'amorce à partir du bruit électromagnétique, celui-ci étant généralement dû à l'émission spontanée dans l'amplificateur.

3.1.2 Régime stationnaire. Intensité et fréquence de l'onde laser

a. Stationnarité du champ

La puissance lumineuse circulant dans la cavité ne peut pas augmenter indéfiniment à chaque tour de cavité. En effet, le gain est presque toujours une fonction décroissante de l'intensité I de l'onde lumineuse circulant dans la cavité à cause des effets de saturation se produisant à haute intensité (voir le paragraphe 2.2.5.7). Il s'ensuit que G est plus petit que $G^{(0)}$.

$$G(I) < G^{(0)} \quad (3.7)$$

Un régime stationnaire stable est atteint lorsque le champ se reproduit identiquement à lui-même après chaque tour de cavité. Notons

$$E_A(t) = E_0 \cos(\omega t + \varphi) \quad (3.8)$$

l'amplitude du champ laser au point A de la figure 3.2. Après *propagation sur un tour*, ce champ devient :

$$E'_A(t) = \sqrt{R(1-A)G(I)} E_0 \cos(\omega t + \varphi - \psi) \quad (3.9a)$$

avec $I = E_0^2/2$. Le facteur précédant E_0 dans (3.9a) rend compte des phénomènes de réflexion, atténuation et amplification subis par l'onde lors de sa propagation sur un tour. Le déphasage accumulé par le champ sur un tour est noté ψ . Il est égal à

$$\psi = 2\pi \frac{L_{\text{cav}}}{\lambda} \quad (3.9b)$$

pour un champ dont la longueur d'onde dans le vide est $\lambda = 2\pi c/\omega$ et pour une cavité de *longueur optique* égale à L_{cav} . En régime *stationnaire*, le champ doit se reproduire identiquement à lui-même au bout d'un tour. On a donc pour tout t :

$$\Rightarrow E'_A(t) = E_A(t) \quad (3.10)$$

Cette condition implique l'identité de l'*amplitude* et de la *phase* des champs $E_A(t)$ et $E'_A(t)$. Nous examinons successivement les conséquences de ces deux égalités dans les paragraphes suivants.

Remarques

(i) La longueur optique L_{cav} de la cavité, qui intervient dans la formule (3.9b), diffère de sa longueur géométrique pour plusieurs raisons. D'abord, les indices de réfraction⁵ du milieu amplificateur et des éléments optiques éventuellement présents dans la cavité sont généralement différents de 1. D'autre part, on inclut dans cette longueur optique les déphasages causés par la réflexion du faisceau sur les miroirs de la cavité ainsi que les variations de phase autour des points de focalisation dans la cavité⁶.

(ii) Nous avons implicitement supposé que les champs dans la cavité sont des ondes planes de sorte que leurs amplitudes et leurs phases ne varient pas dans le plan perpendiculaire à la direction de propagation. Dans la pratique, les champs ont toujours une extension spatiale finie et il faudrait tenir compte de la distribution *transverse* du champ. Ce problème est étudié dans le Complément III.2. Notons cependant que ces effets peuvent être décrits par la même approche que celle conduisant à l'équation (3.10). Supposons que le champ au point A soit $E_A(x, y, t)$ où x et y sont les coordonnées dans le plan perpendiculaire à la propagation. Après un tour dans la cavité, le champ est égal à $E'_A(x, y, t)$. La condition de stationnarité impose que l'on ait :

$$E'_A(x, y, t) = E_A(x, y, t) \quad \text{pour tout } x \text{ et } y.$$

Puisqu'une onde limitée spatialement tend à s'étaler à cause des effets de *diffraction*, la condition précédente ne peut être satisfaite que si des éléments focalisants (lentilles, miroirs sphériques) sont placés dans la cavité pour compenser les effets de la diffraction.

La diffraction est aussi à l'origine de la *cohérence spatiale* de l'onde laser. En effet, d'après le principe de Huyghens-Fresnel⁷, le champ en un point (x', y') résulte de l'interférence entre les champs provenant de tous les points (x, y) au tour précédent. Ceci fixe la *phase relative* du champ entre deux points du plan perpendiculaire à la propagation et est à l'origine de la cohérence spatiale de l'onde laser. Cette propriété de cohérence différencie la lumière émise par un laser de celle issue de sources classiques, qui sont des sources « étendues », pour lesquelles la phase relative du champ entre deux points distants du plan transverse est une variable aléatoire.

De même que la distribution transverse de la phase, la distribution transverse d'amplitude est également fixée (voir le cas typique des modes gaussiens dans le complément III.2).

b. Intensité stationnaire du laser

D'après (3.9a) et (3.10), la condition sur l'intensité s'écrit :

$$R(1 - A)G(I) = 1 \tag{3.11}$$

⁵Nous supposons pour simplifier que les indices et, par conséquent, la longueur optique sont indépendants de l'intensité circulant dans la cavité (pour les effets associés à une éventuelle dépendance de l'indice avec l'intensité, voir le complément III.5).

⁶Ce dernier effet est généralement connu sous le nom de « déphasage de π au foyer ». La phase de Gouy présentée dans le complément III.2 permet de trouver ce résultat.

⁷Voir plus précisément la théorie de Kirchhoff de la diffraction (M. Born and E. Wolf « Principles of Optics » §§ 8.2 et 8.3, J.P. Jackson « Classical Electrodynamics » § 9.8, Wiley).

ou bien

$$\Rightarrow G(I) = \frac{1}{(1-T)(1-A)} \quad (3.12)$$

L'intensité du champ circulant dans la cavité s'ajuste à une valeur telle que **le gain de l'amplificateur compense exactement les pertes de la cavité**. Notons que l'équation (3.12) *détermine l'intensité du champ circulant dans la cavité* (puisque G est une fonction de l'intensité).

Remarques

(i) Le gain de l'amplificateur résulte de la différence entre l'émission stimulée à partir du niveau supérieur et de l'absorption à partir du niveau inférieur de la transition amplifiante. Dans les formules (3.5), (3.6) ou (3.12) (tout comme dans les formules donnant le coefficient d'amplification dans le chapitre II et le complément II.2), les pertes dues à l'absorption dans le milieu amplificateur sont incluses dans la valeur du gain ; plus précisément elles sont responsables du terme de saturation inférieur à 1.

(ii) Dans la cavité en anneau de la figure 3.1 on a considéré que la lumière tourne dans la cavité selon une certaine direction ($M_S M_1 M_2$) dans le cas de la figure 3.1). La propagation dans l'autre sens de rotation est également possible. Pour éviter une compétition entre les deux sens de rotation, on ajoute généralement dans la cavité un *élément non réciproque* introduisant des *perles différentes selon le sens de rotation*. Le seuil est ainsi plus bas pour un sens de rotation qui se trouve favorisé.

(iii) On peut concevoir bien sûr des cavités en anneau ayant plus de trois miroirs. Il n'est pas rare de trouver des cavités à quatre ou cinq miroirs.

(iv) Beaucoup de lasers utilisent des cavités linéaires (Fig. A.3) plutôt que des cavités en anneau. Si l'on considère une cavité composée d'un miroir totalement réfléchissant M_1 et d'un miroir partiellement transparent M_S , il y aura deux traversées du milieu amplificateur (une fois de M_S vers M_1 , une autre fois de M_1 vers M_S) pour une réflexion sur le miroir M_S de sorte que la condition (3.12) doit a priori être remplacée par :

$$G^2(I) = \frac{1}{(1-T)(1-A)^2} . \quad (3.13)$$

En fait la situation est généralement plus subtile. Ainsi, si la polarisation est identique à l'aller et au retour, on a une onde stationnaire, avec une intensité modulée spatialement. La modulation spatiale du gain qui en découle (du fait de la saturation) est responsable de comportements subtils des lasers à cavité linéaire.

Quant à la longueur optique pour un tour de cavité, elle est de l'ordre de deux fois la distance L_0 entre miroirs (la valeur exacte dépend de l'indice de réfraction du milieu amplificateur et de déphasage sur les miroirs).

c. Condition sur la phase

L'équation pour la phase résultant de la condition de stationnarité (3.10) s'écrit :

$$\psi = 2p\pi \quad (3.14)$$

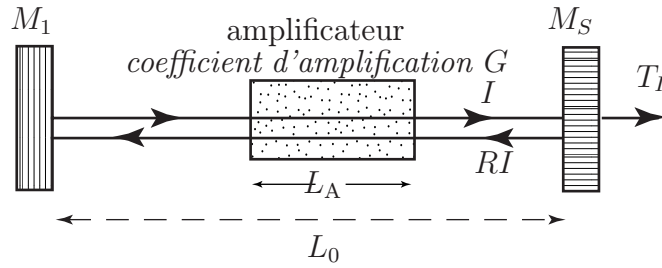


FIG. 3.3: **Laser à cavité linéaire.** La longueur optique pour un tour L_{cav} est de l'ordre de $2L_0$.

où p est un entier. Cette condition s'exprime encore, d'après (3.9b), sous la forme :

$$L_{\text{cav}} = p\lambda \quad (3.15)$$

ou puisque $\lambda = 2\pi c/\omega$

$$\Rightarrow \quad \frac{\omega}{2\pi} = p \frac{c}{L_{\text{cav}}} \quad (3.16)$$

La condition (3.15) exprime que *la longueur optique de la cavité doit être un multiple entier de la longueur d'onde*, ce multiple étant en général très grand (de l'ordre de 10^6 pour des cavités typiquement de l'ordre du mètre et des longueurs d'onde voisines du micron)⁸.

Remarque

La condition de stationnarité (3.10) n'impose aucune contrainte sur la phase φ du champ laser, qui peut donc prendre n'importe quelle valeur. Le complément III.6 montre que cette propriété permet le phénomène de *diffusion de la phase* lui-même responsable d'une largeur spectrale finie, souvent faible, de l'émission laser.

Il importe de bien noter la différence entre la phase absolue du champ qui n'est pas déterminée par la condition de stationnarité, et la fréquence qui est la *dérivée de la phase* et qui est soumise à la contrainte (3.16).

d. Fréquence d'oscillation

Le coefficient d'amplification non saturé $G^{(0)} = \exp\{g^{(0)}L_A\}$ est une fonction de la fréquence⁹ que l'on peut souvent représenter par une courbe en cloche présentant un maximum pour $\omega = \omega_M$. (Il varie généralement comme la section efficace laser : voir par exemple la loi Lorentzienne 2.2.167). Si l'on suppose que $G^{(0)}$ évalué en $\omega = \omega_M$ est plus grand que $1 + T + A$, il existe une bande de fréquence limitée $[\omega', \omega'']$ (voir figure 3.4) pour laquelle la condition de démarrage (3.5) est vérifiée.

⁸Les lasers à semiconducteurs (voir § 3.2.3) font exception, à cause de la très petite taille de leurs cavités.

⁹Pour éviter d'alourdir le texte et pour nous conformer à un usage qui n'est pas sans danger, nous désignons parfois ω sous le terme fréquence, alors qu'il s'agit d'une pulsation ayant pour unité le rad s^{-1} ; la fréquence, au sens propre, est égale à $\omega/2\pi$ et a pour unité le Hz.

En fait, le laser n'émet pas de la lumière dans toute la plage $[\omega', \omega'']$. Parmi toutes les fréquences comprises dans l'intervalle $[\omega', \omega'']$, seules celles dont les valeurs sont données par l'équation (3.16) sont susceptibles d'être émises par le laser (voir figure 3.5). Ces fréquences ω_p *discrètes* sont repérées par l'indice p et correspondent à ce qu'on appelle les **modes longitudinaux d'oscillation**. L'écart en fréquence entre deux modes consécutifs est égal à c/L_{cav} . (Il vaut typiquement 5×10^8 Hz pour une cavité de longueur $L_{\text{cav}} = 0,6$ m ; cet intervalle est donc très petit devant la fréquence de la lumière qui est, pour le visible, de l'ordre de 5×10^{14} Hz).

Le nombre de modes longitudinaux pour lesquels la condition d'oscillation (3.5) est vérifiée peut varier entre 1 et 10^5 selon la nature du milieu amplificateur. Nous reviendrons sur ce point dans la partie C et nous discuterons également le comportement relatif des divers modes : oscillent-ils simultanément ou y a-t-il compétition entre eux ?

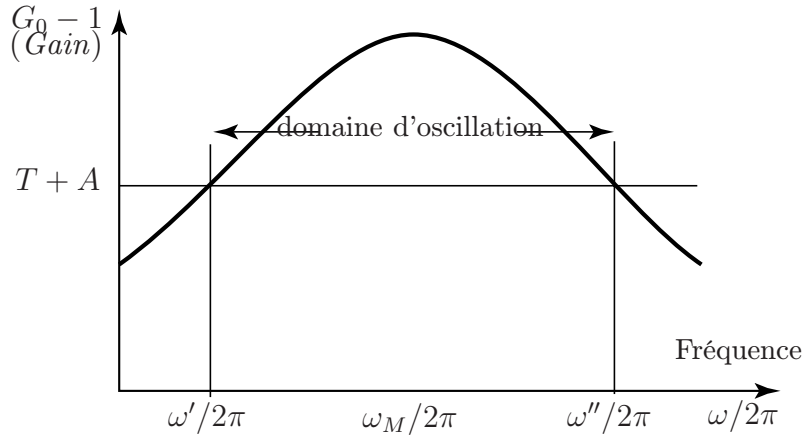


FIG. 3.4: **Courbe d'amplification en fonction de la fréquence.** Le laser est susceptible d'osciller sur la plage $[\omega', \omega'']$ pour laquelle le gain non saturé est plus grand que les pertes.

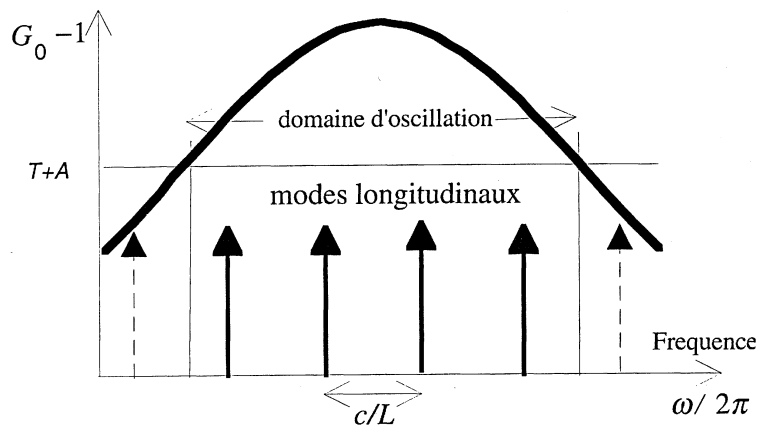


FIG. 3.5: **Fréquences d'oscillation du laser.** Le laser peut osciller sur les modes longitudinaux de la cavité tombant dans le domaine de fréquence $[\omega', \omega'']$ où le gain non saturé est supérieur aux pertes.

Remarques

(i) La *longueur optique* de la cavité L_{cav} dépend des *indices* du milieu amplificateur et des éléments optiques. Si ces indices sont pratiquement constants dans la gamme $[\omega', \omega'']$, L_{cav} est le même pour toute fréquence. La distance entre deux modes longitudinaux d'oscillation donnée par (3.16) est alors indépendante de la longueur d'onde.

(ii) L'intervalle entre modes est c/L_{cav} où L_{cav} est la longueur parcourue par la lumière sur un tour de cavité. Dans le cas d'une cavité linéaire de longueur L_0 (figure 3.3), le trajet est bouclé quand la lumière a parcouru un aller-retour de sorte que $L_{\text{cav}} = 2L_0$ si l'on peut identifier longueur géométrique et longueur optique. Dans ce cas, l'intervalle entre modes d'une cavité linéaire est donc $c/2L_0$.

3.2 Description des milieux amplificateurs

3.2.1 Nécessité d'une inversion de population

Nous avons présenté au chapitre II (Partie E) un premier exemple de milieu susceptible de réaliser l'amplification d'une onde lumineuse : il s'agit d'un ensemble d'atomes dont on a considéré deux niveaux particuliers b et a d'énergies E_b et E_a et tels que l'on excite préférentiellement le niveau d'énergie le plus élevé (b). Dans ces conditions, l'émission induite (ou émission stimulée) est le processus dominant (comparé à l'absorption) ce qui permet l'amplification. Nous nous proposons ici de généraliser ces résultats à des systèmes plus compliqués, mais plus réalistes.

Pour ces systèmes, la plupart des résultats obtenus dans le chapitre II restent valables et notamment les comportements symétriques de l'absorption et de l'émission induite. Si on considère un milieu d'épaisseur dz , avec une densité¹⁰ (N_a/V_A) d'atomes ou de molécules dans le niveau inférieur, et (N_b/V_A) dans le niveau excité, une onde de fréquence ω se propageant vers les z positifs voit son intensité $I(z)$ modifiée suivant la loi

$$\frac{1}{I(z)} \frac{dI(z)}{dz} = -2k'' \quad (3.17)$$

l'absorption (ou le gain) par unité de longueur $2k''$ étant proportionnelle à la différence des populations

$$2k'' = \frac{N_a - N_b}{V} \sigma_L(\omega) \quad (3.18)$$

Cette équation résulte du bilan des processus d'absorption et d'émission induite, l'absorption étant proportionnelle à N_a et l'émission induite à N_b . La section efficace d'absorption (et d'émission stimulée) $\sigma_L(\omega)$ est un nombre positif (ayant la dimension d'une surface) qui a généralement un caractère résonnant, c'est-à-dire qu'elle passe par un maximum pour une certaine fréquence ω_M .

¹⁰On pourra considérer V_A comme le volume commun au milieu amplificateur et à la cavité.

Pour un milieu homogène et non saturé, k'' est indépendant de l'intensité. Celle-ci varie donc suivant la loi

$$I(z) = I(0) \exp(-2k''z) \quad (3.19)$$

Si k'' est négatif, l'onde est amplifiée : l'émission stimulée l'emporte sur l'absorption. C'est évidemment ce cas qui est intéressant pour les lasers. Il se produit lorsqu'il y a *inversion de population*, c'est-à-dire pour

$$N_b > N_a \quad (3.20)$$

Une telle situation est opposée à la situation d'équilibre thermodynamique, caractérisée par la loi de Boltzmann

$$\left(\frac{N_b}{N_a}\right)_{\text{éq.th.}} = \exp\left(-\frac{E_b - E_a}{k_B T}\right) \quad (3.21)$$

La réalisation d'un état hors d'équilibre ne peut être obtenue qu'en jouant sur la **cinétique**. Le cas idéal est celui où l'on alimente énergiquement le niveau supérieur b et où le niveau inférieur a se vide rapidement. On se heurte bien sûr aux *processus de relaxation* qui vident le niveau b et qui remplissent le niveau a . Les découvertes de méthodes permettant de surmonter ces difficultés et d'obtenir *l'inversion de population* ont jalonné les progrès de lasers. Nous présentons dans la suite de ce paragraphe quelques systèmes caractéristiques.

Remarques

(i) Parmi les processus de relaxation néfastes, l'émission spontanée de b vers a joue un rôle à part par son caractère inévitable. Or on sait depuis Einstein que l'émission spontanée est d'autant plus intense (comparée à l'émission stimulée) que la fréquence est plus élevée. Cela explique qu'il soit relativement facile d'obtenir l'effet laser dans l'infrarouge, très difficile dans l'ultraviolet et que le laser à rayons X ait dû attendre près de quatre décennies pour faire ses débuts.

(ii) Dans l'approche utilisée ci-dessus, l'existence d'une inversion de population est une condition nécessaire à l'amplification de la lumière. Des physiciens se sont demandés si tous les processus d'amplification de la lumière impliquent nécessairement une inversion de population. En fait, il existe quelques processus exotiques pour lesquels une inversion de population n'est pas nécessaire, mais il faut alors préparer le système émetteur dans une superposition linéaire d'états (on dit que l'on a créé des « cohérences »)¹¹. Bien qu'ayant un grand intérêt théorique, les très rares systèmes de ce type ayant réellement fonctionné se sont, jusqu'à présent, révélés sans portée au niveau des applications.

Par ailleurs, d'autres processus de gain peuvent être compris en terme d'effets macroscopiques sans invoquer nécessairement le phénomène d'inversion de population. Il existe ainsi un système intéressant, la *laser à électrons libres*, où le processus d'amplification résulte du groupement de paquets ordonnés d'électrons rayonnant en phase (voir remarque du § 3.2.5.b).

¹¹Voir par exemple O. Kocharovskaya, Phys. Rep. **219**, 175 (1992), M.O. Scully ibid p. 191 ; W. Gawlik, Comm. At. Mol. Phys. **29**, 189 (1993), P. Mandel, Contemporary Physics **34**, 235 (1993).

3.2.2 Inversion de population dans les systèmes à quatre niveaux

a. Principe

Le modèle simple d'amplificateur du chapitre 2 supposait l'existence d'au moins trois niveaux, un niveau fondamental et deux niveaux excités b et a entre lesquels on induisait une inversion de population. Dans la pratique, il est apparu que de très nombreux systèmes lasers nécessitaient au moins un niveau supplémentaire, de sorte que les systèmes à quatre niveaux sont maintenant considérés comme les systèmes modèles pour la mise au point de lasers. Nous en décrivons brièvement plusieurs exemples : le laser à néodyme, le laser hélium-néon (et plus généralement les lasers à gaz rares), les lasers accordables à colorant ou à saphir dopé au titane, et les lasers moléculaires.

Le schéma des niveaux d'énergie d'un laser à quatre niveaux est présenté sur la figure 3.6. Un mécanisme de *pompage* (pompage optique sous l'effet de la lumière émise par une lampe ou par un autre laser, excitation par une décharge électrique etc . . .) fait passer les atomes de leur niveau fondamental f vers un niveau excité e. Un phénomène de *relaxation rapide* transfère alors les atomes de ce niveau e au niveau b, niveau supérieur de la transition laser. La transition radiative spontanée couplant b à a en l'absence d'émission laser est lente, comparée à la relaxation de e vers b. Enfin, les atomes dans le niveau inférieur a de la transition laser sont ramenés vers le niveau fondamental f par un processus de *relaxation rapide*. Les temps caractéristiques des divers processus de désexcitation sont notés τ_e, τ_b, τ_a et les hypothèses précédentes correspondent aux conditions : $\tau_e, \tau_a \ll \tau_b$. Le pompage du niveau e se fait avec un taux W_p . À l'aide de ces taux, on peut écrire des *équations cinétiques* (voir § 2.2.5.6 ou § C.2 du complément II.1) pour les populations des divers niveaux, et calculer la différence de populations $N_b - N_a$ qui conditionne l'amplification.

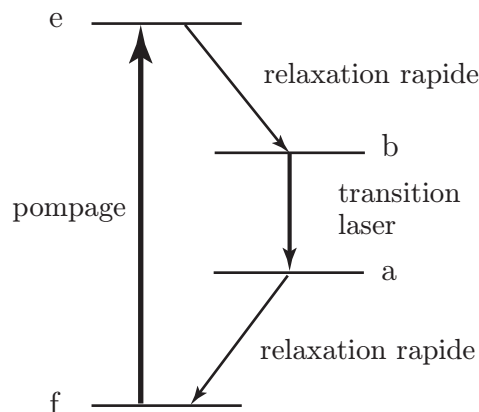


FIG. 3.6: **Schéma du laser à 4 niveaux.** La transition laser se fait sur la transition $b \rightarrow a$. Les niveaux e et a se désexcitent rapidement vers les niveaux b et f (respectivement). Un mécanisme de pompage peuplé e à partir de f.

Quand l'effet laser est suffisamment faible pour pouvoir négliger l'émission stimulée ou l'absorption sur la transition $b \rightarrow a$, c'est-à-dire en régime non saturé, les équations de pompage¹² s'écrivent :

$$\frac{d}{dt}N_e = W_p(N_f - N_e) - \frac{N_e}{\tau_e} \quad (3.22a)$$

$$\frac{d}{dt}N_b = \frac{N_e}{\tau_e} - \frac{N_b}{\tau_b} \quad (3.22b)$$

$$\frac{d}{dt}N_a = \frac{N_b}{\tau_b} - \frac{N_a}{\tau_a} \quad (3.22c)$$

Ces équations assurent la conservation de la population totale :

$$N_a + N_b + N_e + N_f = N \quad (3.23)$$

En régime stationnaire ($dN_i/dt = 0$ avec $i = f, a, b, e$), on déduit de (3.22c) :

$$\frac{N_b}{N_a} = \frac{\tau_b}{\tau_a} \quad (3.24)$$

ce qui, compte tenu des différences d'ordre de grandeur entre les divers taux

$$\tau_e, \tau_a \ll \tau_b \quad (3.25)$$

montre que $N_b > N_a$. Dans un système à quatre niveaux pour lequel la condition (3.25) est vérifiée, *l'inversion de population est automatiquement réalisée*. Dans l'hypothèse (généralement correcte) d'un pompage faible :

$$w\tau_e \ll 1 \quad (3.26)$$

et en l'absence d'émission laser, la fraction des atomes participant à l'inversion de population est égale à

$$\frac{N_b - N_a}{N} = \frac{w\tau_b}{1 + w\tau_b} \quad (3.27)$$

Cette fraction peut être importante si le niveau b a une longue durée de vie. Un tel système à quatre niveaux est donc bien adapté à la réalisation d'amplificateurs laser comme nous le montrons dans les paragraphes ci-dessous.

Remarque

¹²Dans le cas où le pompage est dû à la lumière, ces équations découlent des équations de la matrice densité présentées dans le complément II.2 à la limite où l'évolution des cohérences se fait sur une échelle de temps bien plus courte que l'évolution des populations (c'est le cas lorsque le temps de relaxation des cohérences est très petit devant le temps de relaxation des populations). Cette situation se rencontre fréquemment parce que les cohérences sont plus fragiles que les populations et donc plus sensibles aux perturbations. Lorsque le pompage se fait par collisions isotropes comme dans une décharge électrique, il s'agit d'un processus incohérent donnant directement des équations cinétiques du type (3.22a).

En pratique, il existe des pertes supplémentaires non incluses dans le modèle précédent et l'inversion de population n'est atteinte que si le taux de pompage est suffisant. Les systèmes à quatre niveaux restent cependant privilégiés pour l'obtention d'une inversion de population.

b. Laser à néodyme

La figure 3.7 montre les niveaux d'énergie de l'ion Nd^{3+} présent sous une faible concentration dans un *verre* (verre dopé au néodyme) ou dans un *cristal* $\text{Y}_3\text{Al}_5\text{O}_{12}$ (appelé YAG dans la littérature pour « Yttrium Aluminium Garnet » en anglais, c'est-à-dire grenat d'yttrium aluminium)¹³. Le pompage se fait par une source extérieure de lumière, lampe à arc ou diode laser, qui permet d'exciter les ions dans des bandes d'énergie élevées d'où ils retombent vers le niveau supérieur de la transition laser par un processus de relaxation rapide¹⁴.

L'émission laser se produit essentiellement sur la transition à $1,06\mu\text{m}$ dans l'infrarouge (d'autres transitions sont susceptibles de donner lieu à un effet laser, mais l'émission étant moins intense, elles sont moins utilisées). En fonctionnement *impulsionnel*, les énergies obtenues peuvent atteindre quelques Joules pour des impulsions de l'ordre de 1 ps ou 1 ns. En fonctionnement *continu*, on obtient des puissances qui peuvent dépasser 100 W. Pour de nombreuses applications, on place un *cristal non linéaire* à la sortie du laser de façon à obtenir un faisceau cohérent dans le vert à 532 nm (deuxième harmonique) ou dans l'ultraviolet à 355 nm (troisième harmonique).

Les lasers à néodyme sont très efficaces (le rendement est de l'ordre de 1 % pour un pompage par lampe et peut atteindre 50 % pour un pompage par laser à semiconducteur) et ils ont une gamme très vaste d'applications. En raison de leur très forte puissance en impulsion, ils sont utilisés, pompés par des lampes flash, dans des expériences visant à réaliser la fusion thermonucléaire (voir Complément III.3, § E). Ils servent souvent aussi de *pompes* pour les lasers accordables (voir paragraphe *d* ci-dessous). Les lasers à néodyme sont aussi utilisés en régime continu. On peut utiliser un pompage par lampe à arc ou par laser à semi-conducteur (voir paragraphe 3 ci-dessous) au voisinage de $0,8\mu\text{m}$. Ces derniers dispositifs sont extrêmement intéressants car ils permettent d'avoir des sources compactes, d'excellent rendement et de coût modéré.

c. Lasers à gaz rares

L'ancêtre (toujours en activité) des lasers à gaz rares est le laser à *hélium-néon*, au schéma de pompage très astucieux. Une décharge électrique continue dans un mélange d'hélium et de néon excite l'hélium dans des niveaux *métastables*¹⁵. Lors des collisions

¹³L'utilisation du verre dopé au néodyme ou du cristal néodyme-YAG dépend des applications. Par exemple, les impulsions très énergétiques des très gros lasers utilisés dans la fusion par confinement inertiel (voir Complément III.3) sont obtenues avec des verres dont on fabrique plus facilement des disques de grands diamètres.

¹⁴Ce processus est dû au couplage avec les vibrations du matériau. N'impliquant pas d'émission de lumière, il est appelé *non-radiatif*.

¹⁵Un niveau métastable est un niveau excité qui n'est pas couplé au niveau fondamental par l'hamiltonien dipolaire électrique. Sa désexcitation fait appel à d'autres mécanismes conduisant à des probabilités faibles et donc à des durées de vie très longues.

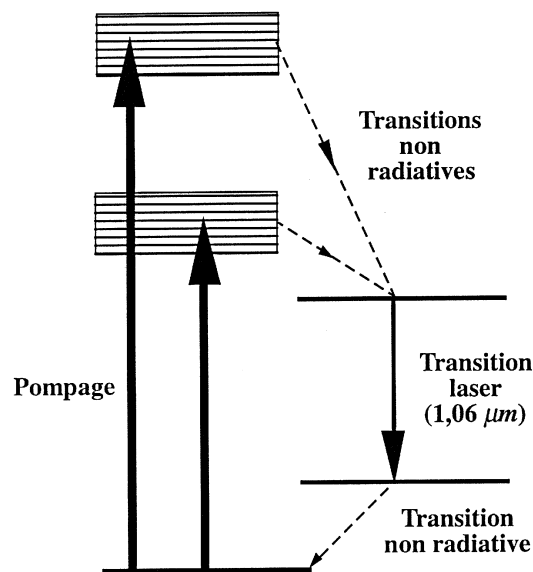


FIG. 3.7: **Schéma des niveaux (et bandes) d'énergie de l'ion néodyme Nd^{3+} dans un solide** (cristal YAG ou verre). Les longueurs d'onde utiles pour le pompage sont situées au voisinage de $0,5\mu m$ et de $0,8\mu m$. L'émission laser a lieu à $1,06\mu m$.

entre hélium métastable et néon dans l'état fondamental, l'énergie interne des atomes d'hélium métastables peut être communiquée aux atomes de néon qui se trouvent portés dans des niveaux excités dont l'énergie est voisine de celle des niveaux métastables de l'hélium. La figure 3.9 montre que ces niveaux excités du néon sont les niveaux supérieurs des raies laser à $3,39\mu m$, $1,15\mu m$ ainsi que de la célèbre raie rouge à $633\mu m$.

La figure 3.9 présente un schéma possible de construction de laser hélium-néon. On remarque que le gaz est enfermé dans un tube fermé par des fenêtres inclinées à l'*angle de Brewster* qui ne provoquent aucune perte par réflexion pour la polarisation dans le plan de la figure. La polarisation orthogonale subissant des pertes, le seuil d'oscillation est très différent pour les deux polarisations ce qui entraîne que la lumière laser ainsi obtenue est *polarisée*. La transition laser est sélectionnée par les miroirs qui n'ont une bonne réflectivité que pour la longueur d'onde que l'on veut favoriser. La puissance d'un laser hélium-néon est typiquement de l'ordre de quelques mW. Dans certaines versions industrielles, les miroirs sont à l'intérieur du tube laser et sont orthogonaux à l'axe du tube : il s'ensuit que le faisceau laser n'est pas polarisé.

Dans les lasers à *argon ionisé* ou à *krypton ionisé*, l'inversion de population est directement obtenue par une décharge électrique intense ($10^3 A/cm^2$). On peut obtenir une puissance de sortie supérieure à 10 W continus sur la raie verte de l'argon ($514,5 nm$), quelques watts sur sa raie bleue, un peu moins dans l'ultra violet. Le laser à krypton ionisé peut pour sa part fournir des puissances du même ordre respectivement dans le rouge, le jaune, ou le violet. Ces lasers ont un *rendement* extrêmement *faible* (quelques watts de lumière pour des dizaines de kilowatts électriques), ils sont chers, et leur fiabilité n'est pas toujours exemplaire. Ils constituent pourtant des outils encore difficilement remplaçables

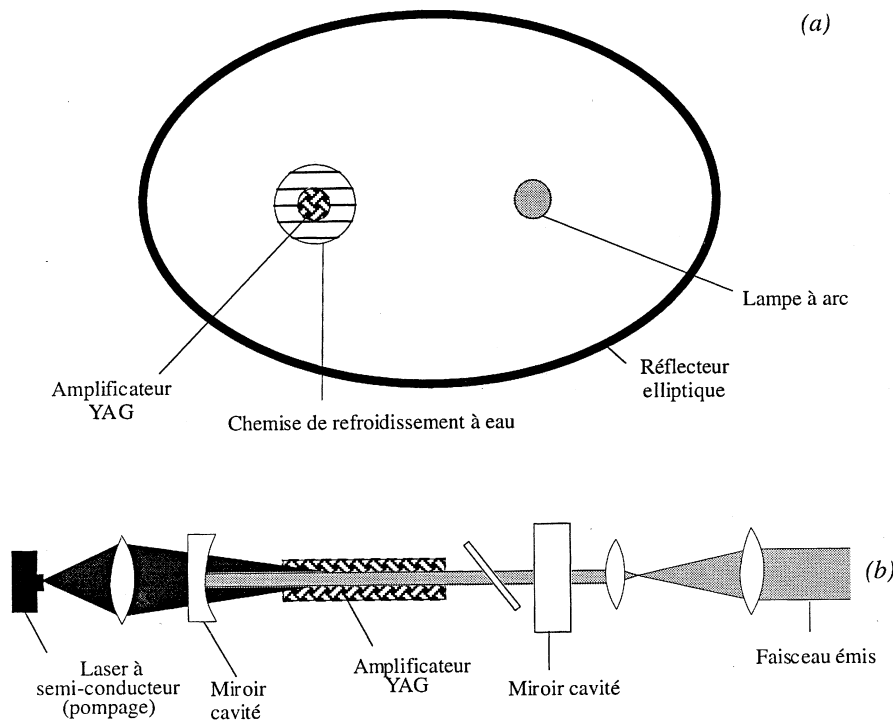


FIG. 3.8: **a)** Section de la zone d'amplification d'un **laser à Néodyme YAG à pompage par lampe à arc**. La lampe et l'amplificateur sont situés aux deux foyers d'une ellipse. Un tel laser peut délivrer de 1 à 100 W dans l'infrarouge. Il consomme de 0,1 à 10 kW électriques. **b)** **Laser à Néodyme YAG pompé par un laser à semi-conducteur**. Un tel laser peut avoir un rendement de l'ordre de 50 %.

dans de nombreuses applications scientifiques ou médicales. Les discothèques sont aussi d'excellents clients pour ce type de laser (« light show »).

Tous ces lasers ont une *courbe de gain très étroite* dont la largeur est de l'ordre de quelques gigahertz. Il est hors de question d'ajuster la longueur d'onde émise dans une large plage (le visible s'étend sur 3×10^{14} Hz!). C'est pour répondre au besoin d'accordabilité sur une plage étendue de longueur d'onde qu'ont été inventés les lasers accordables.

d. Lasers accordables

Ces lasers offrent la possibilité de choisir la longueur d'onde de fonctionnement dans une plage relativement vaste. Cette accordabilité est liée à la *grande largeur de la courbe de gain* qui résulte du fait que le *niveau inférieur de la transition laser appartient à un continuum* (ou à un quasi-continuum dense) *de niveaux*. Un exemple typique de laser accordable est le laser à *colorant* dont le milieu actif est constitué de molécules de colorant (Fig. 3.11) en phase liquide. Ces lasers sont, en général, pompés optiquement par lampe ou mieux par un laser intense mais non accordable comme le laser néodyme-YAG ou le laser à Argon ionisé.

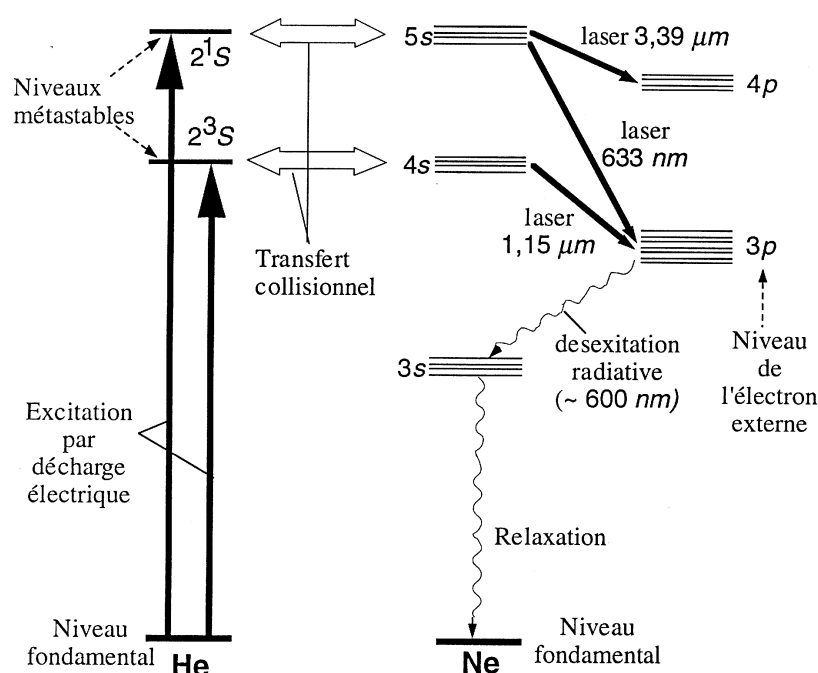


FIG. 3.9: **Niveaux de l'hélium et du néon** impliqués dans les transitions du laser à Hélium-Néon. L'excitation électronique de l'hélium dans un niveau métastable est transférée aux niveaux d'énergie voisine du néon lors des collisions entre atomes.

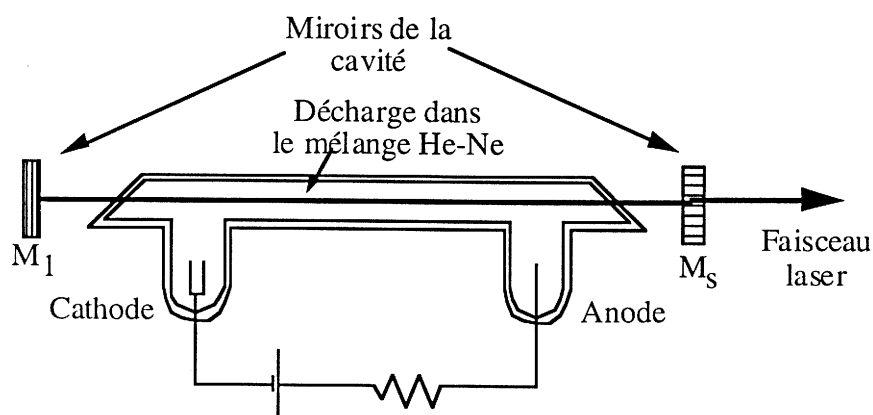


FIG. 3.10 : **Schéma d'un laser Hélium-Néon** conduisant à une émission laser polarisée.

L'accord en longueur d'onde se fait en insérant dans la cavité des *éléments sélectifs* (filtres) réglables qui favorisent une longueur d'onde au détriment des autres (voir figure 3.12). Un laser utilisant le colorant *rhodamine 6G*, par exemple, pompé par la raie verte du laser à argon ionisé, constitue une source accordable sur la plage [565 nm, 595 nm] (soit une plage de 25 000 GHz en fréquence, ce qui est considérable par rapport à la précision de 1 MHz, ou mieux, avec laquelle on peut choisir la fréquence lumineuse). Pour une puissance de pompe de 8 W, la puissance émise par le laser à rhodamine 6G est de 1 W environ. En changeant de colorant, il est possible de couvrir toute la gamme visible. Ces

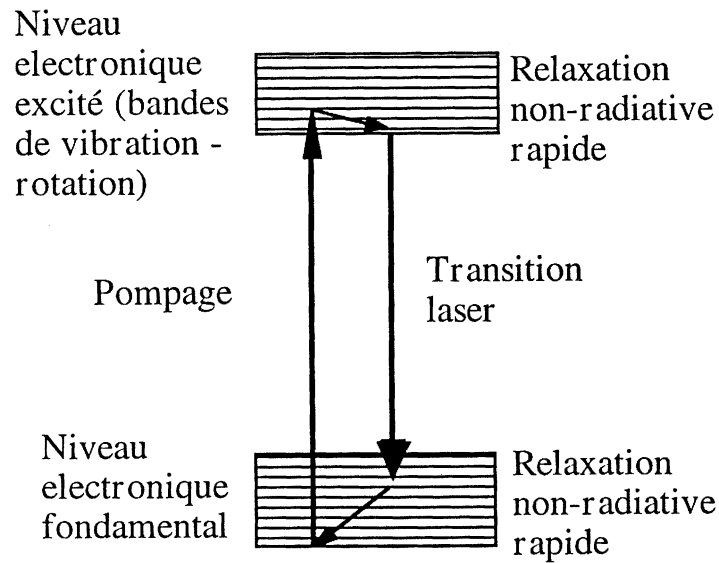


FIG. 3.11: **Schéma du mécanisme de pompage d'un laser à colorant.** Les niveaux d'énergie d'une molécule de colorant forment des bandes de niveaux d'énergie correspondant à la vibration et à la rotation de la molécule. La relaxation dans une bande donnée se fait vers les niveaux de bas de bande par des processus non-radiatifs rapides (temps typique : 1 ps). Deux bandes sont séparées par une énergie électronique. La gamme d'accordabilité du laser est déterminée par la largeur de la bande inférieure.

lasers sont chers et compliqués, mais ils sont parfois indispensables.

Dans les années 1980 est apparu un milieu amplificateur dont les niveaux d'énergie sont très analogues à ceux présentés sur la figure 3.11 : il s'agit d'ions *titane* dans une matrice de *saphir*. Ce milieu donne une amplification entre $0,7\mu\text{m}$ et $1,1\mu\text{m}$ lorsqu'il est pompé par un laser à argon ionisé. Un tel laser peut délivrer en continu une puissance de plusieurs watts, ce qui combiné avec l'existence d'excellents doubleurs de fréquence dans cette gamme, permet d'obtenir des sources accordables continues de puissance supérieure à 100 mW dans le violet et le proche ultra-violet. À cause de sa plus grande fiabilité et de sa relative simplicité d'utilisation, le laser titane-saphir a, au début des années 1990, détrôné les lasers à colorant dans les domaines de longueur d'onde qu'il permet de couvrir.

e. Lasers moléculaires

Le laser à dioxyde de carbone, appelé laser à CO_2 , se rapproche des lasers à gaz du § 3.2.2.c puisque l'excitation est aussi une décharge électrique. Mais il s'agit ici de lasers pouvant être très puissants (10 kW ou plus) aux applications industrielles et militaires nombreuses. La transition laser a lieu entre deux états moléculaires *vibrationnels* distincts¹⁶, dans l'*infrarouge lointain* à $10,6\mu\text{m}$.

¹⁶C'est-à-dire deux états d'excitation différents pour le mouvement relatif des centres de masse des atomes constituant la molécule.

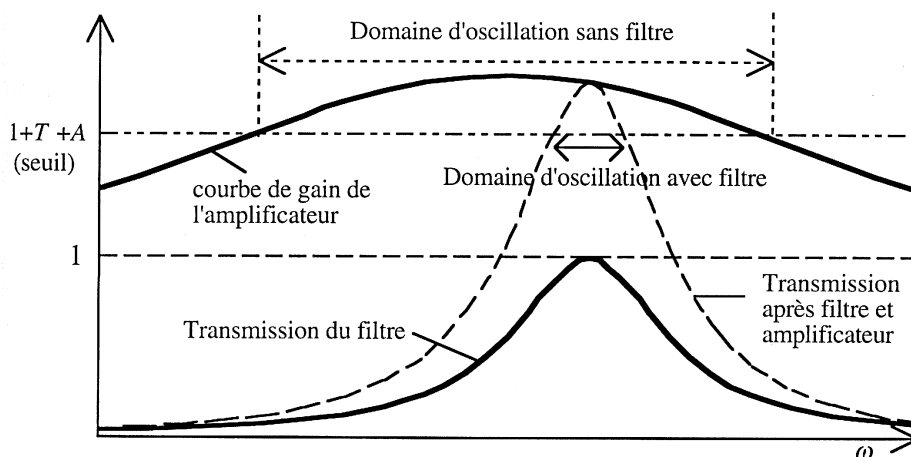


FIG. 3.12: **Sélection d'une longueur d'onde au moyen d'un filtre.** La transmission du filtre est plus étroite que la largeur de la courbe de gain et est voisine de 1 à son maximum. La transmission de la lumière après passage dans le milieu amplificateur et le filtre est le produit du coefficient d'amplification par la transmission du filtre. L'oscillation laser est possible quand ce produit est supérieur à la valeur seuil.

Il existe de très nombreux autres lasers moléculaires émettant dans l'infrarouge. Pour certains d'entre eux (lasers chimiques), la molécule est créée directement *par une réaction chimique dans un état excité*, ce qui assure *automatiquement* l'existence d'une inversion de population. Ainsi pour le laser HF, on a les réactions¹⁷ :



Un tel laser peut produire des impulsions géantes : 4 kJ en 20 ns, soit une puissance crête de 200 GW. Notons que l'énergie étant stockée sous forme chimique il n'est pas nécessaire de se connecter à une ligne électrique de puissance ; il suffit de disposer de cartouches de gaz réactifs.

Si la transition laser a lieu entre niveaux moléculaires *électroniques* distincts, la longueur d'onde de l'émission laser est souvent dans l'*ultra-violet*, ce qui confère un intérêt certain à ces sources. Les lasers les plus utilisés sont les laser à « excimères »¹⁸ comme ArF ou KrF qui émettent dans l'ultra-violet respectivement à 195 nm et 248 nm. Dans ces derniers lasers, le niveau inférieur de la transition (Fig. 3.13) est *instable* parce qu'il n'existe pas d'état lié de l'argon et du fluor (par exemple) dans le niveau fondamental (il en existe en revanche dans les niveaux excités). On obtient ainsi de façon élégante un dépeuplement efficace du niveau inférieur de la transition puisque la molécule dans cet état disparaît très rapidement en libérant les deux atomes.

¹⁷Une faible source d'énergie peut être requise pour initialiser la réaction chimique, c'est-à-dire produire les atomes libres permettant à la réaction de démarrer.

¹⁸Ce nom provient d'une abréviation de l'anglais « excited dimers ». Les lasers à « exciplexes », dont l'origine est l'anglais « excited complex », sont très voisins des lasers à excimères.

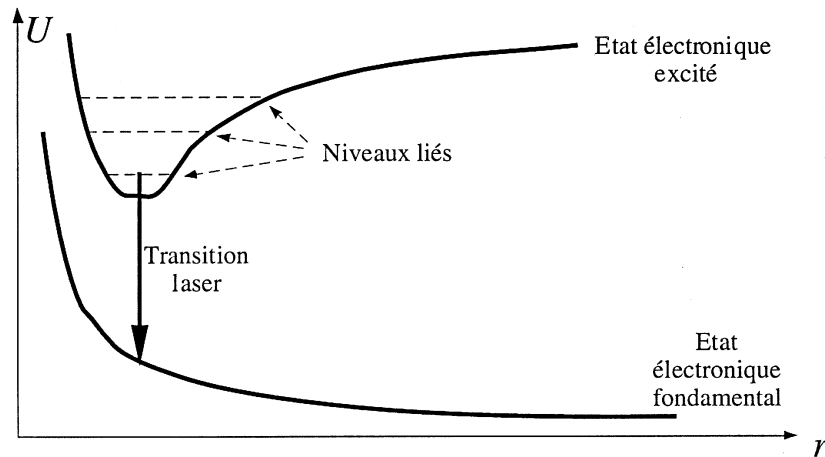


FIG. 3.13: **Schéma des niveaux d'énergie pour un laser à excimère.** On représente l'énergie d'interaction entre les deux atomes de l'excimère en fonction de leur distance r et cela pour deux niveaux d'excitation électronique. Dans le niveau supérieur, le potentiel d'interaction présente un minimum dans lequel des niveaux liés de vibration-rotation peuvent être trouvés. Une décharge dans un mélange contenant un gaz rare et un halogène (par exemple Ar et F_2) conduit à la création de molécules ArF stables dans l'état excité. Le niveau supérieur de la transition est ainsi un état lié de ArF . Le niveau inférieur de la transition est un niveau où la molécule se dissocie spontanément en ses deux éléments Ar et F.

