

Cours laser
Licence/Master de chimie

Table des matières

1	Rappels et compléments...	5
1.1	Statistiques quantiques	5
1.1.1	Particules identiques en mécanique quantique: bosons et fermions	5
1.1.2	Applications	8
1.1.3	Espace de Fock	10
1.2	Théorie quantique...	12
1.2.1	Brefs rappels d'électromagnétisme	12
1.2.2	Champ électromagnétique quantique	14
1.2.3	Électrodynamique quantique en cavité.	15
1.2.4	Coefficients d'Einstein.	18
1.3	Appendice technique...	27
1.3.1	Position du problème	27
1.3.2	Théorie des perturbations	28
1.3.3	Couplage stationnaire à un continuum, Règle d'or de Fermi	29
2	Principe de fonctionnement d'un laser	35
2.1	Laser continu	35
2.1.1	Cavité Fabry-Pérot	36
2.1.2	Amplification optique et inversion de population	41
2.1.3	Dynamique d'un laser à 3 niveaux pompé optiquement	43
2.1.4	"Laserologie"	46
2.2	Propriétés de cohérence	53
2.2.1	Modes longitudinaux d'un laser	54
2.2.2	Quelques applications	57
2.3	Lasers à impulsion	72
2.3.1	Impulsions déclenchées	76

2.3.2	Blocage de mode	76
2.3.3	Applications à la spectroscopie résolue en temps	80
2.4	Complément : sources non-classiques	83
2.4.1	Sources à photons uniques	83
2.4.2	Application à la cryptographie quantique	85
3	Propagation lumineuse	89
3.1	Optique gaussienne	89
3.1.1	Aspects qualitatifs de la propagation lumineuse	89
3.1.2	Approximation de l'enveloppe lentement variable	90
3.1.3	Faisceau gaussien	92
3.1.4	Optique gaussienne	94
3.1.5	Microscopie à haute résolution	98
3.2	Optique non-linéaire	101
3.2.1	Retour sur la théorie des perturbations	101
3.2.2	Applications	105
4	Forces radiatives	115
4.1	Force dipolaire et pression de radiation	115
4.1.1	Force dipolaire	115
4.1.2	Pression de radiation	121
4.2	Applications	122
4.2.1	Pinces optiques	122
4.2.2	Refroidissement d'atomes par laser	125
4.2.3	Fusion inertielle	126

Chapitre 1

Rappels et compléments de mécanique quantique

Le principe de fonctionnement du laser repose sur des phénomènes d'origine purement quantique. Dans ce premier chapitre, nous développerons les outils nécessaires au reste du cours : dans un premier temps, nous établirons la théorie des perturbations qui permet de déterminer (de manière approchée) l'évolution d'un système quantique puis nous nous intéresserons aux particules identiques en mécanique quantique en nous focalisant plus particulièrement sur le cas des bosons, dont font partie les photons. Ceci nous permettra enfin d'établir une théorie quantique de l'électromagnétisme et introduira les concepts d'émission spontanée et d'émission stimulée qui sont au cœur du fonctionnement d'un laser.

1.1 Statistiques quantiques

1.1.1 Particules identiques en mécanique quantique : bosons et fermions

Peut-on différencier deux électrons? Cette question, qui peut sembler au premier abord purement philosophique, joue un rôle central en mécanique quantique. En effet, considérons le dispositif expérimental de la figure 1.1 dans lequel on étudie la collision de deux électrons numérotés 1 et 2 et venant respectivement de la droite et la gauche. Après la collision, supposons que l'on observe qu'un électron est parti vers le haut et l'autre vers le bas. Une

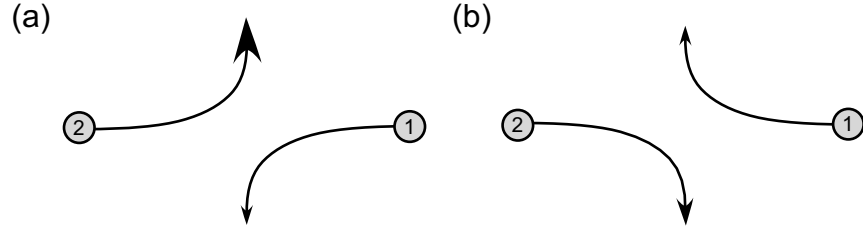


FIG. 1.1 – Indiscernabilité de deux électrons : lors d'une collision entre les électrons 1 et 2, les états finaux des processus (a) et (b) sont indiscernables.

question se pose : est-ce l'électron 1 ou 2 qui est parti vers le haut ? Si l'on observe pas le système durant la collision, il est impossible de départager les deux éventualités et, comme de coutume en mécanique quantique, cela signifie qu'à la fin de la collision, chaque électron se trouve dans une superposition des états diffusés vers le haut et vers le bas.

L'exemple de la diffusion de deux électrons peut se formaliser et se généraliser en introduisant l'opérateur d'échange. Considérons un système à 2 particules, baptisées 1 et 2. Si la particule 1 (resp. la particule 2) est dans l'état $|\psi_a\rangle$ (resp. $|\psi_b\rangle$), on note $|1 : \psi_a; 2 : \psi_b\rangle$ l'état du système complet $1 \oplus 2$. Par définition, l'action opérateur d'échange, $\hat{P}$ sur un tel état revient à échanger les états dans lesquels sont les particules, autrement dit :

$$\hat{P}|1 : \psi_a; 2 : \psi_b\rangle = |1 : \psi_b; 2 : \psi_a\rangle$$

En appliquant deux fois l'opérateur d'échange, on voit que $\hat{P}^2 = 1$, ce qui traduit mathématiquement que l'échange de deux particules est une *symétrie*. Une conséquence importante de cette propriété est que les valeurs propres de $\hat{P}$ sont ± 1 . Les états propres $|\psi_{\pm}\rangle$ correspondant sont symétrique ou antisymétrique par échange des particules et peuvent donc se mettre sous la forme

$$\begin{aligned} |\psi_+\rangle &= \lambda_+ (|1 : \psi_a; 2 : \psi_b\rangle + |1 : \psi_b; 2 : \psi_a\rangle) \\ |\psi_-\rangle &= \lambda_- (|1 : \psi_a; 2 : \psi_b\rangle - |1 : \psi_b; 2 : \psi_a\rangle), \end{aligned}$$

les constantes $\lambda_{\pm}$ s'obtenant en normalisant $|\psi_{\pm}\rangle$.