3.2.3 Lasers à semi-conducteurs

Il s'agit des lasers de très loin les moins chers et les plus répandus en 1995. Leur domaine spectral se situe essentiellement dans le rouge et l'infrarouge. Chaque lecteur de disque compact en contient plusieurs (émission vers $0,8\mu\text{m}$), et les télécommunications à fibre optique stimulent fortement la production des lasers à semi-conducteurs (encore appelés « diodes laser ») à $1,3\mu\text{m}$ et $1,5\mu\text{m}$ (zones de dispersion minimale et de transparence maximale des fibres optiques, voir § D.3 du complément III.4).

L'émission de lumière se produit dans la zone de jonction d'une diode à semi-conducteurs formée de matériaux fortement dopés et polarisée dans le sens passant (Figure 3.14). La *recombinaison électron-trou* peut se faire avec libération de l'énergie sous forme d'un photon d'énergie $\hbar\omega \approx E_g$ où E_g est l'énergie séparant le haut de la bande de valence du bas de la bande de conduction (intervalle ou « gap », de l'ordre de 1,4 eV pour l'Ar-séniure de Gallium AsGa, soit une longueur d'onde entre $0,8\mu\text{m}$ et $0,9\mu\text{m}$). En régime d'émission spontanée, on a une diode électroluminescente (LED en anglais), composant de base de nombreux afficheurs. Mais si le courant injecté dans la jonction augmente, on peut atteindre le régime où l'émission stimulée est prédominante : le système devient alors un amplificateur optique de très grand gain.

Les premiers lasers de ce type fonctionnèrent dès 1962, mais leur emploi était réservé à des laboratoires spécialisés : il fallait les maintenir à basse température (77 K, température de

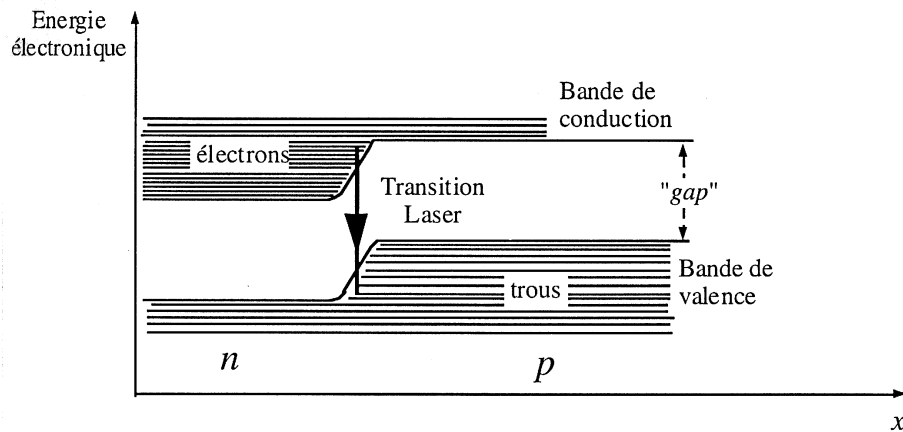


FIG. 3.14: *Schéma des niveaux d'énergie d'un laser à semiconducteur en fonction de la position. L'émission de lumière est associée à la recombinaison électron-trou, induite lorsqu'on applique une tension positive entre la zone p et la zone n.*

l'azote liquide) pour limiter les relaxations non radiatives ; la densité de courant nécessaire pour dépasser le seuil laser était très élevée ($5 \times 10^3 \text{ A cm}^{-2}$ à 77 K).

Dans les années 70, les technologies des composants semi-conducteurs, comme l'arséniure de gallium et les produits dérivés, ont connu des progrès considérables. Il est devenu possible, avec des structures à plusieurs couches de composition différente (hétérostructures, voir Fig. 3.15), de confiner la zone de recombinaison à une épaisseur de $0,1 \mu\text{m}$. Le seuil laser est alors atteint avec des courants beaucoup plus faibles.

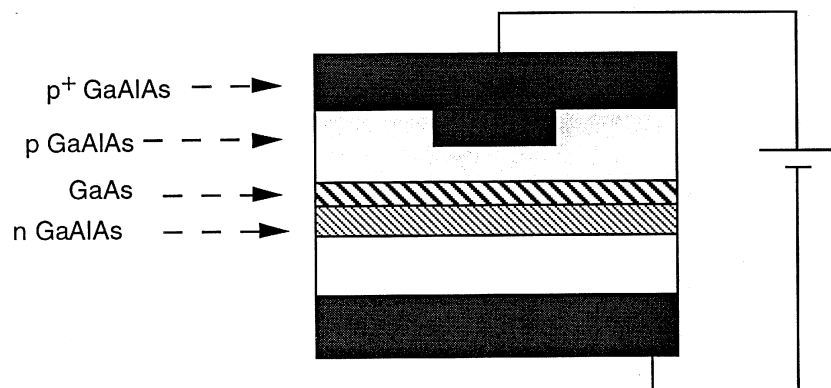


FIG. 3.15: *Schéma de la structure des couches dans un laser à semi-conducteur (hétérostructure). Plusieurs couches, présentant des « gaps » différents, sont superposées.*

L'émission stimulée se produit dans la mince couche active, au sein de laquelle les ondes lumineuses sont guidées comme dans une fibre optique (l'indice de réfraction y est plus élevé). On clive deux facettes perpendiculaires à cette nappe et parallèles entre elles ; elles servent de miroirs pour une cavité résonnante linéaire monolithique (longueur typique : $400 \mu\text{m}$). Enfin un

confinement latéral (largeur de l'ordre de $10\mu\text{m}$) permet d'abaisser encore le courant nécessaire à l'obtention de l'inversion de population (Fig. 3.15). Le rendement de ces lasers étant typiquement de l'ordre de 1 Watt de puissance lumineuse émise pour un 1 Ampère de courant, on dispose de lasers commerciaux alimentés par des courants de l'ordre de l'ampère fournissant en continu quelques centaines de milliwatts à température ambiante.

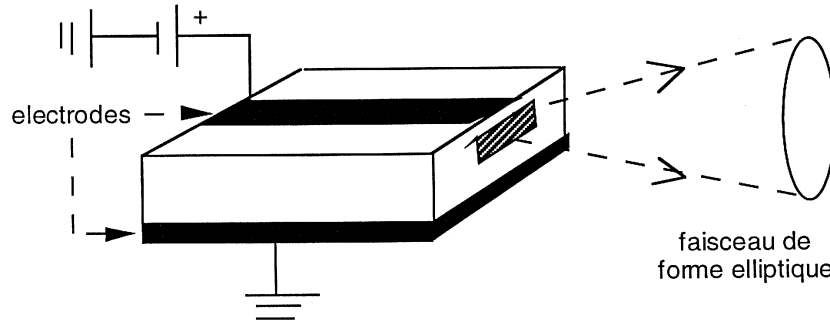


FIG. 3.16: **L'émission du laser à semi-conducteur** se fait par les facettes dont les dimensions sont de l'ordre de $1\mu\text{m} \times 5\mu\text{m}$. Le faisceau émis est elliptique avec une divergence verticale de l'ordre de 30° et une divergence latérale inférieure à 10° .

Du fait de la diffraction, le faisceau obtenu est divergent et elliptique (voir figure 3.16) : on doit le remettre en forme avec une optique élaborée. Ces lasers n'en restent pas moins extraordinairement intéressants par leur simplicité de mise en œuvre et leur rendement supérieur à 50 %. Les progrès dans les performances sont constants, avec les tendances suivantes :

- raccourcissement des longueurs d'onde, de l'infrarouge vers le bleu grâce à des compositions judicieuses des semi-conducteurs ;
- augmentation des puissances disponibles (1 W continu à $0,8\mu\text{m}$, avec un courant de quelques ampères) ;
- pureté spectrale améliorée grâce aux lasers monomodes à réaction répartie (Distributed FeedBack) et aux lasers à miroir de Bragg incorporé (Distributed Bragg Reflector encore appelé laser à émission surfacique).

Remarque

Suivant une règle fréquente dans l'industrie des semi-conducteurs, l'investissement initial pour la mise au point d'un laser de type donné, à une longueur d'onde spécifiée, est très élevé. Ce n'est donc que pour les longueurs d'onde associées à un marché important que l'on trouve des lasers à semi-conducteurs de faible coût.

3.2.4 Transition laser aboutissant sur le niveau fondamental : systèmes à trois niveaux

Pour la plupart des lasers, le niveau inférieur de la transition ne coïncide pas avec le niveau fondamental. Il est, en effet, difficile d'obtenir une population plus élevée dans un niveau excité que dans le niveau fondamental. Par un curieux hasard de l'histoire,

il se trouve pourtant que le premier laser à avoir fonctionné (laser à rubis) déroge à la règle commune. Plus récemment, un autre système à trois niveaux a pris une importance considérable dans les amplificateurs à fibre optique utilisés en télécommunications : il s'agit de l'ion erbium dans une matrice de silice. Nous nous proposons donc d'étudier la répartition des populations dans les systèmes à deux et trois niveaux, et d'évaluer dans quelles conditions une émission laser aboutissant dans le niveau fondamental peut être obtenue.

Montrons d'abord que dans un système à deux niveaux fermés (a étant le niveau inférieur et b le niveau supérieur), il n'est pas possible de créer une inversion permanente de population par un mécanisme de pompage pouvant être décrit par des équations cinétiques d'une forme analogue à (3.22a) :

$$\frac{d}{dt}N_b = w(N_a - N_b) - \frac{N_b}{\tau_b} \quad (3.28a)$$

$$\frac{d}{dt}(N_a + N_b) = 0 \quad (3.28b)$$

En régime permanent, on trouve, d'après (3.28a)

$$\frac{N_b}{N_a} = \frac{w\tau_b}{1 + w\tau_b} \quad (3.29)$$

ce qui montre que la population du niveau excité est toujours plus petite que la population du niveau fondamental.

Remarques

(i) Nous avons vu au chapitre précédent (§ II.C.2) que lorsqu'une onde monochromatique cohérente et résonnante interagit avec un système à deux niveaux initialement dans l'état fondamental, il est possible au bout d'un certain temps de trouver avec certitude le système dans l'état excité (précession de Rabi). On réalise ainsi, *de façon transitoire*, une inversion de population. Cependant, cette inversion nécessite l'existence préalable d'une onde cohérente identique à celle que l'on veut créer par émission laser.

Cet exemple montre par ailleurs la différence existant entre un *pompage incohérent* réalisé, par exemple, au moyen d'une lampe (voir exercice E.3) et conduisant à des *équations cinétiques pour les populations*, et un pompage réalisé par une *source cohérente*. Dans le premier cas, il n'y a jamais inversion de population pour le système à deux niveaux, comme le montre l'équation (3.29). Dans le second cas, l'évolution, obtenue en résolvant l'équation de Schrödinger ou l'équation d'évolution de la matrice densité (voir Complément II.2), peut conduire pour des temps d'interaction déterminés à une inversion de population.

(ii) Dans les masers¹⁹ à NH_3 ou à hydrogène, la transition sur laquelle se produit l'effet maser s'apparente à un système à deux niveaux. Dans ces dispositifs, on prépare un jet d'atomes dans le niveau excité de la transition maser par une méthode du type Stern-Gerlach. Ce jet d'atomes traverse ensuite une cavité résonnante analogue à la cavité optique du laser. Une fraction importante d'atomes est alors transférée vers le niveau fondamental, à cause

¹⁹Voir BD § VI.5 ou C. Townes et A.L. Schawlow « Microwave Spectroscopy » § 15.10 et 17.7, (Mc Graw Hill).

du processus d'émission stimulée. Les atomes sortent ensuite de la cavité. Les atomes dans le niveau inférieur de la transition sont donc mécaniquement extraits de la cavité ce qui maintient automatiquement l'inversion de population. Un tel système est en fait bien décrit par les équations de la partie D du chapitre II, à condition d'interpréter Γ_D comme l'inverse du temps moyen de séjour des atomes dans la cavité résonnante.

La figure 3.17.a présente un schéma de système à trois niveaux pouvant, comme nous le montrons dans la suite, conduire à une émission laser sur la transition joignant le niveau intermédiaire b au niveau fondamental a . Le pompage se fait du niveau a vers un niveau excité e qui se vide rapidement avec une constante de temps τ_e vers le niveau b . La désexcitation radiative spontanée du niveau b vers le niveau a est, en l'absence d'émission stimulée, beaucoup plus lente (les constantes de temps vérifient $\tau_e \ll \tau_b$).

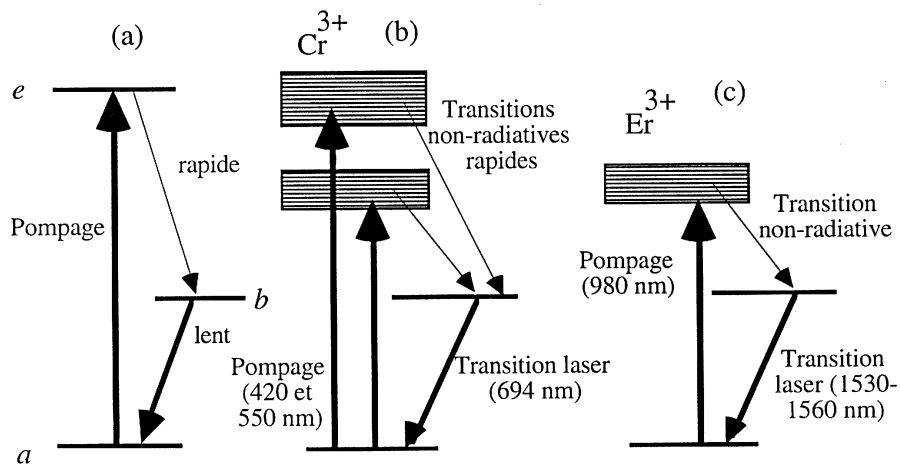


FIG. 3.17: **Schéma de principe des niveaux d'énergie pour un laser à trois niveaux** (a), et sa réalisation (b) dans le cas de l'ion Cr^{3+} dans un cristal de rubis (c) et de l'ion Er^{3+} .

Dans le cas du laser à *rubis*, l'élément actif est l'ion Cr^{3+} situé en substitution de l'ion Al^{3+} dans une matrice d'alumine. La figure 3.17.b montre les niveaux d'énergie de l'ion Cr^{3+} dans le cristal. On note l'existence de deux bandes d'absorption pouvant être excitées par absorption de la lumière émise par une lampe flash. Dans le premier laser réalisé par Maiman, les faces du barreau de rubis étaient polies parallèlement l'une à l'autre et munies d'une argenture semi-réfléchissante de façon à constituer une cavité laser linéaire.

Un autre laser dont le schéma de principe est identique à celui présenté sur la figure 3.17.a est le laser à *erbium* (voir Fig. 3.17.c). Ce laser, dont la longueur d'onde d'émission est $1,5\mu\text{m}$, a donné une nouvelle jeunesse aux schémas à trois niveaux. L'intérêt essentiel de ce laser est que sa longueur d'onde se situe exactement au minimum d'absorption des fibres optiques et il présente donc un avantage évident pour les *télécommunications optiques* (voir Complément III.4). De plus, les ions d'erbium peuvent être insérés dans la silice pour obtenir des *fibres optiques dopées à l'erbium*. Sous l'action d'un faisceau obtenu par exemple à partir d'un laser à semi-conducteur, ces fibres constituent de remarquables amplificateurs de lumière permettant de compenser, à intervalles réguliers, les pertes des impulsions lumineuses se propageant dans les fibres et transportant l'information. De tels

systèmes fonctionnent, entre autres, sur des cables téléphoniques transocéaniques à fibre optique.

Remarque

Une oscillation laser peut également être obtenue avec une fibre dopée à l'erbium en plaçant la fibre entre deux miroirs. Ces miroirs peuvent d'ailleurs être intégrés aux extrémités de la fibre.

Les équations cinétiques pour le pompage des lasers à trois niveaux sont les suivantes :

$$\frac{d}{dt}N_e = W_p(N_a - N_e) - \frac{N_e}{\tau_e} \quad (3.30a)$$

$$\frac{d}{dt}N_b = \frac{N_e}{\tau_e} - \frac{N_b}{\tau_b} \quad (3.30b)$$

$$N_a + N_b + N_e = N \quad (3.30c)$$

où w décrit le taux de pompage du niveau a au niveau e . Ces équations, qui négligent l'absorption et l'émission stimulée, sont valables en régime non saturé. Dans la limite (réalisée en pratique) où $W_p\tau_e \ll 1$, le rapport des populations en régime permanent (c'est-à-dire lorsque les dérivées par rapport au temps sont nulles dans les équations (3.30a),(3.30b),(3.30c)) est :

$$\frac{N_b}{N_e} = \frac{\tau_b}{\tau_e} \gg 1 \quad (3.31a)$$

$$\frac{N_e}{N_a} = W_p\tau_e \ll 1 \quad (3.31b)$$

L'inversion de population est donnée par :

$$\frac{N_b - N_a}{N} = \frac{W_p\tau_b - 1}{W_p\tau_b + 1} \quad (3.32)$$

Ce résultat peut être comparé à celui obtenu pour la laser à quatre niveaux (formule (3.27)). Alors que l'inversion de population est obtenue très facilement pour un laser à quatre niveaux, il faut, dans le cas du laser à trois niveaux que *le pompage soit suffisamment intense* pour atteindre

$$W_p \gg \frac{1}{\tau_b} \quad (3.33)$$

Cette nécessité d'avoir de forts taux de pompage est un premier inconvénient des lasers à trois niveaux. Un second inconvénient est associé à la difficulté de faire fonctionner certains de ces lasers en régime continu. En effet, dès que l'émission laser se produit, il faut rajouter dans les équations cinétiques le terme d'émission stimulée de b vers a qui va conduire à un raccourcissement de la durée de vie du niveau b . De ce fait, l'inversion de population donnée par (3.32) va diminuer jusqu'au moment où le gain ne sera plus suffisant pour assurer l'oscillation laser qui va s'arrêter. Le processus de pompage peut

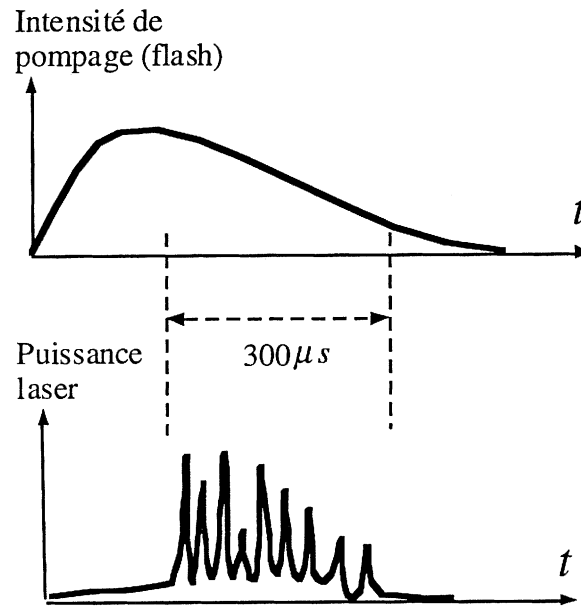


FIG. 3.18: **Représentation schématique de l'évolution temporelle** du pompage par lampe flash (a) et de la puissance de sortie du laser à rubis (b).

alors accumuler à nouveau des atomes dans le niveau b ce qui conduit à l'émission d'un train d'*impulsions* (oscillations de relaxation) par le laser (voir figure 3.18 et § 3.4.2).

Remarque

Il existe un autre système laser à trois niveaux mais dans lequel l'émission laser n'aboutit pas sur le niveau fondamental. Il s'agit du système étudié dans le chapitre 2 (Partie 2.5) lorsque le niveau supérieur de la transition laser b est alimenté directement par pompage à partir du niveau fondamental. Un certain nombre de lasers dans l'infra-rouge fonctionnent selon ce principe, le pompage étant effectué sur une transition résonnante à l'aide d'un laser de plus petite longueur d'onde.

3.2.5 Lasers Raman

Les spectres optiques des molécules possèdent des séries de raies associées à des changements de niveaux de vibration (ou de rotation) de la molécule. Si un gaz constitué de telles molécules est soumis à un laser incident de fréquence ω , non nécessairement résonnant, il y a une possibilité d'obtenir du gain par effet Raman stimulé (Chapitre II, § A.5) pour une onde de fréquence ω' telle que $\omega - \omega'$ coïncide avec une fréquence de vibration (ou de rotation) de la molécule. Considérons, par exemple, les transitions vibrationnelles. À l'équilibre thermodynamique, le niveau vibrationnel le plus profond (correspondant au nombre quantique $n = 0$) étant le plus peuplé parce qu'il correspond au minimum d'énergie (voir figure 3.19), la transition Raman allant de $n = 0$ à $n = 1$ est plus probable que la transition Raman inverse allant de $n = 1$ à $n = 0$. L'amplification du faisceau

de fréquence ω' peut donner lieu à une oscillation laser si les molécules sont enfermées dans une cavité résonnante pour la fréquence ω' . On obtient ainsi un laser Raman dont la longueur d'onde est décalée vers le rouge par rapport à la longueur d'onde du laser incident. Ce mécanisme est notamment utilisé pour obtenir des lasers dans l'infra-rouge. On peut utiliser plusieurs *lasers Raman* en cascade, fonctionnant éventuellement avec des molécules différentes, pour accéder à l'infra-rouge lointain. On réalise dans le domaine millimétrique la jonction avec les fréquences produites par des oscillateurs électroniques (klystrons, carcinotrons, diodes Gunn etc ...).

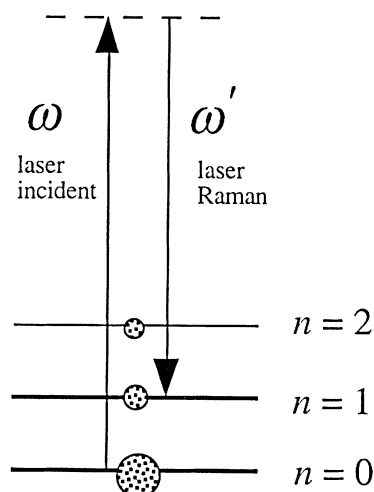


FIG. 3.19: **Schéma des niveaux d'énergie dans un laser Raman.** Les populations sont schématisées par les disques. La transition Raman stimulée se produit d'un niveau vibrationnel plus peuplé vers un niveau vibrationnel moins peuplé. À l'équilibre thermodynamique, le laser Raman a donc une fréquence ω' inférieure à celle ω du laser incident.

Remarque

Le *laser à électrons libres* peut être décrit par un groupement des électrons en paquets comme mentionné dans la remarque (ii) du § 3.2.1. Il est cependant possible d'en avoir une autre description, se rapprochant des lasers Raman, en raisonnant dans l'espace des *impulsions* de l'électron. Considérons un ensemble d'électrons, dont les vitesses sont à l'équilibre thermodynamique, interagissant avec une onde de fréquence ω . Un processus de diffusion stimulée²⁰, avec absorption d'un photon de fréquence ω , émission d'un photon ω' et modification de l'impulsion de l'électron, est susceptible d'amplifier un faisceau de fréquence ω' se propageant en sens opposé au faisceau de fréquence ω . Ce processus, schématisé sur la figure 3.20, est plus probable que le processus inverse (absorption de ω' , émission de ω) lorsque le niveau initial (correspondant à une impulsion nulle sur la figure) est plus peuplé que le niveau final du processus (électron ayant une impulsion $\hbar(\mathbf{k} - \mathbf{k}')$ provenant de l'échange d'impulsion entre champ électromagnétique et électron). Notons qu'ici encore le faisceau amplifié a une fréquence plus petite que le faisceau incident (la différence de fréquences se retrouvent sous forme d'énergie cinétique de l'électron). En fait, dans le laser à

²⁰Ce processus est appelé « diffusion Compton stimulée ». Il s'apparente à la diffusion Raman stimulée étudiée dans le paragraphe (A.5) du chapitre II, la différence essentielle étant que les niveaux initial et final sont discrets dans la diffusion Raman et continus (l'état de l'électron étant défini par son impulsion) dans la diffusion Compton.

électrons libres, les électrons ne sont pas soumis à un champ électromagnétique incident. On envoie un faisceau d'électrons rapides dans une structure magnétique modulée spatialement (« onduleur ») qui donne, par passage dans le référentiel propre de l'électron le champ de fréquence ω induisant le processus de diffusion stimulée. La fréquence de l'onde susceptible d'être amplifiée est obtenue, dans le référentiel du laboratoire, par une transformation de Lorentz. Une cavité résonnante accordée à cette fréquence entoure l'onduleur de façon à pouvoir obtenir une oscillation²¹.

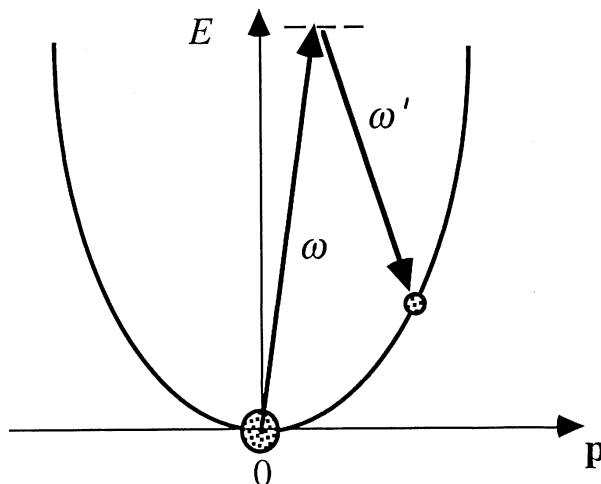


FIG. 3.20: **Laser à électrons libres.** Schéma d'un processus Compton stimulé conduisant à un processus de gain pour une onde électromagnétique de fréquence ω' inférieure à la fréquence incidente ω . La courbe représente la relation entre énergie et impulsion de l'électron. À l'équilibre thermodynamique, les états les plus peuplés correspondent aux énergies les plus basses. La transition stimulée se fait donc du niveau le plus peuplé ($p = 0$) vers le niveau le moins peuplé qui appartient ici à un continuum d'états.

3.3 Propriétés spectrales des lasers

3.3.1 Modes longitudinaux

Nous avons montré au § 3.1.2, que la condition d'oscillation (gain non saturé plus que suffisant pour compenser les pertes) pouvait être vérifiée par plusieurs modes longitudinaux du laser. Le nombre de modes susceptibles d'osciller est associé au rapport entre la largeur de la courbe de gain et l'intervalle c/L_{cav} entre modes longitudinaux (L_{cav} est, rappelons-le, la longueur optique parcourue par la lumière sur un tour de cavité). Par exemple, pour la raie rouge du laser hélium-néon, la largeur de la courbe de gain est 1,2 GHz. Si $L_{\text{cav}} = 0,6$ m, la distance entre modes est 0,5 GHz et le laser hélium-néon pourra osciller sur 2 ou 3 modes. Il faut noter que la courbe de gain est fixe en fréquence alors que le peigne de modes de fréquence ω_p se déplace quand la longueur L varie (par

²¹D.A.G. Deacon, L.R. Elias, J.M.J. Madey, G.J. Ramian, H.A. Schwettman et T.I. Smith, Phys. Rev. Lett. **38**, 892, 1977. Voir également Yariv « Quantum Electronics » 3rd ed. p. 277 (Wiley).

suite de phénomènes de dilatation par exemple). Il s'ensuit que le laser peut osciller à certains instants sur deux modes, et à d'autres instants sur trois.

Le comportement du laser hélium-néon, pour lequel tous les modes susceptibles d'osciller apparaissent simultanément dans l'émission laser, n'est cependant pas universel. Pour d'autres lasers, comme par exemple le laser néodyme-YAG, le nombre de modes oscillant réellement est très inférieur au nombre de modes pouvant apparaître d'après l'analyse précédente. Dans certains cas, il n'y a même qu'un seul mode : le laser est alors dit *mono-mode*. La différence de comportement est directement liée à la nature de l'élargissement de la courbe de gain : l'élargissement est dit « *inhomogène* » dans le cas du laser hélium-néon, alors qu'il est essentiellement « *homogène* » pour le laser néodyme-YAG. Une raie spectrale est inhomogène si deux fréquences différentes dans le profil spectral proviennent de l'émission de deux atomes placés dans des conditions différentes, alors qu'une raie est homogène si la forme de la distribution spectrale est la même pour tous les atomes émetteurs. Une raie inhomogène est en fait la superposition de plusieurs raies homogènes étroites juxtaposées et correspondant à diverses valeurs d'un paramètre extérieur dont dépend la fréquence d'émission.

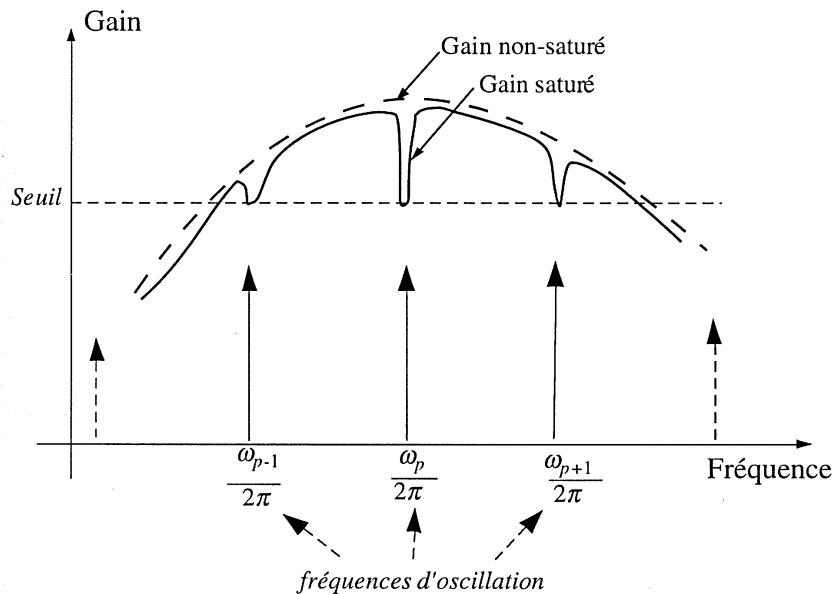


FIG. 3.21: **Effet de la saturation du gain pour un élargissement inhomogène** : il y a creusement de trous dans la courbe de gain (« hole burning ») seulement aux fréquences d'oscillation laser : plusieurs modes peuvent ainsi coexister simultanément.

Ces considérations sur largeur homogène et inhomogène s'appliquent, par exemple, à la transition à $10,6\mu\text{m}$ du laser à CO_2 . La largeur naturelle de cette transition est extrêmement faible (bien inférieure à 1 MHz) mais du fait de l'agitation des molécules dans le gaz, on a un élargissement de la raie par effet Doppler (50 MHz environ à température ambiante). Cet élargissement est inhomogène, car ce ne sont pas les mêmes molécules qui participent à l'aile haute fréquence de la raie (elles se déplacent en sens contraire du faisceau laser), ou à l'aile basse fréquence (elles

se déplacent dans le même sens). Si maintenant on considère du dioxyde de carbone à pression élevée, ce sont les collisions qui deviennent la cause prédominante d'élargissement de la raie.

Ce nouvel élargissement est homogène car toutes les molécules subissent en moyenne les mêmes oscillations et voient donc leurs fréquences de transition affectées de façon identique²².

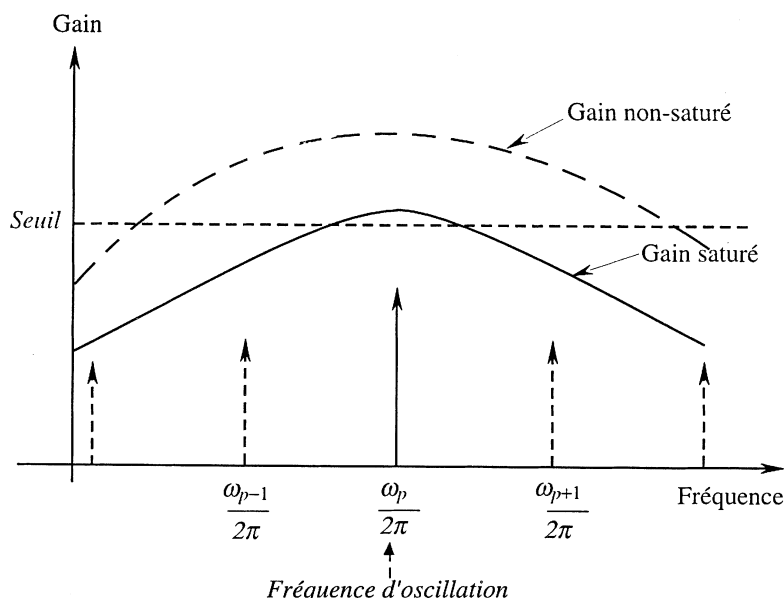


FIG. 3.22: **Saturation du gain pour un élargissement homogène** : toute la courbe de gain est atténuée par la saturation, et un seul mode peut osciller en régime permanent, celui qui est le plus proche du maximum de la courbe du gain.

La distinction entre largeur homogène et inhomogène est fondamentale pour déterminer le nombre de modes oscillants. Rappelons d'abord que le comportement stationnaire du laser est déterminé par la *saturation du gain* qui impose en régime permanent l'égalité entre le gain saturé et les pertes (voir § 3.1.2). Dans le cas d'un *élargissement inhomogène* (Figure 3.21), la *saturation de l'amplification* (c'est-à-dire la diminution du gain lorsque l'intensité lumineuse augmente) n'affecte que les molécules (ou les atomes) du milieu amplificateur dont la fréquence est accordée sur la fréquence lumineuse active : *l'effet laser à une fréquence donnée n'affecte pas le gain à une autre fréquence*. On observera alors l'*oscillation simultanée de tous les modes* autorisés par la condition de seuil. Cette situation est celle du laser hélium-néon où l'élargissement est essentiellement inhomogène (il est dû à l'effet Doppler).

Considérons à présent un laser pour lequel la courbe de gain a un *élargissement homogène* (Figure 3.22). Dans ce cas, la réponse de chaque atome est affectée de la même façon par la saturation du gain. La courbe de gain va donc subir *en bloc* l'effet de la saturation,

²²La distinction entre largeur homogène et largeur inhomogène est également discutée dans le complément III.5 consacré à la spectroscopie laser. Cette distinction joue alors un rôle fondamental puisque la spectroscopie laser peut permettre d'obtenir des signaux résonnants ayant la largeur homogène même quand l'élargissement inhomogène est dominant.

et la condition d'égalité entre le gain saturé et les pertes n'est réalisée que pour *un seul mode* (voir figure 3.22). L'oscillation sur ce mode empêche donc l'émission laser sur les autres modes.

Remarque

Cette dernière situation est celle du laser à néodyme-YAG, où l'élargissement est dû aux vibrations des atomes de la matrice du solide, qui modulent l'environnement de l'atome émetteur et donc ses niveaux d'énergie, comme dans une collision (on parle alors de collision avec les phonons).

En pratique, il est exceptionnel qu'un élargissement soit purement homogène ou purement inhomogène. Pour les lasers fonctionnant en régime continu mentionnés dans le tableau 1, un des types d'élargissement est dominant, mais dans de nombreux autres cas les largeurs homogènes et inhomogènes peuvent avoir des valeurs comparables.

	Largeur de la courbe de gain	Intervalle typique entre modes	Nombre de modes longitudinaux possibles	Nature de l'élargissement dominant
He - Ne	1,2 GHz	500 MHz	3	Inhomogène (Doppler)
Ar ⁺	12 GHz	90 MHz	140	Inhomogène (Doppler, effet Stark)
CO ₂ (haute pression)	0,5 GHz	100 MHz	5	Homogène
Néodyme- YAG	120 GHz	300 MHz	400	Homogène
Colorant (Rhodamine 6G)	25 THz	250 MHz	10 ⁵	Homogène
Saphir dopé au titane	100 THz	250 MHz	4 × 10 ⁵	Homogène
Semi- conducteur	1 THz	100 GHz	10	Homogène

Tableau 1 : Caractéristiques importantes de lasers fonctionnant en régime continu. Ordre de grandeur typique pour le nombre de modes longitudinaux.

3.3.2 Fonctionnement monomode longitudinal

Pour de nombreuses applications (spectroscopie²³, holographie, métrologie) on souhaite avoir une source laser aussi monochromatique que possible. Il faut donc obtenir un laser oscillant sur *un seul* mode longitudinal (fonctionnement monomode longitudinal). Nous venons de voir qu'un tel régime peut être atteint spontanément dans le cas d'une

²³Voir complément III.5.

raie élargie de façon purement homogène, mais cette situation reste exceptionnelle, et il faut en général intervenir pour rendre le laser monomode. Une méthode particulièrement simple consiste à utiliser une cavité assez *courte* pour que l'intervalle entre modes c/L_{cav} soit de l'ordre de la largeur de la courbe de gain : c'est le cas des lasers hélium-néon très courts ($c/L_{\text{cav}} = 1$ GHz pour $L_{\text{cav}} = 0,3$ m) ou de certains lasers à semi-conducteurs ($c/L = 1000$ GHz). Cette méthode n'est cependant possible que dans des cas favorables.

Le moyen le plus sûr pour rendre un laser monomode consiste à introduire à l'intérieur de la cavité un filtre, c'est-à-dire un élément sélectif en fréquence dont la courbe de transmission est assez fine pour que le gain effectif (gain de l'amplificateur multiplié par la transmission du filtre, voir figure 3.12) ne dépasse le seuil laser que pour un seul mode longitudinal. La transmission maximale du filtre doit être aussi voisine de 1 que possible, pour ne pas introduire de pertes supplémentaires à la fréquence laser. La fréquence de ce maximum doit être exactement ajustée sur un mode de la cavité laser (de préférence celui qui est le plus proche du maximum de la courbe de gain), et l'ajustement doit être assuré en permanence grâce à un système d'asservissement électronique.

Remarques

(i) Le système de sélection de fréquence se complique encore lorsqu'il faut utiliser plusieurs éléments sélectifs en cascade, de largeur de filtrage de plus en plus étroite, et dont les maximums doivent coïncider exactement : c'est ainsi qu'un laser à colorant (courbe de gain de 25 THz de large) comporte un triple filtre biréfringent sélectionnant une bande de 300 GHz, un interféromètre de Fabry-Perot²⁴ « mince » (1 mm) sélectionnant une bande de 3 GHz et un Fabry-Perot « épais » (5 mm) sélectionnant un seul mode de la cavité. Un tel ensemble requiert des asservissements imbriqués assez délicats à régler.

(ii) Il est en général plus difficile de tendre de rendre monomode longitudinal un laser à cavité linéaire (Figure 3.3) qu'un laser en anneau. Dans une cavité linéaire, l'onde laser est une onde stationnaire, qui possède donc des nœuds et des ventres alternés, décalés d'un quart de longueur d'onde optique. Aux ventres, l'intensité lumineuse est maximale et le phénomène de saturation diminue le gain. En revanche, l'intensité est nulle aux nœuds et le gain y est donc potentiellement beaucoup plus grand (gain non saturé). La situation est donc tout à fait favorable à l'établissement de l'effet laser pour une deuxième onde stationnaire dont les ventres coïncideraient avec les nœuds de la première. Si le milieu amplificateur est au milieu de la cavité linéaire, il est facile de vérifier que deux modes voisins ω_p et ω_{p+1} sont précisément dans cette situation, ce qui explique la tendance de tels lasers à osciller sur deux modes en l'absence d'éléments sélectifs efficaces. Ce comportement est donc à rapprocher de celui mentionné précédemment dans le cas de la saturation inhomogène du gain (Fig. 3.21), mais les trous créés dans la courbe de gain se situent maintenant dans l'espace réel et non plus dans l'espace des fréquences.

Le phénomène de modulation longitudinale du gain n'existe pas pour un laser en anneau dans lequel la lumière se propage dans un seul sens. En effet, l'onde laser est alors une onde progressive dont l'intensité est constante. Ceci permet de comprendre pourquoi un laser à colorant (dont l'élargissement est essentiellement homogène) fonctionne spontanément en régime monomode avec une cavité en anneau, et en régime multimode avec une cavité linéaire.

Le laser étant monomode, il est généralement souhaitable de contrôler la fréquence émise avec une précision bien meilleure que l'intervalle entre modes c/L_{cav} . La fréquence du

²⁴Voir Complément III.1.

mode p (voir équation (3.16)) est donnée par

$$\frac{\omega_p}{2\pi} = p \frac{c}{L_{\text{cav}}}$$

Elle varie lorsque la longueur L_{cav} de la cavité change. Un changement de fréquence égal à un intervalle entre modes $\Delta = 2\pi c/L_{\text{cav}}$ est obtenu par un changement de longueur égal à une longueur d'onde optique :

$$\delta L_{\text{cav}} = \frac{L_{\text{cav}}}{p} = \lambda \quad (3.34)$$

Pour obtenir une émission dont la largeur en fréquence est inférieure à l'intervalle spectral libre, la longueur de cavité doit donc être contrôlée à beaucoup mieux qu'une longueur d'onde optique. On y parvient en montant un des miroirs de renvoi sur un transducteur piézoélectrique, ce qui permet des déplacements maîtrisés à mieux que 0,1 nm près²⁵.

Remarques

- (i) Dans le cas des lasers à semi-conducteurs, c'est la température qui permet de contrôler la longueur optique de la cavité, ainsi que le courant (qui agit sur l'indice de réfraction).
- (ii) Dans certains lasers à semi-conducteurs, le caractère monomode de l'oscillation est assuré par le mécanisme de contre-réaction répartie. Les propriétés du guide d'onde optique sont modulées longitudinalement avec une période égale à un multiple de la demi-longueur d'onde à favoriser : on sélectionne ainsi un seul mode longitudinal du laser à semi-conducteur.

3.3.3 Largeur de raie laser

Un laser fonctionnant en régime monomode, on peut se demander avec quelle précision sa fréquence est définie. En d'autres termes, quelle est la largeur de la courbe caractérisant sa densité spectrale de puissance ? Cette largeur dépend de plusieurs causes que nous discutons maintenant.

a. Élargissement technique

Nous avons vu ci-dessus que tout changement de longueur de cavité provoque une modification de la fréquence ω_p du mode p sélectionné. Pour une cavité typique ($L_{\text{cav}} = 1$ m) un déplacement d'un miroir de $\lambda/600$ (1 nanomètre pour un laser émettant dans le jaune) entraîne une variation de fréquence de 1 MHz. Les causes de si faibles changements de longueur ne manquent pas : *dilatations* (2×10^{-3} degrés suffisent pour une structure en quartz dont le coefficient de dilatation est pourtant particulièrement faible), *variations de pression* (atmosphérique ou provoquée par une onde acoustique) entraînant un changement de l'indice de réfraction de l'air et donc un changement de longueur optique, etc.

²⁵Notons qu'il n'y a rien de paradoxal à contrôler la position d'un miroir, objet macroscopique, à mieux que la taille caractéristique des molécules de silice qui constituent le miroir : la position macroscopique d'un miroir est associée à un concept de surface moyenne.

L'ensemble de ces phénomènes entraîne une modulation plus ou moins aléatoire de la fréquence laser, que l'on appelle « *jitter* » (« gigue » en français). En fonctionnement libre, la largeur de raie à court terme peut difficilement être inférieure à quelques MHz. À long terme (plus d'une minute), les dérives thermiques provoquent en plus des déplacements au moins égaux à l'intervalle entre modes et la fréquence n'est pas définie à mieux que c/L_{cav} . Pour de nombreuses applications, une telle situation n'est pas satisfaisante.

Ici encore, la solution réside dans l'utilisation d'asservissements : on va comparer la fréquence du laser à une fréquence de référence et modifier la longueur L_{cav} de la cavité laser de façon à minimiser la différence entre la fréquence du laser et la référence. La stabilité obtenue dépend de la stabilité de la référence, et du rapport signal sur bruit dans l'obtention du signal d'erreur (mieux la différence est mesurée et meilleure sera la correction). Cette référence peut être un mode d'une cavité Fabry-Perot (voir Complément III.1) stabilisée en température, placée dans une enceinte à vide, isolée des vibrations. On obtient alors une largeur de raie qui peut approcher 1 Hz à court terme, mais on ne peut malheureusement pas empêcher les dérives lentes (vieillesse des matériaux dont la structure microscopique évolue).

Le contrôle à long terme de la fréquence laser repose sur la comparaison avec une raie atomique ou moléculaire qui constitue une référence absolue. Encore faut-il trouver une raie assez fine, et éliminer toutes les causes d'élargissement : effet Doppler (voir complément III.5), champ électrique ou magnétique parasite. Une stabilité de l'ordre de 0,1 Hz sur des raies optiques, ce qui correspond à une précision de 10^{-16} (soit environ 1 minute sur l'âge de l'univers!) est déjà réalisée au laboratoire²⁶.

b. Limite fondamentale. Largeur de Schawlow-Townes

Supposons que toutes les causes d'origine technique soient corrigées et que la longueur L_{cav} de la cavité laser soit rigoureusement constante. Quelle est alors la largeur de la raie laser ? C'est à Schawlow et Townes que l'on doit l'identification du mécanisme fondamental qui empêche la raie d'être infiniment étroite : il s'agit de l'*émission spontanée*, inséparable de l'émission stimulée dans les processus d'interaction entre matière et photons. En effet, lorsqu'une émission spontanée a lieu dans le mode de la cavité correspondant à l'oscillation laser, on rajoute au champ électromagnétique déjà présent (oscillation laser) une petite contribution ayant une phase aléatoire. Ceci provoque donc une variation d'amplitude et de phase du champ dans la cavité. La variation d'amplitude va être automatiquement corrigée par le phénomène de saturation du gain. En revanche, le fonctionnement du laser n'impose aucune contrainte sur la phase. Contrairement à l'amplitude, la phase n'est pas rappelée vers une valeur déterminée ; elle évolue donc, sous l'effet de l'émission spontanée, selon un processus de diffusion ou de marche au hasard. Cette *diffusion de la phase* entraîne un élargissement spectral de la raie $\Delta\omega_{\text{ST}}$ de l'ordre de l'inverse du temps de corrélation de la phase (temps au bout duquel la mémoire de la phase initiale a été perdue).

²⁶Un tel laser pourrait constituer un étalon de temps 10^3 fois plus précis que les horloges atomiques à Césium actuelles, dont la précision est utilisée couramment pour certaines applications (multiplexages dans les télécommunications, repérage par satellite « GPS », etc...), et dont on peut prévoir qu'elles se révéleront bientôt insuffisamment précises.

Un calcul simplifié de ce processus est fait dans le complément III.6. Qualitativement, la largeur $\Delta\omega_{\text{ST}}$ obtenue pour un laser très au-dessus du seuil est de l'ordre de la *largeur d'un mode de la cavité* $\Delta\omega_{\text{cav}}$ (en l'absence de processus de gain) *divisée par le nombre de photons $\mathcal{N}$ dans la cavité laser*²⁷.

Pour un laser hélium-néon, on obtient typiquement $\Delta\omega_{\text{ST}}/2\pi \approx 10^{-3}$ Hz tandis que pour un laser à semi-conducteur émettant 1 mW, on trouve $\Delta\omega_{\text{ST}}/2\pi \approx 1$ MHz. Par suite du petit nombre de photons présents dans le petit volume de la cavité, la limite de Schawlow-Townes est accessible à l'expérience dans ce dernier cas²⁸ alors que les causes d'élargissement technique sont encore largement dominantes pour le laser hélium-néon.

3.4 Lasers en impulsion

La possibilité offerte par les lasers de fournir l'énergie lumineuse sous forme d'impulsions brèves, tout en conservant une puissance moyenne à peu près constante, permet d'atteindre des puissances instantanées très élevées, et donc de très grands champs électriques donnant accès à une physique nouvelle : *optique non-linéaire, ionisation multiphotonique, création de plasmas laser* etc ... En fait, certains lasers comme le laser à rubis de Maiman fonctionnent spontanément en impulsion, à cause des oscillations de relaxation dues à la structure à trois niveaux du milieu amplificateur (voir § 3.2.4 et figure 3.18). Nous montrons aux paragraphes 2 et 3 de cette partie comment il est possible d'obtenir des impulsions géantes en contrôlant ce phénomène (fonctionnement *déclenché*). Auparavant, nous présentons dans le § 1 la méthode de « *synchronisation des modes* » (mode-locking en anglais) qui permet d'obtenir une suite d'impulsions extrêmement brèves à partir d'un laser présentant une large courbe de gain et pompé de manière continue. C'est en raffinant cette technique que des impulsions de quelques femtosecondes, ne contenant donc qu'un très petit nombre de périodes d'oscillation, ont été obtenues pour des impulsions visibles. Ces impulsions ont permis des avancées remarquables dans l'étude des *phénomènes ultra rapides*.

3.4.1 Laser à modes synchronisés²⁹

La synchronisation des modes est un moyen élégant pour produire des impulsions brèves. Pour en comprendre le principe, considérons d'abord un laser continu multimode. La lumière émise comporte plusieurs composantes monochromatiques correspondant aux

²⁷Rappelons que $\Delta\omega_{\text{cav}}$ et le temps τ_{cav} mis par l'énergie électromagnétique pour disparaître de la cavité sont reliés par une relation de la forme $\Delta\omega_{\text{cav}}\tau_{\text{cav}} \approx 1$ (voir Complément III.1).

²⁸Il existe d'autres causes d'élargissement dans les lasers à semi-conducteur. C'est pourquoi la largeur de raie laser y est généralement supérieure à la largeur Schawlow-Townes, typiquement par un facteur de l'ordre de 10 (facteur de Henry).

²⁹Un tel laser est souvent baptisé « à modes bloqués », traduction approximative de l'anglais « mode locked », dont le sens précis est « à modes verrouillés » (en phase).

divers modes, et dont les *phases sont a priori incorrélées*. Le champ électrique à la sortie d'un tel laser est égal à :

$$E(t) = \sum_{k=0}^{N-1} E_0 \cos(\omega_k t + \varphi_k) \quad (3.35)$$

où N est le nombre de modes oscillants et les φ_k sont des variables aléatoires incorrélées. Nous avons supposé pour simplifier que tous les modes ont même amplitude E_0 . La fréquence ω_k de chaque mode est de la forme

$$\omega_k = \omega_0 + k\Delta \quad (3.36)$$

où $\Delta/2\pi = c/L_{\text{cav}}$ (intervalle entre modes, voir § 3.3.2c) et ω_0 est la fréquence du mode $k = 0$.

L'intensité lumineuse, moyennée sur un temps long devant la période optique, mais court devant $1/N\Delta$, (ce qui est en pratique le signal auquel on peut avoir accès avec les photodétecteurs les plus rapides dont le temps de réponse est supérieur à 10 ps) est égale à :

$$I(t) = \frac{NE_0^2}{2} + E_0^2 \sum_{j>k} \cos[(\omega_j - \omega_k)t + \varphi_j - \varphi_k] \quad (3.37)$$

C'est la somme d'une *intensité moyenne*

$$\bar{I} = \frac{NE_0^2}{2} \quad (3.38)$$

et de *fluctuation* dont l'écart type $\Delta I = \sqrt{\overline{(I(t) - \bar{I})^2}}$, pour des variables aléatoires incorrélées φ_k , est calculable à partie de (3.37) et est égale, dans la limite $N \gg 1$, à

$$\Delta I = \bar{I} \quad (3.39)$$

La variation de l'intensité d'un laser multimode, en fonction du temps, présente donc des fluctuations importantes de l'ordre de $\bar{I}$ comme le schématise la figure 3.23

Quelle est l'origine physique des pics d'intensité observés sur la figure 3.23 ? Ceux-ci résultent du fait qu'à certains instants, plusieurs modes du champ (3.35) *interfèrent constructivement*. Il est clair que si *tous* les modes interféraient constructivement les pics d'intensité seraient encore plus intenses. Comme l'intensité moyenne est toujours donnée par la formule (3.38) (le second terme du membre de droite de (3.37) est de valeur moyenne nulle quelles que soient les hypothèses sur φ_k), les pics d'intensité seraient alors extrêmement brefs.

Pour traduire mathématiquement ces considérations qualitatives, étudions à présent ce qu'il advient lorsque les phases φ_k sont corrélées. Pour simplifier, nous les supposons toutes égales :

$$\varphi_k = \varphi$$

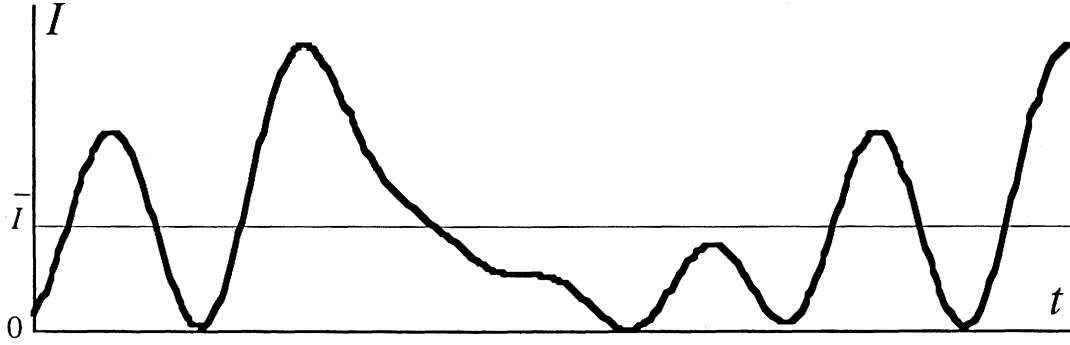


FIG. 3.23: **Variation de l'intensité d'un laser multimode en fonction du temps.** Les phases des divers modes sont des variables aléatoires indépendantes. L'intensité fluctue autour de sa valeur moyenne correspondant à la droite horizontale.

En réécrivant l'intensité lumineuse donnée en (3.37) sous la forme :

$$I(t) = \frac{1}{2} \left| \sum_{k=0}^{N-1} E_0 e^{-i(\omega_k t + \varphi_k)} \right|^2 \quad (3.40)$$

nous trouvons pour $\varphi_k = \varphi$:

$$I(t) = \frac{E_0^2}{2} \left| \sum_{k=0}^{N-1} e^{-i\omega_k t} \right|^2$$

soit en utilisant (3.36) :

$$I(t) = \frac{E_0^2}{2} \left| \sum_{k=0}^{N-1} e^{-ik\Delta t} \right|^2 \quad (3.41)$$

Nous en déduisons l'expression suivante pour l'intensité lumineuse :

$$I(t) = \frac{E_0^2}{2} \left| \frac{\sin \left(N \Delta \frac{t}{2} \right)}{\sin \left(\Delta \frac{t}{2} \right)} \right|^2 \quad (3.42)$$

La figure 3.24 montre l'allure de $I(t)$. Le laser émet des **impulsions lumineuses** dont la période de répétition T est :

$$\Rightarrow T = \frac{2\pi}{\Delta} = \frac{L_{\text{cav}}}{c} \quad (3.43)$$

c'est-à-dire le temps de rebouclage (ou d'aller-retour) dans la cavité. Ceci signifie qu'il n'y a qu'une seule *impulsion* circulant dans la cavité. Par ailleurs, l'intensité de chaque impulsion est égale à

$$I_{\text{max}} = \frac{N^2 E_0^2}{2} \quad (3.44)$$

ou encore en utilisant (3.38)

$$\Rightarrow I_{\max} = N\bar{I} \quad (3.45)$$

La largeur temporelle de chaque impulsion est

$$\Rightarrow \frac{2\pi}{N\Delta} = \frac{T}{N} \quad (3.46)$$

Les impulsions sont d'autant plus *intenses* et *plus courtes* que le nombre de modes oscillant N est *élevé*. Comme nous l'avons vu dans le § 3.1.2.d, ce nombre de modes est proportionnel à la *largeur spectrale* de la courbe de gain. Pour avoir des impulsions très courtes, il faut que cette largeur soit *très grande*.

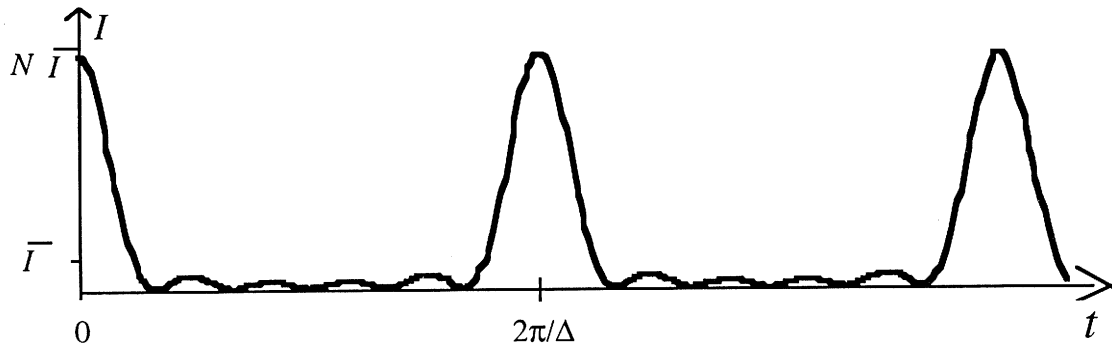


FIG. 3.24: **Intensité d'un laser à modes synchronisés en fonction du temps.** La courbe ci-dessus représente l'intensité résultant de l'addition cohérente des champs de N modes lasers de même amplitude et équidistants (séparation entre modes Δ).

Avec cette méthode, on peut obtenir à partir d'un laser à argon ionisé commercial des impulsions de durée 80 ps, répétées toutes les 12 ns (cf. le tableau du § 3.3.1); avec un laser à colorant ou un laser à saphir dopé au titane, des impulsions plus courtes que 100 fs sont aisément obtenues. Ce sont ces impulsions qui, amplifiées et convenablement traitées, permettent d'obtenir des impulsions de quelques femtosecondes.

Pour synchroniser les modes d'un laser, on peut par exemple *moduler les pertes*³⁰ de la cavité à la fréquence Δ . Considérons, en effet, un mode particulier de fréquence ω_k dont l'amplitude E est modulée à la fréquence Δ :

$$E = E_0[1 + m \cos(\Delta t)]$$

Le champ électrique correspondant s'écrit :

$$E(t) = E_0[1 + m \cos(\Delta t)] \cos(\omega_k t + \varphi_k)$$

soit encore en transformant cette expression pour faire apparaître les fréquences $(\omega_k + \Delta)$ et $(\omega_k - \Delta)$ des deux modes voisins du mode ω_k :

$$E(t) = E_0 \left[\cos(\omega_k t + \varphi_k) + \frac{m}{2} \cos((\omega_k + \Delta)t + \varphi_k) + \frac{m}{2} \cos((\omega_k - \Delta)t + \varphi_k) \right] \quad (3.47)$$

³⁰Pour cela on utilise, par exemple, un modulateur électrooptique ou acousto-optique.

La modulation du mode k crée des composantes aux fréquences de modes $k - 1$ et $k + 1$. Si certaines conditions sont réunies, ces composantes, qui sont en phase avec le mode k , vont *verrouiller la phase* des modes $k - 1$ et $k + 1$ sur celle du mode k . Le processus s'étend de proche en proche à tous les modes, qui sont ainsi synchronisés.

Remarques

(i) Une analyse *temporelle* peut être substituée à cette analyse spectrale. Considérons une impulsion circulant dans le laser : le gain pour cette impulsion sera maximum si elle passe dans le modulateur au moment où celui-ci a le minimum de pertes. La période de modulation doit donc être exactement égale à la période de rotation de l'impulsion dans la cavité donnée par la condition (3.43).

(ii) Le rôle du modulateur n'est pas totalement identique selon que l'amplificateur a une courbe de gain présentant un élargissement homogène ou inhomogène (voir § 3.3.1). Dans le cas d'un élargissement inhomogène, les divers modes oscillent naturellement et le seul but du modulateur est de coupler leurs phases. Dans le cas d'un élargissement homogène, le modulateur provoque, en outre, l'apparition de modes qui, en son absence, n'oscilleraient pas spontanément.

Il est également possible de synchroniser les modes en plaçant dans la cavité un *absorbant saturable*. Celui-ci est un milieu dont l'absorption diminue (et la transmission augmente) lorsque l'intensité augmente. À forte intensité, il est transparent³¹. Dans un laser multimode où l'on insère un absorbant saturable, les pertes de celui-ci *empêchent toute oscillation incorrélée* et seules les fortes impulsions, associées à une synchronisation des modes, peuvent circuler dans la cavité.

3.4.2 Laser déclenché

La figure 3.18 montre l'évolution temporelle de l'intensité lumineuse émise par un laser à rubis lorsque le cristal est excité par une lampe flash. On a un train d'impulsions de largeur $0,1\mu\text{s}$, séparées de quelques microsecondes. Pour une énergie typique de 1 Joule dans le train d'impulsions, on voit que la puissance moyenne est de l'ordre du kilowatt mais que la puissance crête est d'un ordre de grandeur plus élevé. Ce fonctionnement s'interprète comme une oscillation de relaxation (voir § 3.2.4), la dynamique de ces oscillations dépendant des divers temps caractéristiques du problème : temps de relaxation atomique, temps d'amortissement de la cavité. Dans le cas d'un système à trois niveaux, un tel fonctionnement est très fréquent puisque le niveau inférieur de la transition laser (niveau a) ne peut pas se vider vers un autre niveau.

Remarque

Les systèmes à quatre niveaux atteignent au contraire en général un régime de fonctionnement stationnaire. Cependant, certains de ces lasers peuvent aussi présenter des oscillations. Une des raisons fondamentales à l'existence de ces oscillations est que le laser est un système non linéaire parce que les équations couplant les atomes au rayonnement dans la cavité sont des équations différentielles non linéaires. Or, il est connu que de telles équations peuvent conduire à des instabilités voire à un chaos déterministe. Dans le cadre d'un modèle

³¹Un ensemble d'atomes à deux niveaux possède cette propriété (voir chapitre II § D.3) et peut donc être utilisé comme modèle élémentaire d'absorbant saturable.

simplifié³², il est possible de montrer que ces instabilités apparaissent lorsque le temps de relaxation de l'énergie dans la cavité est inférieur aux temps de relaxation atomique (cette condition correspond à ce qui est appelé une « mauvaise cavité »).

Une émission laser relaxée sous forme d'un train d'impulsions séparée par des intervalles de temps plus ou moins aléatoires pose des problèmes d'utilisation car on souhaite souvent disposer d'une impulsion unique. On peut insérer dans la cavité laser un dispositif électro-optique rapide permettant d'extraire sélectivement une seule impulsion. Il est beaucoup plus astucieux d'utiliser un système obturateur rapide (« Q switch » ou commutation des pertes) qui, fermé, *empêche le laser d'osciller pendant la phase d'augmentation de l'inversion de population*, puis qui s'ouvre brutalement : *toute l'énergie accumulée peut alors être émise en une seule impulsion géante* (figure 3.25). La puissance crête est accrue d'un à deux ordres de grandeur et des puissances bien supérieures au mégawatt peuvent ainsi être obtenues, dans ce fonctionnement appelé « déclenché ».

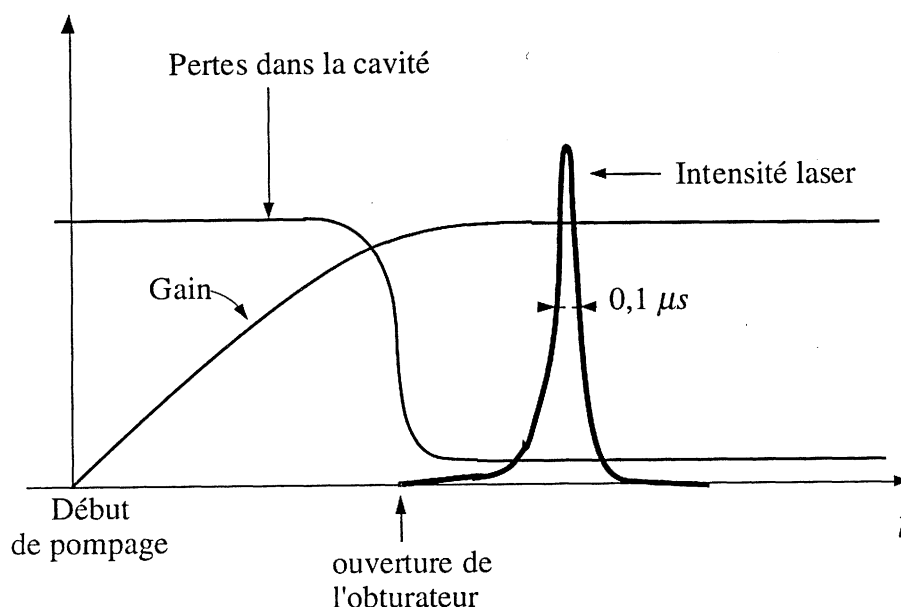


FIG. 3.25: **Fonctionnement d'un laser déclenché.** Les pertes dans la cavité sont maintenues à une valeur élevée pendant la phase de pompage, puis elles sont brutalement diminuées (« Q switch »). Toute l'énergie est alors délivrée en une seule impulsion. La puissance crête peut atteindre $10^7 W$.

Les dispositifs de déclenchement utilisés en pratique sont très variés. Les premiers utilisaient simplement un miroir tournant à l'une des deux extrémités de la cavité laser, qui n'était donc fermée que pendant une durée très courte. Cette méthode a été presque complètement abandonnée à cause de problèmes liés aux vibrations du dispositif tournant et aux difficultés de synchronisation entre la lampe flash et la rotation du miroir. On utilise plutôt aujourd'hui des obturateurs électro-optiques (ouverture en 10 nanosecondes) ou acousto-optiques (ouverture en 1 microseconde).

³²Voir N.B. Abraham, P. Mandel and L.M. Narducci, Prog. Optics **XXV**, 1, 1988 et références incluses.

Une autre méthode mérite une mention particulière, c'est le déclenchement passif induit par la présence d'un absorbant saturable dans la cavité. Tant que l'inversion de population est faible, *l'intensité circulant dans la cavité n'est pas suffisante pour compenser les pertes de l'absorbant saturable*. En revanche, lorsque l'inversion de population atteint son maximum, le démarrage de l'oscillation laser s'accompagne d'une forte décroissance des pertes dues à l'absorbant saturable et une impulsion géante est alors émise. Dans ce cas, c'est l'amorçage spontané du laser qui provoque « l'ouverture » de « l'obturateur » constitué par l'absorbant saturable. Il n'y a plus de problèmes de synchronisation.

Les lasers déclenchés à amplificateur solide (rubis ou néodyme) sont utilisés couramment dans deux types d'applications : (i) les situations où on a besoin de puissances crêtes élevées avec une puissance moyenne modérée (certains effets d'optique non-linéaire) ; (ii) les situations où on a besoin d'une impulsion unique, par exemple les mesures de distance d'objets susceptibles de diffuser la lumière laser (dispositif LIDAR ou « radar lumineux ») ; on peut ainsi déterminer la distance terre-lune à 10 centimètres près, ou évaluer la dérive des continents, par des mesures de temps d'aller-retour de la lumière via un satellite (voir § B.2 du complément III.4).

3.5 Spécificité de la lumière laser

Au début des années 2000, un laser à argon ionisé de 15 watts coûtait environ 50.000 euros et une lampe à incandescence de 150 watts un euro. Quel est donc l'avantage de la lumière laser ? C'est que l'énergie de la lumière laser peut être **concentrée sur des domaines extrêmement étroits de longueur d'onde, d'espace, de temps**. On obtient ainsi des densités d'énergie incomparablement plus élevées qu'à partir d'une source ordinaire.

Pour comprendre pleinement la différence entre la lumière laser et la lumière émise par une source classique incohérente (lampe à incandescence, soleil, lampe à décharge) il faut revenir aux lois de la photométrie classique (voir le complément III.7). Loin d'être circonstancielles, ces *lois qui s'appliquent à toute source incohérente*, et en particulier au rayonnement thermique, *découlent des deux principes de la thermodynamique*. Le résultat essentiel – relié à la non décroissance de l'entropie d'un système isolé – est que la luminance d'un faisceau ne peut augmenter lors de la propagation à travers un instrument d'optique. En conséquence, à partir d'une source dont la luminance L est donnée (il s'agit de la puissance émise par unité de surface et d'angle solide), on ne pourra en aucun cas – quel que soit l'instrument d'optique utilisé – obtenir une puissance par unité de surface E (l'éclairement énergétique) supérieure à πL :

$$E \leq \pi L . \quad (3.48)$$

Sachant qu'un filament de lampe à incandescence à 3000 K possède une luminance d'environ $150 \text{ W cm}^{-2} \text{ sr}^{-1}$, on ne pourra obtenir au mieux qu'un éclairement de l'ordre de 500 W/cm^2 , et cela à condition d'utiliser une optique de très grande ouverture sans aberration – par exemple un miroir elliptique. Avec une lampe à décharge de type arc à haute pression, ou à partir du soleil, on peut au mieux obtenir dix fois plus, soit 5 kW/cm^2 .

Considérons maintenant un faisceau laser de 15 Watts : étant spatialement cohérent, il peut, comme nous le montrons dans le complément III.2, être focalisé sur une tache de diffraction élémentaire, dont la dimension est de l'ordre de la longueur d'onde, soit une surface inférieure à $1\mu\text{m}^2$. Nous disposons donc maintenant de 10^9 W/cm^2 , soit un éclairement énergétique 10^5 à 10^6 fois plus élevé qu'avec la plus intense des sources classiques. Cette propriété est à la base de nombreuses applications des lasers où on recherche des densités d'énergie élevées (complément III.3).

Une autre façon de comprendre le résultat ci-dessus est de remarquer que les 15 Watts du faisceau laser sont émis dans un faisceau très peu divergent (10^{-3} radians pour un diamètre de 1 millimètre), à la différence de la source classique incohérente qui rayonne dans un demi-espace (2π steradians). Cette propriété est mise à profit lorsque l'on cherche à matérialiser une droite, par exemple pour le guidage des engins de chantier, en particulier des tunneliers. On peut réduire encore la divergence en faisant passer le faisceau laser dans un télescope, en sens inverse du sens habituel (de l'oculaire vers le miroir de sortie). Pour un miroir de un mètre de diamètre, la divergence n'est plus que de 10^{-6} radians³³.

On peut en fait exploiter la cohérence suivant la troisième dimension – longitudinale – du faisceau, par exemple en réalisant des impulsions très brèves. La puissance instantanée maximale peut alors être phénoménale. Ainsi, un laser à 10^5 modes verrouillés en phase donne une puissance crête 10^5 fois plus élevée que la puissance moyenne, soit après focalisation 10^{14} W/cm^2 . Par passage d'une seule impulsion dans un amplificateur laser on peut encore gagner plusieurs ordres de grandeur, avec des systèmes accessibles à de petits laboratoires. On atteint aussi le régime où le champ électrique de l'onde lumineuse est supérieur au champ de Coulomb exercé par le noyau de l'atome d'hydrogène dans son état fondamental. On accède ainsi à un nouveau régime de l'interaction entre lumière et matière avec des systèmes de taille modeste, tenant sur une table quelques dizaines de mètres carrés³⁴.

Ici encore, on peut choisir de concentrer non pas dans le temps mais dans sa grandeur conjuguée, la fréquence. On peut ainsi délivrer plusieurs Watts de lumière laser parfaitement monochromatique, avec une largeur de raie inférieure au kiloHertz. L'éclairement (continu) par unité de bande spectrale peut alors atteindre $10^6\text{ W cm}^{-2}\text{ Hz}^{-1}$, valeur dont on ne réalise l'énormité qu'en la comparant à ce qui peut être obtenu à partir du soleil ou même d'une lampe à arc, toujours inférieur à $10^{-10}\text{ W cm}^{-2}\text{ Hz}^{-1}$ (le spectre de la lumière solaire s'étend sur près de 10^{15} Hz).

Ainsi, ce qu'apporte la lumière laser, c'est la **possibilité d'être concentrée, grâce**

³³Le faisceau correspondant a un rayon de seulement 300 mètres lorsqu'il est projeté sur la lune : l'éclairement est assez intense pour que l'on puisse capter sur terre les quelques photons réfléchis par un système catadioptrique déposé par une mission Apollo, et mesurer ainsi la distance terre-lune, par détermination du temps d'aller-retour de la lumière.

³⁴Les très grandes installations laser destinées à créer des milieux thermonucléaires ne cherchent pas à obtenir des densités plus élevées mais elles visent des durées d'impulsion plus longues (des microsecondes au lieu de picosecondes ou femtosecondes) et des volumes d'irradiation plus grands (des millimètres cube au lieu de micromètres cube).

à ses propriétés de cohérence : cohérence temporelle qui permet une concentration d'énergie dans le temps ou en fréquence ; cohérence spatiale permettant soit de focaliser très fortement l'énergie, soit de réaliser des faisceaux extraordinairement bien collimatés. Ce sont ces propriétés qui sont exploitées dans les applications des lasers.

Nous avons déjà indiqué que la différence entre lumière classique et lumière laser est de nature fondamentale, et que la thermodynamique en donne une explication profonde. On peut aller plus loin et analyser cette différence dans le cadre de la physique statistique des photons. Il est alors frappant de constater (voir le complément III.7), que les photons émis par une source classique sont distribués dans un très grand nombre de cellules élémentaires de l'espace des phases – ou modes du champ électromagnétique – et que le nombre moyen de photon par mode est très inférieur à 1. Au contraire, tous les photons d'un laser étant émis dans le même mode (ou dans un petit nombre de modes) du champ électromagnétique, le nombre de photons par mode peut atteindre des valeurs considérables (10^{10} photons par mode est courant)³⁵.

Ainsi, au niveau le plus fondamental, on peut dire que ce qui caractérise la lumière laser – comparée à la lumière ordinaire, c'est que **le nombre de photons par mode est très supérieur à 1**.

³⁵Cette possibilité d'accumuler tous les photons dans un seul mode est évidemment liée à leur nature bosonique, et on peut dans une certaine mesure considérer un faisceau de lumière laser comme un condensat de Bose-Einstein de photons à température nulle. On comprend alors qu'il soit possible d'utiliser la lumière laser pour refroidir des atomes à des températures très basses, proches du zéro absolu (chapitre 7).

Complément III.1

Cavité résonnante Fabry-Perot

Les interféromètres de Fabry et Perot ont une longue et riche histoire. Inventés pour réaliser des mesures spectroscopiques de haute précision, ils se sont révélés être un outil essentiel de la physique des lasers. Un laser n'est souvent, en effet, qu'un milieu amplificateur enfermé dans une cavité résonnante Fabry-Perot. Cette cavité assure d'une part le rebouclage de la lumière issue de l'amplificateur sur celui-ci, et elle permet d'autre part de réaliser une sélection de longueur d'onde dans la courbe de gain de l'amplificateur (modes de la cavité). Le rôle des cavités Fabry-Perot ne s'arrête cependant pas là : on peut en insérer à l'intérieur de la cavité laser où elles permettent de filtrer une gamme spectrale déterminée dans la courbe de gain. Dans la cavité de certains lasers accordables, il est fréquent de trouver en série plusieurs filtres de ce type, de largeurs spectrales de plus en plus étroites. Enfin, on trouve aussi à l'extérieur du laser des cavités Fabry-Perot qui servent de référence de longueur d'onde pour l'oscillateur laser.

1. Position du problème. Cavité linéaire

Considérons deux miroirs plans M_1 et M_2 sans absorption, partiellement réfléchissants (leurs coefficients de réflexion et de transmission en amplitude sont respectivement r_1, r_2 et t_1, t_2) parallèles entre eux. L'ensemble de ces deux miroirs constitue, comme nous le verrons, une cavité *résonnante* pour la lumière. Une onde électromagnétique $\mathbf{E}_i(\mathbf{r}, t)$ de polarisation $\boldsymbol{\varepsilon}$ se propage suivant la direction Oz perpendiculaire aux miroirs :

$$\mathbf{E}_i(\mathbf{r}, t) = E_i \boldsymbol{\varepsilon} \cos(\omega t - kz + \varphi_i) = \boldsymbol{\varepsilon} \operatorname{Re} [\mathcal{E}_i e^{i(kz - \omega t)}] \quad (1)$$

E_i et $\mathcal{E}_i$ sont respectivement les amplitudes réelle et complexe du champ électrique incident. On notera de la même façon les amplitudes des champs réfléchis (E_r), transmis (E_t) et se propageant entre les deux miroirs (E_c et E'_c). Ces diverses ondes sont représentées sur la figure 1.

Plutôt que l'amplitude des champs, on utilisera souvent les flux lumineux Π par unité de surface qui sont égaux à la valeur moyenne du vecteur de Poynting. Le flux incident

est, par exemple, égal à (voir Éq. II.C.8.c) :

$$\Pi_i = \varepsilon_0 c \frac{E_i^2}{2} = \varepsilon_0 c \frac{|\mathcal{E}_i|^2}{2} \quad (2)$$

On se propose de calculer les coefficients de transmission T et de réflexion R de la cavité pour ces flux¹ :

$$T = \frac{\Pi_t}{\Pi_i} = \left| \frac{E_t}{E_i} \right|^2$$

$$R = \frac{\Pi_r}{\Pi_i} = \left| \frac{E_r}{E_i} \right|^2 \quad (3)$$

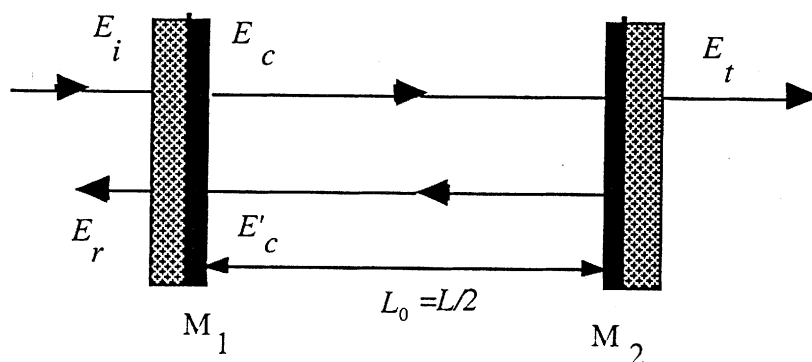


FIG. 1: Les amplitudes des champs électriques (ondes polarisées linéairement) E_i, E_r, E_c et E'_c sont considérés en M_1 et E_t en M_2 .

En choisissant convenablement les origines des phases, on peut écrire les relations entre champs incident et émergent sur le miroir M_1 sous la forme :

$$\mathcal{E}_c = t_1 \mathcal{E}_i - r_1 \mathcal{E}'_c$$

$$\mathcal{E}_r = r_1 \mathcal{E}_i + t_1 \mathcal{E}'_c \quad (4)$$

avec r_1 et t_1 réels et obéissant à la relation :

$$r_1^2 + t_1^2 = R_1 + T_1 = 1 \quad (5)$$

Remarquons que le signe $-$ de la première relation (4) garantit que la matrice de transformation des champs est unitaire, ce qui assure la conservation des flux d'énergie, c'est-à-dire : $E_c^2 + E_r^2 = E_i^2 + E'_c{}^2$ (ou encore $|\mathcal{E}_c|^2 + |\mathcal{E}_r|^2 = |\mathcal{E}_i|^2 + |\mathcal{E}'_c|^2$).

¹Pour éviter d'alourdir le texte, nous omettons de préciser par « unité de surface » quand il n'y a pas ambiguïté

De même, les relations entre les champs s'écrivent au niveau du miroir M_2 :

$$\begin{aligned}\mathcal{E}_t &= t_2 \mathcal{E}_c \exp(ikL_0) \\ \mathcal{E}'_c &= -r_2 \mathcal{E}_c \exp(2ikL_0)\end{aligned}\tag{6}$$

avec r_2 et t_2 réels et tels que

$$r_2^2 + t_2^2 = R_2 + T_2 = 1\tag{7}$$

Ici encore, la conservation des flux est assurée : $E_t^2 + E_c'^2 = E_c^2$ (ou $|\mathcal{E}_t|^2 + |\mathcal{E}'_c|^2 = |\mathcal{E}_c|^2$). Dans la formule (6), L_0 est la distance entre les deux miroirs.

Remarque

Le signe de r_2 dans la seconde équation (6) n'est pas arbitraire. Ce signe, opposé à celui considéré dans l'équation (4) pour le champ se propageant dans la même direction Oz , est lié au fait que la face réfléchissante des miroirs est toujours positionnée vers *l'intérieur* de la cavité.

2. Facteur de transmission et de réflexion de la Cavité. Résonances

Les équations (4) et (6) donnent immédiatement, en introduisant la longueur de rebouclage $L = 2L_0$

$$\frac{\mathcal{E}_t}{\mathcal{E}_i} = \frac{t_1 t_2 \exp(ikL_0)}{1 - r_1 r_2 \exp(ikL)}\tag{8}$$

d'où l'on déduit

$$T = \left| \frac{\mathcal{E}_t}{\mathcal{E}_i} \right|^2 = \frac{T_1 T_2}{1 + R_1 R_2 - 2\sqrt{R_1 R_2} \cos(kL)}\tag{9}$$

La formule (9) montre que T passe par un maximum chaque fois que $kL = 2p\pi$ (voir figure 2), c'est-à-dire pour :

$$\omega = \omega_p = p 2\pi \frac{c}{L} \quad \text{avec } p \text{ entier}\tag{10}$$

Le facteur de transmission vaut alors

$$T_{\max} = \frac{T_1 T_2}{(1 - \sqrt{R_1 R_2})^2}\tag{11}$$

Si $R_1 = R_2$, nous trouvons $T_{\max} = 1$, même si $T_1 (= T_2)$ est arbitrairement proche de zéro. *La cavité transmet alors la totalité du champ incident.* Dans la pratique, les imperfections des miroirs limiteront la valeur de $T_{\max}$.

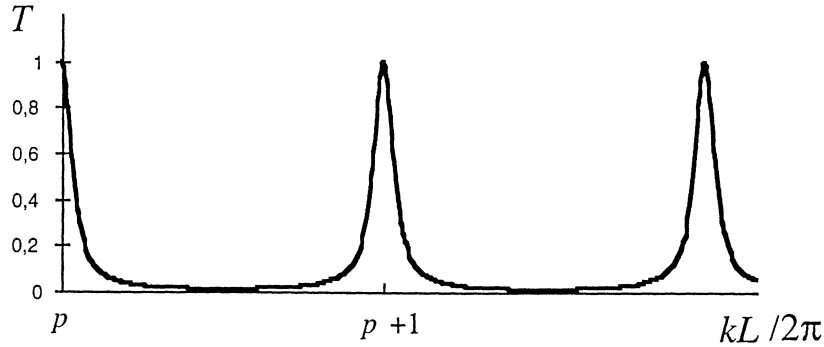


FIG. 2: Transmission de la cavité. Des résonances très marquées apparaissent lorsque le déphasage au bout d'un aller-retour kL est égal à un nombre entier de fois 2π ; (la courbe ci-dessus est obtenue pour $R_1 = R_2 = 0,8$).

Les fréquences ω_p définies par l'équation (10) correspondent à des résonances de la cavité. En effet, puisque $|\mathcal{E}_1|^2 = T_2 |\mathcal{E}_c|^2$, nous trouvons en utilisant (9) :

$$\left| \frac{\mathcal{E}_c}{\mathcal{E}_i} \right|^2 = \frac{T_1}{1 + R_1 R_2 - 2\sqrt{R_1 R_2} \cos(kL)} \quad (12)$$

ce qui montre que pour un champ incident E_i donné, la puissance circulant dans la cavité est maximale pour $\omega = \omega_p$ (c'est-à-dire $kL = 2p\pi$). Par exemple, pour une cavité symétrique, l'intensité circulant dans la cavité est alors plus grande que l'intensité incidente par le facteur $1/T_2$. Ce facteur peut être très grand devant 1 si T_2 est petit.

Notons enfin que le coefficient de réflexion de la cavité peut être déduit de la valeur (9) de T ²

$$R = \left| \frac{\mathcal{E}_r}{\mathcal{E}_i} \right|^2 = 1 - T \quad (13)$$

Il passe par un *minimum à résonance*. En particulier, dans le cas d'une cavité symétrique ($R_1 = R_2$), nous trouvons $R = 0$: le coefficient de réflexion passe par 0 pour les fréquences de résonance.

Remarque

(i) Pour trouver la relation (9) on peut, au lieu de simplement écrire les relations de passage globales (4) et (6) sur les deux miroirs, écrire les séries d'ondes transmises et réfléchies, ce qui met l'accent sur le fait qu'il s'agit d'un phénomène d'interférence à ondes multiples. Le traitement présenté ici est plus simple et tout aussi rigoureux.

(ii) De nombreuses expériences de spectroscopie ou d'optique non-linéaire comme les transitions à deux photons (Complément III.5) ou le doublage de fréquence (Complément III.7)

²On peut aussi l'obtenir directement à partir des équations (4) et (6).

sont réalisées, avec des lasers continus de puissance modérée, à l'intérieur d'une cavité Fabry-Perot parce que l'intensité lumineuse peut y être bien plus grande que l'intensité incidente. Selon les applications, il pourra être plus habile d'utiliser une cavité en anneau (pour n'avoir qu'une onde progressive) ou une cavité linéaire.

(iii) Dans ce complément, on suppose implicitement que l'espace entre les deux miroirs est vide de sorte qu'il y a coïncidence entre longueur optique et longueur géométrique pour la propagation de la lumière entre les deux miroirs. En fait, les résultats ne dépendent que du déphasage. Si la cavité est remplie avec un milieu d'indice n , le déphasage dû à la propagation est égal à $(n\omega/c)L_0$ (où L_0 est la longueur géométrique de la cavité). On pourra écrire ce déphasage sous la forme $(\omega/c)(nL_0)$ où l'on a fait apparaître la longueur *optique* nL_0 séparant les deux miroirs. De manière générale, les résultats de ce complément peuvent être généralisés en considérant L comme la longueur *optique* de rebouclage et en posant $k = \omega/c$.

3. Cavité en anneau à un seul miroir de couplage

La plupart des résultats démontrés pour une cavité linéaire peuvent être étendus au cas d'une cavité en anneau. Considérons l'anneau représenté sur la figure 3. Il consiste en un miroir M partiellement transparent (coefficients r_1, t_1) et deux miroirs totalement réfléchissants.

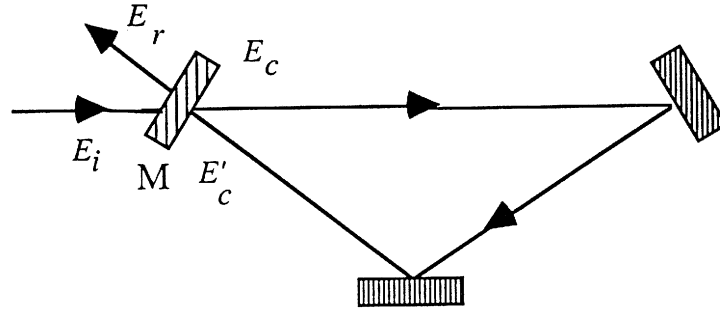


FIG. 3: Cavité de Fabry-Perot en anneau. Les champs E_i, E_r, E_c et E'_c sont évalués sur le miroir M partiellement transparent. La longueur de l'anneau est L .

Les relations entre le champ incident E_i , le champ réfléchi E_r et le champ dans la cavité E_c sont analogues aux équations (4) :

$$\mathcal{E}_c = t_1 \mathcal{E}_i + r_1 \mathcal{E}'_c \quad (14.a)$$

$$\mathcal{E}_r = r_1 \mathcal{E}_i + t_1 \mathcal{E}'_c \quad (14.b)$$

$$\mathcal{E}'_c = \mathcal{E}_c \exp(ikL) \quad (14.c)$$

où L est la longueur de l'anneau. En combinant la première et la dernière des équations (14), nous trouvons :

$$\mathcal{E}_c = \frac{t_1 \mathcal{E}_i}{1 - r_1 \exp(ikL)} \quad (15)$$

soit :

$$\left| \frac{\mathcal{E}_c}{\mathcal{E}_i} \right|^2 = \frac{T_1}{1 + R_1 - 2\sqrt{R_1} \cos(kL)} \quad (16)$$

L'équation (6) coïncide avec l'équation (12) lorsque l'on pose $R_2 = 1$. La valeur de l'intensité à l'intérieur de la cavité en anneau à un seul miroir de couplage est donc *identique* à celle obtenue dans une cavité *linéaire* pour lequel *un des miroirs a un coefficient de réflexion égal à 100%*. Les résultats établis pour les cavités linéaires peuvent donc être facilement adaptés aux cavités en anneau en faisant $R_2 = 1$. En particulier, les phénomènes de résonance sont identiques et se produisent pour les mêmes valeurs de la pulsation, données par l'équation (10).

En ce qui concerne le champ réfléchi, nous trouvons en combinant les équations (14.b) et (14.c) avec (15) :

$$\frac{\mathcal{E}_r}{\mathcal{E}_i} = \frac{1 - r_1 \exp(-ikL)}{1 - r_1 \exp(ikL)} \exp(ikL) \quad (17)$$

ce qui implique que le *coefficient de réflexion* $|\mathcal{E}_r|^2/|\mathcal{E}_i|^2$ *est égal à 1* quelque soit la longueur de la cavité. Dans la cavité en anneau de la figure 2, toute l'intensité incidente se retrouve nécessairement (en l'absence de pertes dans la cavité) dans le faisceau réfléchi. Ceci ne veut pas dire pour autant que les résonances de la cavité ne se manifestent pas sur le champ réfléchi mais les *résonances apparaissent uniquement sur sa phase*.

4. Finesse. Facteur de qualité

Il est intéressant (et traditionnel) de transformer le dénominateur résonnant D des formules (9) et (12), de la façon suivante :

$$\begin{aligned} D &= 1 + R_1 R_2 - 2\sqrt{R_1 R_2} \cos(kL) = 1 + R_1 R_2 - 2\sqrt{R_1 R_2} + 4\sqrt{R_1 R_2} \sin^2\left(k\frac{L}{2}\right) \\ &= (1 - \sqrt{R_1 R_2})^2 \left(1 + m \sin^2\left(k\frac{L}{2}\right)\right) \end{aligned} \quad (18)$$

où nous avons introduit le paramètre m :

$$m = \frac{4\sqrt{R_1 R_2}}{(1 - \sqrt{R_1 R_2})^2} \quad (19)$$

La valeur maximale de $1/D$ est d'autant plus grande que $R_1 R_2$ est plus proche de 1. Par ailleurs, $1/D$ varie avec le déphasage kL et est maximum lorsque la condition de résonance (10) est vérifiée. La largeur de la résonance peut-être caractérisée en cherchant les valeurs de kL telles que $1/D$ soit égal à la moitié de sa valeur maximum : on trouve ainsi pour $m \gg 1$

$$k\frac{L}{2} \approx p\pi \pm \frac{1}{\sqrt{m}}$$

On définit la « finesse » $\mathcal{F}$ comme le rapport entre l'intervalle de fréquence entre deux pics et la largeur des résonances :

$$\mathcal{F} = \pi \frac{\sqrt{m}}{2} = \pi \frac{(R_1 R_2)^{1/4}}{1 - \sqrt{R_1 R_2}} \quad (20)$$

Si on balaie la fréquence de l'onde incidente, les pics de résonance ont une largeur

$$\frac{\Delta\omega_{\text{cav}}}{2\pi} = \frac{1}{\mathcal{F}} \frac{c}{L} \quad (21)$$

Il est clair que le système présente des résonances d'autant plus fines que $\mathcal{F}$ est plus grand.

Une façon habituelle de caractériser un système résonnant en électricité ou en mécanique est d'étudier la façon dont il s'amortit lorsque l'excitation cesse. L'énergie stockée dans une cavité linéaire de surface transverse S est :

$$W = \frac{L_0}{c} (\Pi_c + \Pi'_c) S \quad (22)$$

où Π_c et Π'_c sont les flux circulants dans la cavité linéaire respectivement dans la direction des z croissants et décroissants (voir figure 1). La puissance perdue aux extrémités est

$$-\frac{d}{dt}W = (T_2\Pi_c + T_1\Pi'_c)S \quad (23.a)$$

En tenant compte de $\Pi'_c = R_2\Pi_c$ on obtient

$$-\frac{d}{dt}W = \frac{T_2 + T_1 R_2}{1 + R_2} \frac{c}{L_0} W \quad (23.b)$$

Pour un système résonnant de fréquence propre ω , on définit le *facteur de qualité* Q par

$$\frac{dW}{dt} = -\frac{\omega}{Q} W \quad (24)$$

En comparant les équations (23.a) et (24), et en utilisant $\omega = 2\pi c/\lambda$, nous trouvons :

$$Q = \pi \frac{1 + R_2}{T_2 + T_1 R_2} \frac{L}{\lambda} \quad (25)$$

À cause du facteur L/λ qui dépasse couramment 10^5 pour des longueurs d'onde optiques et pour R_1 et R_2 proche de 1, le champ intracavité effectue de nombreuses oscillations avant de s'amortir.

5. Exemple. Cavité de grande finesse

Il est intéressant de considérer le cas où les miroirs ont des coefficients de réflexion élevés. Des développements limités en T_1 et T_2 sont alors possibles. Ils conduisent à :

$$\mathcal{F} \approx \frac{2\pi}{T_1 + T_2} \quad (26.a)$$

$$Q = \mathcal{F} \frac{L}{\lambda} \quad (26.b)$$

On sait couramment faire des miroirs sans absorption de coefficient de transmission de l'ordre de 1%, ce qui donne des finesesses supérieures à 100. Il est possible d'atteindre des valeurs supérieures à 10^5 avec des couches diélectriques de très haute qualité.

Dans cette limite (R_1 et R_2 voisins de 1), l'équation (12) donnant l'intensité du champ dans la cavité s'écrit en utilisant (18), (19) et (20)

$$\left| \frac{\mathcal{E}_c}{\mathcal{E}_i} \right|^2 = \frac{4T_1}{(T_1 + T_2)^2 \left(1 + \frac{4\mathcal{F}^2}{\pi^2} \sin^2 \left(\frac{kL}{2} \right) \right)} \quad (27)$$

Donnons pour conclure quelques chiffres, relatifs à une cavité linéaire de longueur $L_0 = 5$ cm, et de finesse $\mathcal{F} = 300$ (obtenue pour $T_1 = T_2 \approx 1\%$) :

- l'intervalle spectral libre (écart entre deux modes successifs) vaut $c/L = 3$ GHz
- la largeur d'une résonance vaut $\Delta\omega_{\text{cav}}/2\pi = 10$ MHz
- dans le visible ($\lambda = 0,5\mu\text{m}$), le facteur de qualité vaut $Q = 6 \times 10^7$.

Une telle cavité peut être utilisée comme un spectromètre de très haute résolution : si on modifie sa longueur L_0 , les fréquences susceptibles d'être transmises changent. En modifiant L_0 de $\lambda/2$, on balaie un intervalle spectral libre de 3 GHz : compte-tenu de la petite valeur de $\Delta\omega_{\text{cav}}$, on peut ainsi analyser la structure de la raie émise par un laser Hélium-Néon (3 modes séparés de 500 MHz).

Remarque

(i) Les pertes par transmission ne sont pas les seules pertes d'une cavité Fabry-Perot. Il faut aussi tenir compte des pertes des miroirs et éventuellement des pertes dues à l'absorption ou à la diffusion lorsque un milieu matériel est placé à l'intérieur de la cavité. Considérons la cavité linéaire de la Figure 1 et supposons que les pertes en intensité pour une propagation d'un miroir à l'autre soient égales à A_1 de sorte que les équations (6) deviennent :

$$\mathcal{E}_t = t_2 \mathcal{E}_c \sqrt{1 - A_1} \exp(ikL_0) \quad (28.a)$$

$$\mathcal{E}'_c = -r_2 \mathcal{E}_c (1 - A_1) \exp(2ikL_0) \quad (28.b)$$

alors que les équations (4) restent inchangées. Les développements précédents peuvent être facilement adaptés pour inclure ce terme supplémentaire. On trouve par exemple que dans la limite $T_1, T_2, A_1 \ll 1$, les coefficients de transmission et de réflexion à résonance et la finesse sont respectivement égaux à :

$$T_{\text{max}} = \frac{4T_1T_2}{(4 + T_1 + T_2)^2} \quad (29.a)$$

$$R_{\text{res}} = \frac{(4 - T_1/T_2)^2}{(A + T_1 + T_2)^2} \quad (29.b)$$

$$\mathcal{F} = \frac{2\pi}{A + T_1 + T_2} \quad (29.c)$$

où $A = 2A_1$ correspond aux pertes sur un aller-retour.

(ii) Le cas d'un laser (milieu amplificateur) correspond à $A < 0$. Il est alors possible d'obtenir $T_{\text{max}} = \infty$ ce qui signifie qu'une onde peut être émise par le laser en absence d'onde incidente.

6. Cavit  laser lin aire

L'analyse pr c dente peut- tre appliqu e   une cavit  laser lin aire, en faisant $T_2 = 0$ ($R_2 = 1$), ce qui correspondrait   un faisceau laser  mis dans la direction des z d croissants sur la figure 1. Dans ce cas, la transmission est bien s r nulle pour la cavit  ($T = 0$), mais on peut v rifier que les formules (12) et suivantes restent valables. En particulier, la r sonance pour la puissance circulant dans la cavit  subsiste. La finesse vaut (cas o  $T_1 \ll 1$, et $T_2 = 0$) :

$$\mathcal{F} = \frac{2\pi}{T_1} \quad (30)$$

Les flux circulant dans la cavit  Π_c et Π'_c et le flux incident sont reli s par

$$\Pi_c = \Pi'_c \quad (31.a)$$

$$\frac{\Pi_c}{\Pi_i} = \frac{2\mathcal{F}}{\pi \left(1 + \frac{4\mathcal{F}^2}{\pi^2} \sin^2\left(\frac{kL}{2}\right)\right)} \quad (31.b)$$

L' nergie stock e dans la cavit  est, d'apr s les  quations (22) et (31.a),  gale   :

$$W = \frac{L}{c} \Pi_c S \quad (32)$$

soit en utilisant les  quations (21) et (31.a) :

$$W = \frac{4\Pi_i}{\Delta\omega_{\text{cav}} \left(1 + \frac{4\mathcal{F}^2}{\pi^2} \sin^2\left(\frac{kL}{2}\right)\right)} S \quad (33)$$

ce qui montre qu'  r sonance, l' nergie dans la cavit  est d'autant plus grande que la largeur est  troite,   flux incident constant.