Considérons à présent le cas de deux particules identiques. Puisque l'on ne peut plus distinguer quelles particules occupent les états $|\psi_a\rangle$ et $|\psi_b\rangle$, on

note $|\psi_a, \psi_b\rangle$ l'état du système $1 \oplus 2$. Comme nous l'avons vu dans le cas de la collision de deux électrons, l'état $|\psi_a, \psi_b\rangle$ peut s'exprimer comme une combinaison linéaire des états $|1 : \psi_a; 2 : \psi_b\rangle$ et $|1 : \psi_b; 2 : \psi_a\rangle$, soit

$$|\psi_a, \psi_b\rangle = \lambda |1 : \psi_a; 2 : \psi_b\rangle + \mu |1 : \psi_b; 2 : \psi_a\rangle.$$

Examinons à présent l'action de l'opérateur d'échange sur l'état $|\psi_a, \psi_b\rangle$. Nous savons déjà que l'échange des particules 1 et 2 doit laisser invariant le système. L'état d'un système étant défini en mécanique quantique à une phase θ près, ceci signifie que

$$\hat{P}|\psi_a, \psi_b\rangle = e^{i\theta}|\psi_a, \psi_b\rangle,$$

autrement dit l'état $|\psi_a, \psi_b\rangle$ est vecteur propre de $\hat{P}$ pour la valeur propre $e^{i\theta}$. Or, nous avons vu que les seules valeurs propres admissibles pour $\hat{P}$ sont ± 1 , ce qui signifie que $|\psi_a, \psi_b\rangle$ est nécessairement symétrique ou antisymétrique sous l'échange des deux particules, c'est-à-dire

$$|\psi_a, \psi_b\rangle = \lambda_{\pm} (|1 : \psi_a; 2 : \psi_b\rangle \pm |1 : \psi_b; 2 : \psi_a\rangle).$$

Si on généralise le raisonnement à N particules, on est amené à énoncer le postulat de symétrisation :

Un ensemble de N particules identiques est décrit en mécanique quantique par un état complètement symétrique ou antisymétrique par échange de deux particules. Le caractère symétrique ou antisymétrique de l'état ne dépend que de la nature de la particule. On appelle bosons (respectivement fermions)¹ les particules décrites par les états symétriques (resp. antisymétriques).

Le postulat de symétrisation est complété par le théorème spin-statistique².

Le caractère bosonique ou fermionique d'une particule dépend uniquement de son spin : les particules de spin entier sont des bosons, les particules de spin demi-entier sont des fermions. Ainsi, les photons de spin 1 sont des bosons alors que les électrons de spin 1/2 sont des fermions.

1. Nommés ainsi respectivement d'après les physiciens S.N. Bose et E. Fermi.

2. La démonstration du théorème spin-statistique nécessite l'utilisation de la théorie de la relativité restreinte et sortirait du cadre de ce cours.

1.1.2 Applications

Dans la suite, nous nous intéresserons aux propriétés des photons, qui comme nous l'avons vu, sont des particules bosoniques. Nous n'aborderons donc pas les propriétés générales des particules fermionique. Signalons juste que le caractère fermionique de l'électron est responsable du Principe de Pauli, lui même à l'origine du remplissage du tableau périodique et de toute la chimie. En effet, considérons un état où deux fermions sont dans le même état $|\psi_a\rangle$. D'après le postulat de symétrisation, l'état du système est donné par

$$|\psi_a, \psi_a\rangle = \lambda_- (|1 : \psi_a; 2; \psi_a\rangle - |1 : \psi_a; 2; \psi_a\rangle) = 0,$$

qui n'est pas normalisable et donc non-physique.

Au contraire des fermions, les bosons ont une tendance “naturelle” à occuper le même état quantique, ce que l'on appelle le groupement, ou *bunching* en anglais. Une expérience très simple³ permettant d'observer ce phénomène consiste à envoyer deux photons (notés comme toujours 1 et 2) à angle droit sur les deux faces opposées d'une lame semi-réfléchissante (Fig. 1.2).

Commençons tout d'abord par préciser le comportement du système lorsqu'un seul photon est présent. Appelons $|a\rangle$ et $|b\rangle$ les états d'un photon arrivant “horizontalement” et “verticalement” sur la figure. Ces deux états sont orthogonaux - au sens quantique - et après passage dans la lame, un photon préparé dans l'état $|a\rangle$ se retrouve dans une superposition à poids égaux de $|c\rangle$ et $|d\rangle$, et de même s'il était dans l'état $|b\rangle$.

Supposons que le photon dans l'état $|a\rangle$ soit amené sur l'état $(|c\rangle + |d\rangle)/\sqrt{2}$. L'évolution d'un système quantique étant unitaire, deux systèmes identiques préparés dans des états orthogonaux doivent rester dans des états orthogonaux à tout instant. Après traversée de la lame, le photon préparé dans l'état $|b\rangle$ doit donc être envoyé sur un état orthogonal à $(|c\rangle + |d\rangle)/\sqrt{2}$, soit $(|c\rangle - |d\rangle)/\sqrt{2}$.

Considérons maintenant la situation où deux photons sont présents, l'un arrivant dans l'état $|a\rangle$, l'autre dans l'état $|b\rangle$. D'après le postulat de symétrisation, cet état noté $|ab\rangle$ se met sous la forme

$$|ab\rangle = \frac{|1 : a, 2 : b\rangle + |1 : b, 2 : a\rangle}{\sqrt{2}}.$$

3. Dans le principe, mais pas dans la réalisation !

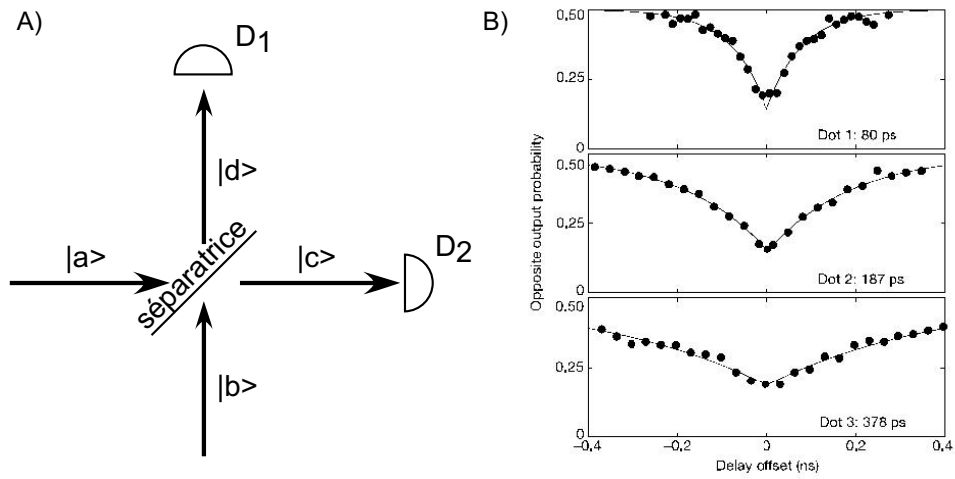


FIG. 1.2 – A) Principe de l'expérience de groupement de photons. B) Résultats expérimentaux tirés de C. Santori et al, *Nature* **419**, 594-597 (2002). On a représenté la probabilité d'observer un photon dans chaque détecteur en fonction de l'écart de temps entre les arrivées de chaque photon sur la lame. Lorsque cet écart est nul, la probabilité présente un minimum correspondant à des photons sortant tous les deux par la même voie. Ce minimum n'est cependant pas nul, du fait de la durée non-nulle du paquet d'onde des photons, introduisant une incertitude quantique irréductible sur le temps d'arrivée de chacun des photons. Les trois courbes correspondent aux trois "boîtes quantiques" différentes ayant servies de sources de photons uniques.