Remarque

(i) Dans le cas d'un laser, la situation est diff rente dans la mesure o  il n'y a pas de faisceau incident ($\Pi_i = 0$), mais il y a cependant de l' nergie dans la cavit  ($\Pi_c \neq 0$)   cause de la pr sence du milieu amplificateur. Pour cette situation, l' quation (32) reste valable et le flux lumineux Π_s  mis par le laser est  gal  

$$\Pi_s = T_1 \Pi_c \quad (34)$$

soit en utilisant les  quations (21), (31) et (32) :

$$W = \frac{\Pi_s}{\Delta\omega_{\text{cav}} S} \quad (35)$$

formule qui permet de relier la puissance de sortie   l' nergie dans la cavit .

(ii) Les r sultats d montr s dans ce paragraphe 6, sont imm diatement applicables au cas d'une cavit  en anneau   cause de la similitude entre une cavit  lin aire dont le miroir M_2 a un coefficient de r flexion  gal   1 et une cavit  en anneau (voir paragraphe 3).

Complément III.2

Faisceaux gaussiens. Modes transverses d'un laser

La description élémentaire du fonctionnement d'un laser, présentée dans la partie A du chapitre III, repose sur l'hypothèse implicite que l'onde lumineuse circulant dans la cavité laser est une onde plane homogène, illimitée dans la direction perpendiculaire à sa direction de propagation. Cette hypothèse est manifestement irréaliste : les divers composants d'un laser (miroirs, milieu amplificateur) ont des dimensions transverses limitées, souvent petites (de l'ordre du centimètre). Si l'onde laser était réellement plane et uniforme, la diffraction au niveau de l'un de ces composants suffirait à provoquer une divergence de l'onde l'empêchant de revenir identique à elle-même au bout d'un tour de cavité. On conçoit qu'il puisse être possible de compenser cette divergence, qui provoque des pertes, en utilisant des miroirs convergents, mais le traitement théorique par ondes planes est alors inadapté.

Une description plus correcte du champ dans la cavité consiste à partir d'une distribution transversale non uniforme pour l'onde et à se poser le problème de la stabilité d'une telle distribution lorsque le faisceau s'est propagé sur un tour de cavité. La description de cette propagation doit naturellement prendre en compte les phénomènes de diffraction, ainsi que la réflexion sur les miroirs. Une telle structure stable, si elle existe, s'appelle un *mode transverse* de la cavité.

La recherche des modes transverses d'une cavité est un problème en général très compliqué. Heureusement, pour les cavités habituellement utilisées dans les lasers (et plus particulièrement dans le cas d'une cavité linéaire formée de deux miroirs concaves) il existe une classe de solutions simples, les *modes transverses gaussiens*, qui sont une excellente approximation pour la plupart des lasers continus.

1. Faisceau gaussien

Considérons le champ électrique suivant (dont on montrera dans le paragraphe 3 qu'il est une solution approchée de l'équation de propagation dans le vide) correspondant à une onde de pulsation ω se propageant dans la direction Oz :

$$\mathbf{E}(\mathbf{r}, t) = \boldsymbol{\varepsilon} E(\mathbf{r}, t) \quad (1.a)$$

avec

$$E(\mathbf{r}, t) = \text{Re} \left\{ E_0 \exp \left(-\frac{x^2 + y^2}{w^2} \right) \exp \left(ik \frac{x^2 + y^2}{2R} \right) \exp[i(kz - \omega t - \varphi)] \right\} \quad (1.b)$$

Dans l'équation (1.b), l'amplitude du champ E_0 , la taille du faisceau w , le rayon de courbure R et la phase du champ φ sont des fonctions réelles de z . Le vecteur unitaire $\boldsymbol{\varepsilon}$, perpendiculaire à Oz , est la polarisation de l'onde. L'amplitude du champ a une symétrie cylindrique autour de Oz et son extension transverse est de l'ordre de w . Il est facile de montrer que la seconde exponentielle de la formule (1.b) décrit une variation transverse de phase caractéristique, pour les petites valeurs de x et y , d'une onde sphérique de rayon de courbure R . Pour une propagation vers les z positifs, l'onde est *divergente* si R est positif comme on peut le constater en considérant les surfaces où la phase est constante. Elle est *convergente* dans le cas contraire ($R < 0$).

En prenant pour origine de l'axe Oz le point où la taille du faisceau est minimum (ce point de focalisation, ou d'étranglement, est appelé col du faisceau ou « waist » en anglais), on trouve (cela sera démontré au paragraphe 3) :

$$w(z) = w_0 \sqrt{1 + \frac{z^2}{z_R^2}} \quad (2.a)$$

$$E_0(z) = E_0(0) \frac{w_0}{w(z)} \quad (2.b)$$

$$R(z) = z + \frac{z_R^2}{z} \quad (3)$$

$$\varphi = \tan^{-1} \left(\frac{z}{z_R} \right) \quad (4)$$

Le rayon au col w_0 caractérise la dimension transversale¹ du faisceau dans le plan $z \approx 0$. La longueur de Rayleigh z_R , égale à

$$z_R = \pi \frac{w_0^2}{\lambda} \quad (5)$$

caractérise l'échelle longitudinale de la variation de la section. En effet, la dimension transversale du faisceau est multipliée par $\sqrt{2}$ entre $z = 0$ et $z = z_R$ (formule (2.a)).

¹On emploie parfois aussi la dénomination « rayon à $1/e^2$ » ce qui fait référence à l'intensité.

Pour des distances z à l'origine petites devant z_R , le faisceau lumineux a une *section pratiquement constante* comme le montre la formule (2.a). Comme par ailleurs, le rayon de courbure (formule (3)) est pratiquement infini dans cette zone, on peut considérer que l'on a un faisceau cylindrique de rayons lumineux *parallèles* à l'axe Oz dans un domaine de longueur z_R de part et d'autre du col.

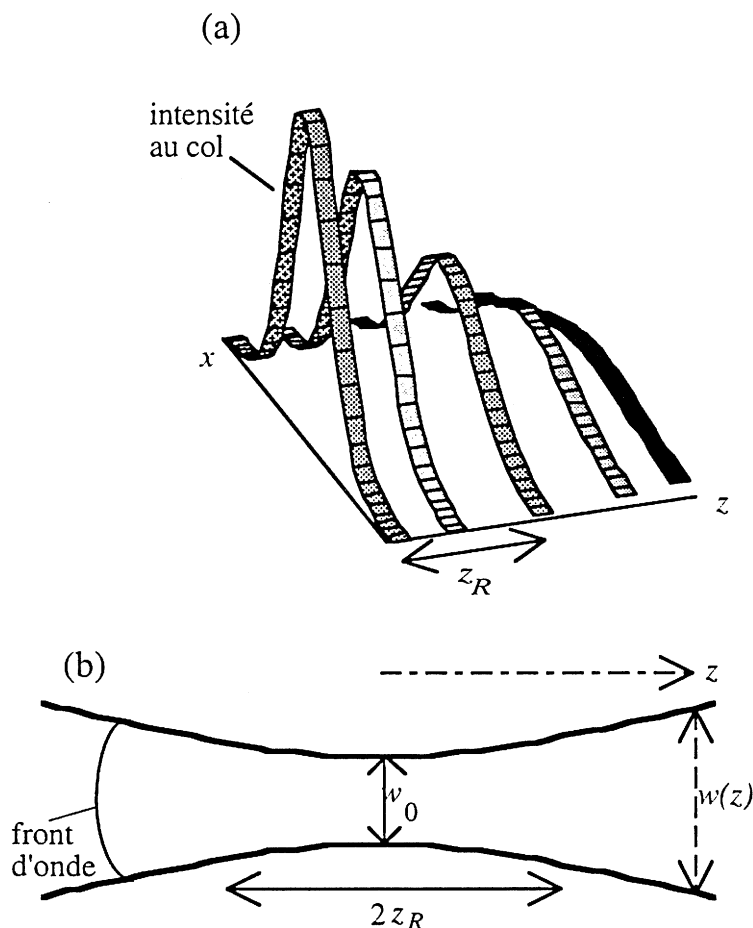


FIG. 1: (a) Variation de l'intensité d'un faisceau gaussien dans la direction de propagation z et dans une direction transverse x . (b) Évolution d'un faisceau gaussien. On a représenté schématiquement le profil. Autour du col, l'onde est proche d'un faisceau cylindrique. Au delà de la zone de Rayleigh le profil est divergent avec des surfaces d'onde sphériques centrées sur le col.

Conformément aux lois de la diffraction, le domaine où le faisceau est pratiquement cylindrique est d'autant plus petit que le rayon au col du faisceau à l'origine est petit puisque z_R est proportionnel à w_0^2 (formule 5). Ainsi pour $\lambda = 0,633 \mu\text{m}$, on a $z_R = 5 \text{ m}$ pour $w_0 = 1 \text{ mm}$ mais seulement $z_R = 0,5 \text{ mm}$ pour $w_0 = 0,01 \text{ mm}$.

À des distances du col grandes devant z_R , nous trouvons

$$w(z) \approx \frac{\lambda}{\pi w_0} z \quad (6.a)$$

$$R(z) \approx z \quad (6.b)$$

On a alors un faisceau divergent à partir du point $x = y = z = 0$ (voir figure 1.b) limité par un cône de demi-angle au sommet $\lambda/\pi w_0$ (qui vaut respectivement 2×10^{-4} et 2×10^{-2} dans les exemples ci-dessus). Cette *divergence*, qui résulte simplement de la diffraction, est d'autant plus grande que la taille est petite. C'est pourquoi la taille des faisceaux doit être augmentée pour diminuer leur divergence dans certaines applications (voir complément III.4, § A.1). La figure 1 donne une représentation schématique d'un faisceau gaussien.

La formule (4) montre que le champ subit une variation de phase de $-\pi/2$ à $\pi/2$ lorsque z varie entre $-\infty$ et $+\infty$. Ce déphasage, appelé déphasage de Gouy, apparaît de part et d'autre de la région focale pour un faisceau focalisé (cet effet, connu en optique classique, est parfois appelé *variation de π au passage par un « foyer »*).

Il est possible de relier E_0 à la puissance Φ transportée dans l'onde et à la section du faisceau w puisque cette puissance est égale au flux du vecteur de Poynting à travers n'importe quel plan d'abscisse z . De la relation

$$\frac{\varepsilon_0 c}{2} \iint dx dy |E(x, y, z, t)|^2 = \Phi \quad (7)$$

on déduit

$$E_0(z) = \frac{1}{w(z)} \sqrt{\frac{4\Phi}{\varepsilon_0 \pi c}} \quad (8)$$

qui montre que E_0 et w sont bien reliés par l'équation (2.b).

2. Mode gaussien fondamental d'une cavité stable

Considérons une cavité résonnante linéaire, formée d'un miroir plan M_1 et d'un miroir concave M_2 de rayon de courbure R_2 distants de L_0 (Fig.2).

S'il existe un faisceau gaussien dont les surfaces d'onde coïncident exactement avec chaque miroir, ce faisceau est un mode stable de la cavité. Un tel faisceau aura donc son col sur le miroir plan et il sera tel que

$$R(L_0) = R_2 \quad (9)$$

d'où l'on déduit à l'aide de (3) :

$$z_R = \sqrt{(R_2 - L_0)L_0} \quad (10)$$

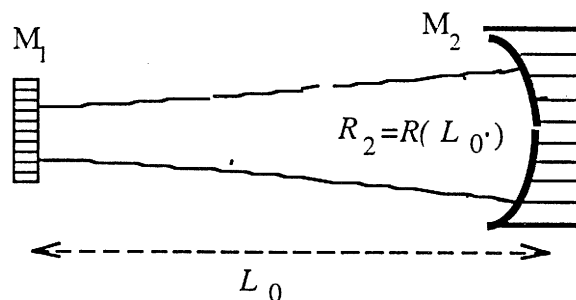


FIG. 2: Cavity stable composée d'un miroir de rayon de courbure R_2 et d'un miroir plan séparé par une distance L_0 inférieure à R_2 . Mode gaussien fondamental de cette cavité.

L'équation (10) n'a de solution que si

$$R_2 \geq L_0 \quad (11)$$

La relation (11) est la *condition de stabilité* de la cavité représentée sur la figure 2. Si elle est vérifiée, *il existe un mode gaussien dont les rayons de courbure sont adaptés aux deux miroirs de la cavité*. Notons en outre que la relation (10) jointe à l'équation (5) permet de déterminer la taille w_0 du faisceau en fonction de R_2 et de L_0 .

Le raisonnement ci-dessus se généralise directement à toute cavité linéaire de longueur L_0 fermée par deux miroirs concaves de rayons de courbure R_1 et R_2 . La condition de stabilité s'écrit alors

$$0 \leq \left(1 - \frac{L_0}{R_1}\right) \left(1 - \frac{L_0}{R_2}\right) \leq 1 \quad (12)$$

Le paramètre z_R (ou w_0) caractérisant le mode s'obtient de façon univoque en imposant aux surfaces d'onde de coïncider exactement avec les miroirs.

3. Modes gaussiens

La solution donnée par les équations (1) et (2) est une solution particulière des équations de Maxwell pour le champ électromagnétique entre les deux miroirs de la cavité. Nous nous proposons maintenant de trouver un ensemble de solutions respectant les mêmes conditions aux limites. L'équation de propagation du champ électrique dans le vide se déduit des équations de Maxwell présentées dans le chapitre II (équations (??,??,??,??)) :

$$\Delta \mathbf{E} - \frac{1}{c^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0 \quad (13)$$

Nous nous plaçons dans l'approximation paraxiale, c'est-à-dire que nous supposons que l'onde électromagnétique se propage approximativement suivant Oz (les normales aux surfaces d'onde forment un angle petit avec Oz) et son amplitude n'est notable qu'au

voisinage de l'axe Oz (à des distances petites comparées au rayon de courbure de la surface d'onde mais grandes par rapport à la longueur d'onde). La condition de transversalité du champ électrique permet alors de négliger² la composante de $\mathbf{E}$ le long de Oz . L'équation (13) se décompose alors en deux équations de propagation indépendantes pour chaque polarisation linéaire $\mathbf{e}_x$ et $\mathbf{e}_y$ et on peut chercher des solutions de la forme

$$\mathbf{E}(\mathbf{r}, t) = E(\mathbf{r}, t)\boldsymbol{\varepsilon} = \text{Re} \left[u(x, y, z) \exp i \frac{k(x^2 + y^2)}{2R(z)} \exp i(kz - \omega t - \varphi(z)) \right] \boldsymbol{\varepsilon} \quad (14)$$

où $\boldsymbol{\varepsilon}$ est la polarisation dans le plan xOy . Pour toute polarisation $\boldsymbol{\varepsilon}$, on se ramène ainsi à la même équation *scalaire* pour E . L'expression (14) prend en compte la variation rapide de la phase due à la propagation le long de l'axe Oz , et le fait que les surfaces d'onde sont tangentes à des sphères centrées sur l'axe Oz (à cause de la symétrie du problème). Les facteurs de phase dépendant de x et y d'une part et de z d'autre part étant intégrés dans les deux exponentielles de (14), la fonction u est réelle. Comme nous le montrons ci-dessous, les lois de variations de $R(z)$ et $\varphi(z)$ sont données par les équations (3) et (21) et les fonctions $u(x, y, z)$ sont de la forme :

$$u_{nm}(x, y, z) = \frac{C}{w(z)} \exp \left\{ -\frac{[x^2 + y^2]}{[w(z)]^2} \right\} H_n \left(\frac{x\sqrt{2}}{w(z)} \right) H_m \left(\frac{y\sqrt{2}}{w(z)} \right) \quad (15)$$

Dans cette expression $w(z)$ est donné par l'équation (2.a), C est une constante de normalisation, H_n et H_m sont les *polynômes d'Hermite* de degrés n et m respectivement (n et m sont des entiers positifs ou nuls).

Pour obtenir la solution (15), nous utilisons l'analogie mathématique avec le problème de l'oscillateur harmonique en Mécanique Quantique. Portons l'expression de $\mathbf{E}$ donnée en (14) dans l'équation scalaire déduite de (13). Nous négligeons dans l'expression du laplacien les termes du type d^2u/dz^2 , $d^2\varphi/dz^2$, $(d\varphi/dz)(\partial u/\partial z)$ ou $(dR/dz)^2$ devant ceux en $ik \partial u/\partial z$ ou $ik d\varphi/dz$ puisque les fonctions u , R et φ doivent varier lentement à l'échelle de la longueur d'onde (cette simplification est connue sous le nom de *l'approximation de l'enveloppe lentement variable*). Nous obtenons alors les deux équations suivantes relatives aux composantes en phase et en quadrature de E (ou encore aux parties réelle et imaginaire de l'amplitude complexe) :

$$-\frac{1}{2} \left[\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} \right] u + \frac{k^2(1 - \frac{dR}{dz})}{R^2} \frac{(x^2 + y^2)}{2} u = k \left(\frac{d\varphi}{dz} \right) u \quad (16)$$

$$\frac{\partial u}{\partial z} + \frac{1}{R} \left(x \frac{\partial u}{\partial x} + y \frac{\partial u}{\partial y} \right) + \frac{u}{R} = 0 \quad (17)$$

L'équation (16) est, en chaque point z , une équation aux valeurs propres analogue à *une équation de Schrödinger indépendante du temps pour un oscillateur harmonique* isotrope à deux dimensions (rappelons que u doit être une fonction décroissante de x et de y). Le rôle de la fonction d'onde est tenu par u et l'énergie est proportionnelle à $kd\varphi/dz$. En appliquant les résultats bien

²Cette composante E_z n'est pas en toute rigueur nulle. En effet, l'équation de Maxwell $\vec{\nabla} \cdot \mathbf{E} = 0$ entraîne $\partial E_z/\partial z = -\partial E_x/\partial x$ pour un champ dont la composante sur Oy est nulle. Pour un champ de la forme (1.b), E_z est environ λ/w_0 fois plus petit que E_x .

connus de Mécanique Quantique³, nous trouvons que les solutions de (16) sont données par des fonctions de la forme (15) où $w(z)$ est défini par la relation :

$$\frac{1}{w^4} = \frac{k^2(1 - \frac{dR}{dz})}{4R^2} \quad (18)$$

De manière analogue au résultat quantique, la *valeur propre* de l'équation de Schrödinger associée à u_{nm} est proportionnelle à $(n + 1/2) + (m + 1/2)$. On en déduit que la phase φ_{nm} associée à u_{nm} vérifie l'équation :

$$\frac{k w^2}{2} \frac{d\varphi_{nm}}{dz} = n + m + 1 \quad (19)$$

En utilisant (18) ainsi que l'équation (17) appliquée à la solution (15), on peut retrouver les équations (2) et (3) donnant respectivement la taille et le rayon de courbure du faisceau et montrer de surcroît que les fonctions $w(z)$ et $R(z)$ sont *indépendantes* de n et m .

Nous savons (notamment par le cours de Mécanique Quantique) que l'ensemble des solutions u_{nm} de « l'équation de Schrödinger »⁴ forme une *base de l'espace des états*. De la même façon pour le problème de la distribution transverse du champ, nous venons de montrer que, pour des ondes transverses écrites sous la forme (14), une base naturelle est donnée par les fonctions u_{nm} .

À chaque vecteur de base u_{nm} est associé une phase φ_{nm} , de la même façon qu'une énergie quantifiée est associée à chaque état propre de l'oscillateur harmonique en Mécanique Quantique. Cette phase φ_{nm} est, d'après les équations (2) et (19), solution de l'équation

$$\frac{d\varphi_{nm}}{dz} = \frac{n + m + 1}{z_R + \frac{z^2}{z_R}} \quad (20)$$

La solution de l'équation (20) est :

$$\varphi_{nm}(z) = (n + m + 1) \tan^{-1} \left(\frac{z}{z_R} \right) \quad (21)$$

formule qui généralise à un mode transverse quelconque l'équation (4) valable pour le mode gaussien fondamental.

En résumé, nous avons trouvé un ensemble orthonormé de solutions, données par l'équation (15), de l'équation d'onde à l'approximation paraxiale. Ces solutions sont appelées *modes gaussiens* et notées TEM_{nm} . Les quantités $R(z)$ et $w(z)$ évoluent en fonction de z comme le *mode fondamental* (qui n'est autre que le mode TEM_{00}). La différence essentielle avec le mode fondamental est donc la dépendance transversale de l'amplitude des modes qui est le produit d'une gaussienne par des polynômes d'Hermite. La conséquence physique est que l'éclairement dans un plan orthogonal au faisceau, au lieu d'être uniformément décroissant du centre vers les bords, possède ici des lignes nodales et qu'il décroît en moyenne moins vite. En utilisant les expressions⁵ des polynômes d'Hermite, il

³Voir CDL Complément BV ou A. Messiah chapitre XII (Dunod, Paris).

⁴Le champ électromagnétique dans cette partie est classique et non quantique. L'analogie avec l'équation de Schrödinger est, rappelons le, purement mathématique.

⁵Voir A. Messiah Appendice B (Dunod, Paris) ou CDL Complément BV.

est possible de tracer les variations transverses d'intensité pour les modes TEM_{nm} . Par exemple, partant de :

$$\begin{aligned} H_0(x) &= 1 \\ H_1(x) &= 2x \\ H_2(x) &= 4x^2 - 2 \end{aligned}$$

nous avons schématisé sur la figure 3 l'allure du profil d'intensité d'un faisceau laser dans quelques modes gaussiens d'ordre peu élevé.

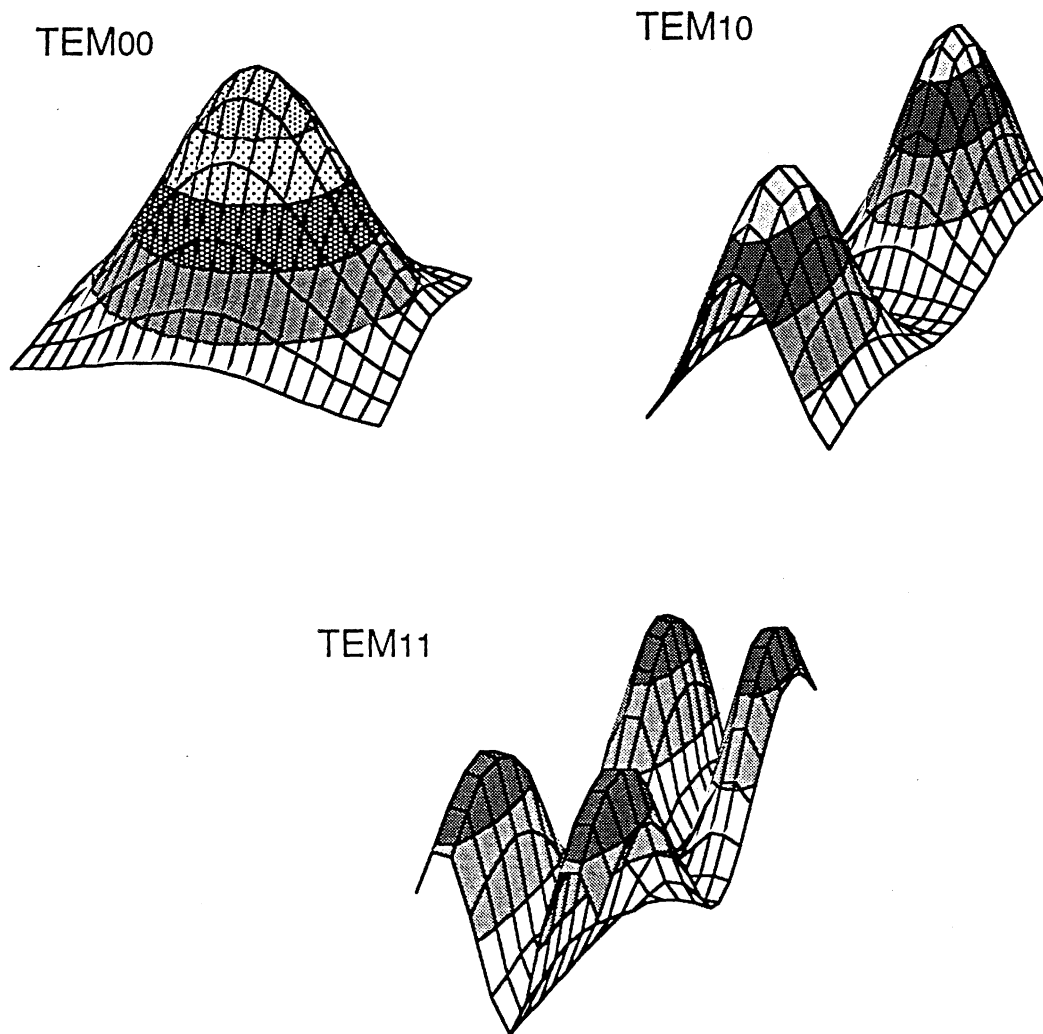


FIG. 3: Distribution transversale d'intensité pour quelques modes gaussiens. On a représenté l'intensité du faisceau en fonction de x et de y pour $-w < x < w$ et $-w < y < w$. Alors que le mode TEM_{00} est pratiquement entièrement contenu dans cette portion du plan, le mode TEM_{11} est bien plus étalé comme le montre la section du profil d'intensité dans les plans $x = \pm w$ et $y = \pm w$.

Remarque

En Mécanique Quantique, il est bien connu que les états propres de l'oscillateur harmonique à deux dimensions peuvent s'exprimer simplement aussi bien en coordonnées cartésiennes qu'en coordonnées cylindriques⁶ (à cause de la dégénérescence en $n + m$ des niveaux d'énergie). De la même façon, il est possible de résoudre l'équation (16) en coordonnées cylindriques (r, θ) . On obtient alors des solutions $u_{\{pl\}}$ (p entier positif et l entier relatif) de la forme⁷

$$u_{\{pl\}}(r, \theta, z) = \left(\frac{r\sqrt{2}}{w(z)} \right)^{|l|} L^{|l|}_p \left\{ \frac{2r^2}{[w(z)]^2} \right\} \exp \left\{ -\frac{r^2}{[w(z)]^2} \right\} \exp(-il\theta) \quad (22)$$

où $L_p^{|l|}$ est un polynôme de Laguerre généralisé. L'évolution de la phase $\varphi_{\{pl\}}$ est donnée par la formule suivante remplaçant (21) :

$$\varphi_{\{pl\}}(z) = (2p + |l| + 1) \tan^{-1} \left(\frac{z}{z_R} \right) \quad (23)$$

4. Modes longitudinaux et transverses d'un laser

Considérons une cavité linéaire stable, par exemple la cavité plan-concave de la figure 2. Nous avons vu dans le chapitre III (§ A.2) que le laser est susceptible d'osciller sur toute fréquence propre de la cavité située dans la partie de la courbe de gain où l'amplification l'emporte sur les pertes. Nous étendons à présent aux modes transverses la discussion du paragraphe III.A.2. *Le mode oscillant est à priori caractérisé par trois entiers : m et n sont relatifs à la détermination du mode transverse u_{nm} , et p est associé à la condition debouclage sur un tour de cavité.* Plus précisément, la phase devant se reproduire identiquement à elle-même après un aller-retour on en déduit la condition suivante qui généralise l'équation (A.14) du chapitre III :

$$\omega_{mnp} \frac{L}{c} - 2[\varphi_{mn}(L_0) - \varphi_{mn}(0)] = 2p\pi \quad (24)$$

avec $L = 2L_0$. Au terme $\omega_{mnp}L/c$ déjà rencontré dans le chapitre III, se rajoute en plus le déphasage de l'onde transverse entre les extrémités de la cavité de la figure 2, ce déphasage étant multiplié par deux puisque la lumière effectue un aller-retour.

En utilisant l'équation (21) adaptée au cas de la cavité plan-concave (la section du faisceau est minimum au niveau du miroir plan), nous pouvons récrire l'équation (24) sous la forme :

$$\omega_{mnp} = \frac{c}{L} \left[2p\pi + 2(n + m + 1) \tan^{-1} \left(\frac{L_0}{z_R} \right) \right] \quad (25)$$

soit en utilisant la valeur de z_R donnée par l'équation (10) :

$$\frac{\omega_{mnp}}{2\pi} = \frac{c}{L} \left[p + \frac{1}{2\pi} (n + m + 1) \cos^{-1} \left(1 - \frac{L}{R_2} \right) \right] \quad (26)$$

⁶Voir CDL Complément D VI.

⁷Pour plus de détails voir par exemple A.E. Siegman, Lasers (University Science Books) ou H. Kogelnik et T. Li, Applied Optics **5**, 1550 (1966).

qui ne dépend que de la longueur de la cavité $L_0 = L/2$ et du rayon de courbure R_2 du miroir concave. L'équation (26) montre que les *modes transverses* correspondant à des valeurs différentes de $(m + n)$ oscillent en général à des *fréquences différentes*.

La discussion du chapitre III doit donc être généralisée pour tenir compte des degrés de liberté supplémentaires introduits par les modes transverses. À priori, divers modes transverse correspondant à des valeurs différentes de n et m peuvent osciller simultanément. Cette oscillation simultanée est gênante, d'abord parce que le laser n'est pas monofréquence, mais aussi parce que la répartition transversale d'éclairement est accidentée, à cause des phénomènes d'interférence entre plusieurs modes transverses alors qu'un éclairement aussi uniforme que possible, est souvent souhaitable. En pratique, le mode TEM_{00} est automatiquement privilégié (à condition que le laser soit bien aligné). Il est en effet clair sur la figure 3 que ce mode est le plus concentré transversalement. Or la cavité laser est toujours *diaphragmée* (ne serait-ce que par le diamètre limité des miroirs ou du milieu amplificateur), ce qui provoque des *pertes, d'autant plus importantes que le mode est plus étalé*. Pour un gain faible, seul le mode TEM_{00} pourra osciller. Si le gain est plus fort, plusieurs modes transverses vont osciller. Lorsqu'un fonctionnement monomode transverse est difficile à obtenir, il faut ajouter un diaphragme à l'intérieur de la cavité. Celui-ci introduit alors des pertes supplémentaires favorisant le mode fondamental mais fait baisser la puissance de sortie. Le compromis entre qualité de faisceau et puissance dépend de l'application envisagée.

Remarque

Les solutions décrites ci-dessus (modes gaussiens) sont celles qui sont généralement rencontrées. Ces formes de faisceau sont associées aux modes d'une cavité *vide* fermée par des miroirs sphériques. L'addition de masques dans la cavité en présence du milieu amplificateur (qui par ses propriétés non linéaires peut coupler des modes gaussiens différents) peut permettre l'oscillation de modes ayant une structure spatiale très différente des modes gaussien⁸.

⁸Voir B. Colombeau, M. Vampouille, V. Kermene, A. Desfarges et C. Froehly, Pure Appl. Opt. **3**, 757, 1994.

Complément III.3

Le laser source d'énergie

Nous avons vu que le laser est une source lumineuse dont les propriétés sont radicalement différentes de celles des sources classiques. C'est la raison pour laquelle, dès sa découverte dans les années 60, on a cherché à tirer profit de ce nouvel « outil ». Le laser est ainsi très rapidement sorti des laboratoires de recherche pour entrer dans le domaine de la production industrielle et de l'activité économique. De nos jours, ses applications sont innombrables, et concernent des domaines aussi variés que la médecine, la métallurgie, les télécommunications etc... Le laser fait partie de la vie quotidienne avec les lecteurs de disques compacts, les lecteurs de code barre ou encore les imprimantes laser.

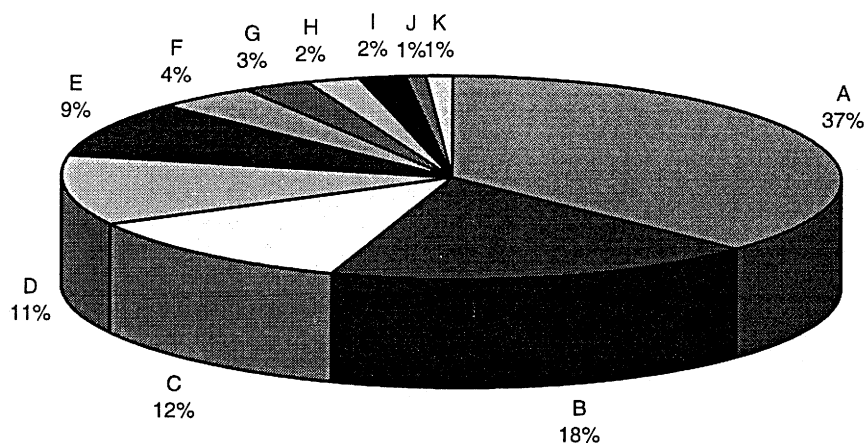


FIG. 1: Ventes annuelles de lasers sur le marché mondial en 1996 (en valeur marchande). A : traitement des matériaux ; B : médecine ; C : télécommunications ; D : mémoires optiques ; E : recherche-développement ; F : instrumentation ; G : imprimantes ; H : LIDAR ; I : spectacles ; J : mesure, contrôle ; K : lecture codes-barre.

La figure 1 donne un aperçu de l'importance relative des différentes applications, mesurée par leur chiffres de vente respectifs. On constate l'importance des applications industrielles de traitement de matériaux, ainsi que des applications médicales. La large partie des ventes consacrée à la recherche et au développement démontre la jeunesse du domaine, qui est loin d'avoir atteint un régime de croisière. De nouvelles applications sont

découvertes chaque année, dont certaines susceptibles d'avoir de profondes implications sur la vie économique.

Il n'est pas possible ici de faire une description exhaustive de ces applications. Nous nous contenterons de quelques exemples significatifs choisis dans différents domaines. Le présent chapitre est consacré aux applications dans lesquelles c'est uniquement l'énergie apportée par le faisceau laser qui est utilisée. Le complément III.4 donnera des exemples où d'autres propriétés, comme la directivité ou la monochromaticité, sont plus spécifiquement mises à profit. Enfin, le complément III.5 est consacré aux applications des lasers en spectroscopie.

1 Effets de l'irradiation laser sur la matière

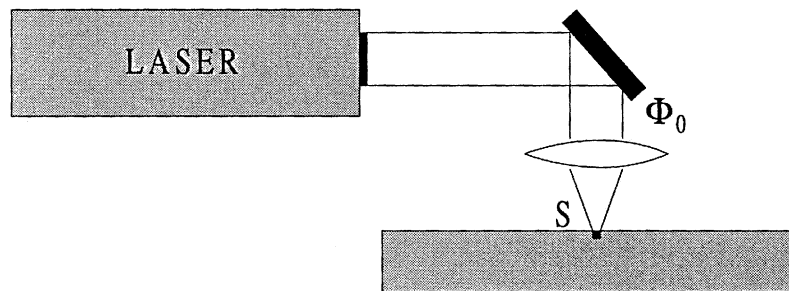


FIG. 2: Un système optique focalise le faisceau d'un laser de puissance Φ_0 sur une surface de matériau d'aire S .

Considérons le dispositif de la figure 2, qui est commun à un grand nombre d'applications : une impulsion laser d'énergie W_0 , de longueur d'onde λ et de durée τ est focalisée à l'aide d'une lentille sur la surface d'un matériau. On obtient ainsi une *irradiation localisée* qui affecte une surface S de petites dimensions. Le flux lumineux par unité de surface $\Pi_0 = W_0/\tau S$ peut être très élevé, et l'amplitude locale du champ électromagnétique atteint alors des valeurs considérables, ce qui peut conduire à de profondes modifications du matériau : échauffement, liquéfaction, vaporisation, voire ionisation complète et création de plasma. Ces différentes phases dépendent évidemment de manière compliquée de la nature du matériau et des caractéristiques de l'impulsion laser¹. On peut néanmoins dégager un certain nombre d'idées générales communes à l'ensemble de ces processus d'interaction matière-rayonnement. Nous noterons x et y les coordonnées de la surface du matériau, et z la coordonnée caractérisant la profondeur de pénétration.

¹Pour une revue sur ce sujet : voir par exemple J. L. Bouinois « Travail des matériaux par laser » dans « Le laser, principes et techniques d'application » p. 107 édité par H. Maillet (Technique et documentation Lavoisier, Paris 1990), ou bien M. Bass « Laser-material interactions » dans « Encyclopedia of lasers and optical technology », Academic Press, San Diego 1991.

1.1 Couplage laser-matériau

L'énergie du faisceau incident sur la surface peut être séparée en trois parties : une fraction R est réfléchiée, ou plus généralement diffusée vers l'arrière, une fraction T est transmise ou diffusée dans les zones plus profondes du matériau, la fraction restante $A = 1 - R - T$ étant absorbée à la surface du matériau. Les deux premières fractions ne contribuent évidemment pas au couplage avec le matériau et doivent être minimisées.

Les coefficients R et T varient considérablement en fonction du matériau, de son état de surface, et de la longueur d'onde du rayonnement utilisé. Pour les grandes longueurs d'onde les métaux sont des conducteurs pratiquement parfaits. Ils réfléchissent donc une grande partie du rayonnement dans le domaine infrarouge et micro-onde. Le coefficient R est supérieur à 90% lorsque $\lambda > 1 \mu\text{m}$ pour la plupart des métaux. Il diminue ensuite et devient inférieur à 50% dans le visible et l'ultraviolet. En revanche, des matériaux non-métalliques, organiques ou non, absorbent fortement le rayonnement laser à toute longueur d'onde. Le coefficient de transmission T est négligeable pour les corps opaques comme les métaux, qui absorbent le rayonnement sur une épaisseur de moins d'un micromètre. Il joue par contre un rôle crucial dans l'interaction avec les *matériaux biologiques* pour lesquels la profondeur de pénétration peut être de l'ordre du centimètre. Ces considérations ne sont en fait valables que pour la phase initiale d'échauffement, car le couplage au rayonnement varie en fonction de l'état atteint par le matériau. Lorsque celui-ci se trouve au voisinage de son point de fusion, il se comporte approximativement comme un *corps noir*, ayant un coefficient d'absorption proche de 1. Le transfert énergétique devient alors beaucoup plus efficace et indépendant du matériau. Lorsque la surface du corps est complètement ionisée, c'est le *plasma* résultant, en expansion au dessus de la région de focalisation, qui va absorber le rayonnement incident. Les charges libérées (en particulier les électrons) vont se mettre à osciller sous l'effet du champ électromagnétique de l'onde laser et à rayonner à leur tour. Ce processus peut conduire à une réflexion totale du rayonnement laser si la densité électronique N dépasse une valeur seuil N_s , pour laquelle la fréquence propre d'oscillation plasma² $\omega_p(Nq^2/m\varepsilon_0)^{1/2}$ est égale à la fréquence du laser. Cette densité vaut donc $N_s = 4\pi^2\varepsilon_0 mc^2/q^2\lambda^2$. Il faut donc éviter à tout prix la formation de cette « barrière plasma » qui diminue fortement l'interaction du matériau avec le champ laser, et peut même dans certaines circonstances la supprimer complètement. L'expression de N_s montre qu'on a intérêt à utiliser des lasers de *courte longueur d'onde*³, qui assurent une valeur plus grande à la densité seuil.

²Voir par exemple R.P. Feynman, R.B. Leighton et M. Sands « The Feynman lectures on physics » Vol II § 32-7

³Parmi les premiers, E. Fabre et ses collaborateurs du Laboratoire pour l'Utilisation des Lasers Intenses de l'École Polytechnique ont étudié les possibilités offertes par les lasers très intenses de courte longueur d'onde pour créer des plasmas lasers très chauds, en particulier en vue d'applications à la fusion inertielle (voir partie E).

1.2 Transfert thermique

La puissance $A\pi_0(x, y)$ déposée par le laser au niveau de la tâche focale, en un point de coordonnées (x, y) de la surface du matériau, provoque une variation de la température $T(x, y, z, t)$ du matériau. Dans le cas où la profondeur d'absorption de la lumière laser est extrêmement faible, l'équation de diffusion de la chaleur dans le matériau irradié s'écrit :

$$\rho C \frac{\partial T}{\partial t} = K \Delta T + A\pi_0(x, y)\delta(z) \quad (1)$$

où ρ est la densité du matériau, C sa capacité calorifique, et K sa conductivité thermique. La solution de cette équation dépend de la géométrie du matériau, des caractéristiques spatiales et temporelles de $\Phi_0(x, y, t)$, et d'un paramètre caractéristique du comportement thermique du matériau, appelé *diffusivité thermique*, $\kappa = K/\rho C$ (κ vaut 0,05 cm²/s pour l'acier, 1 cm²/s pour le cuivre).

Envisageons le cas d'un matériau s'étendant dans tout le demi-espace des z positifs, irradié par une brève impulsion⁴ de forme gaussienne de diamètre w (distribution d'énergie en $\exp\{-(x^2 + y^2)/w^2\}$) et d'énergie totale W_0 . La solution de l'équation (1) s'écrit dans ce cas :

$$T(x, y, z, t) = \frac{AW_0}{\pi\rho C\sqrt{\pi\kappa t}(4\kappa t + w^2)} \exp\left(-\frac{z^2}{4\kappa t} - \frac{x^2 + y^2}{4\kappa t + w^2}\right) \quad (2)$$

Cette expression montre que la chaleur diffuse, aussi bien en profondeur que latéralement, sur une distance de l'ordre de $\sqrt{\kappa t}$, caractéristique d'un phénomène de marche au hasard. La diffusivité κ qui intervient dans cette distance caractérise donc la capacité du matériau à évacuer l'énergie introduite localement.

La dépendance spatiale et temporelle de la température du matériau pendant la durée de l'irradiation laser prend en général des formes compliquées. On obtient cependant un résultat simple pour la température $T(x = 0, y = 0, z = 0, t)$ à la surface du matériau et au centre de la tache gaussienne, dans le cas où la puissance instantanée Φ_0 du laser est supposée constante pendant toute la durée de l'impulsion :

$$T(x = 0, y = 0, z = 0, t) = \frac{A\Phi_0}{\pi^{3/2}Kw} \tan^{-1}\left(\frac{\sqrt{4\kappa t}}{w}\right) \quad (3)$$

La formule (3) montre qu'aux temps longs ($\tau \gg w^2/4\kappa$), on atteint un régime stationnaire dans lequel la puissance déposée sur la surface du matériau par le laser équilibre le flux de chaleur diffusé. La température d'équilibre est égale à $T_{eq} = A\Phi_0/2\pi^{1/2}Kw$. Par exemple, la température de fusion de l'acier (1430°C) est atteinte avec une puissance de 100 W environ avec un laser à CO₂ ($\lambda = 10,6 \mu\text{m}$, $A \approx 10^{-2}$) focalisé au maximum ($w \approx \lambda$), et une puissance de 1 W environ si l'on utilise un laser Nd:YAG ($\lambda = 1,06 \mu\text{m}$, $A \approx 10^{-1}$). Pour

⁴L'impulsion laser sera qualifiée de « courte » si la chaleur n'a pas le temps de diffuser sur une distance supérieure à la taille du faisceau pendant la durée de l'impulsion

ces deux exemples ⁵ le temps mis pour atteindre le régime d'équilibre est respectivement de 5 μs et 500 μs .

En revanche, aux temps courts, la formule (3) montre que la température de surface croît extrêmement rapidement avant d'atteindre le régime d'équilibre. L'accroissement initial de température, en $\sqrt{t}$, est proportionnel au flux instantané absorbé par la surface, $A\Phi_0/\pi w^2$. Dans les conditions de focalisation précédentes, un laser à CO_2 de 1 kW porte l'acier à son point de fusion en 300 ns. On peut ainsi utiliser des impulsions extrêmement brèves et d'énergie relativement faible pour porter localement un métal à sa température de fusion. L'utilisation d'impulsions brèves présente en outre l'avantage de limiter au maximum la zone d'échauffement, puisque la diffusion en profondeur de la chaleur, en $\sqrt{\kappa t}$, n'a pas eu le temps de se produire.

1.3 Effet mécanique

Lorsque la température de vaporisation du matériau est dépassée, l'irradiation laser produit un effet nouveau : il y a *mise en mouvement de matière*. Des fragments solides ou vaporisés, ou encore des ions dans le cas de la formation de plasma, s'éloignent de la surface avec des vitesses très importantes. Si l'éjection se produit à une vitesse supérieure à celle du son dans le milieu, il y a création d'une *onde de choc* (détonation) qui peut se révéler extrêmement utile pour « nettoyer » complètement la zone chauffée de ses différentes impuretés.

Enfin, par conservation de l'impulsion dans le processus, le substrat non évaporé subit une force normale à la surface qui peut être considérable, puisqu'il y a éjection de matière dans un temps extrêmement court. *Cette force est bien supérieure à celle due à la simple pression de radiation*. Nous verrons dans le paragraphe E de ce complément que l'on peut tirer profit de cet effet inertiel.

1.4 Effets photochimiques, photoablation

Le laser peut aussi produire des effets de nature chimique, et aboutir à la *photodissociation*, ou à la *photofragmentation*, des molécules composant le matériau. Les lasers de *courte longueur d'onde*, pour lesquels le quantum d'énergie disponible est plus grand, seront les plus efficaces pour produire de tels effets avec des flux énergétiques relativement modestes. Dans d'autres circonstances, un flux lumineux de puissance modérée pourra induire des réactions entre plusieurs molécules présentes sur la surface, par exemple des réactions de *polymérisation*.

⁵Ces deux exemples correspondent à des cas limites idéalisés. Dans la pratique, les faisceaux des lasers ne peuvent être aussi fortement focalisés, le coefficient d'absorption dépend de l'état de surface du matériau, et les puissances nécessaires pour atteindre le point de fusion sont nettement supérieures, tout en restant accessibles aux lasers usuels.

Un autre effet apparenté aux effets photochimiques, et appelé « *photoablation* », a été observé dans l'interaction entre le rayonnement ultraviolet, produit par exemple par les lasers à excimère, et certains matériaux comme les chaînes polymérisées : on a constaté que si l'énergie déposée par unité de surface dépasse une valeur seuil (de l'ordre de 100 mJ/cm^2), il y a « ablation » du matériau sur la surface illuminée sur une épaisseur bien définie ($0,3 \mu\text{m}$ environ). La figure 3 montre à titre d'exemple des trous de $300 \mu\text{m}$ de diamètre percés dans un substrat de polyamide à l'aide d'un laser Nd :YAG à $1,06 \mu\text{m}$ (a), un laser CO_2 à $10,6 \mu\text{m}$ (b) et un laser excimère KrF à $0,248 \mu\text{m}$ (c). On constate dans le cas (c), correspondant à la photoablation, l'extrême netteté de la découpe, due à l'effet de seuil, et l'absence de dommage sur les bords des zones découpées, indiquant qu'à la différence des cas (a) et (b) il n'y a pratiquement pas d'effets thermiques.

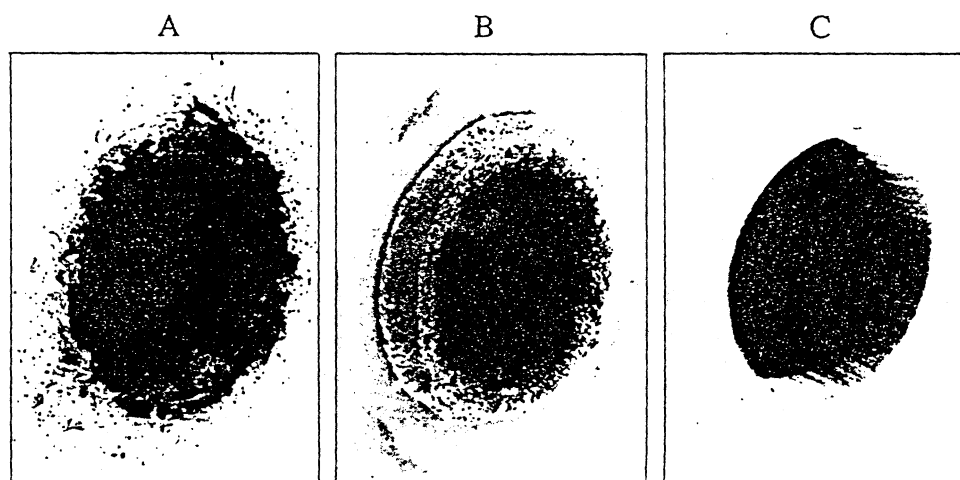


FIG. 3: Trous de $300 \mu\text{m}$ de diamètre percés dans une feuille de polyamide de $75 \mu\text{m}$ d'épaisseur à l'aide de différents lasers : (a) : Nd : YAG à $1,06 \mu\text{m}$; (b) CO_2 à $10,6 \mu\text{m}$; (c) excimère KrF à 248 nm . On remarque dans le cas (c) la netteté des bords du trou, percé par effet de photoablation.

On peut donner de ce phénomène l'explication simplifiée suivante : les photons ultraviolets du faisceau brisent de nombreuses liaisons chimiques dans les polymères du matériau irradié. Les monomères et les produits de dissociation ainsi formés occupant un volume plus grand que le polymère original, il y a explosion locale avec expulsion rapide des produits de réaction. L'énergie apportée par le laser, ayant principalement servi à dissocier les polymères, diffuse très peu dans le substrat, dont la température ne varie pratiquement pas pendant l'irradiation. C'est ce qui explique l'absence de traces de « brûlure » sur le matériau qui n'a pas été enlevé.