En supposant que chaque photon évolue indépendamment l'un de l'autre et en utilisant la linéarité de l'équation de Schrödinger, l'état du système après passage dans la lame séparatrice est décrit par

$$\frac{1}{2\sqrt{2}} ((|1 : c\rangle + |1 : d\rangle) \otimes (|2 : c\rangle - |2 : d\rangle) + (|1 : c\rangle - |1 : d\rangle) \otimes (|2 : c\rangle + |2 : d\rangle)),$$

soit, en développant

$$\frac{|1 : c, 2 : c\rangle - |1 : d, 2 : d\rangle}{\sqrt{2}}.$$

Cet état est particulièrement remarquable. En effet, on constate que si un photon est détecté dans un état de sortie donné, l'état $|c\rangle$ par exemple, le deuxième photon sera alors détecté dans ce même état $|c\rangle$ avec une probabilité 100 %. Autrement dit, dans cette expérience de traversée d'une lame semi-réfléchissante, les deux photons "s'attirent" vers la même voie de sortie.

1.1.3 Espace de Fock

Dans le cas où le nombre de particules devient important, les notations utilisées au paragraphe précédent deviennent assez peu commode. En effet, si l'on considère un système contenant N bosons occupant les états $|\psi_1\rangle, |\psi_2\rangle \dots |\psi_N\rangle$, l'état à N corps, noté comme précédemment $|\psi_1, \psi_2, \dots, \psi_N\rangle$ décrivant le système complet s'écrira en fait

$$|\psi_1, \psi_2, \dots, \psi_N\rangle = A \sum_{\sigma \in \mathcal{S}_N} |1 : \psi_{\sigma(1)}, 2 : \psi_{\sigma(2)} \dots N : \psi_{\sigma(N)}\rangle,$$

où $\mathcal{S}_N$ désigne l'ensemble des permutations de $\{1, 2, \dots, N\}$ et A est une constante de normalisation. La somme et la présence du groupe des perturbations rend le formalisme très lourd. Le nombre de particules étant, pour cet état, fixé à N , il ne peut donc pas décrire la création ou la disparition de photon, phénomène que l'on doit cependant prendre en compte dans la description de l'interaction matière rayonnement.

Pour toutes ces raisons, on est amené à développer un formalisme alternatif, mais mathématiquement équivalent que l'on appelle formalisme de seconde quantification (car il consiste à considérer les champs, tels que le champ de Schrödinger, comme des opérateurs).

Dans ce nouveau formalisme, on se place dans un espace de Hilbert (dit espace de Fock) où le nombre de particules est indéfini. Si on note $|\psi_\alpha\rangle$ une base d'états à *une particule*, une base de l'espace de Fock est constitué des états notés

$$|N_1, N_2, \dots, N_\alpha, \dots\rangle = |\{N_\alpha\}\rangle,$$

où N_α désigne le nombre de bosons occupant l'état $|\psi_\alpha\rangle$.

Ainsi, l'état à deux particules que nous notions anciennement $|\psi_\alpha, \psi_\beta\rangle$ devient à présent

$$|0, 0, \dots, 0, \underset{(N_\alpha)}{1}, 0, \dots, 0, \underset{(N_\beta)}{1}, 0, \dots\rangle$$

Dans le cas où les particules sont indépendantes et que les états $|\psi_\alpha\rangle$ sont des états propres du hamiltonien à une particule, pour l'énergie propre E_α , on se convainc sans difficulté que l'état $|\{N_\alpha\}\rangle$ possède l'énergie propre $E(\{N_\alpha\})$ donnée par

$$E(\{N_\alpha\}) = \sum_{\alpha} N_\alpha E_\alpha. \quad (1.1)$$

Dans la plupart des phénomènes physiques faisant intervenir la lumière, des photons peuvent être émis ou absorbés. Ceci nous amène à introduire les opérateurs $\hat{a}_\alpha$ d'annihilation d'un photon dans l'état à une particule $|\psi_\alpha\rangle$ que l'on définit par

$$\hat{a}_\alpha |N_1, N_2, \dots, N_\alpha, \dots\rangle = \sqrt{N_\alpha} |N_1, N_2, \dots, N_\alpha - 1, \dots\rangle.$$

et dont on montre ensuite que l'action de l'opérateur adjoint s'écrit

$$\hat{a}_\alpha^\dagger |N_1, N_2, \dots, N_\alpha, \dots\rangle = \sqrt{N_\alpha + 1} |N_1, N_2, \dots, N_\alpha + 1, \dots\rangle.$$

L'opérateur $\hat{a}_\alpha^\dagger$ correspond donc à la création d'un photon dans le mode α et porte le nom d'opérateur *création d'un photon dans le mode α* . D'après la définition de ces deux opérateurs, il s'ensuit la relation de commutation $[\hat{a}_\alpha, \hat{a}_\beta] = \delta_{\alpha, \beta}$.

Les préfacteurs $\sqrt{N_\alpha}$ et $\sqrt{N_\alpha + 1}$ ne sont pas purement conventionnels. D'une part, ils soulignent l'analogie existant entre les opérateurs création et annihilations de photons, avec les opérateurs création et annihilation d'une

excitation d'un oscillateur harmonique. Par ailleurs, l'opérateur $\hat{N}_\alpha = \hat{a}_\alpha^\dagger \hat{a}_\alpha$ possède une interprétation physique très simple puisque

$$\hat{a}_\alpha^\dagger \hat{a}_\alpha |N_1, N_2, \dots, N_\alpha, \dots\rangle = N_\alpha |N_1, N_2, \dots, N_\alpha, \dots\rangle.$$

L'opérateur $\hat{N}_\alpha$ est donc l'opérateur *nombre de photons dans le mode α* .

L'introduction des opérateurs $\hat{a}_\alpha$ et $\hat{a}_\alpha^\dagger$ nous permet d'exprimer simplement le hamiltonien décrivant le champ électromagnétique. En effet, en notant que, d'après l'équation (1.1) l'énergie de l'état $|\{N_\alpha\}\rangle$ vaut $\sum_\alpha \hbar\omega_\alpha N_\alpha$, on voit qu'en passant aux opérateurs, on obtient

$$\hat{H} = \sum_\alpha \hbar\omega_\alpha \hat{a}_\alpha^\dagger \hat{a}_\alpha. \quad (1.2)$$

1.2 Théorie quantique élémentaire du rayonnement et coefficients d'Einstein

Nous allons appliquer les notions introduites dans les deux sections précédentes à l'élaboration d'une théorie quantique du rayonnement, l'électrodynamique quantique. Au contraire du modèle utilisé précédemment pour décrire de façon perturbative l'interaction d'un atome avec une onde électromagnétique *classique*, la théorie que nous allons développer permettra une description complètement quantique de l'interaction matière-rayonnement. Ceci nous amènera naturellement à l'introduction de l'émission stimulée qui est au cœur du fonctionnement d'un laser.

1.2.1 Brefs rappels d'électromagnétisme

Rappelons brièvement qu'en électromagnétisme classique, l'évolution dans le vide des champs électriques et magnétiques $\mathbf{E}$ et $\mathbf{B}$ est donnée par les équations de Maxwell

$$\begin{aligned} \operatorname{div}(\mathbf{E}) &= 0 \\ \operatorname{div}(\mathbf{B}) &= 0 \\ \operatorname{rot}(\mathbf{E}) &= -\frac{\partial \mathbf{B}}{\partial t} \\ \operatorname{rot}(\mathbf{B}) &= \frac{1}{c^2} \frac{\partial \mathbf{E}}{\partial t}. \end{aligned}$$

La solution générale de ces équations peut se décomposer sous la forme d'ondes planes transverses pour lesquelles le champ électrique se met sous la forme

$$\begin{aligned}\mathbf{E}(\mathbf{r}, t) &= \sum_{\alpha} \mathcal{E}_{\alpha}(t) \mathbf{v}_{\alpha} e^{i\mathbf{k}_{\alpha} \cdot \mathbf{r}} + \mathcal{E}_{\alpha}^*(t) \mathbf{v}_{\alpha}^* e^{-i\mathbf{k}_{\alpha} \cdot \mathbf{r}} \\ \mathbf{B}(\mathbf{r}, t) &= \frac{1}{c} \sum_{\alpha} \mathcal{E}_{\alpha}(t) \boldsymbol{\kappa}_{\alpha} \times \mathbf{v}_{\alpha} e^{i\mathbf{k}_{\alpha} \cdot \mathbf{r}} + \mathcal{E}_{\alpha}^*(t) \boldsymbol{\kappa}_{\alpha} \times \mathbf{v}_{\alpha}^* e^{-i\mathbf{k}_{\alpha} \cdot \mathbf{r}}\end{aligned}$$

où l'indice $\alpha = (\mathbf{k}, \sigma)$ indique le vecteur d'onde $\mathbf{k}$ et la polarisation σ du mode. Les vecteurs $\mathbf{v}_{\alpha}$ sont des vecteurs unitaires, orthogonaux à $\mathbf{k}$ formant une base de l'espace (bidimensionnel) des polarisations⁴. Le vecteur $\boldsymbol{\kappa} = \mathbf{k}/k$ garantit que le champ magnétique reste orthogonal à $\mathbf{E}$ et $\mathbf{k}$. Les coefficients $\mathcal{E}_{\alpha}$ sont donnés par les équations de Maxwell et sont de la forme

$$\mathcal{E}_{\alpha}(t) = \mathcal{E}_{\alpha}^0 e^{-i\omega_{\alpha} t},$$

la pulsation ω_{α} étant donnée par la relation de dispersion linéaire $\omega_{\alpha} = ck_{\alpha}$. Dans l'espace libre, les vecteurs d'ondes $\mathbf{k}$ forment un continuum et les la somme sur les α devrait être en toute rigueur une intégrale. Pour garder une somme discrète, on est souvent amené à introduire un *volume de quantification*. On suppose dans ce cas que le champ électromagnétique est confiné dans une cavité de volume V très grand devant la taille caractéristique de l'expérience. Les bords de la cavité étant placés très loin du lieu de l'expérience, la forme exacte de la boîte ainsi que les conditions aux limites imposées par ses parois importent peu. Dans la suite, nous choisirons donc de travailler avec les conditions aux limites de Born-von Karman, où l'on suppose la cavité de forme cubique, de côté $L = V^{1/3}$ et où l'on impose des conditions aux limites périodiques, c'est à dire que l'on identifie les champs en $x_i = 0$ et $x_i = L$.

L'utilisation des conditions de Born - von Karman impose la quantification des vecteurs d'ondes. En effet, pour assurer la périodicité de chaque mode α , il faut que les coordonnées $k_{\alpha,i}$ des vecteurs d'ondes $\mathbf{k}_{\alpha}$ se mettent sous la forme

$$k_{\alpha,i} = n_{\alpha,i} \frac{2\pi}{L},$$

4. Les $\mathbf{v}_{\alpha}$ sont éventuellement complexes si l'on travaille dans une base de polarisations circulaires.

où $n_{\alpha,i} \in \mathbb{Z}$. Autrement dit, la somme sur α est en fait une somme *discrète* sur les entiers $n_{\alpha,i}$.

Pour retrouver un résultat physiquement acceptable et donc indépendant du volume de quantification, on sera très souvent amené à prendre la limite $V \rightarrow \infty$. Dans ce cas, la somme discrète sur les n_i se transforme en une intégrale sur $\mathbf{k}$. Chaque mode occupant un volume $(2\pi/L)^3$ dans l'espace des impulsions, on voit que l'on pourra faire la transformation

$$\sum_{\alpha} \longleftrightarrow \int \frac{d^3 \mathbf{k} V}{(2\pi)^3}.$$

Pour finir, on sait que la densité volumique d'énergie électromagnétique s'écrit

$$\varpi = \frac{\epsilon_0 E^2}{2} + \frac{B^2}{2\mu_0}.$$

Intégrons cette densité sur tout le volume de quantification, en décomposant les champs $\mathbf{E}$ et $\mathbf{B}$ selon les modes α . On constate qu'il va apparaître une somme de termes du type

$$\int_V e^{i(\mathbf{k}_{\alpha} - \mathbf{k}_{\alpha'}) \cdot \mathbf{r}} d^3 \mathbf{r}.$$

En utilisant la périodicité imposée par les conditions aux limites, cette intégrale vaut 0, à moins d'avoir $\mathbf{k}_{\alpha} = \mathbf{k}_{\alpha'}$ auquel cas elle vaut V . Les modes se découpent donc complètement et l'on obtient

$$U_{\text{em}} = \int_V \varpi d^3 r = 2\epsilon_0 V \sum_{\alpha} |\mathcal{E}_{\alpha}|^2. \quad (1.3)$$

1.2.2 Champ électromagnétique quantique

Le champ électromagnétique étant constitué de bosons, les photons, nous avons vu que l'espace des états caractérisant le champ électromagnétique quantique aura pour base les états $|\{N_{\mathbf{k},\sigma}\}\rangle$, où $N_{\mathbf{k},\sigma}$ désigne le nombre de photons dans le mode de vecteur d'onde $\mathbf{k}$ et de polarisation σ .

Classiquement, le champ électromagnétique n'est pas décrit en terme de photon, mais plutôt du champ électrique $\mathbf{E}$. Pour faire le lien entre physique classique et physique quantique, il nous reste à définir *l'observable* champ électrique $\hat{\mathbf{E}}$.

Quantiquement, on va remplacer $\mathcal{E}_\alpha$ par un opérateur, ce qui va nous permettre d'écrire l'observable champ électrique sous la forme

$$\hat{\mathbf{E}}(\mathbf{r}) = \sum_{\alpha} \hat{\mathcal{E}}_{\alpha} \mathbf{v}_{\alpha} e^{i\mathbf{k}_{\alpha} \cdot \mathbf{r}} + \hat{\mathcal{E}}_{\alpha}^{\dagger} \mathbf{v}_{\alpha}^* e^{-i\mathbf{k}_{\alpha} \cdot \mathbf{r}}.$$

Et de même, le hamiltonien (c'est-à-dire l'énergie) du champ électromagnétique s'écrit d'après l'équation (1.3)

$$\hat{H} = \sum_{\alpha} 2\epsilon_0 V \hat{\mathcal{E}}_{\alpha}^{\dagger} \hat{\mathcal{E}}_{\alpha}.$$

Comparons cette expression à la relation (1.2) que nous avons obtenu pour un système de bosons sans interaction. On constate que les deux descriptions coïncident si l'on choisit pour l'opérateur champ $\hat{\mathcal{E}}_{\alpha}$

$$\hat{\mathcal{E}}_{\alpha} = \sqrt{\frac{\hbar\omega_{\alpha}}{2\epsilon_0 V}} \hat{a}_{\alpha}$$

Cette dernière relation nous permet donc d'écrire l'opérateur champ électrique à l'aide des opérateurs création et annihilation de photons

$$\hat{\mathbf{E}}(\mathbf{r}) = \sum_{\alpha} \sqrt{\frac{\hbar\omega_{\alpha}}{2\epsilon_0 V}} (\hat{a}_{\alpha} \mathbf{v}_{\alpha} e^{i\mathbf{k}_{\alpha} \cdot \mathbf{r}} + \hat{a}_{\alpha}^{\dagger} \mathbf{v}_{\alpha}^* e^{-i\mathbf{k}_{\alpha} \cdot \mathbf{r}}).$$

1.2.3 Électrodynamique quantique en cavité.