2 Usinage et traitement de surface des matériaux

2.1 Effets thermiques

Le paragraphe précédent nous a montré qu'en jouant sur la longueur d'onde, l'intensité, la durée, la focalisation d'un faisceau laser, on est en mesure de porter avec une assez bonne précision une région de surface et de profondeur définies à une température élevée déterminée. L'effet est donc comparable à celui d'un chalumeau, mais il a l'avantage sur ce dernier de n'affecter pratiquement pas les régions avoisinantes, ce qui limite l'énergie totale consommée et évite les coûteux réusinages de pièces après traitement. Signalons aussi que le faisceau laser est transporté à l'aide de simples miroirs, que l'on peut facilement positionner de manière automatique : un système laser se prête particulièrement bien à une gestion assistée par ordinateur. Bien sûr, dans le choix d'un système de traitement industriel, les considérations de coût et de rendement sont primordiales. En particulier, le rendement du laser employé, sa facilité d'emploi et son coût d'entretien entrent en ligne de compte. Malheureusement, les lasers qui ont les meilleurs rendements émettent le plus souvent dans l'infrarouge, où le couplage matériau-rayonnement est peu efficace ! Par exemple, pour le traitement de l'acier, il était en 1994 plus intéressant d'utiliser un laser à CO_2 , d'excellent rendement, mais émettant dans l'infra-rouge lointain ($10,6 \mu\text{m}$) plutôt qu'un laser à YAG ($1,06 \mu\text{m}$) dont le rayonnement est pourtant dix fois mieux absorbé.

En fonction de la température T localement atteinte, on obtiendra sur le matériau irradié différents effets qui sont, dans l'ordre croissant :

- $T < T_{\text{fusion}}$: *traitement de surface*. Il peut être de nature chimique (*cémentation*, c'est-à-dire alliage avec un autre métal déposé sur la surface), ou bien structurale (*trempe* ou *vitriification de surface*). On obtient un matériau qui résiste mieux aux chocs, aux agressions thermiques et chimiques. Ce type de traitement est par exemple utilisé dans l'industrie automobile, pour augmenter la résistance de la surface des pistons, des soupapes ou des vilebrequins.
- $T_{\text{fusion}} < T < T_{\text{vaporisation}}$: *soudage*. Il se fait sans apport de métal extérieur, sans contact direct avec la pièce et jusque dans des endroits difficilement accessibles à des chalumeaux classiques (intérieurs de tuyaux ...). Un laser à CO_2 de 10 kW, par exemple, est capable de souder des tôles d'acier de 1 cm d'épaisseur à la vitesse de 1 m/mn.
- $T > T_{\text{vaporisation}}$: *gravure, perçage, découpage*, en fonction de la profondeur et de l'extension latérale plus ou moins grande sur laquelle le matériau est volatilisé. Il est possible en particulier d'usiner des matériaux réfractaires, à haut point de fusion : céramiques, rubis... Les exemples d'utilisation sont innombrables, et le pilotage par ordinateur joue un rôle essentiel : gravure de logos de marques sur divers substrats (métaux ou polymères), modification ou réparation de pistes sur les « puces » des microprocesseurs, perçage des tétines de biberons, découpe des tableaux de bord des automobiles grâce à un système de bras articulé permettant d'atteindre des endroits d'accès difficile, découpe « sur mesure » et sans effilochage de tissu pour la confection de vêtements...

2.2 Transfert de matière

Il est possible de tirer profit de l'éjection de matière qui se produit lorsqu'une surface est illuminée par un faisceau laser intense. La méthode permet en effet de vaporiser de manière contrôlable des matériaux éventuellement très réfractaires, qui se déposent ensuite sur une autre surface placée à proximité de la surface irradiée. C'est de cette manière par exemple que l'on dépose des couches minces de céramiques supraconductrices de type Y-Ba-Cu-O. Ces couches ont une excellente qualité cristalline, comme l'atteste la valeur élevée du courant critique que l'on mesure dans de tels dispositifs.

2.3 Effets photochimiques, photoablation

Il existe une application intéressante des effets photochimiques liés à l'irradiation laser, appelée « stéréolithographie ». Elle permet de fabriquer des prototypes de pièces de formes compliquées conçues sur ordinateur : la surface d'une cuve contenant une résine acrylique liquide photopolymérisable est éclairée par un faisceau laser piloté par ordinateur. La résine durcit au point d'irradiation, et le faisceau est balayé jusqu'à matérialisation complète d'une « tranche » de la pièce à créer. On élève alors de 0,1 mm environ le niveau du liquide et on recommence l'opération pour la couche immédiatement supérieure. Il faut entre 4 et 30 heures pour fabriquer ainsi une pièce d'une dizaine de cm, mais le processus est complètement automatique et il peut être dans certains cas moins coûteux qu'un usinage classique.

D'autre part, les propriétés spécifiques de la photoablation commencent à être utilisées dans un certain nombre d'applications, en dépit du coût plus élevé des lasers UV : perçage de précision dans des films de polymères, confection de masques pour la micro-électronique, dénudage de fils électriques de très faible diamètre.

3 Application médicales

Dans la plupart des applications médicales, il s'agit comme dans le paragraphe précédent d'apporter de l'énergie sous forme contrôlée. Il y a tout de même des différences, liées au fait que le matériau à traiter ici est vivant⁶ :

- L'énergie nécessaire au traitement est beaucoup plus faible et ne nécessite pas de lasers de forte puissance.
- Il faut limiter à la zone malade la zone d'interaction par diffusion thermique. Pour cela, l'utilisation d'impulsions brèves est indispensable.
- La profondeur de pénétration est plus importante que dans le cas des métaux. Elle dépend de la longueur d'onde utilisée : 50 μm pour un laser à CO_2 , 800 μm pour un

⁶Voir par exemple J.M. Brunetaud : « *Les applications médicales du laser* » dans : Les lasers et leurs applications, Cours de l'Ecole d'été de la SFO, Cargèse 1994, Les Éditions de Physique (Paris 1996).

laser à YAG, 200 μm pour un laser à Argon ionisé (0,514 μm). Cette profondeur reste toujours faible, ce qui limite l'usage du laser à des traitements de tissus dans un volume restreint et pouvant être atteint par le faisceau laser. Nous verrons plus loin que cette caractéristique ne présente pas que des désavantages.

- Le tissu traité est vivant, et il réagit au traumatisme subi d'une manière qui peut *annuler* l'effet recherché. Il faut donc limiter au maximum ce traumatisme.

L'effet obtenu dépend ici aussi de la température maximale atteinte. Si celle-ci est très grande, on volatilise localement le tissu, ce qui permet de réaliser des coupes ou incisions, ou de détruire des cellules malades. Si celle-ci est comprise entre la température de coagulation et la température de carbonisation, on obtient un effet de photocoagulation locale très utile en chirurgie. Enfin, la technique de photoablation (voir § B.3) rend possible une élimination des tissus très propre et très précise, en raison de la netteté de bords des zones enlevées et de l'absence de nécrose des tissus environnants. Ces propriétés permettent de comprendre l'intérêt de l'utilisation du *bistouri laser* en chirurgie et en odontologie. En dépit de ses avantages (cautérisation, stérilité), le développement du bistouri laser est ralenti par des problèmes d'encombrement et de maintenance.

La *dermatologie* a été un des premiers champs d'application du laser en médecine : il permet en principe de traiter d'innombrables affections cutanées, tels que les angiomes, les taches de vin... Signalons aussi à titre d'exemple les applications en *ophtalmologie*, où la longueur d'onde du laser utilisé joue un rôle déterminant : les décollements de rétine sont traités par un laser à Ar^+ émettant un faisceau de lumière visible qui n'est absorbé qu'au niveau de la rétine. On fait également des recherches pour traiter la myopie à l'aide de lasers UV, dont le rayonnement, absorbé au niveau de la surface du cristallin, crée de très fines incisions qui en modifient le rayon de courbure.

Le complément naturel du laser en médecine est l'*endoscope*, tube flexible que l'on peut introduire à l'intérieur du corps humain (tube digestif, vaisseaux sanguins...) et qui contient une fibre optique d'illumination, un système optique d'inspection visuelle (l'endoscope proprement dit), la fibre optique transportant le faisceau laser, et un canal auxiliaire permettant d'injecter différents produits au niveau de l'extrémité ou d'aspirer les débris des tissus détruits. Cet ensemble permet de réaliser des interventions extrêmement localisées à l'intérieur du corps humain au prix d'un traumatisme de l'organisme nettement inférieur à celui résultant d'une opération chirurgicale classique. Le coût plus élevé de l'équipement laser peut dans certains cas être largement compensé par la réduction des dépenses d'hospitalisation. Un exemple particulièrement médiatique est fourni par les recherches actuelles de traitement laser de certaines obstructions coronariennes (angioplastie) : la fibre est introduite dans l'artère coronaire jusqu'au niveau de l'athérome, responsable du rétrécissement ou de l'occlusion du vaisseau. Celui-ci est alors détruit grâce au laser par effet thermique ou mieux, par photoablation. Ce type de traitement n'est cependant pas dénué d'effets secondaires. Les effets thermiques, par exemple, conduisent ainsi souvent à des complications sérieuses : brûlure des tissus proches ou réocclusion précoce. Ce traitement n'est donc pas, pour l'instant, la panacée. Néanmoins, les pro-

grès offerts par la photoablation⁷ tout comme une meilleure intégration des lasers dans la panoplie de soins des maladies artérielles, offrent un espoir raisonnable à terme.

Les effets thermiques ou de photoablation présentent le désavantage d'être non sélectifs médicalement parlant, puisqu'ils s'attaquent aussi bien aux tissus malades qu'aux parties saines. Il est possible de tirer profit de l'absorption sélective en longueur d'onde de certaines substances. Des recherches sont actuellement en cours sur l'utilisation de plusieurs molécules, et en particulier des dérivés de l'hématoporphyrine. Ces molécules présentent la particularité de se fixer préférentiellement dans les *cellules cancéreuses* et *athéromateuses*. Si l'on irradie les tissus par un faisceau laser de longueur d'onde correspondant à un pic d'absorption de cette molécule (635 nm), on la décompose en produisant in situ un poison qui tue uniquement la cellule malade. Cette technique, a priori très séduisante, présente néanmoins des inconvénients importants : l'ingestion de ce produit provoque une photosensibilisation globale du malade et la faiblesse de la profondeur de pénétration du rayonnement lumineux limite le volume de cellules malades détruites.

Le bilan de l'utilisation des lasers en médecine et en chirurgie doit, pour l'instant, être relativement nuancé. Après une phase de très forte expansion, le début de la décennie 90 a vu une stabilisation, parfois une régression, de l'emploi des lasers. Si pour certaines indications (en ophtalmologie par exemple), l'utilisation des lasers semble désormais incontournable, dans d'autres cas, le laser a été vécu comme une mode parfois discutable (utilisation du « soft-laser », source laser continue de quelques mW, pour soigner les processus inflammatoires des sportifs par exemple). Cette pause, probablement salutaire, devrait permettre de mieux définir les sources nécessaires de façon à adapter les lasers à la médecine plutôt que de suivre la stratégie inverse. À plus long terme, il ne fait guère de doutes que la miniaturisation des sources lasers, l'augmentation de leur fiabilité, la possibilité de délivrer des impulsions dont la forme spatiale et temporelle et l'intensité peuvent être programmées tout comme les progrès dans le domaine des fibres optiques, conduiront à de nouveaux développements positifs.

Anticipant quelque peu sur le complément suivant, il nous faut signaler enfin que le champ des applications du laser en médecine comprend le traitement mais aussi le diagnostic des maladies : contrôle de l'écoulement sanguin par vélocimétrie Doppler (§ B.3), dosage de protéines dans le sang, ou bien détermination de la dimension de cellules par diffractométrie (§ B.4).

4 Application militaires

Dans les dessins animés, les films de science-fiction et dans l'inconscient collectif, l'image communément associée au laser est celle de « rayon de la mort ». Cet aspect des applications du laser n'a pas échappé aux militaires. L'arme laser présente en effet de

⁷La photoablation étant faite avec des lasers UV, il faut toutefois s'assurer qu'ils n'induisent pas de mutations cellulaires dans les tissus environnants.

grands avantages potentiels : le faisceau étant dirigé à l'aide de miroirs, l'inertie du système de tir est négligeable comparée à un système classique (canon ou missile). De plus, le « projectile » se déplace à la vitesse de la lumière. L'erreur de pointage, et donc la probabilité de manquer la cible, peut donc être rendue très faible, même pour des distances très importantes. C'est la raison pour laquelle des recherches intensives ont été consacrées à l'arme laser, qui peut avoir des utilisations aussi bien tactiques que stratégiques : on sait que, dans les années 80, les superpuissances ont consacré de gigantesques efforts au projet dit de « guerre des étoiles », qui devait notamment conduire à développer un système de destruction par laser des missiles balistiques adverses à l'aide d'un réseau de satellites⁸.

Ces recherches se heurtent cependant à un grand nombre de problèmes difficiles à résoudre. En effet, pour obtenir un effet destructeur à partir d'un rayonnement lumineux, il faut apporter sur l'objet à endommager un flux de quelques kW/cm² pendant un temps de l'ordre de la seconde. La focalisation du faisceau sur la cible et le suivi de son pointage pendant le temps nécessaire à la destruction doivent donc être réalisés avec une grande précision. De plus, le faisceau subit de nombreux effets perturbateurs sur son trajet : atténuation et diffusion par l'atmosphère, d'une manière qui dépend beaucoup des conditions météorologiques, et, pour les faisceaux de très grandes densités de puissance, effets d'optique non-linéaire comme l'autodéfocalisation (voir complément III.6 § D.1). Le laser à utiliser doit émettre des puissances considérables avec un bon rendement et un taux de répétition important. Il doit aussi être autonome, fiable et facile à mettre en opération même dans des conditions difficiles. Enfin le système formé des lasers et de leurs miroirs de renvoi doit être le plus possible invulnérable vis à vis des contre-mesures de l'adversaire. De nombreuses recherches sont actuellement consacrées au développement de *lasers chimiques*, comme le laser à HF (voir chapitre III § B.2.e), qui peuvent produire des puissances moyennes très importantes, supérieures à 100 kW, sans aucun apport extérieur d'énergie sous forme électrique, et peuvent par exemple fonctionner dans l'espace. Ils présentent par contre l'inconvénient de produire un rayonnement infrarouge, qui, nous le savons, est fortement réfléchi par les matériaux usuels. D'autres types de laser sont aussi étudiés à l'heure actuelle dans le même but, comme les lasers à excimères ou les lasers à électrons libres.

Les problèmes techniques sont donc loin d'être tous résolus à l'heure actuelle. C'est la raison pour laquelle l'intérêt stratégique de l'arme laser reste encore à démontrer. Le laser est en revanche couramment utilisé à l'heure actuelle pour neutraliser par éblouissement les systèmes de visualisation adverses, ce qui nécessite évidemment des flux lumineux beaucoup plus faibles. La majeure partie des applications militaires actuelles des lasers concerne en fait des domaines que nous allons traiter dans le Complément III.4 : la distance de cibles est couramment déterminée par télémétrie laser impulsionnelle (paragraphe B.2), et leur vitesse par vélocimétrie Doppler (paragraphe B.3). Enfin, la directivité du laser est utilisée pour la désignation d'objectifs : le projectile est guidé sur la tache lumineuse diffusée par la cible que l'on a illuminée par un faisceau laser.

⁸Voir Review of Modern Physics, **59**, pp S1-S201 (Juillet 1987).

5 Fusion inertielle

Les recherches concernant la fusion thermonucléaire contrôlée par laser sont activement poursuivies dans le monde entier. Même si en 1995, on n'a pas démontré la validité de cette technique, nous la décrivons ici en raison de son importance potentielle⁹. Elle est basée sur l'utilisation de la réaction de fusion entre noyaux de deutérium et tritium, qui produit une particule α et un neutron (n) de grande énergie :



Le deutérium se trouve en abondance dans l'eau de mer, et le tritium se fabrique à partir de neutrons thermiques et de lithium, provenant aussi de l'eau de mer. L'humanité dispose donc en principe d'un énorme gisement d'énergie. Le problème est, que pour réaliser la réaction de fusion, il faut être capable de porter les noyaux de deutérium et de tritium à une température suffisante pour leur faire franchir la barrière de répulsion électrostatique. Celle-ci est de l'ordre de 10 keV, soit environ $10^8 K$. Il faut aussi que le nombre de chocs entre ces noyaux soit suffisamment grand pour que la combustion puisse s'entretenir. On montre qu'il faut pour cela que le produit de la densité N par le temps de confinement τ soit au moins de l'ordre de 10^4 particules $\text{cm}^{-3} \times s$. Deux méthodes concurrentes pour obtenir ces conditions extrêmes sont à l'heure actuelle étudiées à grande échelle : le confinement magnétique dans les tokamaks pour lequel τ est de l'ordre de la seconde et le confinement inertiel par laser, ou par faisceaux d'ions lourds, pour lequel τ est beaucoup plus court, de l'ordre de la nanoseconde. Pour que la réaction s'amorce, il faut dans ce dernier cas une densité beaucoup plus élevée, équivalente à 1000 fois celle du solide, c'est-à-dire des conditions de température et de densité proches de celles de l'intérieur des étoiles ou des engins thermonucléaires.

Le principe de la fusion inertielle est le suivant (Figure 4) : une cible sphérique de quelques millimètres de diamètre contenant un mélange deutérium-tritium est illuminée par des faisceaux lasers convergents de très grande énergie. La cible est ionisée superficiellement, et s'entoure d'une couronne de plasma où le rayonnement est absorbé. L'énergie laser déposée chauffe en surface la cible et en volatilise les couches superficielles. Le milieu est alors soumis à une force très intense dirigée vers le centre, par l'effet mécanique de réaction dont nous avons parlé au début de ce chapitre (paragraphe A.3). Il est fortement comprimé, et un point chaud est créé au centre, où des réactions de fusion commencent à se produire. Celles-ci, servant en quelque sorte « d'allumette », produisent des particules α énergétiques qui vont initier les réactions de fusion dans le reste de la cible et conduire à son inflammation. Le terme « inertiel » qui qualifie cette technique fait référence au fait que c'est l'inertie mécanique qui maintient la cible comprimée avant qu'elle n'explose, pendant un temps de confinement τ satisfaisant au critère de Lawson. Pendant la réaction, des neutrons énergétiques s'échappent du milieu et fournissent aux parois de l'enceinte, sous forme thermique, l'énergie utilisable.

Pour réaliser une centrale électrique à fusion thermonucléaire, il faut pouvoir disposer

⁹Cette partie s'inspire étroitement d'un texte rédigé par A. Migus, du laboratoire LULI de l'École polytechnique.

d'un laser fournissant 10 fois par seconde des impulsions dont l'énergie est d'une dizaine de Mégajoules, et la durée 10 ns environ¹⁰. Dans ces conditions, l'inflammation de la cible doit libérer à chaque impulsion du laser une énergie de l'ordre de 1 GJ (soit un gain de 100 environ), qui est ensuite convertie en chaleur puis en électricité. Ce dispositif conduit à une production d'énergie « utile » de 100 MJ à chaque impulsion du laser, permettant de fournir en continu 1 GW de puissance électrique.

Pour que la description précédente soit autre chose qu'un scénario de science-fiction, il faut résoudre un certain nombre de problèmes physiques et technologiques délicats. Tout d'abord, le couplage avec la cible doit être le plus efficace possible. Il s'avère que l'essentiel de l'énergie du laser est absorbé par le plasma créé lorsqu'il a une densité légèrement inférieure à la densité critique N_s (voir paragraphe A.1). Le couplage est donc meilleur avec des lasers de courte longueur d'onde, comme nous l'avons vu au début de ce chapitre. Il peut atteindre des valeurs de 90 % vers 300 nm. Il faut ensuite éviter que l'énergie cédée à la cible se couple aux modes d'oscillation du plasma, ce qui conduirait à un chauffage prématuré du centre de la cible, et donc à une diminution de l'effet de compression. La cible doit avoir une structure en couches très étudiée, et l'impulsion laser une dépendance temporelle appropriée, de manière à maximiser le rendement de compression. Il faut aussi veiller à l'uniformité de l'éclairement afin d'éviter l'apparition d'instabilités hydrodynamiques : pour cela, on utilise un grand nombre de faisceaux convergents très bien équilibrés. D'autres schémas sont aussi envisagés, où les phases de compression et d'allumage seraient découplées.

Les lasers ayant à la fois les performances requises en énergie, durée d'impulsion, en taux de répétition et en rendement, ne sont prêts ni technologiquement, ni surtout économiquement. De gros efforts sont engagés pour atteindre ce but. Les concepts physiques et les techniques du procédé ont été étudiés à l'aide d'énormes installations comme Nova à Livermore aux USA, Gekko-XII à Osaka au Japon, et Phébus à Limeil en France, mais aussi grâce à des installations plus petites du type LULI à l'École Polytechnique. Tous sont insuffisants pour permettre l'inflammation de la cible. Les systèmes laser les plus importants ont la taille d'un immeuble. Ils utilisent un laser au Néodyme, à 1,06 μm , dont le faisceau de sortie est subdivisé en plusieurs sous-faisceaux (10 dans le cas de Nova). Ceux-ci sont ensuite amplifiés par toute une série de disques de verre au Néodyme de section croissante pompés par des batteries de lampes flash de grande puissance. Des cristaux de KDP de grandes dimensions produisent, par effet d'optique non-linéaire (voir complément III.7 § B.I), les harmoniques 2 ou 3 dont les courtes longueurs d'onde sont mieux adaptées à l'interaction avec la cible. Les différents faisceaux, dont le diamètre pour Nova ou Phébus est alors de 70 cm, sont concentrés dans une enceinte à vide et focalisés sur la cible de quelques centaines de microns de diamètre. Ces dispositifs permettent d'obtenir des énergies par impulsion considérables : Pour Nova, 100 kJ en 1 ns dans l'infrarouge, 40 kJ pour l'harmonique 3 dans l'ultra-violet. Les résultats obtenus sont encourageants, puisque des neutrons thermonucléaires ont été obtenus dans des cibles comprimées, et que la physique associée commence à être bien maîtrisée. On a maintenant la quasi-certitude

¹⁰Soit une puissance crête de 1000 TW, supérieure à la puissance électrique consommée en moyenne dans le monde, mais une puissance moyenne de 100 MW.

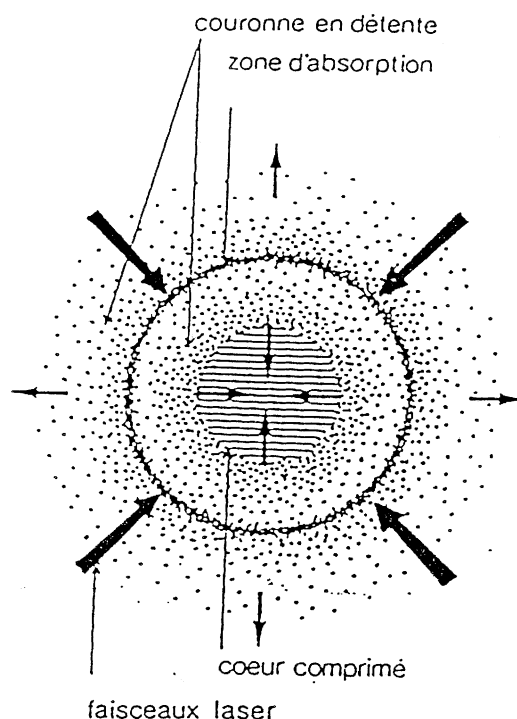


FIG. 4: *Compression inertielle : la bille contenant le combustible est irradiée par des faisceaux lasers convergents qui en volatilisent la couche superficielle. La compression résultante porte le cœur à des températures et des densités très élevées.*

qu'un laser de 10 MJ d'énergie devrait induire un gain convenable pour un réacteur.

Les problèmes à résoudre sont donc essentiellement d'ordre technologique et économique. La technologie actuelle des lasers à verre-Néodyme, qui devrait permettre de produire une telle énergie par impulsion, ne permet toutefois pas de dépasser une cadence de tir d'un coup toutes les heures et a un rendement extrêmement faible. Le pompage par lasers à semi-conducteurs de grande puissance devrait permettre d'améliorer le rendement et les performances du laser à verre-Néodyme. D'autres lasers sont étudiés, comme le laser KrF à excimère qui a l'avantage d'émettre directement dans l'UV et qui a un bon rendement. À long terme, le remplacement des lasers par des faisceaux d'ions semble une approche prometteuse.

Signalons aussi un autre intérêt de ce type d'installations : il permet de recréer les conditions physiques existant au cœur d'une bombe thermonucléaire, et donc de réaliser des tests qui ne présentent pas les mêmes inconvénients que les essais thermonucléaires en vraie grandeur. C'est la raison pour laquelle les USA, puis la France ont programmé la construction de lasers « Mégajoules ». L'installation française, prévue dans la région de Bordeaux et comprenant plus de 250 faisceaux, devrait fournir vers 2010 une puissance de 2 MJ environ dans l'UV. Ce type d'installation devrait permettre de réaliser à moyen terme la première expérience d'ignition de la cible, puis la démonstration d'un gain intrinsèque supérieur à 1. Il faudra donc encore plusieurs décennies de recherche avant de savoir si

les centrales électriques à fusion inertielle résoudront nos problèmes d'approvisionnement énergétique¹¹.

¹¹Pour plus de détails, on consultera par exemple *Physics Today* **45**, pp 32 et 42 (1992).

Complément III.4

Le laser source de lumière cohérente

Dans les applications du complément précédent, on utilisait uniquement l'énergie du faisceau laser. Ces applications ne sont pas les seules possibles. Le présent complément montre comment on peut tirer profit des propriétés de cohérence (spatiale ou temporelle) du rayonnement émis par un laser.¹

1 Les atouts du laser

1.1 Propriétés géométriques

Comme nous l'avons expliqué dans le complément III.2, la *cohérence spatiale* d'un faisceau laser lui assure une grande directivité. Il est de ce fait une excellente approximation du *rayon lumineux* de l'optique géométrique, et matérialise la droite, ou bien, focalisé à l'aide d'une lentille, le point. Nous verrons dans le paragraphe suivant tout le parti que l'on peut tirer de cette propriété. Bien entendu, à cause des lois de la diffraction, il ne s'agit pas vraiment de droite ou de point au sens mathématique : le faisceau lumineux diverge nécessairement, et le point lumineux a une extension minimale non nulle. Plus précisément, si le faisceau issu du laser est gaussien (voir Complément III.2.), caractérisé par son rayon au col w_0 , son demi angle de divergence α est donné par :

$$\alpha = \lambda / \pi w_0 \quad (1)$$

où λ est la longueur d'onde du laser. Donnons quelques ordres de grandeur pour fixer les idées : un laser Hélium-Néon standard, souvent utilisé dans ce genre d'application ($w_0 \approx 2$ mm, $\lambda = 0,6\mu$ m) a une divergence de l'ordre de 4×10^{-4} radians. Cette valeur peut paraître faible, mais elle donne tout de même une tache de 20 cm de diamètre au bout d'une distance de 300 m, ce qui est insuffisant pour certaines applications. On voit

¹En fait, ces diverses propriétés de la lumière laser ne sont pas indépendantes. Comme on l'a montré dans la partie E du chapitre III, la possibilité de concentrer l'énergie d'un faisceau laser est liée à la cohérence.

sur la formule (1) que pour réduire la divergence, il faut utiliser soit un laser de plus courte longueur d'onde, soit un faisceau de plus grand diamètre. Le rayon w_0 est initialement imposé par la géométrie de la cavité laser, mais on peut ensuite le modifier par un système optique, par exemple un système afocal, composé de deux lentilles de distances focales f_1 et f_2 ayant un foyer commun F (Figure 1). Le diamètre du faisceau est multiplié par le rapport f_2/f_1 . Pour des valeurs de ce rapport de l'ordre de 10 à 25, on obtient des faisceaux bien collimatés, utilisables jusqu'à des distances excédant le kilomètre. Sur de telles distances, d'autres facteurs commencent à entrer en jeu, liés à la variation de l'indice de l'air : courbure sous l'effet d'un gradient d'indice, ou bien diffusion par les turbulences atmosphériques.

Réciproquement, lorsque le faisceau laser, supposé parallèle, est focalisé à l'aide d'une lentille de distance focale f , le rayon au col w_0 du faisceau focalisé est donné par l'expression² :

$$w_0 = \lambda f / \pi w \quad (2)$$

En utilisant des lentilles dépourvues d'aberration et de grande ouverture numérique, comme par exemple les objectifs de microscope, éclairées sur toute leur surface utile par un faisceau laser élargi, on obtiendra une *tache focale qui aura des dimensions voisines de la longueur d'onde*. Ceci n'est vrai que pour des lasers fonctionnant sur le mode transverse fondamental TEM₀₀, ce qui n'est pas toujours facile à obtenir, surtout avec des lasers en impulsion.

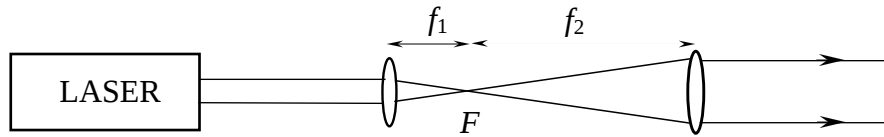


FIG. 1: Système afocal de réduction de la divergence du faisceau (par exemple, il peut consister en un télescope habituel, utilisé à l'envers).

1.2 Propriétés spectrales et temporelles

Dans le chapitre III, nous avons montré qu'un laser monomode est une source de rayonnement extrêmement monochromatique. On dit aussi qu'il a une grande *cohérence temporelle*. Le rayon laser est donc aussi la réalisation la plus parfaite de l'*onde lumineuse* monochromatique de l'optique ondulatoire.

On peut écrire tout champ électrique monochromatique $\mathbf{E}(\mathbf{r}, t)$ sous la forme :

$$\mathbf{E}(\mathbf{r}, t) = \mathbf{E}_0(\mathbf{r}) \cos(\omega t + \Phi(\mathbf{r}) + \varphi(t)) \quad (3)$$

²Cette expression n'est valable que si f est plus grand que la longueur de Rayleigh $\pi w^2/\lambda$ du faisceau initial. Pour les formules exactes de focalisation d'un faisceau Gaussien, voir par exemple A. Siegman *Lasers*, p. 675 (Oxford University Press 1986).

dans lequel $\Phi(\mathbf{r})$ caractérise la dépendance spatiale du front d'onde ($\Phi(\mathbf{r}) = -\mathbf{k} \cdot \mathbf{r}$ pour une onde plane, par exemple). Pour une source lumineuse quelconque, la phase $\varphi(t)$ varie de manière aléatoire, et ne reste constante que sur un temps τ , caractéristique de la cohérence temporelle de la source, qui a pour ordre de grandeur l'inverse de la largeur spectrale du rayonnement. Pour des sources lumineuses classiques, ce temps τ est très court. Pour les lasers, il peut atteindre la microseconde, ou plus. Les trains d'onde issus d'un laser peuvent ainsi avoir une *longueur de cohérence* très supérieure au kilomètre. En conséquence, il va être extrêmement facile avec des lasers d'observer des phénomènes d'optique ondulatoire, interférences ou figures de diffraction, même sur de grandes surfaces ou avec de *grandes différences de marche*. On reconnaît aisément un faisceau laser aux granulations (« tavelures ») que l'on observe lorsqu'il éclaire une surface quelconque, et qui ne sont autres que la figure d'interférence entre les différents rayons diffusés par les irrégularités de la surface.

De même, l'*holographie*, imaginée en 1948 par D. Gabor, ne s'est vraiment développée qu'après l'avènement du laser. Un hologramme permet de stocker à la fois les informations de phase et d'intensité de l'onde lumineuse (voir figure 2). Il est obtenu par enregistrement photographique de la *figure d'interférence résultant de la superposition de l'onde diffusée par un objet et d'une onde de référence* E_R . Lorsqu'on éclaire l'hologramme par un faisceau identique à l'onde de référence, l'onde diffractée par l'hologramme est identique, en amplitude et en phase, à l'onde initiale qui provenait de l'objet lui-même. On obtient ainsi une reproduction en *trois dimensions* de l'objet initial. L'holographie laser a de multiples applications, dont nous mentionnerons quelques unes dans la suite de ce complément.

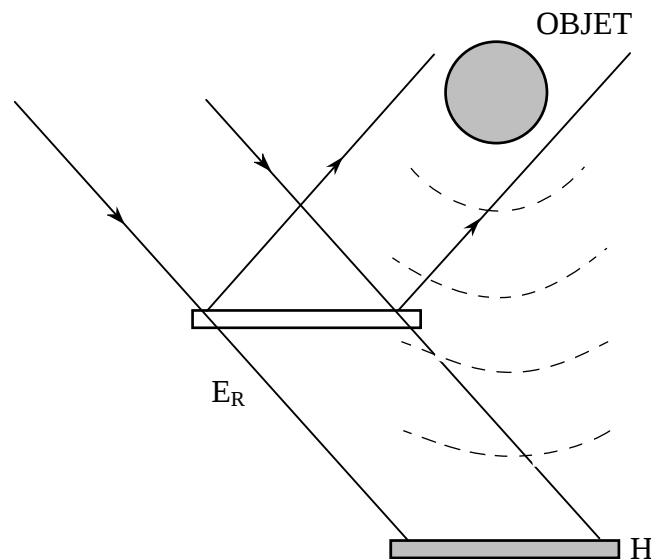


FIG. 2: Enregistrement d'un hologramme : la plaque H photographique est illuminée à la fois par l'onde de référence E_R et la lumière diffusée par l'objet, qui produisent un système d'interférences.

La pureté spectrale du laser a de nombreuses applications directes, soit pour discriminer deux systèmes dont les fréquences d'absorption sont très proches l'une de l'autre,

comme dans la séparation isotopique par laser (§ 4), soit pour mesurer un phénomène lié à la variation de fréquence de l'onde (vélocimétrie Doppler : § 2.3, gyrométrie : § 3). La monochromaticité est enfin utilisée en spectroscopie (voir Complément III.5).

Le laser peut aussi fonctionner en mode impulsionnel (voir Chapitre III.D) et produire des impulsions très courtes, jusqu'à la dizaine de femtosecondes. On dispose ainsi de « tops » extrêmement précis, utilisables pour mesurer des intervalles de temps très courts, ou bien pour transporter des « bits » d'information avec des débits très importants.

1.3 Manipulation du faisceau

À cause de son petit diamètre et de sa bonne collimation, un faisceau laser peut aisément être manipulé. Cet avantage, même s'il est moins fondamental que les précédents, est cependant à l'origine d'un grand nombre d'applications, à commencer par les spectacles lasers : pour un même éclairage, on peut utiliser des manipulateurs de taille plus faible qu'avec des sources classiques.

Les systèmes mécaniques, comme les miroirs mobiles ou tournants, sont à la fois les plus simples et les plus lents (fréquences caractéristiques jusqu'au kiloHertz). On peut aussi faire diffracter le faisceau laser sur un réseau de phase créé par une onde acoustique se propageant dans un cristal transparent. Ce sont les systèmes *acousto-optiques*, qui agissent sur le faisceau jusqu'à des fréquences de plusieurs centaines de MHz. On peut enfin utiliser la modification de l'indice d'un cristal sous l'effet d'un champ électrique oscillant (effet *électro-optique*), qui est efficace jusqu'à une dizaine de GHz. Ces dispositifs sont susceptibles de modifier à volonté les différents paramètres de l'onde laser : direction ou intensité, mais aussi phase ou polarisation. Ils sont le complément indispensable du laser dans la plupart des dispositifs que nous allons décrire ci-après.

Il faut enfin signaler que dans le cas particulier des lasers à semi-conducteur, il est possible de moduler l'intensité de sortie, même à des fréquences très élevées, simplement en faisant varier le courant d'alimentation de la diode. La grande commodité de cette méthode est un avantage supplémentaire des diodes lasers, extrêmement précieux par exemple dans le cas des télécommunications (voir paragraphe 4.3).

1.4 Conclusion

Nous avons tenté de donner dans ce premier paragraphe quelques indications sur les atouts du laser, afin que le lecteur puisse lui-même replacer dans ce contexte les diverses applications dont il pourra avoir connaissance. En ce sens, les quelques exemples qui suivent ne doivent être considérés que comme des illustrations. Il semble clair que les applications utilisant les propriétés géométriques et spectrales des faisceaux laser et leur facilité de manipulation vont continuer à se développer, dans une gamme allant des techniques les plus rudimentaires aux plus sophistiquées. Signalons par exemple que le projet

VIRGO de détection des ondes gravitationnelles par un interféromètre optique de plusieurs kilomètres de long emploie les lasers, non seulement dans son principe même, mais pour de nombreuses tâches annexes allant du contrôle des optiques en cours de fabrication à leur positionnement.

2 Mesures par laser

2.1 Mesure de position

Les lasers d'alignement (principalement des lasers Hélium-Néon) sont utilisés dans de nombreuses applications. La méthode la plus simple consiste à s'aligner de visu sur le faisceau lumineux, éventuellement élargi pour en réduire la divergence comme indiqué dans le paragraphe 1.1. La précision du pointé est considérablement augmentée si l'on utilise un photodétecteur à quadrants (voir figure 3), qui permet de déterminer le centre de la tache lumineuse à $10\mu\text{m}$ près. On peut ainsi asservir la position d'une machine sur le centre du faisceau. Cette technique est idéale pour les travaux publics (perçage de tunnels, construction d'autoroutes) où elle assure une conduite automatique des engins sur des distances de l'ordre de la dizaine de kilomètres, en ligne droite évidemment, mais aussi selon un tracé légèrement courbe, en déplaçant progressivement le point d'impact du faisceau sur la machine.

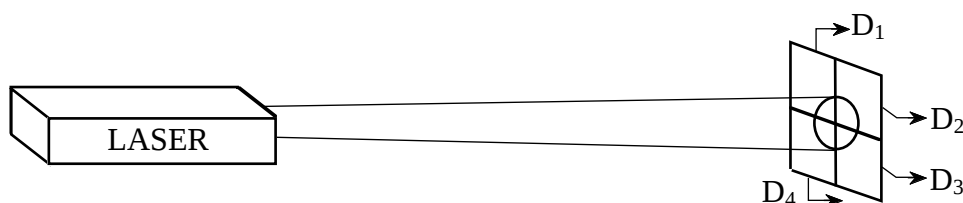


FIG. 3: Détecteur à quadrants : le faisceau est centré lorsque les signaux D_1 , D_2 , D_3 et D_4 sont égaux.

En utilisant un miroir de renvoi à 90° susceptible de tourner autour du faisceau initial, on peut matérialiser non plus une droite, mais un plan. Ce type d'appareillage est utilisé en construction navale, par exemple pour le contrôle des cuves de méthanier.

2.2 Mesure de distances

Les méthodes sont ici très variées, en fonction de l'ordre de grandeur de la distance de l'objet à l'observateur et de la précision exigée pour la mesure. Elles permettent de mesurer cette distance, ou bien des déplacements de l'objet.

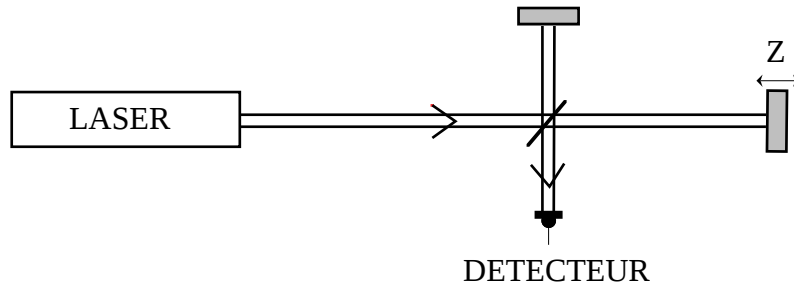


FIG. 4: Mesure interférométrique de déplacements : le détecteur compte les franges d'interférence liées au déplacement Z du miroir.

Pour la détermination de faibles déplacements avec une très grande précision, les techniques *interférométriques* sont sans rivales. L'objet dont on veut mesurer le déplacement porte un miroir qui ferme un des deux bras d'un interféromètre de Michelson (voir figure 4). On compte ensuite le nombre de franges qui défilent sur le détecteur D . Sur les machines-outil utilisées en mécanique de précision, on peut ainsi contrôler les déplacements de l'outil à $0,3\mu\text{m}$ près sur des distances de quelques dizaines de centimètres.

Si on demande une précision moindre, toujours sur de faibles distances, on peut faire appel à l'optique géométrique. On peut par exemple déterminer par triangulation la position de la tache lumineuse qu'un laser crée sur l'objet dont on veut déterminer la distance (figure 5.a). On peut aussi utiliser un faisceau laser focalisé sur l'objet à mesurer. On détecte la lumière rétrodiffusée, qui passe par un maximum lorsque l'objet est exactement dans le plan focal de la lentille (figure 5.b). Ce type d'appareillage présente l'avantage d'être simple, rapide, et de ne pas nécessiter de contacts avec l'objet, évitant ainsi toute déformation ou contrainte.

Pour des distances plus importantes, de l'ordre du kilomètre, on utilise couramment des méthodes de modulation³. Ce procédé, schématisé sur la figure 6, consiste à diriger vers l'objet le faisceau d'un laser continu dont l'intensité est modulée sinusoïdalement, et à comparer les phases des modulations des faisceaux aller et retour. Si l'intensité I de l'onde émise est de la forme :

$$I = I_0(1 + a \cos \Omega t) \quad (4)$$

le photodétecteur reçoit un signal I' égal à :

$$I' = I_1[1 + a \cos \Omega(t - 2L/c)] \quad (5)$$

où L est la distance à mesurer. Un comparateur de phase mesure la différence de phase $\Delta\varphi$ entre les modulations de I et I' :

$$\Delta\varphi = (2L\Omega/c) - 2n\pi \quad (n \text{ nombre entier}) \quad (6)$$

³Ces méthodes ne sont en définitive qu'une mesure du temps d'aller-retour de la lumière. Il s'agit de versions modernes de la roue dentée de Foucault.

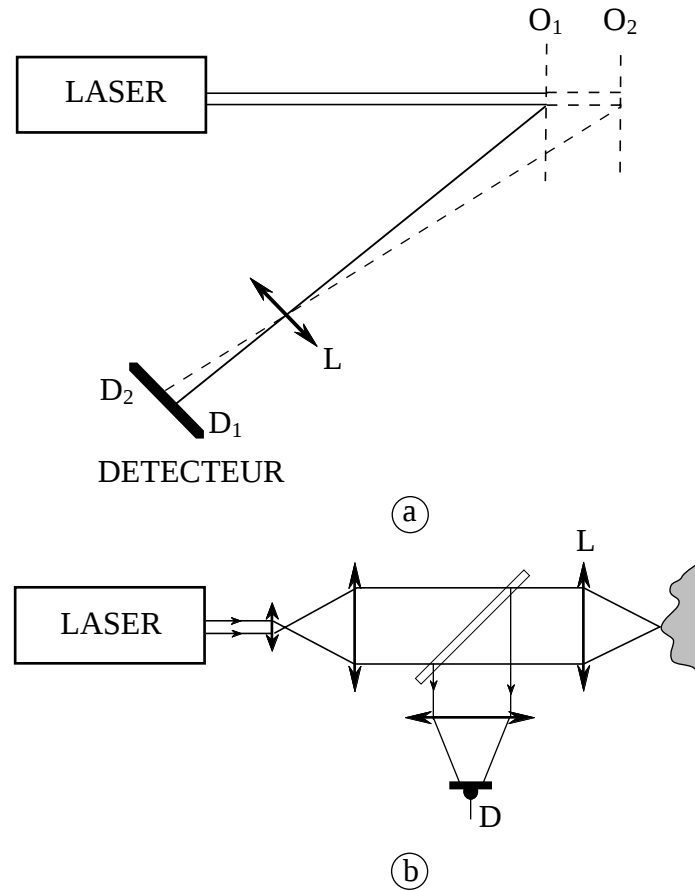


FIG. 5: Mesure de distance par optique géométrique : a) les positions O_1 et O_2 de l'objet donnent des images D_1 et D_2 dans le plan focal de la lentille L ; b) la lumière en D passe par un maximum lorsque l'objet est dans le plan de focalisation de la lentille L .

On obtient donc L par la formule :

$$L = \frac{\pi c}{\Omega} \left(\frac{\Delta\varphi}{2\pi} + n \right) \quad (7)$$

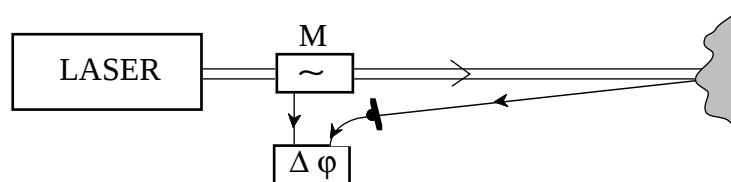


FIG. 6: Mesure de distance par modulation : on mesure le déphasage $\Delta\varphi$ entre la phase de l'oscillation imprimée par le modulateur M sur le faisceau laser et celle du signal rétrodiffusé par l'objet.

On détermine n en faisant deux mesures successives à deux fréquences de modulation

différentes. La précision d'une mesure de phase étant excellente (meilleure que le milliradian), de tels appareils atteignent des précisions de l'ordre de quelques mm par km pour des fréquences de modulation typiques de 10 MHz.

Il est enfin possible d'utiliser des méthodes impulsionnelles, dans lesquelles on mesure le temps d'aller-retour d'une impulsion laser très brève (≤ 10 ns) qui a été rétrodiffusée par l'objet dont on veut connaître la distance. Par rapport au radar habituel, basé sur le même principe mais utilisant des microondes, la télémétrie laser utilise un émetteur de taille plus réduite pour une précision identique, ce qui intéresse notamment les militaires pour connaître la position de cibles (précision de l'ordre de 10 m à quelques kilomètres). D'autre part, la divergence du faisceau étant beaucoup plus faible (voir équation 1), on peut mesurer des distances beaucoup plus importantes. L'exemple extrême est fourni par la mesure de la distance terre-lune. On envoie le faisceau d'un laser Nd : YAG en impulsion dans un télescope astronomique utilisé « à l'envers », c'est-à-dire par son oculaire. Le faisceau laser sortant par l'ouverture du télescope a alors un diamètre de l'ordre du mètre, et, en vertu de (1), il illumine sur la lune une région dont le diamètre n'est que de l'ordre de la centaine de mètres. La lumière renvoyée sur le télescope par des rétroreflecteurs (installés sur la lune lors de différentes missions spatiales) est très faible (moins de un photon par impulsion), mais détectable. La précision de la détermination de la distance est de quelques dizaines de centimètres. La prise en compte d'un grand nombre d'impulsions au cours des années permet de moyennner à une valeur très faible l'effet des turbulences atmosphériques et donc d'augmenter la précision de la mesure. On a pu ainsi déterminer que la lune s'éloignait actuellement de la terre à la vitesse de 3,2 cm par an !

2.3 Mesure de vitesses

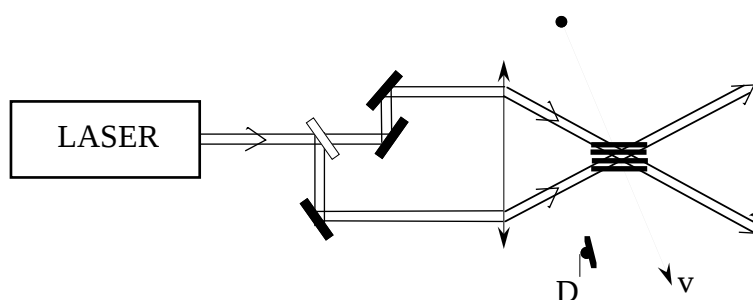


FIG. 7: Vélocimétrie laser : le laser crée un réseau lumineux au foyer de la lentille. Le détecteur D mesure la fréquence du clignotement de la lumière diffusée par l'objet de vitesse v .

Pour mesurer des vitesses perpendiculaires à la ligne de visée, on utilise le dispositif de la figure 7. Le faisceau d'un laser est séparé en deux parties puis recombinaé à l'aide d'une lentille dans la région où se déplace l'objet (de petites dimensions) dont on veut déterminer la vitesse. Dans la région de recouvrement des deux faisceaux, on obtient une figure d'interférences formée de zones alternativement sombres et éclairées parallèles à l'axe de la lentille. Une particule traversant ce réseau va diffuser la lumière laser de

manière périodique. On détermine la vitesse v_{perp} perpendiculaire aux plans d'interférence à partir de la fréquence ν_c de ce « clignotement » et de la valeur d de l'interfrange :

$$v_{\text{perp}} = d\nu_c \quad (8)$$

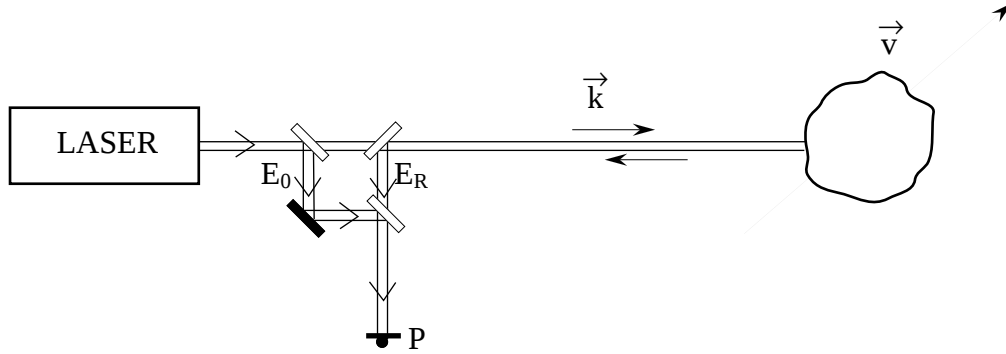


FIG. 8: Vélocimétrie Doppler : le photodétecteur P enregistre le battement entre l'onde directe E_0 et l'onde E_R rétrodiffusée par l'objet de vitesse $\mathbf{v}$. On mesure ainsi la composante de la vitesse suivant la ligne de visée.