L'expérience la plus simple permettant de vérifier la validité du formalisme développé ci-dessus consiste à placer un atome unique dans une cavité Fabry-Pérot résonnante avec une transition atomique. La présence de la cavité permet de n'avoir à considérer qu'un volume V très petit et ainsi limiter le nombre de modes du champ électromagnétique pertinents dans l'expérience. Nous allons voir que ceci permet notamment de se ramener à la situation très simple d'un système à deux niveaux.

Dans l'approximation dipolaire, L'interaction d'un atome placé en $\mathbf{R}_0$ avec le rayonnement se traduit par le hamiltonien

$$\hat{V}_{\text{dip}} = -\hat{\mathbf{D}} \cdot \hat{\mathbf{E}}(\mathbf{R}_0) = -\hat{\mathbf{D}} \cdot \hat{\mathbf{E}}(0),$$

si l'on suppose l'atome localisé en $\mathbf{R}_0 = 0$. Dans cette approximation, le hamiltonien complet décrivant le système atome + champ s'écrit⁵

$$\hat{H} = \hbar\omega_a (|e\rangle\langle e| + \hat{a}^\dagger\hat{a}) - \sqrt{\frac{\hbar\omega_\alpha}{2\epsilon_0 V}} \mathbf{d} \cdot \mathbf{v} (|e\rangle\langle f| + |f\rangle\langle e|) (\hat{a} + \hat{a}^\dagger),$$

où nous avons posé l'énergie de l'état fondamental atomique $|f\rangle$ égale à zéro et $\hbar\omega_a$ est l'énergie de l'état excité atomique. On suppose que le mode électromagnétique de la cavité résonnant avec cette transition. et on suppose pour simplifier que le mode de la cavité est polarisé linéairement, on peut supposer $\mathbf{v}$ réel. Enfin, nous avons noté $\mathbf{d} = \langle f|\hat{\mathbf{D}}|e\rangle = \langle e|\hat{\mathbf{D}}|f\rangle$ et nous avons comme précédemment supposé $\langle e|\mathbf{D}|e\rangle = \langle f|\mathbf{D}|f\rangle$

Dans le terme du hamiltonien correspondant à l'interaction dipolaire entre l'atome et le champ, les termes du types $|e\rangle\langle f|\hat{a}^\dagger$ ou $|f\rangle\langle e|\hat{a}$ sont non résonnants, puisqu'ils correspondent à la création (resp. l'absorption) d'un photon par un atome dans l'état fondamental (resp. excité). On peut donc les négliger et le hamiltonien se simplifie donc en

$$\hat{H} \sim \hbar\omega_a (|e\rangle\langle e| + \hat{a}^\dagger\hat{a}) - \sqrt{\frac{\hbar\omega_\alpha}{2\epsilon_0 V}} \mathbf{d} \cdot \mathbf{v} (|e\rangle\langle f|\hat{a} + |f\rangle\langle e|\hat{a}^\dagger).$$

Préparons l'atome dans l'état fondamental avec N photons dans la cavité. On voit que cet état initial, noté $|f, N\rangle$ est couplé par le hamiltonien à l'état $|e, (N-1)\rangle$. L'atome étant initialement dans l'état fondamental, il ne peut passer dans l'état excité qu'en absorbant un photon. Dans le sous-espace engendré par les deux états $|f, N\rangle$ et $|e, (N-1)\rangle$, le hamiltonien a par conséquent pour éléments de matrice

$$\hat{H} = \begin{pmatrix} N\hbar\omega_a & -\sqrt{N\frac{\hbar\omega_\alpha}{2\epsilon_0 V}} \mathbf{d} \cdot \mathbf{v} \\ -\sqrt{N\frac{\hbar\omega_\alpha}{2\epsilon_0 V}} \mathbf{d} \cdot \mathbf{v} & N\hbar\omega_a \end{pmatrix},$$

les coefficient N et $\sqrt{N}$ provenant de l'action des opérateurs création et annihilation de photons dans le mode de la cavité. On est donc ramené à l'étude d'un système quantique à deux niveaux. On montre sans difficulté que les états propres $|\psi_\pm\rangle$ et les valeurs propres $E_\pm$ de $\hat{H}$ s'écrivent

5. Un seul mode du champ électromagnétique étant résonnant, nous omettrons l'indice α dans ce paragraphe.

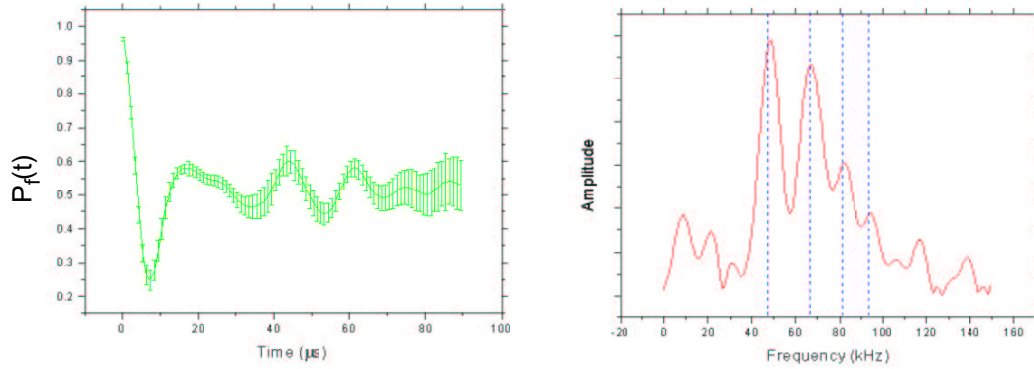


FIG. 1.3 – *Oscillations de Rabi en électrodynamique en cavité. À gauche, probabilité de trouver l'atome dans l'état fondamental à l'instant t . À droite, transformée de Fourier du signal précédent. Les pics à 47 kHz, 67 kHz $\sim \sqrt{2} \times 47$ kHz, 82 kHz $\sim \sqrt{3} \times 47$ kHz et 94 kHz $\sim \sqrt{4} \times 47$ kHz sont bien dans les rapports prédits par la théorie quantique.*

$$|\psi_{\pm}\rangle = \frac{1}{\sqrt{2}} (|f, N\rangle \pm |e, N-1\rangle)$$

$$E_{\pm} = \hbar\omega_a \pm \hbar\Omega\sqrt{N}$$

où la fréquence de Rabi Ω est définie par $\hbar\Omega = -\sqrt{\frac{\hbar\omega_a}{2\epsilon_0 V}} \mathbf{d} \cdot \mathbf{v}$.