Les vitesses parallèles à l'axe de visée peuvent être déterminées par vélocimétrie Doppler (voir figure 8). On utilise le fait que la lumière diffusée par un objet fixe a la même fréquence que la lumière incidente (diffusion élastique voir Chapitre II § A.4 et Chapitre VI § D). Si l'objet est animé d'une vitesse $\mathbf{v}$ par rapport à la source, la diffusion est élastique dans le repère lié à l'objet. La pulsation de l'onde incidente est décalée dans ce repère par *effet Doppler* d'une quantité $\delta\omega$ qui vaut :

$$\delta\omega = \mathbf{k} \cdot \mathbf{v} \quad (9)$$

où $\mathbf{k}$ est le vecteur d'onde de la lumière incidente. Revenant au repère fixe de l'observateur, on constate que la lumière diffusée vers l'arrière voit sa pulsation décalée une deuxième fois de la même quantité. Le photodétecteur P reçoit en même temps cette lumière rétrodiffusée $E_R \cos[(\omega + 2\delta\omega)t + \varphi']$ et une partie du faisceau initial $E_0 \cos(\omega t + \varphi)$. Il produit un signal électrique proportionnel à l'intensité totale I du champ électromagnétique incident :

$$I = \overline{\{E_0 \cos(\omega t + \varphi) + E_R \cos[(\omega + 2\delta\omega)t + \varphi']\}^2} \quad (10)$$

(la barre supérieure dénotant la moyenne temporelle sur quelques périodes optiques). Dans ce signal, les termes oscillant à des fréquences optiques sont moyennés par le détecteur, et le terme de battement $E_0 E_R \cos(2\delta\omega t + \varphi' - \varphi)$ oscille à une fréquence suffisamment basse pour tomber dans la bande passante du détecteur ($2\delta\omega/2\pi \approx 10$ MHz pour $v = 6$ m/s si on utilise un laser Hélium-Néon). La mesure de cette fréquence permet donc, à partir de (9), de déterminer la valeur de la projection de la vitesse de l'objet sur la ligne de visée. Cette technique nécessite un laser monomode, et sera d'autant plus précise que les fluctuations de phase du laser seront plus réduites. Elle est utilisée notamment en *médecine* pour déterminer la *vitesse de déplacement des globules sanguins*. Nous retrouverons dans le paragraphe consacré au gyrolaser un schéma de détection analogue pour déterminer des vitesses angulaires et non plus linéaires.

2.4 Contrôle non destructif

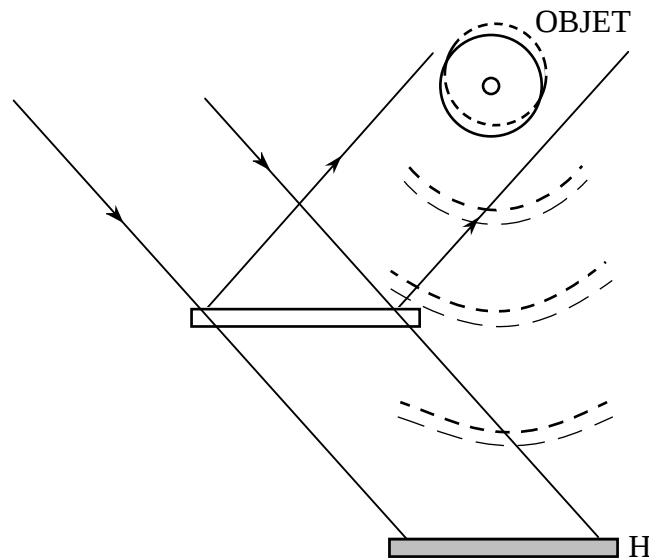


FIG. 9: Mesure de déformation par holographie : on observe à travers l'hologramme H l'interférence entre l'objet O déformé et l'enregistrement holographique du même objet non déformé.

Outre les positions et les vitesses, la mesure de la lumière diffusée par un objet illuminé par un laser permet de remonter à un grand nombre de paramètres intéressants, et cela de manière non destructive et non perturbatrice de l'objet. Ce type de mesure est utilisé sur un grand nombre de chaînes de fabrication (papeteries, sidérurgie ...) pour contrôler en ligne la qualité de fabrication. Citons par exemple :

- le contrôle de l'état de surface par inspection de la lumière diffusée ou diffractée ;
- la détermination de la dimension de grains (*granulométrie*) par analyse de la figure de diffraction.

Il est aussi possible de mettre en évidence en temps réel des déformations ou des contraintes par *interférométrie holographique* : le procédé est schématisé sur la figure 9. On enregistre l'hologramme H d'un objet O hors contraintes. On reproduit ensuite exactement le montage qui a permis cet enregistrement (faisceau laser, objet, hologramme). Toute déformation de l'objet O donne alors naissance à une *figure d'interférence entre l'onde diffusée par l'objet réel et l'onde diffractée par l'hologramme* (qui reproduit l'onde diffusée par l'objet dans son état antérieur). Des variations de géométrie de l'ordre de la longueur d'onde sont ainsi aisément détectables. On contrôle par cette méthode des pneumatiques, des structures d'avion, des maquettes d'éléments de construction dans les ponts et chaussées ...

2.5 Analyse chimique

En utilisant des lasers accordables, émettant à certaines fréquences d'absorption spécifiques de différentes molécules, on peut faire à distance des mesures de composition chimique⁴. Une méthode couramment utilisée consiste à faire de l'absorption différentielle, c'est-à-dire à comparer les intensités de deux faisceaux lasers L et L' transmis à travers un échantillon. Les fréquences légèrement différentes ν et ν' des faisceaux L et L' sont choisies de telle sorte que ν est une fréquence d'absorption de la molécule que l'on cherche à doser, et ν' est en dehors de la zone d'absorption.

Une autre méthode consiste à analyser la fluorescence du milieu en fonction de la fréquence de l'onde incidente. Compte-tenu de la spécificité des raies atomiques et moléculaires et de leur faible largeur spectrale dans le cas des atomes et des petites molécules, il est possible de détecter des éléments à l'état de *traces*. Ce type de technique spectrométrique est utilisé en usine pour le contrôle de fabrication, mais aussi à l'extérieur, sous le nom de LIDAR (Light Detection And Ranging), pour faire de la *téledétection dans l'atmosphère*. Le laser utilisé est en impulsion, et la mesure du temps d'aller-retour des impulsions permet de mesurer la distance des molécules qui ont rétrodiffusé sélectivement la lumière laser. On peut ainsi mesurer des concentrations de molécules inférieures à 10^{-8} sur une ligne de visée d'1 km. En déplaçant le faisceau et en utilisant plusieurs couples de fréquence, on peut ainsi faire la carte de la répartition de différentes espèces chimiques (H_2O , SO_2 , O_2 , $O_3 \dots$) dans une zone donnée. Le dispositif complet (laser et son alimentation, télescope de visée et de détection, traitement du signal rétrodiffusé), peut être installé à bord d'un camion, et ainsi être complètement autonome. Il permet de réaliser in situ de précieuses études de *pollution*. Par une analyse plus fine de la position et de la forme des raies de diffusion, on peut aussi déterminer la température locale et la vitesse de déplacement des nuages moléculaires. On a ainsi accès à des paramètres atmosphériques importants pour les études *météorologiques* ou *climatiques*.

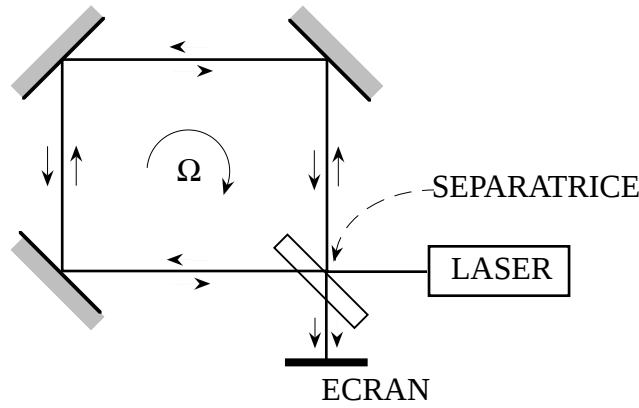
3 Le gyroscope laser ou « gyrolaser »

Nous décrivons plus en détail dans ce paragraphe une mesure laser particulière. Il s'agit de la mesure des *rotations par rapport à un référentiel d'inertie*.

3.1 Effet Sagnac

Cet effet, découvert par Sagnac en 1914, permet de mesurer par une méthode d'interférométrie optique la rotation d'un référentiel par rapport à tout référentiel d'inertie.

⁴Pour plus de détails, voir A. Briand, B. Fleurot, P. Mauchien, C. Moulin et B. Remy, Ann. Physique Colloque 2, Vol. 16, pp. 63-68 (1991), ou M.L. Chanin, « Les lasers et leurs applications scientifiques et médicales », p. 507, Les Éditions de Physique (1996).

FIG. 10 : *Interféromètre de Sagnac.*

Considérons l'interféromètre en anneau de la figure 10. Lorsqu'il est animé d'un mouvement de rotation de vitesse angulaire Ω par rapport à un référentiel d'inertie, il existe une différence de marche δx entre les deux faisceaux se propageant dans des sens opposés dans l'interféromètre, qui vaut :

$$\delta x = c\delta t = 4A\Omega/c \quad (11)$$

où δt est la différence entre les temps de propagation de l'onde dans les deux sens de rotation, et A la surface sous-tendue par le trajet des faisceaux lumineux dans l'interféromètre. Donnons de cette formule générale une démonstration simplifiée dans le cas d'un interféromètre de forme circulaire, animé d'un mouvement de rotation autour de son centre (voir figure 11). Un tel dispositif peut être réalisé avec une fibre optique.

Considérons un rayon lumineux entrant dans l'interféromètre à l'instant 0 au point C (symbolisant la lame séparatrice de la figure 10). Il se sépare en deux parties, l'une (1) tournant dans le sens direct, l'autre (2) dans le sens rétrograde. Au bout d'un tour, ces faisceaux sortent en des points respectivement appelés C' et C'' , différents pour l'observateur immobile puisque l'anneau, et donc la lame séparatrice, a tourné pendant le temps de propagation. Si R est le rayon du cercle, on peut exprimer de deux manières différentes les temps de révolution t_1 et t_2 des deux faisceaux avant leur sortie de l'interféromètre :

$$\begin{cases} t_1 = (2\pi R - CC')/c = CC'/R\Omega \\ t_2 = (2\pi R + CC'')/c = CC''/R\Omega \end{cases} \quad (12)$$

En éliminant les valeurs des arcs de cercle CC' et CC'' entre ces relations, on obtient la différence $t_2 - t_1 = \delta t$ des temps de propagation des deux faisceaux (1) et (2), mesurée dans le repère d'inertie :

$$\delta t = 4\pi R^2\Omega / (c^2 - (R\Omega)^2) \quad (13)$$

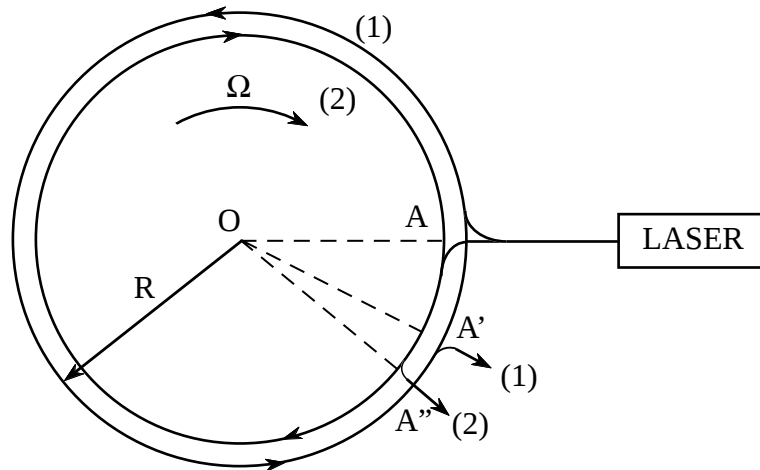


FIG. 11: Lors de la rotation à vitesse Ω de l'interféromètre, la lumière est réfléchiée par la lame séparatrice en des points C' et C'' différents selon que la lumière et l'interféromètre tournent dans le même sens ou dans des sens opposés.

À l'approximation non-relativiste⁵, les intervalles de temps sont les mêmes pour le référentiel en rotation et le référentiel fixe (il faut remarquer que la mesure se fait généralement dans le référentiel en rotation. Par exemple, dans le cas des figures 10 et 12, l'écran et le photodétecteur sont solidaires de l'interféromètre). On obtient dans cette limite :

$$\delta t = 4A\Omega/c^2 \quad (14)$$

puisque $A = \pi R^2$, ce qui redonne bien l'expression (11). On obtient un effet proportionnel à la vitesse angulaire, dont on peut montrer qu'il est indépendant de la forme de l'anneau et de la position du centre de rotation. Il en résulte un déplacement de frange égal à $c\delta t/\lambda$. Pour une surface A de 1m^2 , une longueur d'onde de $1\mu\text{m}$, la rotation terrestre ($15^\circ/h$) donne lieu à un déplacement relatif de frange de l'ordre de 10^{-6} . Pour mesurer une vitesse de rotation aussi faible, on peut augmenter la surface A , par exemple en utilisant une fibre optique bobinée sur un grand nombre de tours. Mais on peut aussi utiliser un système actif, le *gyrolaser*.

3.2 Le gyrolaser

Il s'agit d'un laser ayant une cavité en anneau dans lequel l'oscillation se produit en même temps sur deux modes (1) et (2) tournant dans la cavité dans des sens opposés (voir figure 12). On superpose les deux faisceaux de sortie sur un même photodétecteur⁶.

⁵Pour les effets relativistes voir par exemple W. Schleich and M.O. Scully « General Relativity and Modern Optics » dans « New Trends in Atomic Physics ». Les Houches XXXVIII édité par G. Grynberg et R. Stora (North-Holland, Amsterdam, 1984).

⁶Pour plus de détails sur les gyrolasers, voir par exemple W.W. Chow, J. Gea-Banacloche, L.M. Pedrotti, V.E. Sanders, W. Schleich et M.O. Scully Rev. Mod. Phys. **57**, 61 (1985).

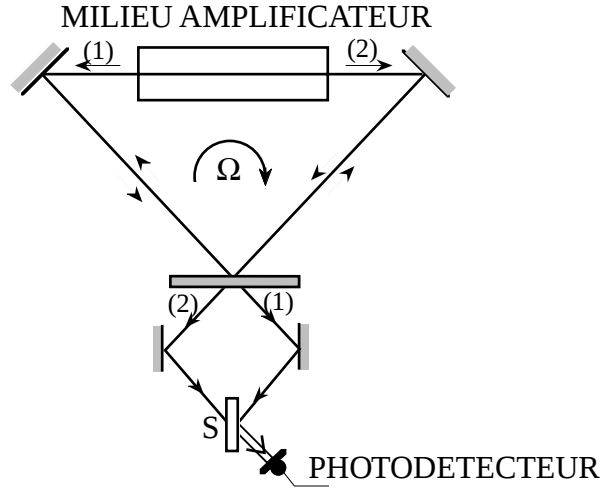


FIG. 12 : Schéma de gyrolaser.

La fréquence d'émission d'un laser étant un nombre entier de fois c/L , et la longueur effective L différant selon le sens de la rotation de la quantité $c\delta t$ calculée précédemment, la différence de fréquence $\delta\omega$ entre les deux faisceaux (1) et (2) est telle que :

$$\delta\omega/\omega = c\delta t/L \quad (15)$$

D'où :

$$\frac{\delta\omega}{2\pi} = \frac{4A}{\lambda L}\Omega \quad (16)$$

La fréquence de battement est proportionnelle à la vitesse de rotation Ω , et le facteur de proportionnalité ne dépend que des caractéristiques géométriques de la cavité laser : surface A et périmètre L . La quantité $4A/L\lambda$ est appelée « facteur d'échelle ». Son ordre de grandeur est environ L/λ . En prenant une cavité de l'ordre de 1 m, on trouve ainsi que $\delta\omega/2\pi$ est plus grand que Ω par un facteur de l'ordre de 10^6 . C'est ce facteur multiplicatif qui permet de mesurer les très faibles vitesses de rotation. Ainsi pour un gyrolaser carré de 1 m de côté fonctionnant à $0,6 \mu\text{m}$, la rotation terrestre provoque un écart en fréquence d'environ 100 Hz.

La mesure est effectuée par un photodétecteur qui mesure l'intensité lumineuse totale $(E_1(t) + E_2(t))^2$ obtenue par combinaison sur la séparatrice S des deux ondes émises par le laser (voir figure 12). Comme dans le cas de la vélocimétrie Doppler (paragraphe B.3), celle-ci est égale à la somme des intensités des deux champs, plus un terme de battement à la fréquence $\delta\omega$, fréquence qui peut être mesurée électroniquement avec une très bonne précision. On obtient ainsi un gyroscope entièrement optique. De tels dispositifs sont fabriqués industriellement et équipent un certain nombre d'avions, fusées ou missiles. Par rapport aux centrales à inertie classiques, utilisant des gyroscopes mécaniques, ils présentent l'avantage de ne comporter aucune pièce mécanique, donc d'être insensibles à tout problème de frottement et d'accélération, d'avoir une grande dynamique de mesure

et une bonne précision dans la détermination de Ω . En effet, celle-ci n'est limitée que par les fluctuations intrinsèques affectant la grandeur $\delta\omega$: les phénomènes de vibration, de dilatation, ou autres qui modifient la longueur L du laser affectent de manière *identique* les fréquences d'émission ω_1 et ω_2 . En définitive, la précision de la mesure n'est limitée que par la largeur ultime de la raie laser (largeur Schawlow-Townes, voir complément III.9). La précision de mesure peut ainsi atteindre le MHz, rendant possible une précision de détermination de Ω meilleure que $10^{-3}^\circ/h$.

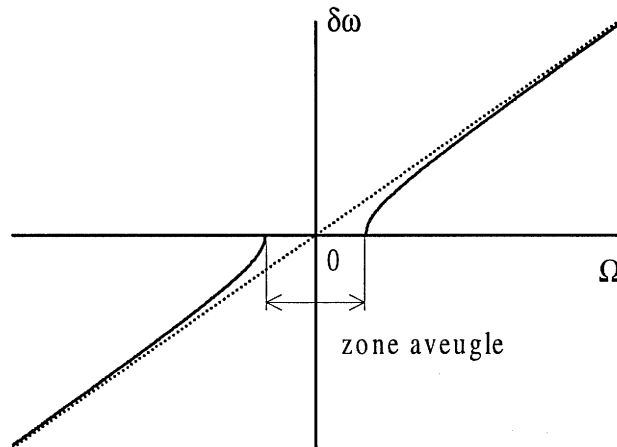


FIG. 13: Variation de la fréquence du battement $\delta\omega$ en fonction de la vitesse de rotation Ω aux faibles vitesses de rotation.

L'utilisation du gyrolaser est cependant contrariée par un phénomène qui l'empêche de mesurer les très faibles vitesses de rotation : il s'agit du *verrouillage de fréquence* entre les modes tournant en sens opposés dans la cavité. À cause de l'existence d'une très faible rétrodiffusion de la lumière sur les miroirs de la cavité, il existe un couplage mutuel entre les deux modes qui peuvent se verrouiller l'un sur l'autre, et osciller à la même fréquence (voir figure 13). Il existe donc une « zone aveugle », dans laquelle $\delta\omega$ reste nul même si Ω n'est pas nul. Une solution à ce problème consiste à moduler Ω (par exemple en faisant vibrer le support du laser) de manière que le gyrolaser passe un temps minimum dans la zone aveugle. On arrive ainsi à mesurer des rotations inférieures à $10^{-2}^\circ/h$, permettant d'utiliser le gyrolaser pour la navigation inertielle en aéronautique⁷.

4 Le laser porteur d'information

On peut utiliser le laser comme *moyen de transport d'une information*, qui est codée en modulant un paramètre du faisceau (son intensité la plupart du temps). Comme nous allons le voir par la suite, grâce aux lasers, on peut lire, écrire, ou transporter de l'information, voire même traiter cette information.

⁷Notons cependant qu'à la fin des années 90, les gyroscopes mécaniques ont encore une stabilité à très long terme (plusieurs mois) inégalée, et sont toujours utilisés pour les sous-marins en plongée.

4.1 Lecture et reproduction

Nous avons indiqué dans le paragraphe A.3 comment on pouvait défléchir et moduler un faisceau laser. Ce dernier peut donc être utilisé de la même manière que le faisceau électronique d'un tube cathodique. Comparé à cette technique largement répandue, il a certains avantages : il fournit une meilleure image, à cause d'un rapport signal sur bruit supérieur sur l'intensité du faisceau, et surtout à cause d'une bien meilleure résolution (on peut focaliser le faisceau laser sur une surface bien plus petite qu'un faisceau d'électrons) ; il se propage dans l'air et non pas dans le vide. En revanche, il a l'inconvénient d'un prix de revient plus élevé, surtout si on veut rendre les couleurs, car on a alors besoin de 3 lasers de couleurs différentes (Rouge, Vert, Bleu).

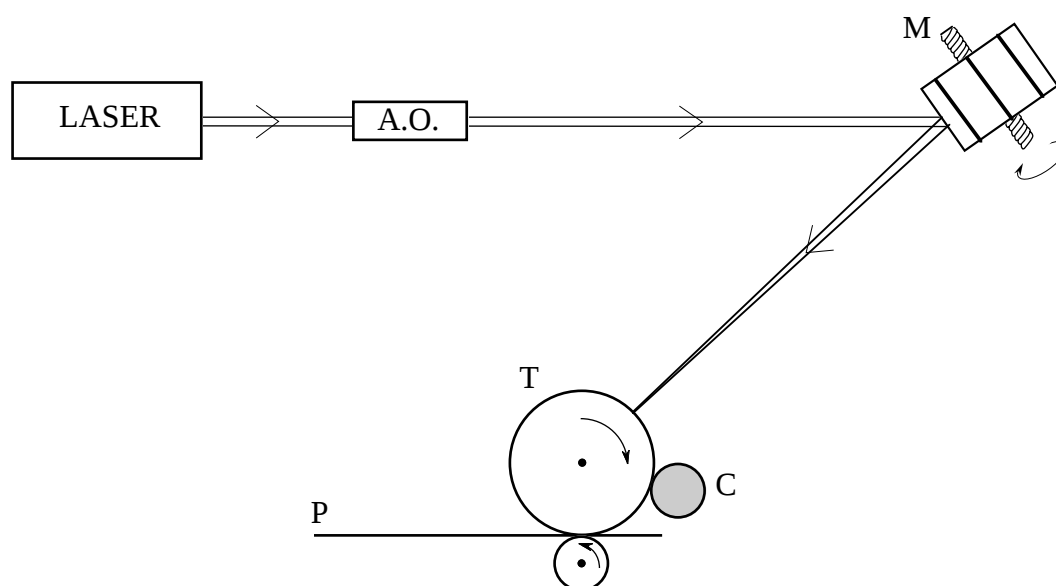


FIG. 14: Schéma d'une imprimante laser : le faisceau laser modulé par l'acousto-optique AO est balayé par le miroir tournant M sur le tambour photoconducteur T. L'image latente est révélée par la poudre de carbone C et transférée sur la feuille de papier P.

En ce qui concerne la reproduction de documents, nous citerons *l'imprimante laser*, qui est un produit de grande diffusion. Son principe est schématisé sur la figure 14. Le faisceau d'une diode laser, modulé en intensité par un système acousto-optique (AO), et balayé par un miroir polygonal tournant (M), est focalisé sur un tambour photoconducteur (T). L'irradiation laser donne naissance, par un processus de photoconduction, à une modification de la charge électrique du tambour sur une surface de $100\mu\text{m}$ de diamètre environ. On crée ainsi point par point, et sous forme électrostatique, une « image latente » sur le tambour T. Le reste du processus est commun avec les photocopieuses ordinaires (pour lesquelles l'image latente est créée par imagerie ordinaire à partir du document à reproduire) : l'image est « révélée » par une poudre de carbone (C) attirée par les zones chargées, puis transférée sur du papier (P) et fixée. Ces imprimantes sont maintenant universellement adoptées en raison de leur rapidité (jusqu'à 4 000 lignes balayées par

seconde!) et de leur excellente résolution, qui leur permet d'imprimer à la fois des textes avec une grande variété de polices de caractères, des graphismes et des images.

En ce qui concerne les opérations de lecture, il faut d'abord citer la lecture des « codes-barre ». Un faisceau laser est balayé à grande vitesse ($2m/s$) sur l'étiquette, et la lumière rétrodiffusée mesurée sur un photodétecteur fournit une succession de niveaux haut et bas qui constitue un code permettant d'identifier le produit. Cette application simple du laser lui fournit un de ses marchés les plus importants : 150 000 lasers Hélium-Néon étaient utilisés en 1990 dans de tels systèmes. Ceux-ci sont progressivement remplacés par des diodes-lasers fonctionnant dans le visible.

Le laser permet de lire un support particulier, le « disque optique », qui est capable de stocker un grand nombre d'informations dans un petit volume. Nous consacrons le paragraphe suivant à une description plus détaillée de cette application importante.

4.2 Mémoires optiques

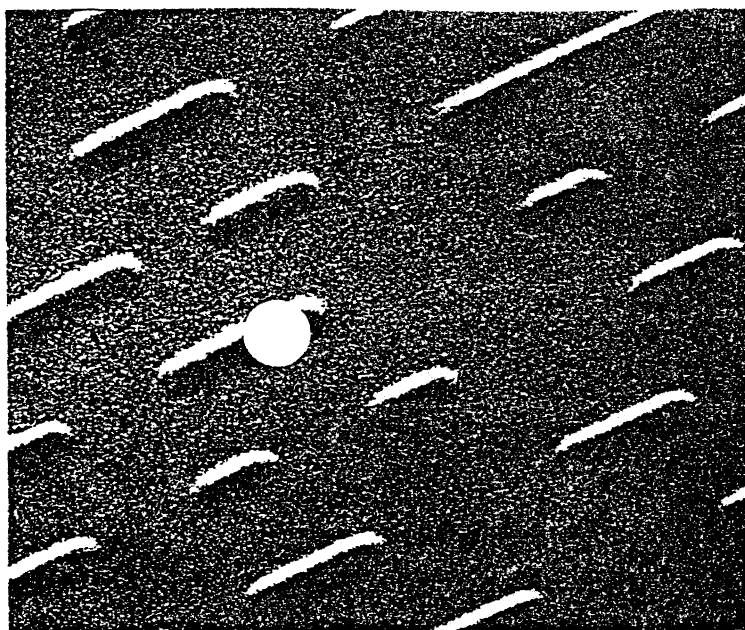


FIG. 15: Image de la surface d'un disque compact au microscope à balayage. La tache circulaire représente la tache focale du faisceau laser.

Comme il a été expliqué dans le paragraphe 1.1, le faisceau d'un laser peut être focalisé sur une surface dont la dimension est de l'ordre du micromètre. De plus, comme nous le verrons plus loin, la lecture de l'information se fait par un système optique capable de suivre avec une grande précision les pistes, qui peuvent donc être très resserrées. On obtient ainsi un moyen de stocker l'information avec une très grande densité, jusqu'à 100 Mbit/cm^2 , à comparer avec la densité de 1 Mbit/cm^2 obtenue sur le « disque dur »

magnétique des micro-ordinateurs ($1,86 \text{ Mbit/cm}^2$ pour les meilleurs systèmes magnétiques au début des années 90). Les disques optiques commerciaux ont une densité de 53 Mbit/cm^2 , et peuvent contenir l'équivalent de 130 000 pages de texte ou 54 000 images de télévision. De plus, grâce à la faible inertie du système de lecture, la lecture optique de cette information peut se faire avec des temps d'accès très courts.

Sur les *disques compacts* actuellement commercialisés, l'information est stockée sous la forme de « cuvettes » de largeur $0,6 \mu\text{m}$, de profondeur $0,1 \mu\text{m}$, avec un espace interspire de $1,6 \mu\text{m}$ (Figure 15). Pour la lecture de l'information, on utilise le faisceau d'une diode laser focalisé sur la surface du disque à l'aide d'une lentille de très courte focale. On mesure sur une photodiode l'intensité réfléchie par la surface du disque (Figure 16). La profondeur de champ au foyer de la lentille étant de l'ordre de $2 \mu\text{m}$, cette intensité varie fortement lorsque le faisceau laser arrive sur le bord d'une cuvette creusée dans le disque. C'est cette variation brutale sur le bord des cuvettes qui fournit le bit « 1 », alors qu'un niveau constant d'intensité réfléchie correspond au bit « 0 ».

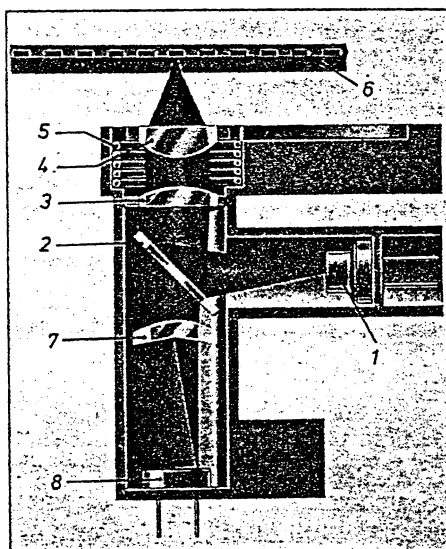


FIG. 16: Schéma du système optique d'un lecteur de compact disque : 1) Diode laser ; 2) lame séparatrice ; 3) 4) Lentilles de focalisation ; 5) Système de déplacement de la tache focale ; 6) Disque optique ; 7) Système optique astigmatique ; 8) Photo détecteur à quadrant.

Pour obtenir une lecture fiable de l'information, le positionnement du système optique par rapport au disque doit être réalisé avec une précision de $0,2 \mu\text{m}$, ce qui exclut tout dispositif mécanique passif, même de grande précision. On utilise donc un système électronique actif qui, grâce à des miroirs montés sur des électroaimants, asservit les trois coordonnées du point de focalisation du laser. En ce qui concerne la distance système optique-surface du disque, on utilise une lentille cylindrique qui rend le système astigmatique, c'est-à-dire que la forme de la tache focale n'est pas circulaire et varie très rapidement en fonction de cette distance. La lumière réfléchie est alors collectée par un photodétecteur à quadrants (analogue à celui décrit dans le paragraphe 2.1) qui est sensible à cette variation de forme et fournit le signal d'erreur nécessaire à l'asservissement. L'asservissement radial sur le « sillon » du disque est souvent

assuré grâce à l'adjonction de deux pinceaux de lumière latéraux. La mesure différentielle des signaux rétrodiffusés à partir de ces derniers fournit le signal d'erreur permettant de centrer le pinceau principal avec la précision requise. Quant au positionnement tangentiel, il est obtenu par asservissement de la vitesse de rotation du disque sur les tops de synchronisation du signal enregistrés sur le disque. C'est donc la vitesse linéaire de défilement, et non la vitesse angulaire de rotation du disque, qui est maintenue constante (et égale à 1,25 m/s dans le cas du disque compact). Le débit brut obtenu est de 4,2 Mbit/s.

Un tel système ne permet que la lecture de l'information stockée au préalable sur le disque (ROM : Read Only Memory). Il existe aussi des systèmes « WORM » (Write Once, Read Many), dans lesquels l'information est inscrite sur le disque grâce au laser lui-même, par ablation d'une mince couche superficielle par exemple. Des disques « WMRM » (Write Many, Read Many) sont maintenant commercialisés par différents constructeurs. Ils sont basés sur l'utilisation d'effets magnéto-optiques : pour écrire l'information, le laser provoque localement une transition du matériau magnétique au-dessus de son point de Curie, ce qui permet, lors du refroidissement dans un champ magnétique extérieur, d'inverser l'aimantation de cette région (bit « 1 »). Pour effacer, il suffit de changer le sens du champ magnétique et de maintenir constante l'intensité du laser. Quant à la lecture, elle est assurée par la mesure de la rotation de la polarisation de la lumière réfléchie par le matériau aimanté (effet Kerr magnéto-optique), en utilisant pour la lecture un faisceau laser polarisé moins puissant que le faisceau laser d'écriture. Cette rotation est très faible (de l'ordre de 1°). Elle peut néanmoins être repérée avec un excellent rapport signal sur bruit, qui permet d'avoir un taux d'erreur de lecture (probabilité de mesurer un bit 0 alors qu'un bit 1 avait été inscrit sur le disque) extrêmement faible. Un tel système marie la souplesse d'utilisation des systèmes magnétiques avec la grande capacité potentielle des systèmes optiques (4 Gbits environ, soit 400 Mégaoctets sur un disque de 5"1/4).

4.3 Télécommunications

Les télécommunications par voie optique sont basées sur l'utilisation de quatre composants essentiels :

- le laser à semi-conducteur, décrit dans le chapitre III § B.3, qui produit un rayonnement de faible largeur spectrale qui peut être facilement modulé en intensité, par une simple variation du courant d'alimentation ;
- la fibre optique en silice à faible taux d'absorption (inférieur à 0,2 dB, soit 5 %, par kilomètre de fibre aux alentours de $1,5\mu\text{m}$ de longueur d'onde), qui peut être produite en tronçons de très grande longueur (plusieurs kilomètres et plus) ;
- l'amplificateur tout optique, qui permet de régénérer le signal affaibli par les pertes au cours de sa propagation ;
- la photodiode (silicium ou germanium), qui transforme les photons incidents en paires électrons-trous, donc en courant électrique, et cela avec un excellent rendement et un bruit de fond minime.

La mise au point de ces éléments, produits à faible coût et en grande quantité, a permis le développement rapide de ce type de communication. Par rapport aux systèmes classiques de transmission (câble ou faisceau hertzien), les avantages sont la plus grande

fréquence de la porteuse, qui permet de lui faire transporter beaucoup plus de canaux ou d'en augmenter la bande passante, et l'insensibilité de la liaison aux perturbations extérieures et aux parasites d'origines diverses.

À l'heure actuelle, toutes les liaisons mises en service sont basées sur un codage en amplitude de l'information : par exemple, le bit « 1 » correspond à une intensité donnée, et le bit « 0 » à une intensité nulle. Le débit de l'information, lié à l'intervalle de temps minimum entre bits successifs, est limité par la *dispersion chromatique de la fibre*. En effet la source lumineuse émettrice n'étant pas rigoureusement monochromatique, ses différentes composantes spectrales se propagent à des vitesses différentes, ce qui entraîne un élargissement temporel des impulsions. Pour minimiser cet effet, il faut, ou bien se placer dans la région spectrale où la dispersion chromatique de la silice est minimale, c'est-à-dire à $1,3\mu\text{m}$, ou bien utiliser une source extrêmement monochromatique⁸. À titre d'exemple, la liaison optique sous-marine transatlantique installée en Mars 1992 utilise des diodes lasers à $1,55\mu\text{m}$, a un débit d'information par fibre de 560 Mbit/s (7680 voies téléphoniques). La puissance initiale de la source de 1 mW et la faible absorption de la fibre permettent une portée de propagation sans répéteur de 130 km (1,5 km avec les câbles coaxiaux classiques).

Les techniques sont en développement rapide, et le record de capacité des liaisons optiques expérimentales est battu régulièrement. Diverses voies sont explorées :

- Le multiplexage spectral permet l'injection sur une même fibre de signaux à des longueurs d'onde distinctes.
- L'utilisation de diodes lasers monomodes, éventuellement asservies sur des cavités extérieures, permet de réaliser des débits de plusieurs *Gbit/s* sans être gênés par la dispersion. De telles sources extrêmement monochromatiques rendent aussi possibles des techniques de transmission « cohérentes », dans lesquelles l'information est codée sur la phase de l'onde émise et non l'intensité. La lecture de l'information est assurée en mélangeant le signal avec le faisceau d'un autre laser, qui joue le rôle d'oscillateur local (détection homodyne ou hétérodyne). On gagne ainsi un facteur 100 sur la capacité, au prix d'une complication importante des dispositifs d'émission et de réception.
- La généralisation de l'utilisation de répéteurs purement optiques en remplacement des répéteurs « classiques » (qui comprennent une photodiode, un amplificateur électronique, et une diode laser). On utilise pour cela le phénomène d'amplification optique par émission stimulée. Des *amplificateurs à fibre dopée à l'Erbium* (voir chapitre III, § B.4), pompés par un rayonnement à 950 nm injecté dans la fibre, présentent un gain optique allant jusqu'à 10^3 à la longueur d'onde de $1,5\mu\text{m}$.
- L'utilisation d'*impulsions solitons*, qui présentent l'avantage de ne pas s'élargir lorsqu'elles se propagent. La tendance à l'élargissement par dispersion chromatique est en effet contrebalancée par l'effet non-linéaire d'automodulation de phase (voir Complément III.6, § D.3).

⁸Dans l'état actuel de la technologie, les impulsions de codage n'ont pas une durée Δt suffisamment brève, et une largeur spectrale $\Delta\nu$ suffisamment faible pour atteindre la « limite de Fourier » $\Delta\nu\Delta t \approx 1$.

D'une manière plus générale, tous les développements récents de l'optique non-linéaire et de l'optique quantique (voir Compléments III.6, III.7, VI.1 et Chapitre V) trouvent des applications naturelles dans le domaine des télécommunications optiques. De nombreux laboratoires de recherche dans le monde consacrent d'importants efforts à l'étude de ces applications. On a annoncé en 1999 des débits de 1 Terahertz (10^{12} bits/s) sur une distance de 1000 km. On se rapproche ainsi de la limite théorique, de l'ordre de la fréquence de la porteuse (2×10^{14} Hz).

4.4 Vers l'ordinateur optique ?

Nous venons de voir que le laser apporte des solutions intéressantes aux problèmes de stockage et de transmission de l'information. On peut alors se poser la question du traitement optique de l'information : est-il possible de réaliser un ordinateur « tout-optique », dans lequel les photons remplaceraient les électrons ?

Le traitement optique de l'information a un grand avantage potentiel : la facilité de réaliser des *architectures massivement parallèles*. À la différence des courants électriques, les faisceaux lumineux n'ont pas forcément besoin de fils pour se propager, et de plus ils n'interagissent pas entre eux lorsqu'ils se croisent. Il est alors possible d'envisager des adressages bi-dimensionnels entre un réseau plan de diodes-lasers et un réseau plan de photodétecteurs par des techniques d'imagerie optique, holographiques en particulier. De tels dispositifs sont appelés à révolutionner de nombreux concepts en informatique, qui s'est pour l'instant développée à partir d'un traitement élémentaire de l'information de type séquentiel, imposé par la technologie même des composants électroniques.

D'autre part, il est possible de réaliser des *portes logiques* purement optiques grâce à des techniques d'optique non-linéaire, comme la *bistabilité optique* (voir complément III.6 § B). En l'état actuel de la technologie, leur fonctionnement consomme une énergie non négligeable, ce qui rend prohibitif le fonctionnement en parallèle d'un très grand nombre de ces dispositifs. Les démonstrations en laboratoire d'« ordinateurs optiques » ont ainsi mis en jeu, jusqu'à présent, un très faible nombre de portes logiques. De plus, le gain en temps de commutation, donc en vitesse de calcul, n'est pas considérable par rapport aux composants purement électroniques les plus performants.

L'ordinateur purement optique aura donc du mal à s'imposer à court terme. Par contre les interconnexions optiques dans l'ordinateur semblent une voie prometteuse.

5 Séparation isotopique par laser

En raison de sa grande monochromaticité, le laser apporte une solution élégante au problème de la séparation isotopique (essentiellement celle de l'Uranium 235 et de l'Uranium 238). Pour comprendre cette technologie, il faut tout d'abord expliquer comment

la nature isotopique d'un noyau influe sur la position des raies spectrales de l'élément correspondant.

5.1 Effets isotopiques sur les niveaux atomiques

Ils ont trois origines différentes, que nous allons décrire successivement.

a. Effet de masse

Considérons tout d'abord le cas simple de l'atome d'hydrogène, comportant un noyau de masse M et un électron de masse m . Le problème est traité dans la plupart des ouvrages de Mécanique Quantique⁹. Les niveaux d'énergie s'expriment en fonction de la constante de Rydberg R_M donnée par :

$$R_M = \frac{\mu q^4}{32\pi^2 \varepsilon_0^2 \hbar^2} \quad (17)$$

où μ est la masse réduite, définie par :

$$\mu = \frac{Mm}{M+m} \quad (18)$$

Lorsqu'on passe d'un isotope à un autre, la masse M varie. Il en résulte une variation de la constante de Rydberg, et donc de l'ensemble des fréquences de transition, que l'on peut écrire sous la forme suivante, sachant que $m/M \ll 1$:

$$\frac{\Delta\nu}{\nu} = \frac{\Delta\mu}{\mu} \approx \frac{m}{M} \frac{\Delta M}{M} \quad (19)$$

Ainsi, lorsqu'on passe de l'hydrogène au deutérium, la variation relative des fréquences de transition est de 2×10^{-4} , ce qui est loin d'être négligeable.

Pour les atomes à plusieurs électrons, et en particulier pour l'uranium, l'effet isotopique de masse est beaucoup plus compliqué, mais l'ordre de grandeur donné par la formule précédente demeure. Cette formule indique que la variation relative des fréquences de transition décroît rapidement avec la masse du noyau : pour l'uranium, cette contribution au déplacement isotopique est négligeable devant l'effet de volume.

b. Effet de volume

Dans un atome, ce n'est qu'en première approximation que le noyau peut être considéré comme *ponctuel*. Son volume *non nul* influe en fait légèrement sur la position des niveaux d'énergie. À cause du théorème de Gauss, cette non-ponctualité ne modifie la valeur du champ par rapport au champ Coulombien qu'à l'intérieur du volume du noyau. Seuls les états stationnaires de l'atome dont la probabilité de présence sur le noyau est importante,

⁹Voir par exemple BD Chapitre X ou CDL Complément A.VII. L'effet de « masse réduite » existe aussi en mécanique classique.

c'est-à-dire les états S , seront sensibles à cet effet correctif, et ce d'autant plus que le noyau est plus gros, donc plus lourd. Pour l'Uranium, l'état S fondamental subit un déplacement d'environ 10 GHz entre les isotopes 235 et 238. Toutes les transitions optiques partant de ce niveau seront décalées de cette grandeur entre les deux isotopes.

c. Effet magnétique

Les différents isotopes d'un élément ont en général des spins nucléaires I différents, et donc des moments magnétiques différents. Il en résulte que leur *structure hyperfine*¹⁰ est différente. En particulier l'uranium 238 a un spin $I = 0$, donc pas de structure hyperfine, tandis que l'uranium 235, de spin $I = 7/2$, voit son niveau fondamental décomposé en plusieurs sous-niveaux.

5.2 Spectre de l'Uranium

À cause de la présence de 92 électrons, le spectre de l'uranium est beaucoup plus compliqué que celui de l'atome d'hydrogène. On a pu ainsi y répertorier près de 300 000 raies ! La figure 17.a montre une petite partie de ce spectre, mesuré avec une résolution spectrale moyenne. À cette échelle les spectres des deux isotopes ne montrent aucune différence. Il en va tout autrement si l'on effectue une spectroscopie à haute résolution d'une petite partie du spectre précédent (Figure 17.b) : on observe une raie unique pour l'isotope 238, séparée d'environ 0,01 nm de l'ensemble des raies hyperfines de l'isotope 235. On voit donc qu'il est possible, avec un laser de fréquence judicieusement choisie et dont la largeur spectrale n'excède pas quelques GHz, de *porter sélectivement un isotope dans un état excité*, l'autre restant dans l'état fondamental. La sélectivité est ici de 100 %, et non pas quelques %, comme pour les autres effets physiques utilisés pour la séparation isotopique.

5.3 Méthode de séparation

Elle est basée sur *l'ionisation sélective par échelons* de l'atome d'Uranium (voir figure 18). La première étape est l'excitation sélective décrite précédemment. Elle porte uniquement l'isotope 235 dans l'état excité a . La deuxième étape porte l'atome dans un niveau b plus proche de la limite d'ionisation. La troisième étape permet de franchir cette limite et d'ioniser l'isotope choisi. Les deux premiers échelons, ayant un caractère résonnant, ont une très forte section efficace (de l'ordre du carré de la longueur d'onde utilisée, soit 10^{-14} cm², voir Complément IV.1) et demandent une faible puissance lumineuse. La troisième n'a pas de caractère résonnant ; elle est donc nettement moins efficace (10^{-18} cm², voir Complément II.3), mais son rendement peut être nettement augmenté (10^{-15} cm²) si l'on choisit la longueur d'onde du laser de façon que l'état final soit un *niveau autoionisant* (voir Complément I.1), qui conserve de nombreuses caractéristiques d'un état lié.

¹⁰Voir par exemple BD Chapitre XIII ou CDL Chapitre XII.

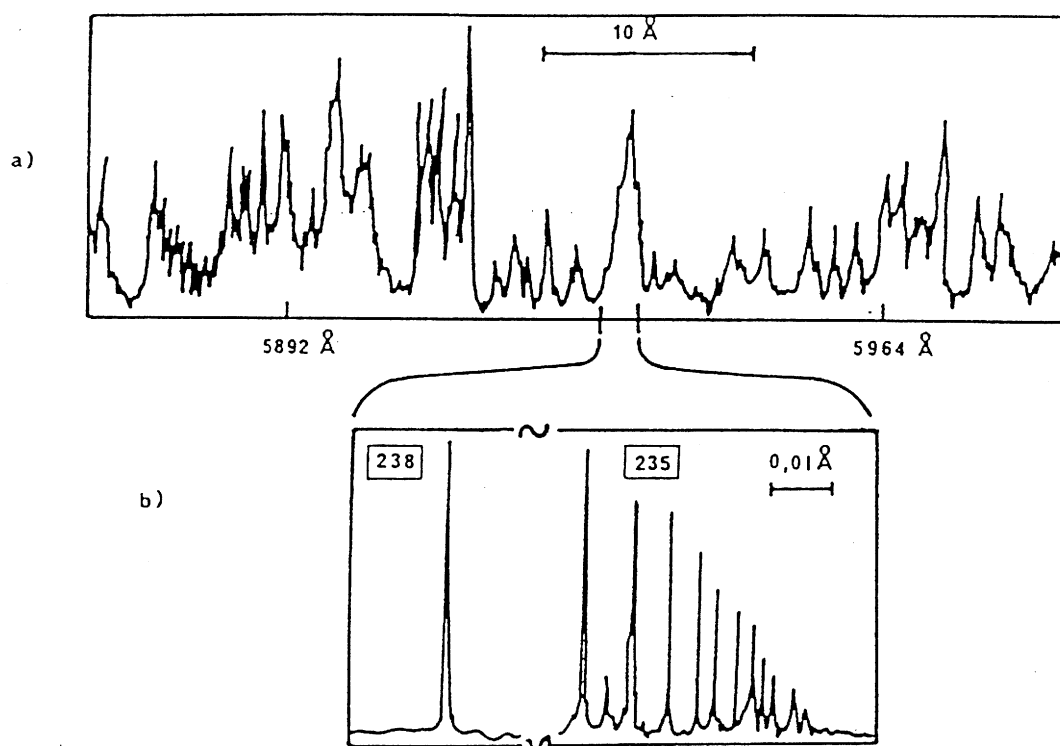


FIG. 17: Spectre de l'uranium : a) à faible résolution; b) à haute résolution. On remarque la présence de nombreuses raies de structure hyperfine pour l'isotope 235.

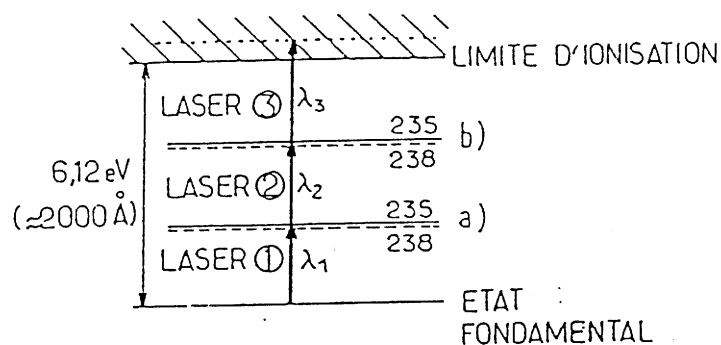


FIG. 18 : Ionisation sélective en trois échelons de l'atome d'uranium.