Partant de l'état initial

$$|\psi(t=0)\rangle = |f, N\rangle = \frac{1}{\sqrt{2}} (|\psi_+\rangle + |\psi_-\rangle),$$

on sait que la solution de l'équation de Shrödinger s'écrit simplement

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}} (e^{-iE_+t/\hbar} |\psi_+\rangle + e^{-iE_-t/\hbar} |\psi_-\rangle),$$

soit, en repassant dans la base $\{|f, N\rangle, |e, N-1\rangle\}$

$$|\psi(t)\rangle = e^{-i\omega_a t} \left(\cos(\sqrt{N}\Omega t) |f, N\rangle - i \sin(\sqrt{N}\Omega t) |e, N-1\rangle \right).$$

La probabilité d'observer l'atome dans l'état fondamental est donc :

$$P_f(t) = |\langle f, N | \psi(t) \rangle|^2 = \cos^2(\sqrt{N}\Omega t). \quad (1.4)$$

Autrement dit, l'atome effectue des oscillations entre l'état fondamental et l'état excité à une pulsation $2\sqrt{N}\Omega$.

La figure (1.3) représente les oscillations de Rabi d'un atome placé dans une cavité. Dans cette expérience réalisée au Laboratoire Kastler Brossel de l'ENS⁶, le système est préparé dans un état

$$|\psi(t=0)\rangle = \sum_N c_N |f, N\rangle,$$

où le nombre de photons dans la cavité est indéterminé et $|c_N|^2$ désigne la probabilité d'observer N photons. En utilisant la linéarité de l'équation de Schrödinger ainsi que les résultats que nous venons d'obtenir, la probabilité de trouver l'atome dans l'état excité s'écrit comme une somme de termes du type (1.4) oscillant aux fréquences 2Ω , $2\Omega\sqrt{2}$, $2\Omega\sqrt{3}$..., fréquences que l'on retrouve sur la transformée de Fourier de $P_f(t)$.

1.2.4 Coefficients d'Einstein.

Pour académique qu'il soit, le problème de l'électrodynamique en cavité illustre deux propriétés fondamentales de l'interaction lumière-matière :

1. L'absorption et l'émission d'un photon par un atome sont des processus symétriques, comme le prouvent les oscillations de Rabi discutées au paragraphe précédent.
2. Le couplage de l'atome à un mode du champ électromagnétique est d'autant plus grand que ce mode est fortement peuplé. Cette propriété, qui se manifeste par le coefficient $\sqrt{N}$ dans les éléments de matrices non-diagonaux, est une conséquence de la présence des opérateurs $\hat{a}$ et $\hat{a}^\dagger$ dans le hamiltonien.

Nous allons retrouver ces propriétés dans le cas plus général qui suit où nous allons considérer que plusieurs modes du champ électromagnétique sont peuplés suivant une distribution continue $\{N_{k,\sigma}\}$, situation que l'on rencontre naturellement si l'on fait tendre vers l'infini le volume de la cavité étudiée à la question précédente.

6. <http://www.lkb.ens.fr/recherche/qedcav/french/frenchframes.html>

Absorption

Supposons dans un premier temps l'atome initialement dans l'état fondamental $|f\rangle$. Le hamiltonien dipolaire couple cet état à l'état $|e\rangle$ en absorbant un photon dans un mode $\alpha = (\mathbf{k}, \sigma)$. Ce mode α étant *a priori* quelconque, l'état initial est couplé à un continuum d'état quelconque caractérisés par le vecteur d'onde $\mathbf{k}$ du photon absorbé. Le taux de transfert $\Gamma_{e \rightarrow f}$ de l'état fondamental à l'état initial est donc donné par la Règle d'or de Fermi (1.19) et vaut

$$\Gamma_{f \rightarrow e} = \frac{2\pi}{\hbar} \sum_{\alpha=\mathbf{k},\sigma} \delta(\hbar\omega_a - \hbar\omega_\alpha) \frac{\hbar\omega_\alpha}{2\epsilon_0 V} |\mathbf{d} \cdot \mathbf{v}_\alpha|^2 N_\alpha. \quad (1.5)$$

Supposons dans un premier temps le rayonnement polarisé dans la direction z . Le vecteur $\mathbf{v}_\alpha$ est alors constamment égal à $\mathbf{u}_z$ et le taux d'absorption s'écrit

$$\Gamma_{f \rightarrow e} = \frac{2\pi d_z^2}{\hbar^2} \sum_{\alpha=\mathbf{k},\sigma} \delta(\omega_a - \omega_\alpha) \frac{\hbar\omega_\alpha}{2\epsilon_0 V} N_\alpha. \quad (1.6)$$

La somme intervenant dans l'équation précédente est en fait la densité volumique spectrale d'énergie électromagnétique à la fréquence ω_a . Par définition, l'énergie électromagnétique U contenue dans un volume V et associée aux modes de pulsation comprise entre ω_1 et ω_2 s'écrit

$$U = V \int_{\omega_1}^{\omega_2} u(\omega) d\omega, \quad (1.7)$$

où u est par définition la densité spectrale d'énergie. En terme de photons, cette énergie U est aussi la somme de l'énergie $N_\alpha \hbar\omega_\alpha$ sur tous les modes de pulsation $\omega_\alpha \in [\omega_1, \omega_2]$, soit

$$U = \sum_{\omega_\alpha \in [\omega_1, \omega_2]} \hbar\omega_\alpha N_\alpha.$$

Mais cette somme peut se récrire comme

$$U = \int_{\omega_1}^{\omega_2} \sum_{\alpha} \hbar\omega_\alpha N_\alpha \delta(\omega - \omega_\alpha) d\omega. \quad (1.8)$$

En effet, l'intégrale de la fonction δ donne 1 si ω_α est contenu dans l'intervalle d'intégration, et 0 sinon.

En identifiant les expressions (1.7) et (1.8), on voit que

$$u(\omega) = \frac{1}{V} \sum_{\alpha} \hbar \omega_{\alpha} N_{\alpha} \delta(\omega - \omega_{\alpha}),$$

ce qui permet d'écrire le taux d'absorption comme

$$\Gamma_{e \rightarrow f} = B_{fe} u(\omega_a),$$

où $B_{fe} = \pi d_z^2 / \epsilon_0 V$ est appelé coefficient d'Einstein pour l'absorption.

Émission stimulée et émission spontanée

On prépare dans un second temps l'atome dans l'état excité. L'application de la règle d'or de Fermi nous donne alors pour le taux de désexcitation radiative

$$\Gamma_{e \rightarrow f} = \frac{2\pi}{\hbar} \sum_{\alpha=\mathbf{k},\sigma} \delta(\hbar\omega_a - \hbar\omega_{\alpha}) \frac{\hbar\omega_{\alpha}}{2\epsilon_0 V} |\mathbf{d} \cdot \mathbf{v}_{\alpha}|^2 (N_{\alpha} + 1). \quad (1.9)$$

La ressemblance entre les équations (1.5) et (1.9) saute immédiatement aux yeux et nous permet d'écrire

$$\Gamma_{e \rightarrow f} = A_{ef} + B_{fe} u(\omega_a), \quad (1.10)$$

où le coefficient d'Einstein A_{ef} d'émission spontanée est défini par

$$A_{ef} = \frac{2\pi}{\hbar} \sum_{\alpha=\mathbf{k},\sigma} \delta(\hbar\omega_a - \hbar\omega_{\alpha}) \frac{\hbar\omega_{\alpha}}{2\epsilon_0 V} |\mathbf{d} \cdot \mathbf{v}_{\alpha}|^2.$$

En passant au continu, cette intégrale se récrit comme

$$A_{ef} = \frac{2\pi}{\hbar} \int \frac{d^3k}{(2\pi)^3} \sum_{\sigma} \delta(\hbar\omega_a - \hbar\omega_{\alpha}) \frac{\hbar\omega_{\alpha}}{2\epsilon_0 V} |\mathbf{d} \cdot \mathbf{v}_{\alpha}|^2.$$

Dans l'équation (1.10), on constate que le taux d'émission est d'autant plus grand que l'énergie électromagnétique est importante dans les modes résonnants avec la transition $e \rightarrow f$. Autrement dit, la présence de photons dans un mode accélère la désexcitation dans ce mode. Ce phénomène, prédit pour la première fois par Einstein, porte le nom d'*émission stimulée*. Notons aussi qu'à la limite où $u \rightarrow \infty$, on a même $\Gamma_{e \rightarrow f} \sim \Gamma_{f \rightarrow e}$: l'émission stimulée

est donc une traduction de la symétrie entre émission et absorption observée dans le cas de l'électrodynamique quantique en cavité.

Le coefficient A_{ef} correspond quant à lui à l'émission "spontanée" d'un photon par l'atome. On peut le calculer en intégrant en coordonnées polaires (k, θ, ϕ) , l'axe z étant pris selon la direction de $\mathbf{d}$. Les vecteurs de polarisations $\mathbf{v}_\alpha$ devant rester orthogonaux à $\mathbf{k}$ ne dépendent que des variables angulaires θ et ϕ . Ceci permet de séparer les intégrales et on trouve

$$A_{\text{ef}} = \frac{1}{8\pi^3 \hbar \epsilon_0} \left(\int \sin(\theta) d\theta d\phi \sum_{\sigma} |\mathbf{d} \cdot \mathbf{v}_\sigma|^2 \right) \int k^2 dk \omega_k \delta(\omega_a - \omega_k).$$

En intégrant sur la variable $\omega_k = ck$ plutôt que k , on trouve

$$\begin{aligned} A_{\text{ef}} &= \frac{1}{8\pi^3 \hbar \epsilon_0 c^3} \left(\int \sin(\theta) d\theta d\phi \sum_{\sigma} |\mathbf{d} \cdot \mathbf{v}_\sigma|^2 \right) \int \omega_k^2 d\omega_k \delta(\omega_a - \omega_k) \\ &= \frac{\omega_a^3}{8\pi^3 \hbar \epsilon_0 c^3} \left(\int \sin(\theta) d\theta d\phi \sum_{\sigma} |\mathbf{d} \cdot \mathbf{v}_\sigma|^2 \right). \end{aligned}$$

L'intégrale angulaire se calcule ensuite en prenant pour base de vecteurs polarisation les vecteurs $\mathbf{u}_\theta$ et $\mathbf{u}_\phi$. Puisque $\mathbf{d}$ est parallèle à l'axe z , $\mathbf{d} \cdot \mathbf{u}_\phi = 0$ et seule contribue la polarisation selon $\mathbf{u}_\theta$ qui donne une contribution $\mathbf{d} \cdot \mathbf{u}_\theta = -d \sin(\theta)$.

On obtient donc

$$A_{\text{ef}} = \frac{\omega_a^3 d^2}{3\pi \hbar \epsilon_0 c^3}. \quad (1.11)$$

D'après l'équation (1.3), on constate que A_{ef} est d'autant plus petite que ω_a est faible. Autrement dit, l'émission spontanée est observable sur les transitions optiques, de grande fréquence ($\omega_a \sim 2\pi \times 10^{15}$ Hz), pour lesquelles on trouve $A \sim 10^9 \text{ s}^{-1}$ et pratiquement négligeable pour les transitions micro-ondes ($\omega_a \sim 2\pi \times 1$ GHz).

Aparté technique : absorption par un rayonnement non polarisé

Historiquement, Einstein s'était intéressé au cas d'un rayonnement à l'équilibre thermique avec la matière. Dans ce cas, les photons ne sont pas

polarisés et la distribution de quantité de mouvement est isotrope, ce qui revient à supposer que $N_{\mathbf{k},\sigma}$ ne dépend que de k . Si l'on transforme la somme en une intégrale, l'expression (1.5) du taux d'absorption peut s'écrire

$$\Gamma_{f \rightarrow e} = \frac{2\pi}{\hbar} \int \frac{k^2 dk d^2\Omega}{(2\pi)^3} \sum_{\sigma} \delta(\hbar\omega_a - \hbar\omega_k) \frac{\hbar\omega_k}{2\epsilon_0} |\mathbf{d} \cdot \mathbf{v}_{\mathbf{k},\sigma}|^2 N_k.$$

En notant que la seule dépendance angulaire provient du produit scalaire $\mathbf{d} \cdot \mathbf{v}_{\mathbf{k},\sigma}$, on peut séparer intégrale radiale et angulaire, ce qui nous donne

$$\Gamma_{f \rightarrow e} = \frac{\pi}{\hbar^2 \epsilon_0} \left(\int d^2\Omega \sum_{\sigma} |\mathbf{d} \cdot \mathbf{v}_{\mathbf{k},\sigma}|^2 \right) \int \frac{k^2 dk}{(2\pi)^3} \delta(\omega_a - \omega_k) \hbar\omega_k N_k$$

L'intégrale angulaire se calcule comme précédemment et l'on peut alors écrire en réintégrant $\int d^2\Omega = 4\pi$, le taux d'absorption peut se réécrire comme

$$\Gamma_{f \rightarrow e} = \frac{2\pi d^2}{3\hbar^2 \epsilon_0} \int \frac{d^3k}{(2\pi)^3} \delta(\omega_a - \omega_k) \hbar\omega_k N_k.$$

On en déduit que le taux d'absorption de photons par l'atome peut se mettre sous une forme similaire à celle obtenue pour le rayonnement polarisé, soit

$$\Gamma_{f \rightarrow e} = B'_{fe} u(\omega_a),$$

où le coefficient B'_{fe} donné par

$$B'_{fe} = \frac{\pi d^2}{3\hbar^2 \epsilon_0}.$$

On constate en particulier la proportionnalité des coefficients A et B' ,

$$A_{ef} = \frac{\hbar\omega_a^3}{\pi^2 c^3} B'_{fe}, \quad (1.12)$$

relation démontrée par Einstein dans son article de 1917.

Application à la loi du Corps Noir

Le problème du corps noir, c'est-à-dire l'étude du rayonnement émis par un corps chaud est intimement lié à la naissance de la mécanique quantique.