Les trois échelons sont assurés par des lasers accordables à colorant, en impulsion (voir chapitre III § B.2.d), de largeur spectrale voisine de 1 GHz, pompés par des lasers en impulsion à vapeur de cuivre (dont les raies à 530 nm et 570 nm sont bien adaptées au pompage des colorants), de puissance moyenne de l'ordre de 100 W, qui présentent l'avantage d'avoir un taux de répétition élevé et un excellent rendement. Les faisceaux de ces trois lasers, rendus colinéaires, excitent transversalement un jet atomique d'uranium (voir Figure 19), créé par la vaporisation par bombardement électronique d'un lingot d'uranium. Les atomes d'uranium 238 se propagent en ligne droite et sont recueillis sur le collecteur A, tandis que les atomes d'uranium 235, ionisés,

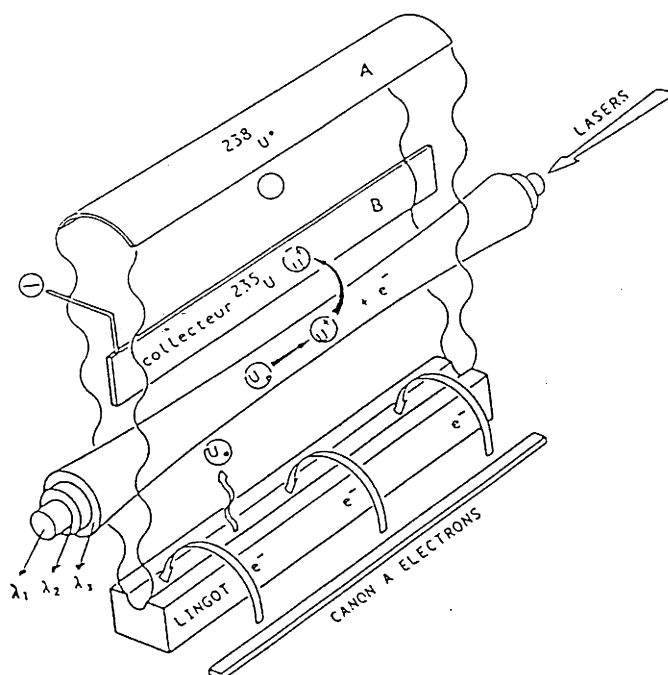


FIG. 19: Schéma du dispositif de séparation isotopique par laser : l'uranium 238 se dépose sur le collecteur A et l'uranium 235 est attiré sur l'électrode B.

sont déviés par un champ électrique et recueillis par la cathode B. L'ensemble est porté à une température de 3000 K, supérieure à la température de fusion de l'uranium, pour que les isotopes ainsi déposés puissent s'écouler.

5.4 Intérêt de la méthode

La technique classique de séparation par diffusion gazeuse présente quelques désavantages : elle est chère (10 % du prix du kWh) et grande consommatrice d'énergie (il faut fournir 02×10^6 eV pour chaque atome d'uranium 238 séparé, soit 3,5 % de l'énergie ultérieurement produite). Dans la Séparation Isotopique par Laser sur Vapeur Atomique (procédé « SILVA »), il ne faut en principe fournir que l'énergie d'ionisation, soit 6 eV, pour séparer un atome d'Uranium 238. En pratique, si l'on tient compte du rendement du laser, de l'efficacité du processus d'ionisation et de l'efficacité de collection, le bilan s'élève à 5×10^4 eV par atome, ce qui donne un avantage à la technique laser. Les études de laboratoire relatives à cette technique ont permis de démontrer sa faisabilité. Le passage à l'échelle industrielle nécessite de progresser dans la résolution des problèmes de fiabilité et de maintenance des dispositifs lasers, ainsi que dans des problèmes plus fondamentaux liés à la propagation des faisceaux en milieux absorbants inhomogènes.¹¹

¹¹Pour plus de détails sur ce sujet, voir par exemple M. Clerc, P. Rigny et O. de Witte « Séparation isotopique par laser » dans « Le laser – Principes et techniques d'application » p. 201 édité par H. Mallet (Technique et documentation Lavoisier, Paris 1990).

Complément III.5

Spectroscopie non-linéaire

Les sources laser ont complètement bouleversé les méthodes et les possibilités de la spectroscopie. Leur très grande monochromaticité permet en effet de résoudre des structures très fines qu'il aurait été difficile de détecter par les méthodes antérieures, mais plus encore, leur densité spectrale de puissance a permis d'imaginer de nouvelles méthodes de spectroscopie bien plus performantes que les techniques classiques. Ces méthodes sont basées sur le fait que la réponse d'un atome à une source lumineuse intense n'est pas linéaire. Ce sont quelques unes de ces méthodes de *spectroscopie non linéaire* que nous décrivons succinctement ici¹.

Nous expliquons d'abord pourquoi la largeur des transitions atomiques observées expérimentalement apparaît plus grande que celle que l'on peut calculer pour un atome isolé au repos. Ceci nous permet d'introduire dans la partie A les notions de *largeur homogène* et de *largeur inhomogène*. Nous expliquons ensuite pourquoi il est possible de se débarrasser de l'élargissement inhomogène en utilisant des méthodes de spectroscopie non-linéaire. Nous nous intéressons plus spécifiquement à l'*élargissement Doppler*, qui est l'élargissement inhomogène le plus étudié et nous présentons deux méthodes permettant de s'affranchir de cet élargissement : la *spectroscopie d'absorption saturée* (Partie B) et les *transitions à deux photons* (Partie C). Nous discutons enfin les progrès faits sur la connaissance du spectre de l'*atome d'hydrogène* (qui est un des problèmes les plus fondamentaux de la physique atomique) en utilisant ces méthodes de spectroscopie sans élargissement Doppler (Partie D).

1 Largeur homogène et inhomogène

L'étude expérimentale de la largeur spectrale d'une raie émise par un atome ou un ion montre que cette largeur est presque toujours plus grande que la largeur prévisible à partir

¹Signalons aussi qu'en utilisant des *impulsions laser*, on peut réaliser une excitation percutielle du système que l'on veut étudier (atome, molécule, solide). En sondant le système au bout d'un temps déterminé, on peut alors avoir accès à la dynamique de son évolution. Nous n'aborderons pas ici ce domaine de la *spectroscopie résolue en temps*.

des premiers principes. Pour un atome isolé, la théorie prévoit une raie Lorentzienne dont la largeur est déterminée par la durée de vie des niveaux considérés (voir Complément II.2 Eq. 29). En fait, les largeurs mesurées expérimentalement sont généralement bien plus grandes. Cela tient souvent au fait que les atomes sont animés d'un mouvement désordonné : lorsque l'atome a une vitesse $\mathbf{v}$, la fréquence de la lumière émise dans la direction $\mathbf{u}$ est décalée de la fréquence naturelle ω_0 d'une quantité égale à $\omega_0 \mathbf{v} \cdot \mathbf{u} / c$, où c est la vitesse de la lumière (effet Doppler). À cause de la distribution maxwellienne des vitesses des atomes dans un gaz², un observateur enregistre donc une raie de forme gaussienne dont la largeur est bien plus grande que la largeur naturelle. Dans le domaine optique, la largeur due à l'effet Doppler, de l'ordre de $\omega_0 \bar{v} / c$ où $\bar{v}$ est la vitesse quadratique moyenne des atomes dans une direction donnée ($\bar{v} = \sqrt{k_B T / M}$, où M est la masse de l'atome et T la température du gaz), est plus grande que la largeur naturelle de la transition, typiquement par un facteur de l'ordre de 100 à température ambiante. Il s'ensuit que la précision sur la mesure de la fréquence de la transition est affectée par cet élargissement Doppler et que, d'autre part, les structures atomiques qui sont plus grandes que la largeur naturelle mais plus petites que la largeur Doppler, sont masquées par cet élargissement.

Un élargissement analogue des raies spectrales existe pour des ions placés dans un verre ou dans une matrice cristalline. Dans ce cas, les ions sont certes immobiles (il n'y a donc pas d'élargissement Doppler³), mais ils sont disposés au hasard de sorte que les déplacements Stark des niveaux d'énergie de l'ion dus au champ électrique créé par les ions voisins varient d'un site à l'autre. Il s'ensuit que la fréquence de la transition pour un ion dans le cristal diffère de la fréquence obtenue pour un ion isolé et que la raie associée à l'ensemble des ions s'étale sur une large plage correspondant aux divers environnements possibles.

Le type d'élargissement décrit plus haut (largeur Doppler, largeur due au champ cristallin) ne correspond pas à une limite ultime puisque des largeurs bien plus fines seraient obtenues s'il était possible d'étudier l'émission d'un seul atome, ou plus exactement d'une seule classe d'atomes (atomes ayant une vitesse donnée, ou bien placés dans un site de type donné). Ce phénomène parasite, dû aux conditions expérimentales, est responsable de la *largeur inhomogène* de la raie. Nous verrons dans les parties B et C que les méthodes de la spectroscopie non linéaire permettent d'éliminer certains de ces élargissements inhomogènes.

Au contraire, la *largeur homogène* correspond à la largeur de la transition que l'on est susceptible d'observer si l'on étudie un atome unique de l'échantillon⁴. La largeur liée à la durée de vie finie des niveaux atomiques excités est un exemple de largeur homogène. Cette largeur, liée à l'émission spontanée, est la plus petite largeur susceptible d'être obtenue pour une transition donnée d'un atome isolé de toute perturbation. La valeur de cette largeur est calculable dans le cadre de l'interaction entre un atome et le

²Voir par exemple R. Balian « Mécanique Statistique » Chapitre VI, § 6.2.

³En fait les atomes ne sont pas immobiles dans le solide, mais vibrent sous l'effet de l'agitation thermique avec une amplitude petite devant la longueur d'onde optique. Ce mouvement n'entraîne pas d'élargissement Doppler (« effet Lamb-Dicke », voir par exemple CDG 2, ex II).

⁴Orrit et al.

champ électromagnétique quantifié. Il existe d'autres causes d'élargissement homogène. Par exemple, on peut montrer que l'élargissement d'une transition due aux collisions que subit un atome dans un gaz est un élargissement homogène, c'est-à-dire qu'il est le même pour tous les atomes. Comme pour l'émission spontanée, cet élargissement est associé à un profil Lorentzien.

2 Spectroscopie d'absorption saturée

Peu de temps après l'avènement des lasers est apparue la première méthode de spectroscopie « sub-Doppler », c'est-à-dire permettant de résoudre des structures inférieures à la largeur inhomogène. Cette méthode, appelée *absorption saturée*, a révolutionné la spectroscopie, et elle est très largement utilisée à l'heure actuelle⁵. Nous en expliquons le principe dans les paragraphes suivants.

2.1 Trous dans une distribution de population

Considérons des ions dans un réseau cristallin. À cause des inhomogénéités du champ électrique dans le cristal, la différence d'énergie entre le niveau excité et le niveau fondamental de l'ion dépend du site où se trouve l'ion. Il est possible d'écrire la fréquence de résonance pour l'ion situé sur le site i sous la forme :

$$\omega_i = \omega_0 + \delta\omega_i \quad (1)$$

où ω_0 est la fréquence de résonance pour un ion isolé et $\delta\omega_i$ est le déplacement de fréquence pour l'ion situé sur le site i . Supposons que ces ions sont soumis à un faisceau laser incident d'amplitude E' et de pulsation ω' dont on mesure l'absorption. Seuls les ions situés sur les sites i tels que :

$$\omega_i = \omega' \quad (2)$$

contribueront au signal. En balayant la fréquence du laser, on explore les différentes fréquences d'absorption ω_i . La forme et la largeur de la raie d'absorption sont donc déterminées par la distribution des atomes sur les différents sites i (élargissement inhomogène, voir figure 1.a).

Un second faisceau laser de pulsation ω , d'amplitude E beaucoup plus grande que E' et capable de *saturer* la transition, est maintenant envoyé sur l'échantillon. Il a pour effet

⁵Le phénomène d'absorption saturée a d'abord été observé dans la courbe de gain d'un laser à gaz à cavité linéaire où existe un creux étroit situé à la fréquence de résonance (« Lamb dip »). La méthode de spectroscopie d'absorption saturée a été ensuite mise au point par C. Bordé et T. Hansch. Cette technique a évolué et s'est affinée en incluant d'autres degrés de liberté comme ceux liés à la polarisation des faisceaux. Pour plus de détails, on pourra consulter l'article de T. Hansch, A. Schawlow et G. Series dans Scientific American. Vol 240 (Mars 1979) p. 72, ou le livre de V. Lethokov et V. Chebotayev : Non Linear Laser Spectroscopy (Springer, 1977).

d'égaliser les populations entre le niveau fondamental et le niveau excité (voir chapitre 11, formule (D.10)) pour les ions situés sur les sites j tels que :

$$\omega_j = \omega \quad (3)$$

Lorsqu'on balaie la pulsation ω' de la source faible E' (faisceau sonde), l'absorption est identique à celle obtenue en l'absence de source E , sauf lorsque $\omega' \approx \omega$ puisqu'alors les deux faisceaux interagissent avec les mêmes ions et que l'action du faisceau E saturant a *diminué le nombre d'ions dans le niveau fondamental* sur les sites j . La forme de la raie d'absorption obtenue dans ces conditions est représentée sur la figure 1.b. La largeur de la courbe fine ainsi obtenue est de l'ordre de $2\sqrt{\gamma^2 + \Omega_1^2\gamma/\Gamma_{sp}}$ où 2γ est la largeur naturelle de la transition (largeur de type homogène), Γ_{sp} l'inverse de la durée de vie du niveau excité, et $\Omega_1 = -dE/\hbar$ la pulsation de Rabi résonnante pour l'onde E (Eq. (24.b) du Complément II.2).

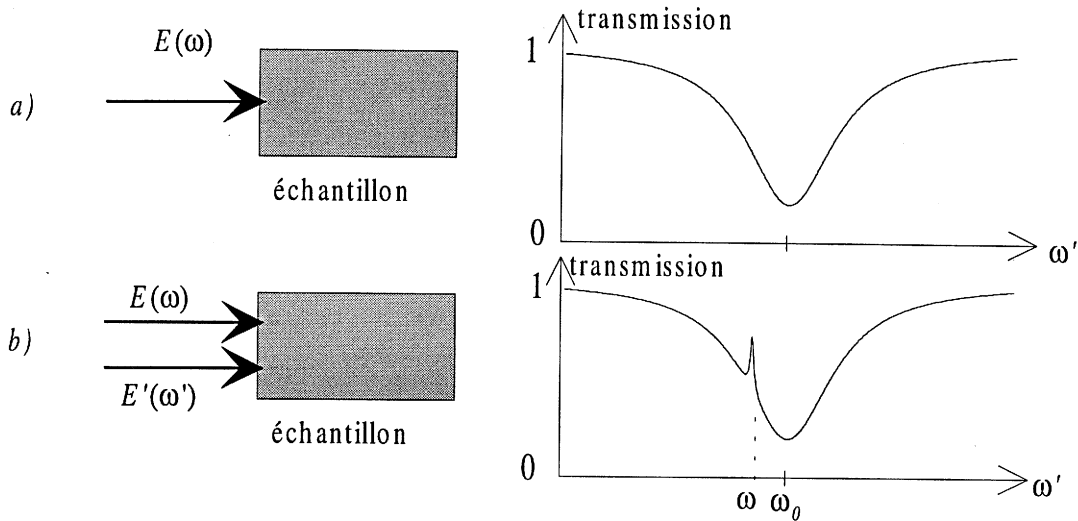


FIG. 1: Absorption d'un faisceau sonde de faible intensité et de pulsation ω' . La courbe d'absorption mesurée avec ce faisceau seul (a) permet de déterminer l'élargissement inhomogène de la transition. La même expérience faite en présence d'un faisceau intense de pulsation ω constante (b) conduit à une courbe identique sauf au voisinage de $\omega = \omega'$ où une diminution de l'absorption due à la saturation par l'onde intense est observée. La largeur de la courbe étroite permet de remonter à la largeur homogène.

Il apparaît clairement sur la figure 1.b que c'est la saturation de la transition due à l'onde intense E qui permet d'observer une courbe étroite sur l'absorption de l'onde E' . En diminuant la pulsation de Rabi Ω_1 , la largeur de la courbe étroite diminue et tend vers la largeur naturelle de la transition. Cette méthode permet donc d'obtenir des raies plus fines que la largeur inhomogène de la transition. Elle présente cependant l'inconvénient que le centre de la raie étroite est déterminé par la pulsation du faisceau ω et non par une valeur spécifiquement liée aux niveaux d'énergie de l'ion. Nous montrons dans le paragraphe suivant qu'il est possible de surmonter ce problème dans le cas de l'élargissement Doppler

et de déterminer avec une grande précision les énergies atomiques en utilisant la saturation de l'absorption.

2.2 Absorption saturée dans un gaz

a. Distribution maxwellienne des vitesses

Considérons un gaz constitué d'atomes à deux niveaux (niveau excité b , niveau fondamental a , $E_b - E_a = \hbar\omega_0$). À l'équilibre thermique à la température T , les atomes sont, dans leur immense majorité, dans le niveau fondamental. Ces atomes, de masse M , sont animés de mouvements rectilignes uniformes dans le gaz et la probabilité que la composante v_x de la vitesse d'un atome soit comprise entre v_x et $v_x + dv_x$ est déterminée par la distribution de Maxwell-Boltzmann :

$$f(v_x) = \frac{1}{\bar{v}\sqrt{2\pi}} \exp\left(-\frac{v_x^2}{2\bar{v}^2}\right) \quad (4)$$

avec $\bar{v} = \sqrt{\frac{k_B T}{M}}$ (vitesse quadratique moyenne à une dimension).

b. Excitation d'une classe de vitesse

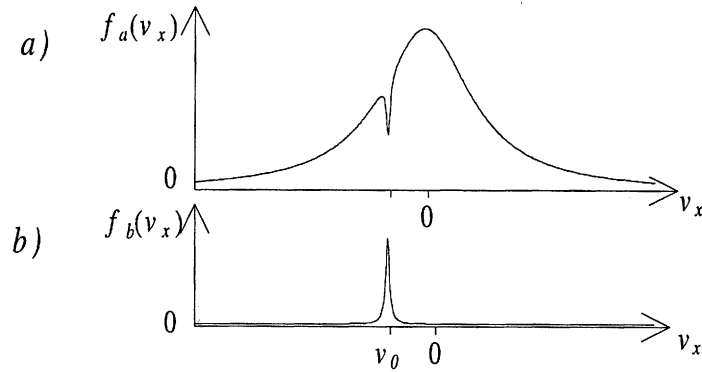


FIG. 2: Proportion d'atomes de vitesse donnée (« distribution de vitesse ») dans le niveau fondamental a et dans le niveau excité b lorsque les atomes sont soumis à une onde incidente progressive de pulsation ω se propageant selon Ox . La vitesse est déterminée par la condition (7).

Supposons que ces atomes interagissent avec une onde monochromatique incidente

$$\mathbf{R}(\mathbf{r}, t) = \mathbf{E} \cos(\omega t - \mathbf{k} \cdot \mathbf{r}) \quad (5)$$

Un atome de vitesse $\mathbf{v}$ et décrivant une trajectoire

$$\mathbf{r} = \mathbf{r}_0 + \mathbf{v}t \quad (6)$$

voit dans son référentiel propre une onde

$$\mathbf{E} \cos[(\omega - \mathbf{k} \cdot \mathbf{v})t - \mathbf{k} \cdot \mathbf{r}_0] \quad (7)$$

dont la fréquence est déplacée par *effet Doppler*. Pour une pulsation incidente ω donnée, seuls les atomes de vitesse $\mathbf{v}$ tels que

$$\hbar(\omega - \mathbf{k} \cdot \mathbf{v}) = \hbar\omega_0 \quad (8)$$

peuvent absorber le rayonnement incident. Si l'onde se propage dans la direction Ox , les atomes excités auront donc une composante v_0 de la vitesse le long de l'axe Ox égale à

$$v_0 = \frac{\omega - \omega_0}{\omega} c \quad (9)$$

(puisque pour un milieu dilué $k \approx \omega/c$).

En présence de l'onde $\mathbf{E}(\mathbf{r}, t)$, la répartition en fonction de v_x des atomes dans le niveau fondamental présente un *trou* autour de la vitesse vérifiant la condition (9). En revanche, les atomes portés dans le niveau excité b ont une *vitesse bien définie* (voir figure 2). Pour des atomes à deux niveaux, la figure (2.b) correspondant aux atomes excités est généralement complémentaire de la figure (2.a) correspondant aux atomes dans le niveau fondamental.

Remarque

La complémentarité des figures (2.a) et (2.b) peut ne plus être respectée lorsque les atomes subissent des *collisions*. En effet, les collisions modifient les vitesses et ont donc tendance à étaler le trou de la figure (2.a) et le pic de la figure (2.b). Cependant comme les *sections efficaces de collision sont généralement différentes dans le niveau fondamental et dans le niveau excité*, les déformations du trou et du pic ne sont généralement pas identiques.

Supposons par exemple que la section efficace de collision avec un gaz « tampon » introduit dans la cellule soit beaucoup plus grande dans le niveau excité que dans le niveau fondamental. Il s'ensuit qu'un trou au voisinage de v_0 subsistera sur la distribution $f_a(v_x)$ des atomes dans le niveau fondamental alors que la distribution $f_b(v_x)$ sera rapidement thermalisée et retrouvera une forme Maxwellienne centrée en $v = 0$. Une conséquence de cette absence de complémentarité est que globalement les atomes acquièrent une *vitesse moyenne non nulle le long de l'axe Ox* (Dans l'exemple discuté ci-dessus, cette vitesse moyenne a un signe opposé à v_0). Le déplacement d'un gaz dû à cet effet, appelé effet de *piston optique*, a été observé dans plusieurs laboratoires⁶. Il convient de noter que pour cet exemple, le mouvement n'est pas dû à un effet mécanique de la lumière lié par exemple à un transfert d'impulsion lors de l'absorption d'un photon par un atome (pression de radiation). Le déplacement peut d'ailleurs se produire dans le sens opposé à celui de la propagation de l'onde si le laser est désaccordé vers les fréquences élevées (Signalons toutefois qu'il existe des situations expérimentales différentes, où l'effet de pression de radiation est prédominant).

c. Principe de la spectroscopie d'absorption saturée

Supposons que le faisceau incident de pulsation ω soit divisé à l'aide d'une lame séparatrice en deux faisceaux : un faisceau intense noté F et un faisceau sonde faible noté F' (voir figure 3). Ces deux faisceaux sont envoyés selon deux directions pratiquement opposées dans le gaz (leur direction faisant un très petit angle juste suffisant pour pouvoir placer un détecteur sur le faisceau F' sans occulter le faisceau F).

⁶Voir par exemple L. Moi « La Recherche » n° 207 (Février 1989) p. 260.

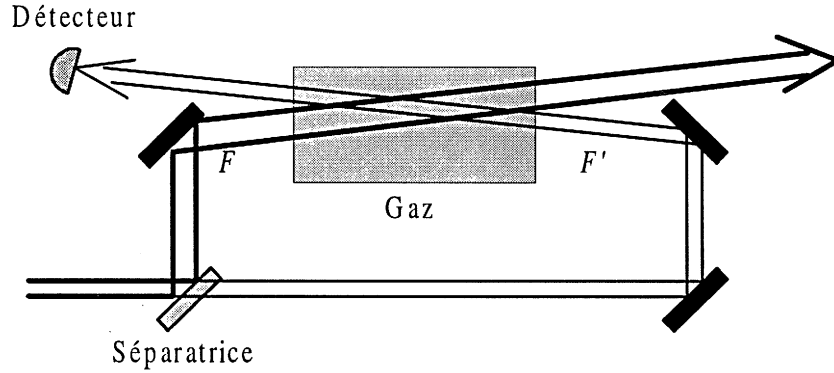


FIG. 3: Schéma d'une expérience d'absorption saturée. Un faisceau incident est divisé en deux faisceaux au moyen d'une lame séparatrice. L'intensité du faisceau le plus faible F' est mesurée en fonction de la pulsation ω après traversée dans le gaz. L'angle entre les deux faisceaux est dans la pratique beaucoup plus petit que l'angle représenté sur la figure.

Le faisceau F interagit avec les classes de vitesse déterminées par la condition (9). Le faisceau F' qui se propage en sens opposé, et dont la phase est donc de la forme $(\omega t + \mathbf{k} \cdot \mathbf{r})$, excite les atomes dont la vitesse $\mathbf{v}'$ vérifie la relation :

$$\hbar(\omega + \mathbf{k} \cdot \mathbf{v}') = \hbar\omega_0 \quad (10)$$

soit encore

$$v'_x = -\frac{\omega - \omega_0}{\omega} c \quad (11)$$

La comparaison des formules (11) et (9) montre que les deux ondes F et F' interagissent avec des *classes de vitesses différentes* (et ont donc une absorption indépendante de leur présence simultanée dans le gaz) tant que ω est différent de ω_0 . En revanche, pour $\omega = \omega_0$, les deux ondes interagissent avec la même classe de vitesse $v_x = 0$ (voir figure 4). Il s'ensuit que si l'onde F sature la transition $a \rightarrow b$, la transmission de l'onde F' sera augmentée lorsque $\omega \approx \omega_0$.

La courbe d'absorption (figure 5) de l'onde faible présente donc une raie large correspondant à la largeur Doppler sur laquelle se superpose (en négatif) une courbe *étroite* correspondant à l'excitation simultanée de la même classe de vitesse ($v_x = 0$) par les deux ondes. Dans le cadre du modèle considéré ici, la largeur de la courbe étroite est de l'ordre de $2\sqrt{\gamma^2 + \Omega_1^2 \gamma / \Gamma_{sp}}$, où Ω_1 est la pulsation de Rabi associée à l'onde saturante F . On conçoit donc qu'en diminuant les intensités des ondes, on puisse déterminer la largeur naturelle de la transition comme valeur asymptotique⁷.

⁷À forte intensité de l'onde F , les populations des niveaux a et b sont égalisées. On pourrait donc s'attendre à ce que l'absorption de l'onde faible F' soit pratiquement nulle lorsque $\omega = \omega_0$. En fait, ce n'est pas le cas : un calcul plus précis tenant compte notamment des déplacements des niveaux d'énergie dans un champ intense montre que l'absorption tout en étant considérablement diminuée ne tend pas vers 0 (pour plus de détails, voir S. Haroche, F. Hartmann, Phys. Rev. **A6** 1280 (1972), ou CDG 2, Exercice 20).

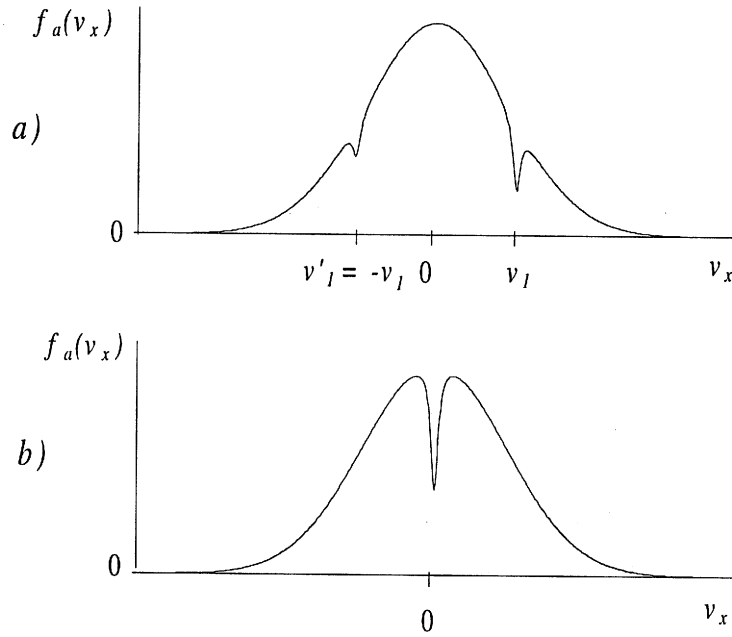


FIG. 4: Distribution de vitesse dans le niveau fondamental. a) Lorsque $\omega \neq \omega_0$, les faisceaux F et F' excitent des classes de vitesse symétriques v_1 et v'_1 . b) Lorsque $\omega = \omega_0$, les faisceaux F et F' interagissent avec la même classe de vitesse $v_x = 0$.

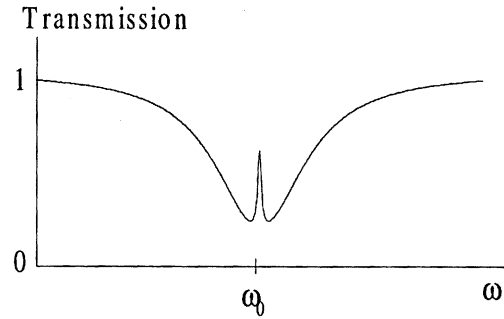


FIG. 5: Forme de raie d'absorption saturée dans une vapeur atomique : transmission du faisceau faible en fonction de la pulsation ω (en pratique le signal étroit est beaucoup plus faible, ce qui nécessite souvent l'utilisation de techniques d'extraction de bruit).

La courbe étroite est donc centrée sur la résonance atomique, et peut avoir une largeur proche de la largeur naturelle. On comprend que cette méthode présente un intérêt spectroscopique considérable. Elle est notamment très utilisée pour obtenir une référence de fréquence absolue sur laquelle asservir la fréquence d'un laser.

Remarque

Lorsque l'atome possède plusieurs sous-niveaux, dans a ou dans b , on peut voir apparaître d'autres signaux résonnants étroits, appelés résonances de « croisement de niveau ». Prenons par exemple un modèle où l'atome a deux sous-niveaux a et a' dans le niveau fondamental,

dont la séparation est inférieure à la largeur Doppler, et un seul niveau excité b . Si le faisceau intense est résonnant sur la transition $a \rightarrow b$, il modifie fortement la population du niveau b , et perturbe aussi bien l'absorption du faisceau sonde sur la transition $a \rightarrow b$ que sur la transition $a' \rightarrow b$. Dans la configuration de la figure 3, on obtiendra une variation résonnante de l'absorption lorsque la même classe de vitesse interagira avec les deux faisceaux sur l'une quelconque des deux transitions. On montre aisément que cela se produit pour trois fréquences ω du champ, égales à :

$$\omega = \omega_{a \rightarrow b'} , \quad \omega = (\omega_{a \rightarrow b} + \omega_{a' \rightarrow b})/2 , \quad \omega = \omega_{a' \rightarrow b} \quad (12)$$

Pour la plupart des atomes et des molécules, le niveau fondamental, aussi bien que les niveaux excités, possèdent une structure hyperfine comprenant plusieurs niveaux hyperfins. Le signal obtenu par le dispositif de la figure 3 est alors complexe. Il est souvent très utile, car il fournit une « grille » de fréquences absolues définies avec une précision limitée seulement par la largeur naturelle des transitions, permettant par exemple de choisir dans une certaine mesure la fréquence sur laquelle le laser est asservi. Notons enfin que l'effet croisé entre les deux faisceaux peut avoir d'autres origines que la saturation. Il peut être lié par exemple à un effet de pompage optique entre les sous-niveaux a et a' (voir Complément II.1).

3 Spectroscopie d'absorption à deux photons sans élargissement doppler

3.1 Transition à deux photons

Nous avons vu au chapitre II paragraphe C.3 qu'un atome interagissant avec deux ondes de pulsations ω_1 et ω_2 peut passer du niveau fondamental a à un niveau excité b en absorbant deux photons, un de pulsation ω_1 , l'autre de pulsation ω_2 (voir figure 6). La condition de résonance pour ce processus est :

$$E_b - E_a = \hbar(\omega_1 + \omega_2) \quad (13)$$

3.2 Élimination de l'élargissement Doppler

Considérons des atomes dans un gaz interagissant avec deux ondes progressives F, F' de même pulsation ω et se propageant dans des sens opposés (figure 7.a). Si dans le référentiel du laboratoire (figure 7.b) un atome de vitesse v interagit avec deux ondes de pulsation ω , l'atome voit dans son référentiel propre deux ondes dont les pulsations ω_1 et ω_2 ont été déplacées par effet Doppler (figure 7.c). Ces pulsations valent, dans le cas où $|v| \ll c$:

$$\omega_1 = \omega \left(1 - \frac{v_x}{c} \right) \quad (C.2.a)$$

$$\omega_2 = \omega \left(1 + \frac{v_x}{c} \right) \quad (C.2.b)$$

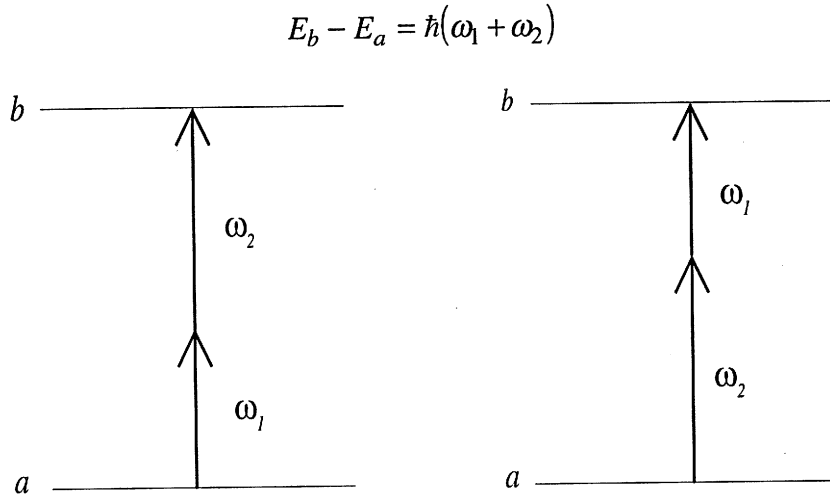


FIG. 6: Transition à deux photons entre les niveaux a et b . L'amplitude de transition est la somme de deux termes, l'un correspondant à l'absorption du photon ω_1 suivi de l'absorption du photon ω_2 , l'autre correspondant à l'ordre inverse. Les deux termes sont résonnants pour la même condition (13).

Étudions la condition d'absorption par l'atome de deux photons se propageant en sens opposé. En utilisant les relations (13) et (C.2.a), nous trouvons :

$$E_b - E_a = \hbar\omega \left(1 - \frac{v_x}{c}\right) + \hbar\omega \left(1 + \frac{v_x}{c}\right) \quad (15)$$

soit

$$E_b - E_a = 2\hbar\omega \quad (16)$$

Nous voyons que les *termes dépendant de la vitesse disparaissent* dans cette condition de résonance. Il s'ensuit que *tous les atomes* sont susceptibles d'être portés dans le niveau excité quand la condition (16) est réalisée.

Il est important de noter que cette condition n'est obtenue que dans la géométrie où deux ondes de même fréquence se propagent en sens opposés. Si nous ne considérons qu'une seule onde progressive, la condition d'absorption de deux photons dans cette onde

$$\omega = \omega_{ba} = \frac{E_b - E_a}{\hbar} \quad (17)$$

dépend de la valeur de la vitesse. Pour une valeur de ω donnée, une seule classe de vitesse (celle vérifiant l'équation (17)) est susceptible de conduire à une absorption à deux photons. Dans ce cas, la largeur de la raie est déterminée par l'effet Doppler.

Comme, en général, l'atome est susceptible d'absorber aussi bien deux photons se propageant en sens opposé que deux photons se propageant dans le même sens, cela entraîne que la forme de raie est la superposition d'une raie large due à l'absorption de

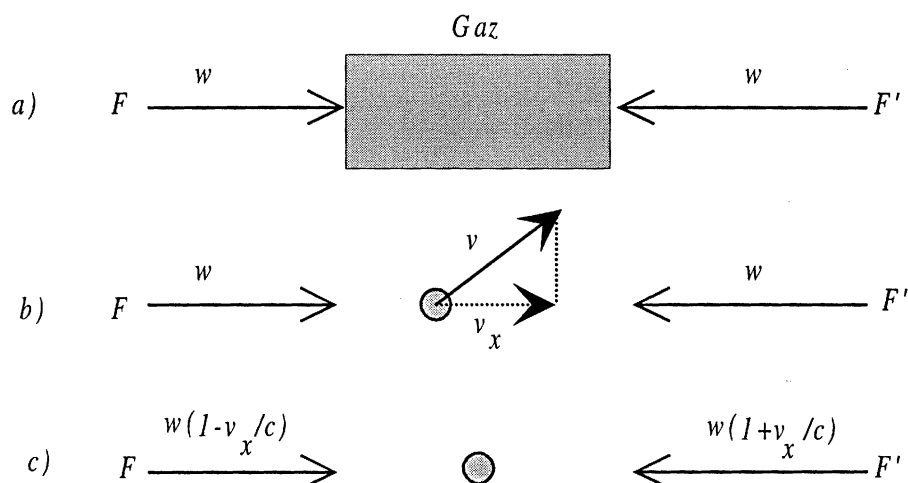


FIG. 7: a) Schéma de l'élimination de l'élargissement Doppler dans une transition à deux photons. b) Dans le référentiel du laboratoire, l'atome de vitesse v interagit avec deux ondes de même pulsation. c) Dans son référentiel propre, l'atome interagit avec deux ondes dont les fréquences ont été déplacées de façon symétrique par effet Doppler.

deux photons se propageant dans le même sens et d'une raie étroite liée à l'absorption de deux photons se propageant en sens opposé (Figure 8). Si les faisceaux F et F' , sont d'intensités comparables, les surfaces des raies large et étroite sont semblables puisqu'elles font toutes deux intervenir toutes les classes de vitesse⁸

Remarque

La possibilité d'éliminer l'élargissement Doppler dans les transitions à deux photons peut s'interpréter en termes de conservation de l'impulsion et de l'énergie. Considérons un atome de masse M et de vitesse $\mathbf{v}$ dans le référentiel du laboratoire. L'impulsion et l'énergie des photons est dans ce même référentiel $\hbar\mathbf{k}$, $\hbar\omega$ pour une onde et $-\hbar\mathbf{k}$, $\hbar\omega$ pour l'autre onde. La conservation de l'impulsion à l'issue du processus d'absorption de deux photons se propageant en sens opposé implique :

$$M\mathbf{v}' = M\mathbf{v} + \hbar\mathbf{k} - \hbar\mathbf{k} \quad (\text{C.6.a})$$

où $\mathbf{v}'$ est la vitesse de l'atome après absorption. L'équation (C.6.a.a) montre que la vitesse ne change pas ($\mathbf{v}' = \mathbf{v}$) dans ce processus.

La conservation de l'énergie entraîne

$$E_b + \frac{1}{2}Mv'^2 = E_a + \frac{1}{2}Mv^2 + 2\hbar\omega \quad (\text{C.6.b})$$

Puisque $v' = v$, on retrouve la condition de résonance (16). Notons que cette démonstration montre de surcroît l'absence de recul de l'atome quand celui-ci absorbe deux photons se propageant en sens opposés.

⁸Cette possibilité d'éliminer l'élargissement Doppler dans une excitation à deux photons a été proposée d'abord par Vasilenko, Chebotayev et Shishaev en 1970. Les premières observations expérimentales datent des années 1973-74 et ont été obtenues à Paris (Biraben, Cagnac, Grynberg), et à Harvard (Bloembergen et Levenson) avec des lasers en impulsion, puis à Stanford (Hansch et Schawlow) avec des lasers continus. Les extensions de la méthode et la généralisation aux processus à plus de deux photons discutée dans la remarque sont présentées dans l'article de G. Grynberg et B. Cagnac (Rep. Prog. Phys. **40**, 791 (1977)) où l'on trouvera également une revue des premières expériences.

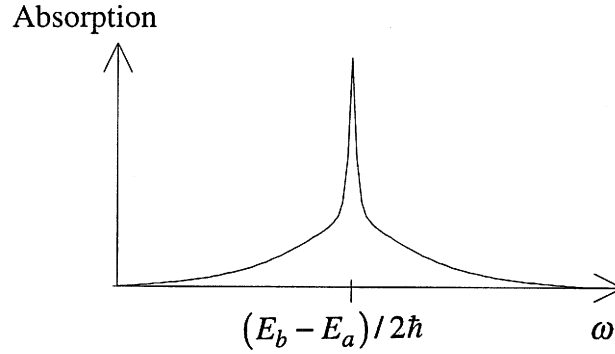


FIG. 8: Variation de l'absorption en fonction de la pulsation. En pratique, la raie étroite a une hauteur plus grande que la raie large dans un rapport (largeur Doppler/largeur naturelle), puisque les deux courbes ont des surfaces voisines. Ce rapport étant fréquemment de l'ordre de 100, on observera essentiellement la raie étroite dans les spectres expérimentaux.

L'approche présentée ici peut aisément être généralisée à des transitions multiphotoniques d'ordre plus élevé. Considérons à titre d'exemple l'absorption de trois photons issus de faisceaux de vecteurs d'onde $\mathbf{k}_1, \mathbf{k}_2$ et $\mathbf{k}_3$ et de pulsation ω_1, ω_2 et ω_3 (dans le référentiel du laboratoire). Si les faisceaux sont disposés dans des directions telles que la somme des impulsions des photons soit nulle :

$$\sum_{i=1}^3 \hbar \mathbf{k}_i = 0 \quad (\text{C.7.a})$$

l'impulsion de l'atome sera la même avant et après absorption et son énergie cinétique ne sera pas modifiée par le processus d'absorption. Il s'ensuit que la condition de résonance :

$$E_b - E_a = \sum_{i=1}^3 \hbar \omega_i \quad (\text{C.7.b})$$

sera la même pour tous les atomes, indépendamment de leur vitesse. En balayant la pulsation de l'un des faisceaux, on obtient alors une raie qui n'est pas élargie par l'effet Doppler.

3.3 Propriétés de la spectroscopie à deux photons sans élargissement Doppler

La spectroscopie à deux photons permet d'obtenir des raies ayant une largeur égale à la largeur naturelle. Il s'agit donc d'une méthode extrêmement puissante dont les propriétés sont complémentaires de celles de l'absorption saturée puisqu'elle permet d'atteindre des niveaux de même parité que le niveau fondamental alors que l'absorption saturée permet d'étudier des transitions reliant le niveau fondamental à un niveau de parité opposé (cf. Complément II.1).

On pourrait craindre qu'une transition à deux photons nécessite des sources intenses occasionnant d'importants déplacements lumineux (voir chapitre II paragraphe C.3.c), ce

qui ferait perdre beaucoup d'intérêt à cette méthode pour la spectroscopie de haute résolution. Souvent, ce n'est pas le cas. En particulier, le fait que tous les atomes contribuent à la résonance (alors que seule la classe de vitesse $v_x = 0$ contribue dans la spectroscopie d'absorption saturée) permet de compenser, au moins partiellement, la faiblesse des probabilités d'absorption à deux photons par rapport à l'absorption à un photon⁹.

4 Spectroscopie de l'atome d'hydrogène

4.1 Le rôle de l'atome d'hydrogène

L'atome d'hydrogène, constitué d'un proton et d'un électron, est l'atome le plus simple. Cette simplicité a permis une confrontation extraordinairement fructueuse entre théorie et expérience. Ainsi, c'est l'étude à la fin du XIX^{ème} siècle du spectre de l'atome d'hydrogène, conduisant à la formule empirique de Balmer, qui a été un des catalyseurs de l'émergence de la *théorie quantique*. Le succès majeur du modèle de Bohr d'abord, de l'équation de Schrödinger ensuite, a été l'explication théorique de la formule de Balmer qui était incompréhensible dans le cadre de la mécanique classique. L'amélioration des techniques de la spectroscopie optique au début du siècle a rapidement montré que les raies de l'hydrogène présentaient une *structure fine*. L'interprétation de ce type de structure a conduit à l'introduction du spin de l'électron, d'abord décrit de façon phénoménologique (Uhlenbeck et Goudsmit), puis interprétée dans un cadre théorique nouveau, celui de la *mécanique quantique relativiste* (équation de Dirac). Plus tard, à la fin des années 40, Lamb et Retherford ont montré, en utilisant la spectroscopie radiofréquence, que les niveaux $2S_{1/2}$ et $2P_{1/2}$ n'étaient pas dégénérés en énergie, contrairement aux prédictions de la théorie de Dirac. L'interprétation et le calcul de cette petite différence d'énergie (moins de 10^{-6} en valeur relative) ont impliqué de nombreux théoriciens (Bethe, Dyson, Feynmann, Schwinger, Tomonaga, Weisskopf...). Ils ont permis d'aboutir à une des théories physiques les plus sophistiquées, et sans doute la plus précise : *l'électrodynamique quantique*.

Ce bref résumé montre clairement que l'atome d'hydrogène a joué un rôle essentiel dans le développement de la physique. Tout progrès expérimental permettant d'améliorer la précision avec laquelle sont connus les niveaux de l'atome d'hydrogène est immédiatement susceptible d'être confronté avec la théorie et ainsi la confirmer à un degré de précision supérieure ou (comme cela a été plusieurs fois le cas dans le passé) de lui imposer une brutale remise en question. À cet égard, les expériences de spectroscopie non-linéaire sans élargissement Doppler menées depuis 1970 sur l'hydrogène présentent une grande importance : elles ont, en effet, permis d'améliorer la précision avec laquelle est connu le spectre de l'atome d'hydrogène dans le domaine optique par plus de trois ordres de

⁹Pour plus de détails sur les transitions à plusieurs photons et sur l'élimination de l'élargissement Doppler dans ces processus, on pourra consulter le chapitre « Multiphoton Resonant Processes in Atoms » par G. Grynberg, B. Cagnac et F. Biraben dans « Coherent Non Linear Optics. Recent Advances » édité par M.S. Feld et V.S. Letokhov (Springer - Verlag, Berlin 1980).

grandeur¹⁰. Notons également que la détermination du spectre de l'« antihydrogène » (formé d'un positron et d'un antiproton) permettrait de tester avec une grande précision la symétrie des lois de la physique vis à vis de l'échange matière anti-matière.

4.2 L'atome d'hydrogène

a. De Balmer à Dirac

La théorie non relativiste de l'atome d'hydrogène présentée dans les cours de Mécanique Quantique élémentaire montre que les niveaux d'énergie d'un électron dans un potentiel coulombien peuvent être repérés par trois nombres quantiques n, ℓ et m et que l'énergie ne dépend en fait¹¹ que de n :

$$E_{n\ell m} = -\frac{R}{n^2} \quad (19)$$

où R est la constante de Rydberg qui s'exprime en fonction de la charge q et de la masse m de l'électron sous la forme¹² :

$$R = \frac{mq^4}{32\pi^2\epsilon_0^2\hbar^2} \quad (20)$$

L'équation de base de la théorie quantique relativiste est l'équation de Dirac¹³. Cette équation permet d'introduire naturellement le spin s de l'électron. Comme le spin s et le moment cinétique orbital ℓ sont couplés par l'interaction spin-orbite, les niveaux d'énergie s'expriment en fonction du moment cinétique total j ($\mathbf{j} = \mathbf{l} + \mathbf{s}$). La résolution de l'équation de Dirac pour l'atome d'hydrogène conduit à la formule suivante pour les niveaux d'énergie :

$$E_{n\ell jm} = -\frac{R}{n^2} \left[1 + \frac{\alpha^2}{n} \left(\frac{1}{j + 1/2} - \frac{3}{4n} \right) \right] \quad (21)$$

Un des résultats remarquables de l'équation de Dirac est que, dans une multiplicité de n donné, deux sous-niveaux de même valeur de j mais de ℓ différent auront même énergie. Ainsi, dans la multiplicité $n = 2$, le niveau $2S_{1/2}(\ell = 0, j = 1/2)$ et le niveau $2P_{1/2}(\ell = 1, j = 1/2)$ sont dégénérés (voir figure 9.a).

¹⁰Un exposé très complet sur nos connaissances actuelles de l'atome d'hydrogène, tant au niveau théorique qu'expérimental, est donné dans le livre « The Spectrum of Atomic Hydrogen-Advances » édité par G.W. Series (World Scientific, Singapour, 1988). Pour une revue des expériences plus récentes, on pourra consulter l'article de B. Cagnac, M. Plimmer, L. Julien et F. Biraben (Rep Prog. Phys. **57**, 853 1994).

¹¹Rappelons que la dégénérescence en ℓ est spécifique au potentiel coulombien.

¹²La constante de Rydberg R correspond à un noyau infiniment lourd. C'est la limite de l'équation (E.1) du complément précédent pour $M \rightarrow \infty$ (Attention à la confusion entre m masse de l'électron, et m nombre quantique magnétique).

¹³Voir par exemple le livre « Quantum mechanics of one and two electron atoms » de H.A. Bethe et E.E. Salpeter (Plenum, New-York. 1977).

b. De Dirac à Lamb

En fait, les expériences menées par Lamb et Retherford ont montré que les niveaux $2S_{1/2}$ et $2P_{1/2}$ n'ont pas exactement la même énergie (Figure 9.b). L'équation de Dirac n'est donc pas la description ultime de l'atome d'hydrogène.

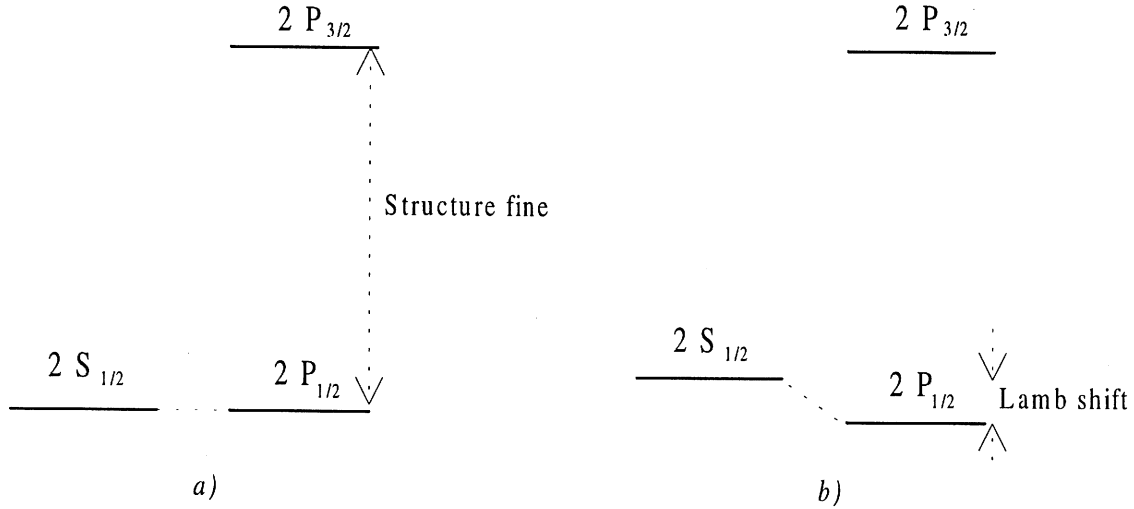


FIG. 9: Structure de la multiplicité $n = 2$ de l'atome d'hydrogène d'après l'équation de Dirac (a) et dans la réalité (b). L'écart entre les niveaux $2S_{1/2}$ et $2P_{1/2}$ est appelé le déplacement de Lamb. La valeur numérique du déplacement de Lamb est (en unité de fréquence) 1,057 GHz. La valeur de l'écart de structure fine entre les niveaux $2P_{1/2}$ et $2P_{3/2}$ est de 9,912 GHz.

L'origine de la déviation par rapport à la formule de Dirac est liée au couplage de l'électron avec le champ électromagnétique. En *théorie quantique du rayonnement* (voir chapitre VI), ce couplage, même en l'absence de tout champ appliqué de l'extérieur, est responsable de la désexcitation par émission spontanée des niveaux excités de l'atome. Mais il a un autre effet : l'électron de l'atome d'hydrogène est toujours susceptible d'émettre puis de réabsorber virtuellement des photons. De tels processus conduisent à de faibles déplacements des niveaux d'énergie, différents pour les niveaux $2S_{1/2}$ et $2P_{1/2}$. Ils permettent donc de comprendre l'origine du déplacement de Lamb.

La mesure de plus en plus précise de ces déplacements permet de confronter la valeur expérimentale avec les prédictions théoriques de l'électrodynamique quantique, suscitant un effort parallèle des théoriciens pour raffiner les calculs dont la difficulté augmente avec la précision recherchée. En 1994, il y a accord entre théorie et expérience avec une précision relative voisine de 10^{-11} sur la valeur des fréquences de transition dans le domaine optique.

4.3 Mesure de la constante de Rydberg

La spectroscopie de l'atome d'hydrogène permet de confronter théorie et expérience en ce qui concerne le détail des écarts énergétiques entre sous-niveaux. Mais elle per-

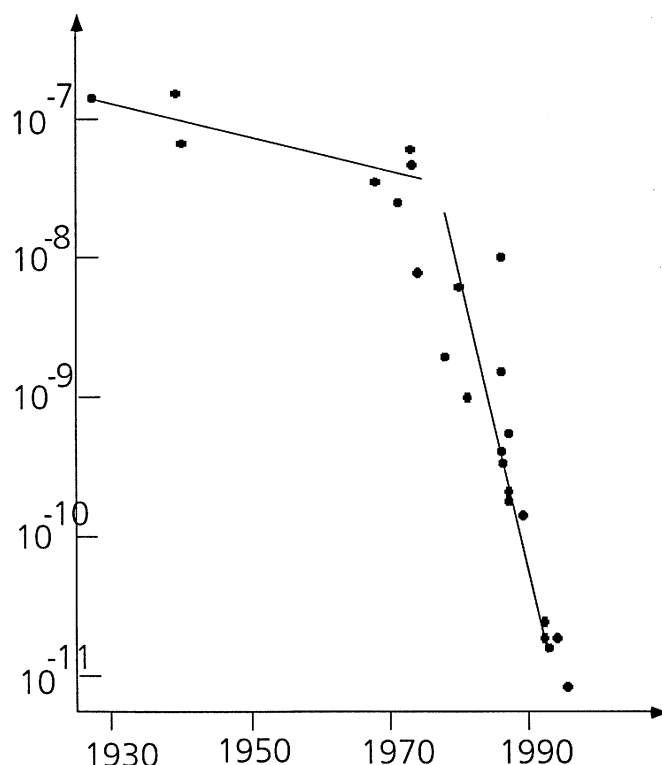


FIG. 10: Évolution de la précision sur la mesure de la constante de Rydberg. La rupture de pente après 1970 est due à l'apparition de la spectroscopie laser (figure extraite de la référence de la note 8).

met aussi une détermination extrêmement précise de la constante de Rydberg elle-même, c'est-à-dire d'une certaine combinaison des quantités $q, m, \hbar \dots$ (Eq. 20). D'autres mesures métrologiques fournissent une valeur pour d'autres combinaisons de ces quantités. La confrontation de tous ces résultats permet en définitive d'attribuer des valeurs extrêmement précises aux constantes fondamentales de la physique, telles que la masse ou la charge de l'électron.

Pour déterminer la constante de Rydberg, la procédure utilisée en 1994 consiste à corriger la mesure faite sur une transition $n\ell j \rightarrow n'\ell'j'$ de l'atome d'hydrogène des effets liés à la masse du noyau, des effets relativistes et de ceux de l'électrodynamique quantique calculés théoriquement, et d'en déduire une mesure de la constante de Rydberg R apparaissant dans l'expression non-relativiste (19) des niveaux d'énergie. La cohérence de la théorie doit être assurée par le fait que la mesure de R doit être *indépendante de la transition considérée*.

Nous avons représenté sur la figure 10 la variation avec le temps de la précision sur la mesure de la constante de Rydberg. La rupture de pente apparaissant après 1970 est due aux développements des méthodes de spectroscopie laser non-linéaire exposées dans ce complément, qui ont donc permis de passer d'une précision de l'ordre de 10^{-7} à une

précision de l'ordre de 10^{-11} . La valeur admise en 1994 est :

$$R = 109\,737,315\,683\,4 \pm 0,000\,002\,4 \text{ cm}^{-1}$$

(soit une précision relative de $2 \cdot 10^{-11}$). Elle a été obtenue grâce à des expériences sur différentes transitions à 2 photons¹⁴, et elle est bien sûr encore susceptible d'améliorations.

¹⁴Pour plus de détails, on pourra consulter : B. Cagnac, M. Plimmer, L. Julien, F. Biraben, Rep. Prog. Physics **57**, 853 (1994).

Complément III.6

Largeur spectrale des lasers : formule de Schawlow-Townes

La largeur spectrale de la plupart des lasers monomode est liée à des problèmes techniques de stabilité de longueur optique de la cavité laser (voir Chapitre III, § C.3.a). Il est néanmoins important d'identifier les facteurs ultimes limitant la pureté spectrale d'un laser. Il est également important de comprendre pourquoi la distribution spectrale du rayonnement d'un laser est bien plus fine que la bande passante de la cavité contenant le milieu amplificateur et que la courbe de gain de cet amplificateur. On se propose donc dans ce complément d'évaluer la largeur spectrale ultime d'un laser très au-dessus du seuil, par une méthode heuristique¹.

L'idée de base est que le milieu amplificateur placé dans la cavité laser émet de temps en temps des *photons spontanés dans le mode laser*. Le champ correspondant à un photon spontané, d'amplitude E_{sp} et de phase aléatoire, vient s'ajouter au champ laser E_L qui subit alors une fluctuation d'amplitude et de phase. La variation d'amplitude est corrigée par le gain qui, en raison du mécanisme de saturation, agit comme une « force de rappel ». En revanche, l'absence de force de rappel sur la phase permet aux variations de phase de persister et s'accumuler. Ainsi, lors des émissions spontanées successives, *la phase subit une diffusion aléatoire* et au bout d'un temps τ_c (temps de corrélation) le champ laser a perdu la mémoire de sa phase initiale : la phase du champ ne peut donc pas être prédite sur une durée supérieure à τ_c et la fréquence qui est la dérivée de la phase par rapport au temps ne peut être définie à mieux que $1/\tau_c$. La largeur de la raie laser auquel conduit ce processus est donc de l'ordre de²

$$\Delta\omega_{ST} \approx \frac{1}{\tau_c} \quad (1)$$

Pour évaluer le temps de corrélation τ_c , nous allons utiliser la représentation du champ

¹Le traitement rigoureux dépasse le cadre de ce complément. Néanmoins, les idées physiques de base sont correctes.

²Ce raisonnement peut être rendu quantitatif à partir du théorème de Wiener-Khintchine, qui relie le spectre de puissance (ici le spectre optique habituel) à la fonction d'autocorrélation de la grandeur aléatoire (ici le champ laser).

laser par son vecteur de Fresnel $\mathcal{E}_L$ (voir Figure 1). Chaque émission spontanée est représentée par un vecteur $\mathcal{E}_{\text{sp}}$ de longueur constante E_{sp} , et de direction aléatoire, que l'on ajoute au champ laser. Les phénomènes de saturation ramènent l'amplitude du champ laser à sa valeur moyenne E_L , mais pas sa phase. Le processus se répétant, l'extrémité du vecteur de Fresnel du champ $\mathcal{E}_L$ effectue *une marche au hasard* avec un pas de longueur $\mathcal{E}_{\text{sp}}$ et décrit donc un cercle de rayon E_L . Au bout de N_{sp} émissions spontanées, l'extrémité du vecteur de Fresnel s'est donc déplacée en moyenne d'un angle $\Delta\varphi$ avec

$$(\Delta\varphi)^2 \approx N_{\text{sp}} \frac{E_{\text{sp}}^2}{2E_L^2} \quad (2)$$

Le facteur 2 au dénominateur de l'équation (2) rappelle que l'émission spontanée contribue à la diffusion de la phase seulement par la composante de $\mathcal{E}_{\text{sp}}$ tangentielle au cercle de rayon E_L .

Le temps de corrélation τ_c est celui au bout duquel $\Delta\varphi$ est de l'ordre de 1, car la phase aura alors diffusé d'environ un radian. Ce temps correspond à un nombre d'émissions spontanées

$$\overline{N_{\text{sp}}} \approx 2 \left(\frac{E_L}{E_{\text{sp}}} \right)^2 \quad (3)$$

L'énergie dans la cavité laser de volume V est égale à $\varepsilon_0 E_L^2 V/2 = \mathcal{N} \hbar \omega$ ($\mathcal{N}$ étant le nombre moyen de photons du mode laser dans la cavité) alors que l'énergie d'un photon spontané est $\hbar \omega = \varepsilon_0 E_{\text{sp}}^2 V/2$. Il s'ensuit que la formule (3) peut encore s'écrire

$$\overline{N_{\text{sp}}} \approx 2\mathcal{N} \quad (4)$$

En appelant Γ_{mod} la probabilité par unité de temps pour qu'un atome dans l'état supérieur b émette un photon spontané dans le mode laser, et N_b le nombre d'atomes dans l'état b , on trouve que le nombre de photons spontanés émis pendant τ_c dans le mode laser est :

$$\overline{N_{\text{sp}}} \approx \Gamma_{\text{mod}} N_b \tau_c \quad (5)$$

En combinant les équations (1), (4) et (5) on obtient une première expression pour la largeur de raie :

$$\Delta\omega_{ST} \approx \frac{1}{\tau_c} \approx \frac{\Gamma_{\text{mod}} N_b}{2\mathcal{N}} \quad (6)$$

Il est d'autre part possible de relier le taux d'émission spontanée dans le mode laser Γ_{mod} à la puissance de sortie du laser Φ_s . En effet, on montre que le taux d'émission stimulé par atome dans l'état b , tout comme le taux d'absorption par atome dans l'état a , est égal à $\Gamma_{\text{mod}} \mathcal{N}$ ($\mathcal{N}$ étant, rappelons-le, le nombre de photons dans le mode laser supposé grand devant 1). Quand le laser fonctionne très au-dessus du seuil, la perte d'énergie due à l'émission spontanée est négligeable et toute l'énergie résultant de la compétition entre les processus d'émission stimulée et d'absorption se retrouve dans l'énergie émise par le laser. De ce bilan d'énergie, nous déduisons donc :

$$\Phi_s = \Gamma_{\text{mod}} \mathcal{N} (N_b - N_a) \hbar \omega \quad (7)$$

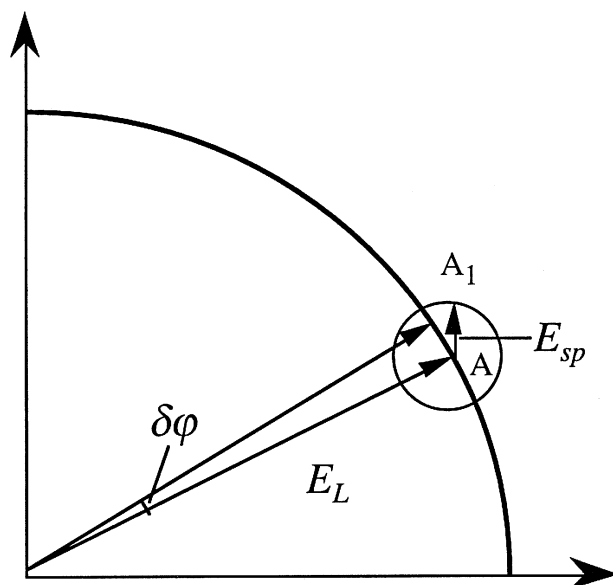


FIG. 1: Diagramme de Fresnel du champ laser E_L . Chaque émission spontanée dans le mode laser rajoute un champ de module E_{sp} et de phase aléatoire (évolution AA_1). L'amplitude du champ est corrigée par le phénomène de saturation du gain mais la phase a varié d'une quantité $\delta\varphi$ de l'ordre de E_{sp}/E_L . Ce processus est à la base de la diffusion de la phase dans le laser.

En éliminant Γ_{mod} entre les équations (6) et (7), nous trouvons

$$\Delta\omega_{ST} \approx \frac{N_b}{2(N_b - N_a)} \frac{\Phi_s}{\hbar\omega} \frac{1}{\mathcal{N}^2} \quad (8)$$

Il est possible de transformer cette expression en remarquant que la puissance de sortie Φ_s est reliée au nombre de photons du mode laser $\mathcal{N}$ dans la cavité. Pour une cavité linéaire avec un miroir de sortie partiellement réfléchissant, nous déduisons de l'équation (35) du complément III.1

$$\mathcal{N}\hbar\omega = \frac{\Phi_s}{\Delta\omega_{\text{cav}}} \quad (9)$$

où $\Delta\omega_{\text{cav}}$ est la largeur à mi-hauteur d'un mode de la cavité laser (voir Complément III.1, Eq. 21). Ceci donne finalement pour largeur de raie :

$$\begin{aligned} \Delta\omega_{ST} &\approx \frac{N_b}{2(N_b - N_a)} \frac{\hbar\omega}{\Phi_s} (\Delta\omega_{\text{cav}})^2 \\ &\approx \frac{N_b}{2(N_b - N_a)} \frac{\Delta\omega_{\text{cav}}}{\mathcal{N}} \end{aligned} \quad (10)$$

Ce résultat coïncide, à un facteur numérique près, avec la formule plus rigoureuse établie par Shawlow et Townes :

$$\Delta\omega_{ST} = \pi \frac{N_b}{2(N_b - N_a)} \frac{\hbar\omega}{\Phi_s} (\Delta\omega_{\text{cav}})^2 \quad (11)$$

Cette formule fait intervenir la puissance extraite du laser (puissance de sortie Φ_s), la largeur $\Delta\omega_{\text{cav}}$ d'un mode de la cavité laser et un facteur dépendant des populations N_b et N_a des niveaux extrêmes de la transition (ce facteur est une fonction décroissante de l'inversion de population qui vaut 1 pour une inversion totale, c'est-à-dire pour $N_a = 0$).

Remarquons que la formule (11) jointe à l'équation (9) permet d'évaluer le rapport entre la largeur ultime de la raie laser et la largeur d'un mode de la cavité :

$$\frac{\Delta\omega_{ST}}{\Delta\omega_{\text{cav}}} = \frac{\pi}{\mathcal{N}} \frac{N_b}{N_b - N_a} \quad (12)$$

Comme le nombre de photons $\mathcal{N}$ est très grand devant 1 et que le facteur d'inversion est de l'ordre de l'unité, cette formule montre que dans la limite Schawlow-Townes, *la raie laser est beaucoup plus fine qu'un mode de la cavité qui est pourtant déjà très étroit*. On comprend également que plus $\mathcal{N}$ est grand, plus l'oscillation laser est stable vis-à-vis des fluctuations. Enfin la présence du facteur $N_b/(N_b - N_a)$ traduit le fait que l'absorption de photons de l'onde laser est une cause de fluctuation d'autant plus faible que l'inversion de population est meilleure.

Pour les applications numériques utilisant la formule (11), il faut estimer $\Delta\omega_{\text{cav}}$. Cela peut-être fait au moyen des formules (21) et (30) du complément III.1 qui conduisent à

$$\Delta\omega_{\text{cav}} = T \frac{c}{L} \quad (13)$$

où $L/2$ est la distance séparant les deux miroirs de la cavité linéaire. Par exemple, pour un laser Hélium-Néon émettant une puissance $P_S = 1\text{mW}$, ayant des miroirs séparés de 1m (soit $L = 2\text{m}$) et avec un miroir de sortie de transmission T égal à 2 %, on a :

$$\frac{\Delta\omega_{\text{cav}}}{2\pi} \approx 5 \times 10^5 \text{ Hz} \quad ; \quad \frac{\Delta\omega_{ST}}{2\pi} \approx 10^{-3} \text{ Hz}$$

Un laser ayant une raie aussi étroite constituerait un oscillateur de facteur de qualité extraordinaire. En fait, la valeur de la largeur donnée par la formule de Schawlow-Townes est très inférieure aux fluctuations « techniques » de fréquence ayant leur origine dans les variations de longueur optique de la cavité laser. L'amélioration de la pureté est donc possible en augmentant la stabilité de la cavité. C'est une démarche de ce type qui est mise en pratique en métrologie de haute précision.

Un deuxième exemple intéressant est celui d'une diode laser émettant 1 mW. La cavité est ici très courte et de petite finesse. Dans ce cas, les grandeurs typiques sont :

$$\frac{\Delta\omega_{\text{cav}}}{2\pi} \approx 2 \times 10^{10} \text{ Hz} \quad ; \quad \frac{\Delta\omega_{ST}}{2\pi} \approx 10^6 \text{ Hz}$$

En pratique, on observe des largeurs de raies supérieures à cette dernière valeur³ d'environ un ordre de grandeur (10 à 30 MHz). Il semble donc que les possibilités de réduire la

³L'origine de cet élargissement supplémentaire est bien identifié : il s'agit des fluctuations d'indice de réfraction, induites par les fluctuations d'intensité de la diode, qui provoquent des variations de la longueur optique de la diode.

largeur soient faibles puisque le laser est proche de la limite de Schawlow-Townes. On peut cependant dépasser cette limite, c'est-à-dire obtenir une raie plus fine, en couplant une telle diode à une cavité Fabry-Perot externe⁴ (voir Figure 2) et ainsi atteindre des largeurs de raies inférieures au kHz.

Le fait que l'on obtienne ainsi des diodes laser de pureté spectrale meilleures que la limite de Schawlow-Townes est un des nombreux exemples de la prudence avec laquelle il faut utiliser les concepts de limite fondamentale et de l'esprit critique qu'il faut avoir à leur égard. La limite de Schawlow-Townes est incontestablement une limite fondamentale et le raisonnement qui y conduit est parfaitement correct. Mais une telle limite peut-être contournée si on se place dans une situation n'obéissant pas aux hypothèses du raisonnement ce qui est le cas lorsque la cavité est couplée à une deuxième cavité beaucoup plus longue.

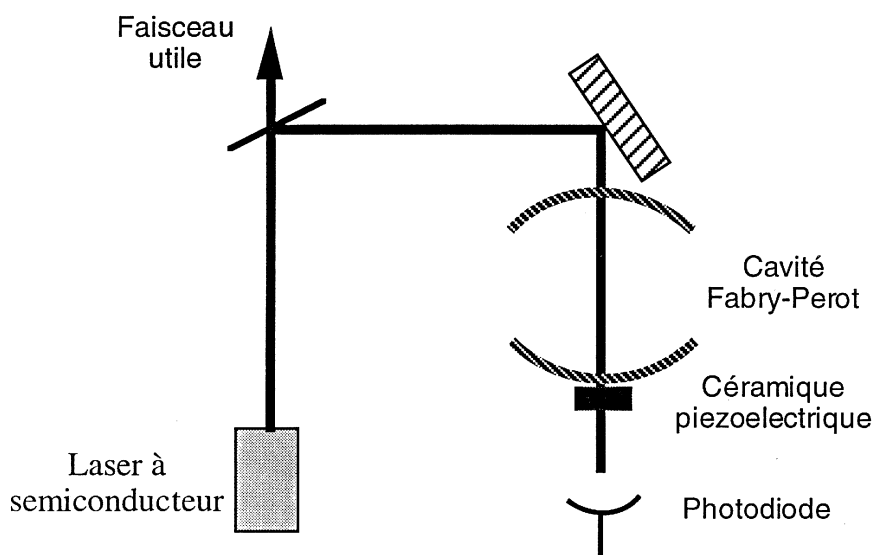


FIG. 2: **Diode laser couplée à une cavité extérieure.** Si la longueur de cette cavité est telle qu'un de ses modes coïncide avec un mode de la diode laser, un grand nombre de photons peut être stocké dans la cavité externe, ce qui permet d'affiner considérablement la largeur spectrale de la diode laser.

⁴Le champ de la cavité est couplé au champ stocké dans la cavité externe ce qui contribue à augmenter son inertie.

Complément III.7

Faisceau de lumière incohérente et lumière laser

La différence entre la lumière laser et la lumière émise par une source incohérente ne peut être pleinement appréciée qu'en rappelant quelques notions de photométrie que nous présentons dans la partie III.7.1. Nous montrons dans la partie 2 qu'elles limitent dramatiquement la densité d'énergie que l'on peut obtenir avec une source classique, à la différence de la lumière laser (partie 3). Loin d'être circonstanciées, les lois de la photométrie des sources classiques sont de nature fondamentale, puisqu'on peut les déduire des principes de la thermodynamique. Une autre façon de relier les propriétés de la lumière aux concepts fondamentaux de la physique est de les examiner dans le cadre de la physique statistique des photons, comme nous le verrons dans les parties 4 et 5.

1 Conservation de la luminance pour une source classique (incohérente)

1.1 Étendue géométrique, luminance

Une source incohérente émet de la lumière dans toutes les directions. Un faisceau lumineux issu de cette source peut être décomposé en pinceaux élémentaires : la lumière étant incohérente, la puissance totale transportée par le faisceau est la somme des puissances transportées par les pinceaux élémentaires.

Un pinceau élémentaire est défini par l'élément de source dS sur lequel il s'appuie, et un deuxième élément de surface dS' (Fig. 1). L'étendue géométrique du pinceau est

$$dU = \frac{dS \cos \theta \, dS' \cos \theta'}{MM'^2} \quad (1)$$

où θ et θ' sont les angles entre la direction moyenne MM' du pinceau et les normales $\mathbf{n}$

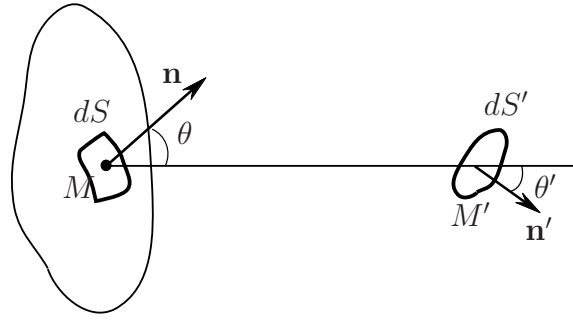


FIG. 1: **Dispositions relatives d'un élément de source dS et d'un élément de surface éclairée dS' .** Ces deux éléments déterminent un pinceau lumineux, ensemble des rayons passant dans dS et dS' .

et $\mathbf{n}'$ aux éléments de surface. On peut évidemment introduire les angles solides

$$d\Omega = \frac{dS' \cos \theta'}{MM'^2} \quad (2a)$$

et

$$d\Omega' = \frac{dS \cos \theta}{MM'^2}, \quad (2b)$$

et récrire

$$dU = dS \cos \theta \, d\Omega = dS' \cos \theta' d\Omega'. \quad (3)$$

Le flux énergétique du rayonnement transporté dans ce pinceau, c'est-à-dire sa puissance, s'exprime par la formule

$$d\phi = L(M, \theta) dU \quad (4)$$

où $L(M, \theta)$ est la luminance au point M . Pour de nombreuses sources (et en particulier le corps noir à une température homogène), la luminance ne dépend ni du point M ni de la direction θ , et nous nous plaçons dans ce cas pour ne pas alourdir les notations.

Un faisceau lumineux s'appuie sur deux diaphragmes (dont l'un peut être confondu avec la source) et la puissance ϕ qu'il transporte est évidemment la somme des contributions $d\phi$. Si la luminance L est homogène et indépendante de la direction, on a alors

$$\phi = LU \quad (5)$$

où l'étendue géométrique U est une quantité purement géométrique.

Les définitions et propriétés ci-dessus ont été implicitement données pour un faisceau monochromatique. Si on a une source polychromatique on définit une luminance spectrale différentielle $\mathcal{L}(\omega)$, et il est entendu que les propriétés ci-dessus s'appliquent à chaque élément spectral $d\omega$ caractérisé par la luminance $\mathcal{L}(\omega)d\omega$. Les divers éléments spectraux d'une source incohérente étant indépendants, leurs contributions s'ajoutent au niveau des grandeurs énergétiques.

1.2 Conservation de la luminance

Si un faisceau lumineux se propage dans un milieu non absorbant, et que les dioptries qu'il rencontre ont reçu un traitement antiréfléchissant, la puissance est conservée. On montre de plus que pour un système optique parfait (stigmatique) l'étendue géométrique est conservée au cours de la propagation. La formule (5) montre alors qu'il y a conservation de la luminance au cours de la propagation : le faisceau reste caractérisé par une luminance inchangée, ce qui permet de calculer aisément les grandeurs photométriques dont on a besoin après un instrument d'optique, comme on va le voir ci-dessous.

Notons que si le faisceau rencontre des milieux absorbants, ou des surfaces partiellement réfléchissantes, la luminance va décroître même si le système optique est stigmatique. De plus, si le système n'est pas stigmatique, l'étendue géométrique ne peut que croître, et ceci conduit également à une diminution de la luminance.

En définitive, on retiendra qu'au cours de la propagation d'un faisceau issu d'une source incohérente, la luminance est conservée si les systèmes optiques traversés sont parfaits ; sinon, elle ne peut que décroître.

Remarques

(i) Cette propriété est étroitement liée au deuxième principe de la thermodynamique, qui serait violé si on pouvait augmenter la luminance au cours de la propagation. Il y a aussi un lien avec la mécanique statistique qui apparaît clairement si on remarque que la conservation de l'étendue géométrique n'est pas autre chose que l'analogue du théorème de Liouville, c'est-à-dire la conservation du volume dans l'espace des phases, lors d'une évolution hamiltonienne.

(ii) Les propriétés vues ci-dessus ne sont plus valables si les milieux traversés ont un comportement non-linéaire. Remarquons par exemple que dans ce cas on peut avoir des changements de fréquence, et qu'il n'y a plus d'indépendance entre les éléments spectraux.

2 Éclairement maximal d'une surface avec une source incohérente

L'éclairement d'une surface est la puissance reçue par unité de surface

$$E = \frac{d\phi}{dS'} . \quad (6)$$

Lorsqu'on dispose d'une source incohérente de luminance L , quelle est l'éclairement maximal que l'on peut obtenir ?

Considérons d'abord le cas de la figure 1, où la surface éclairée dS' est directement en vue de la source. En utilisant (3) et (5), on voit que l'éclairement dû à l'élément de source dS , vue sous l'angle solide $d\Omega'$, vaut

$$dE = L \cos \theta' d\Omega' \quad (7)$$

En intégrant sur toute la source, on trouve l'éclairement total. Si la source est très grande, comparée à la distance MM' , l'angle solide total est quasiment égal à 2π , mais dans l'intégration le terme $\cos \theta'$ donne un facteur $1/2$. On trouve finalement

$$\Rightarrow E \leq \pi L, \quad (8)$$

l'égalité étant obtenue si la source est vue depuis M' sous un angle solide de 2π .

Peut-on, à l'aide d'un instrument d'optique, augmenter cette valeur limite supérieure de l'éclairement ? La réponse est non, car en vertu de la conservation de la luminance, la formule (6) reste vraie même si le pinceau lumineux est passé à travers un instrument parfait. Lorsqu'on intègre cette formule, l'angle solide est limité par l'angle solide sous lequel est vue la pupille de sortie de l'instrument. L'éclairement dépend donc uniquement de l'ouverture de l'instrument, et il est d'autant plus grand que cette ouverture est grande, sans pouvoir dépasser la valeur limite (8).

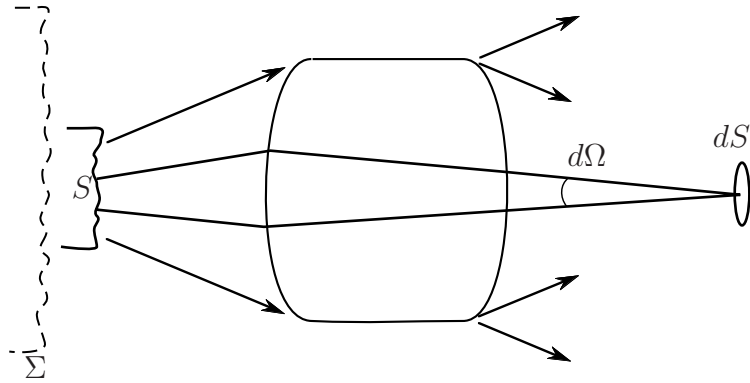


FIG. 2: **L'élément de surface dS' est maintenant éclairé à travers un instrument d'optique.** Si l'instrument est parfait, on peut se ramener à la situation de la figure 1 : l'image Σ de S est vue depuis dS' avec une luminance inchangée, et l'éclairement en dS' ne peut dépasser la valeur πL (instrument très ouvert, dont la pupille de sortie est vue sous un angle solide de 2π stéradian).

Ainsi, quel que soit l'instrument utilisé, la luminance fixe l'ordre de grandeur de l'éclairement maximal que l'on peut obtenir en un point à partir d'une source donnée, c'est-à-dire en définitive le champ électrique maximal, ce qui est la quantité importante pour l'interaction atome rayonnement.

Remarques

(i) Avec des systèmes à miroirs elliptiques ou paraboliques, on peut éclairer le point M' sous un angle solide supérieur à 2π . Dans ce cas la limite (8) doit être multipliée par 2.

(ii) Les considérations ci-dessus permettent de comprendre comment on peut, avec une loupe qui concentre la lumière du soleil, enflammer des brindilles : on augmente le diamètre apparent, qui passe de 2α ($30'$) à $2u$ (qui peut atteindre aisément 60°) (Figure 3). L'éclairement est alors (au facteur $\cos \theta'$ près) multiplié par $\left(\frac{u}{\alpha}\right)^2$ (soit $120^2 = 14400$ dans notre exemple). Il peut sembler surprenant que l'éclairement obtenu ne dépende que de l'ouverture numérique u , et pas de la taille de la loupe (proportionnelle à sa distance focale f , pour une

ouverture u fixée). Il en est pourtant ainsi, et si c'est l'éclairement qui compte, alors une petite loupe fera très bien l'affaire. Mais si l'on veut allumer un feu, on comprend sur la figure l'intérêt de prendre une grande loupe : la surface éclairée Σ (l'image du soleil) croît avec f , et on pourra donc embraser des brindilles d'autant plus grosses que la distance focale (et donc la loupe) sera grande, l'ouverture u restant la même.

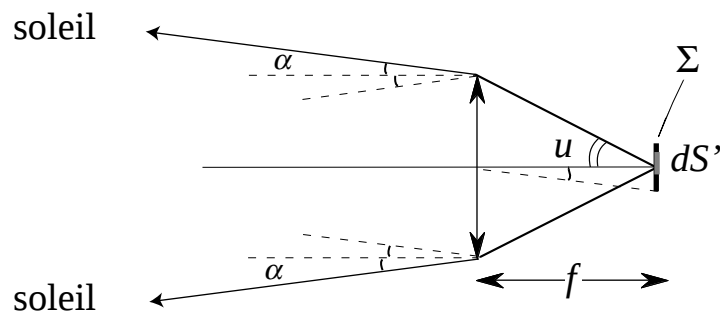


FIG. 3: **Éclairement au foyer d'une loupe placée perpendiculairement aux rayons solaires.** L'augmentation d'éclairement est liée au rapport de l'ouverture u de l'instrument au rayon angulaire α du soleil ($2\alpha = 30'$). L'éclairement ne dépend que de l'angle u .

Ordres de grandeurs

D'après la loi de Planck, un corps noir à 3000 K émet environ 500 W/cm^2 dans un demi espace, sur une bande spectrale s'étendant de l'infrarouge ($1,5\mu\text{m}$) au vert ($0,5\mu\text{m}$) soit de $4 \times 10^{14} \text{ Hz}$. C'est à peu près ce que fournit le filament d'une lampe à incandescence. La luminance est π fois moins grande, et l'éclairement maximal que l'on pourra obtenir est donc de l'ordre de 500 W/cm^2 . Une lampe à arc à haute pression, ou le soleil, donneront au mieux 10 fois plus, c'est-à-dire 5 kW/cm^2 .

La quantité évaluée ici est la densité de puissance totale reçue par unité de surface. Mais une quantité souvent plus importante est la densité de puissance par unité de surface et par unité de bande spectrale : elle vaut environ $10^{-12} \text{ W cm}^{-2} \text{ Hz}^{-1}$ dans notre exemple. Pour comprendre l'intérêt de cette grandeur, supposons par exemple que nous voulions exciter une transition atomique : seule une bande de fréquence très étroite, de l'ordre du MHz, sera utilisée, et l'éclairement utile est de l'ordre de 10^{-6} W/cm^2 . De tels éclairagements sont beaucoup trop faibles pour atteindre le régime des champs forts, et généralement insuffisants pour saturer une transition atomique (les intensités de saturation, conduisant à un paramètre de saturation à résonance s de l'ordre de 1, sont plutôt de l'ordre de 10^{-3} W/cm^2).

3 Éclairement maximal d'une surface avec de la lumière laser

Considérons maintenant un faisceau laser de 15 Watts : étant spatialement cohérent, il peut, comme nous le montrons dans le Complément III.2, être focalisé sur une tache

de diffraction élémentaire, dont la dimension est de l'ordre de la longueur d'onde¹, soit une surface inférieure à $1\mu\text{m}^2$. Nous disposons donc maintenant de 10^9 W/cm^2 , au lieu des 500 W/cm^2 précédents. De plus, cette puissance est fournie sur une bande spectrale qui peut être inférieure au MHz. Si on souhaite exciter une transition atomique étroite, la puissance est utilisable en totalité : une telle source se situe d'emblée dans le régime des champs forts.

Dans le domaine temporel, la technique de synchronisation des modes présentée dans le § 3.4.1 permet de concentrer la puissance laser dans le temps sous forme d'impulsions de durée 0,1 ns répétées toutes les 10ns. La puissance maximale est 100 fois plus grande que la puissance moyenne et atteint 10^{11} W/cm^2 . Avec des lasers à colorant ou des lasers à saphir dopé au titane, il est possible d'atteindre des impulsions encore plus courtes et des puissances crêtes encore plus grandes pouvant dépasser 10^{15} W/cm^2 . Dans ces conditions, le champ électrique de l'onde électromagnétique est *supérieur* au *champ de Coulomb* exercé par le proton sur l'électron d'un atome d'hydrogène dans son état fondamental. On accède ainsi à une nouvelle échelle d'énergie dans l'interaction entre la lumière et la matière.

On constate que **ce qui caractérise la lumière laser, c'est la possibilité d'être concentrée, dans l'espace, et dans le temps** (ou dans le spectre), offrant ainsi la possibilité d'obtenir des densités d'énergie énormes, à la base de la plupart des applications des lasers.

4 Nombre de photons par mode d'une cavité

4.1 Rayonnement thermique dans une cavité

L'énergie du rayonnement thermique enfermé dans une cavité (par exemple un cube d'arête L) est proportionnelle au volume L^3 du cube. Ainsi la loi de Planck nous donne l'énergie comprise dans la bande de fréquence $d\omega$ autour de ω :

$$dE = L^3 \frac{\hbar\omega^3}{\pi^2 c^2} \frac{1}{\exp\left(\frac{\hbar\omega}{k_B T}\right) - 1} d\omega \quad (9)$$

On sait que dans une telle cavité, on peut définir des modes discrets, caractérisés par un vecteur $\mathbf{k}$ (ou par une fréquence $\omega = ck$ et une direction de propagation), et une polarisation. La densité de modes dans l'espace réciproque (espace des $\mathbf{k}$) est, elle aussi, proportionnelle au volume L^3 , et le nombre de modes de fréquence comprise dans la bande

¹Une autre façon de comprendre ce résultat est de remarquer que ces 15 Watts sont émis dans un faisceau très collimaté (très peu divergent), à la différence de la source classique qui rayonne dans 2π stéradians. Cette propriété est mise à profit dans de nombreuses applications allant du guidage des tunneliers à la mesure de distance terre-lune.

$d\omega$ autour de ω vaut²

$$dN = L^3 \frac{\omega^2}{\pi^2 c^2} d\omega \quad (10)$$

On constate alors que le nombre de photons dans un mode de fréquence ω est donné par la distribution de Bose-Einstein³

$$\mathcal{N}(\omega) = \frac{1}{\exp\left(\frac{\hbar\omega}{k_B T}\right) - 1} \quad (11)$$

Ce nombre, indépendant du volume et de la forme de la cavité, caractérise le *degré de dégénérescence* des photons : les photons dans un même mode ne peuvent être distingués ni par leur fréquence, ni par leur direction de propagation, ni par leur polarisation. Pour une source thermique à 3000 K, le nombre $\mathcal{N}(\omega)$ de photons par mode est très inférieur à 1, de l'ordre de 3×10^{-4} au milieu du spectre visible (5×10^{14} Hz, soit $\lambda = 0,6 \mu\text{m}$), et de l'ordre de 10^{-2} pour la lumière solaire (température superficielle : 5 800 K) : dans les deux cas, *il y a beaucoup moins d'un photon par mode*, de fréquence correspondant à de la lumière visible.

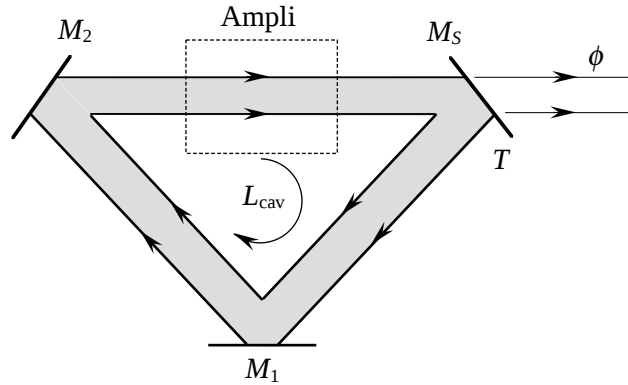


FIG. 4: **Cavité laser de longueur** L_{cav} . Le miroir de sortie a un coefficient de transmissions T , et la puissance sortante est Φ .

4.2 Cavité laser

Considérons maintenant un laser mono-mode émettant une puissance Φ . Si le miroir de sortie a un coefficient de transmission T , et la cavité une longueur totale L_{cav} , (Figure 4) le nombre de photons contenus dans la cavité laser vaut

$$\mathcal{N}_{\text{cav las}} = \frac{\Phi}{\hbar\omega} \frac{1}{T} \frac{L_{\text{cav}}}{c} \quad (12)$$

²Le calcul peut être fait comme au paragraphe A.2 du complément II.3, en n'oubliant pas de multiplier le résultat par 2 pour tenir compte des 2 polarisations associées à chaque $\mathbf{k}$.

³Voir E. Brézin, cours de Physique Statistique de l'École polytechnique.

Si nous prenons un laser Hélium-Néon délivrant 1 mW, et une cavité de longueur 1 m, on trouve environ 2×10^9 photons dans le mode de la cavité laser. Cette dégénérescence très élevée est manifestement une caractéristique essentielle de la lumière laser, par opposition au rayonnement thermique.

5 Nombre de photons par mode pour un faisceau libre

5.1 Mode libre propagatif

Considérons maintenant un pinceau lumineux émis dans l'espace libre, par exemple à partir d'un petit trou percé dans la paroi d'une enceinte (rayonnement du corps noir, Figure 5). Comment généraliser la notion de mode dans ce cas, alors qu'on n'a plus de condition aux limites pour discrétiser le problème ?

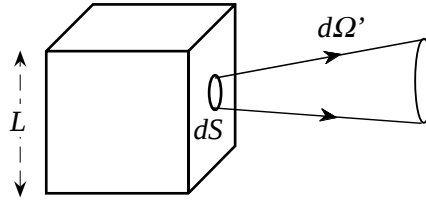


FIG. 5: **Pinceau de rayonnement thermique** obtenu en perçant un petit trou de surface dS dans une cavité contenant du rayonnement à température T . La luminance est celle du corps noir à température T . Un mode libre à la fréquence ω est associé à un pinceau minimal, tel que $dU = dS d\Omega' = \lambda^2$ ($\lambda = 2\pi c/\omega$), et à un paquet d'onde minimal de durée $\Delta t = 2\pi/\Delta\omega$.

Nous définirons un mode libre en considérant d'abord un pinceau d'étendue minimale compatible avec la nature ondulatoire du rayonnement. Un tel pinceau élémentaire, de section dS , a une divergence angulaire (due à la diffraction) $\lambda/\sqrt{S}$, et son étendue géométrique vaut

$$dU_{\text{mini}} = \lambda^2 = 4\pi^2 \frac{c^2}{\omega^2} \quad (13)$$

Cette formule est évidemment la conséquence des relations

$$\Delta x \cdot \Delta k_x \geq 2\pi \quad (14a)$$

$$\Delta y \cdot \Delta k_y \geq 2\pi \quad (14b)$$

qui contraignent les dispersions des variables conjuguées $\{x, k_x\}$ et $\{y, k_y\}$, caractérisant respectivement les tailles transversales du pinceau dans l'espace réel et réciproque.

En ce qui concerne la direction de propagation z (longitudinale), nous considérerons les variables conjuguées {temps, fréquence}, et nous nous intéresserons à un paquet d'onde de durée Δt . Sa largeur spectrale obéit alors à la relation

$$\Delta t \cdot \Delta \omega \geq 2\pi , \quad (15)$$

et pour un paquet d'onde minimal on a

$$\Delta t \cdot \Delta \omega = 2\pi . \quad (16)$$

À la dispersion $\Delta \omega$ est associée une dispersion

$$\Delta k_z = \frac{\Delta \omega}{c} \quad (17)$$

de la composante k_z du vecteur d'onde, tandis qu'à la durée Δt on associe évidemment l'extension longitudinale Δz du paquet d'onde :

$$\Delta z = c \Delta t . \quad (18)$$

On en déduit alors

$$\Delta z \cdot \Delta k_z \geq 2\pi \quad (19)$$

l'égalité étant obtenue pour un paquet d'onde minimal

$$\Delta z \cdot \Delta k_z = 2\pi \quad (20)$$

Un mode de l'espace libre est caractérisé par une extension minimale dans l'espace des phases (produit direct de l'espace réel par l'espace réciproque). Les dispersions en position et vecteur d'onde sont reliées par les égalités (14a), (14b) et (20). (Pour la direction longitudinale, on peut remplacer le couple de variables conjuguées $\{z, k_z\}$ par $\{t, \omega\}$).

Un tel mode est en fait associé à une cellule élémentaire de l'espace des phases. Il est impossible de mieux préciser en position et en vecteur d'onde un paquet d'onde électromagnétique.

Remarque

Si on se restreint aux dimensions transversales, le volume de l'espace des phases à 4 dimensions (x, y, k_x, k_y) est l'étendue U , et la cellule élémentaire a pour volume λ^2 . Le rapport U/λ^2 est le nombre de Fresnel : c'est le nombre de pinceaux minimaux indépendants contenus dans le faisceau.

5.2 Pinceau de rayonnement thermique

La luminance spectrale $\mathcal{L}(\omega)$ du corps noir, ou encore d'un trou percé dans la paroi d'une enceinte, contenant du rayonnement thermique, est égale à

$$\mathcal{L}(\omega) = \frac{c}{4\pi} \frac{1}{L^3} \frac{dE}{d\omega} \quad (21)$$

où la densité volumique d'énergie du rayonnement $\frac{1}{L^3} \frac{dE}{d\omega}$ est donnée par la loi de Planck (9).

L'énergie contenue dans un mode, caractérisé par (14a), (14b) et (20), vaut

$$\begin{aligned} E_{1 \text{ mode}} &= \frac{1}{2} \mathcal{L}(\omega) \cdot dU_{\text{mini}} \Delta\omega \Delta t \\ &= 4\pi^3 \frac{c^2}{\omega^2} \mathcal{L}(\omega) \end{aligned} \quad (22)$$

(Le facteur $\frac{1}{2}$ rend compte de l'existence de deux polarisations orthogonales, et donc deux modes distinguables pour un paquet d'onde minimal). En utilisant (9), on trouve

$$E_{1 \text{ mode}} = \frac{\hbar\omega}{\exp\left(\frac{\hbar\omega}{k_B T}\right) - 1} \quad (23)$$

c'est-à-dire que le *nombre de photons par mode libre du rayonnement du corps noir* est encore donné par la distribution de Bose-Einstein (10). Ce nombre est *petit devant 1* pour la lumière visible émise par toutes les sources traditionnelles.

5.3 Faisceau émis par un laser

La divergence d'un faisceau laser monomode transverse n'est due qu'à la diffraction. C'est donc un pinceau minimal. Pour définir un mode, on considère un paquet d'onde de durée minimale compatible avec la largeur de raie $\Delta\omega$, suivant la relation (20)

$$\Delta t = \frac{2\pi}{\Delta\omega} \quad (24)$$

Le nombre de photons par mode libre est donc

$$\mathcal{N}_{\text{laser}} = \frac{\phi}{\hbar\omega} \frac{2\pi}{\Delta\omega} \quad (25)$$

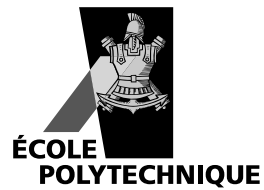
La largeur de raie $\Delta\omega$ est souvent très inférieure à $2\pi c/L_{\text{cav}}$, et en comparant (25) à (13) on voit que le nombre de photons par mode dans un faisceau laser est très grand. Si on prend pour largeur de raie la limite de Schawlow Townes (cf. § 3.3.3.b, et complément III.6), qui vaut typiquement 10^{-3} Hz, le nombre de photons par mode peut dépasser 10^{15} pour un laser de quelques mW. Même si on prend des largeurs de raies plus réalistes de l'ordre du MHz, on trouve des nombres supérieurs à 10^9 photons par mode pour des cavités de longueur de l'ordre du mètre.

Nous retrouvons dans le cas des faisceaux libres, le fait qu'*un faisceau laser a un nombre de photons par mode très grand devant 1, par opposition à un faisceau issu d'une source ordinaire.*

6 Conclusion

Nous avons vu que la lumière laser offre la possibilité d'être concentrée aussi bien dans l'espace que dans le temps, ou de façon complémentaire, angulairement (faisceau très collimaté) et dans le spectre (faisceau très monochromatique). Ces possibilités sont fondamentalement reliées au fait que tous les photons sont dans le même mode : étant des bosons indiscernables ils sont cohérents spatialement et temporellement, et on peut les *concentrer sur les tailles minimales* autorisées par la relation *taille $\times$ divergence* (diffraction), et par la relation *temps $\times$ fréquence*.

Ainsi, on peut dire que ce qui caractérise fondamentalement la lumière laser, comparée à la lumière ordinaire, c'est un nombre de photons par mode très supérieur à 1.



ÉDITION 2005

Achevé d'imprimer le 21 septembre 2005 sur les presses
du Centre Poly-Média de l'École Polytechnique.



Dépôt légal : 3^e trimestre 2005
N° ISBN 2 – 7302 – 1271 – X

TOME I

IMPRIMÉ EN FRANCE