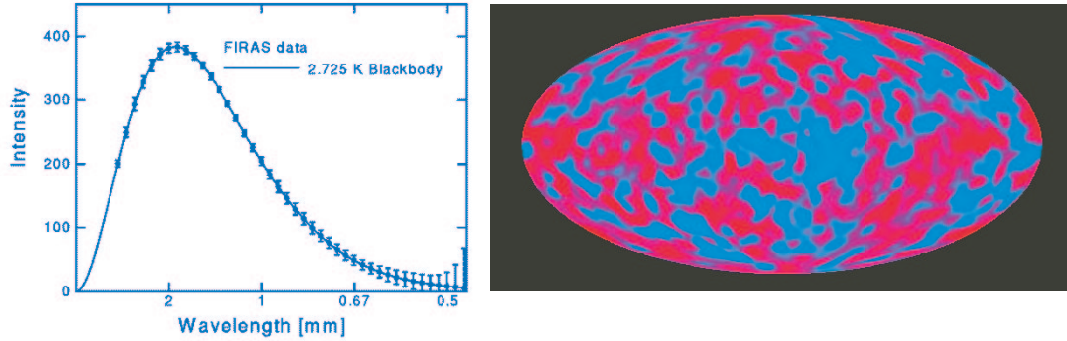


FIG. 1.4 – Spectre du corps noir. *Gauche* : spectre du rayonnement cosmologique. La courbe en trait plein correspond à une loi de corps noir à 2.725 K. *Droite* : fluctuations du spectre du corps noir traduisant les inhomogénéités de l’Univers primordial, responsables de la formation des étoiles et des galaxies.

En effet, c’est par suite de l’incapacité de la physique classique à décrire ce rayonnement que M. Planck fut amené à introduire en 1900 la constante qui porte son nom.

Considérons la situation physique d’une enceinte fermée (un “four”) dont les parois sont supposées parfaitement absorbantes et sont maintenues à une température T . La matière constituant les parois va rayonner des photons sous l’effet de l’agitation thermique qui fait passer certains atomes dans des états excités. Les photons ainsi émis dans l’enceinte du four sont réabsorbés lorsqu’ils atteignent la paroi opposée. Cette succession d’émission et d’absorption de photons par les parois du four permet de placer les parois du four et le rayonnement contenu dans celui-ci à l’équilibre thermique.

Supposons que les parois du four sont constitués des atomes à deux niveaux étudiés au paragraphe précédent et notons N_e et N_f le nombre d’entre eux dans les états $|e\rangle$ et $|f\rangle$ respectivement. On voit sans difficulté que la variation du nombre d’atomes dans l’état fondamental est donné par l’équation de taux

$$\frac{dN_f}{dt} = -\Gamma_{f \rightarrow e} N_f + \Gamma_{e \rightarrow f} N_e.$$

À l’équilibre thermique, le nombre d’atomes dans chaque niveau est constant (en moyenne). On a donc $\dot{N}_f = 0$, ce qui impose la condition

$$\frac{N_f}{N_e} = \frac{\Gamma_{e \rightarrow f}}{\Gamma_{f \rightarrow e}}.$$

Par ailleurs, on sait d'après la relation de Boltzmann que, pour une enceinte à la température T , $N_f/N_e = \exp(\hbar\omega_a/k_B T)$. En écrivant les taux Γ en fonction des coefficients d'Einstein, on voit que la densité spectrale d'énergie u doit satisfaire l'équation

$$\frac{A + B'u}{B'u} = e^{\beta\hbar\omega_a},$$

avec $\beta = 1/k_B T$. On obtient alors

$$u = \frac{\hbar}{\pi^2 c^3} \frac{\omega_a^3}{e^{\beta\hbar\omega_a} - 1}.$$

C'est la *loi de Planck pour le spectre du corps noir*. Comme nous l'avons signalé plus haut, cette relation fait explicitement intervenir la constante de Planck h ce qui explique l'incapacité de la physique classique à expliquer ce spectre.

La manifestation la plus spectaculaire de la distribution du corps noir est le spectre du rayonnement cosmologique à 2.7 K découvert en 1965 par Penzias et Wilson. Ce rayonnement a été émis aux premiers âges de l'univers quand le rayonnement n'était plus assez énergétique pour ioniser la matière (essentiellement les atomes d'hydrogène) : c'est donc un outil très précieux qui nous fournit d'importantes informations sur la structure de l'Univers primordial. Son spectre, présenté sur la figure 1.4, a été récemment mesuré par les satellites COBE (COsmic Background Explorer) et WMAP (Wilkinson Microwave Anisotropic Probe) et s'ajuste parfaitement avec une loi de corps noir à 2.7 K.

Condensation de Bose-Einstein et émission stimulée d'ondes de matière

L'émission stimulée est une propriété commune à toutes les particules bosoniques qui traduit leur tendance à s'accumuler dans un même état quantique. Dans le cas d'atomes et, non plus de photons, une conséquence spectaculaire de cet effet est l'existence de la condensation de Bose-Einstein qui se manifeste par l'accumulation brusque des atomes dans l'état fondamental du système. Cette transition, prédite par Einstein en 1924, est une

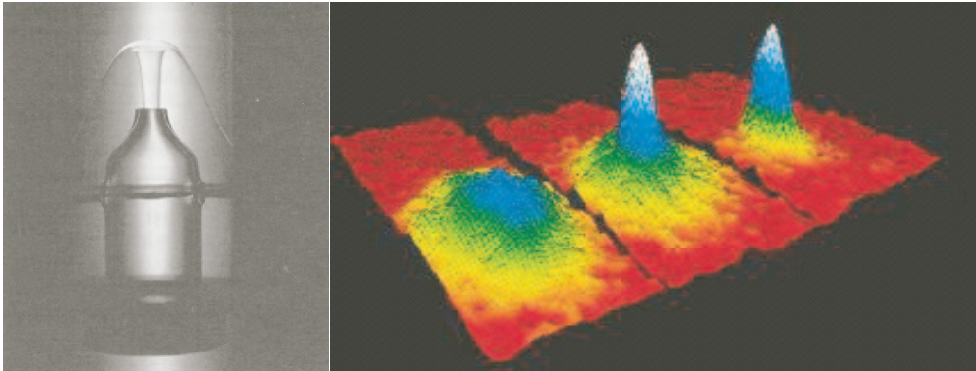


FIG. 1.5 – Condensation de Bose-Einstein. *Gauche : Effet fontaine dans l'hélium superfluide. En raison de l'absence de viscosité, l'hélium superfluide chauffé ne bout pas mais s'échappe du récipient sous forme d'un jet. Droite : Profil de densité d'un condensat de Bose-Einstein gazeux. De gauche à droite on réduit la température. Lorsque l'on passe le seuil de la transition, il apparaît un pic de densité correspondant aux atomes peuplant l'état fondamental.*

conséquence de la “surpopulation” due à la loi de Boltzmann des états de basse énergie à l'équilibre thermique. Lorsque deux atomes entrent en collision, l'effet “d'émission stimulée” peut se manifester et les atomes vont avoir tendance à être diffusés vers un état majoritairement peuplé, c'est-à-dire un état de basse énergie. Lorsque la température est suffisamment basse, ce processus s'emballe et les atomes vont pratiquement tous dans l'état de plus basse énergie.

La condensation de Bose-Einstein a été observée expérimentalement dans l'hélium superfluide où elle se manifeste par une disparition brusque de la viscosité en dessous de 2.17 K. Plus récemment (1995), elle a été obtenue dans les gaz d'atomes ultra-froids (~ 100 nK) en utilisant notamment les techniques de refroidissement que nous discuterons à la fin du cours.

Par opposition à l'hélium, les gaz d'atomes ultra-froids sont des systèmes extrêmement dilués où les interactions sont très faibles et permettent par conséquent de réaliser des expériences modèles permettant de tester finement les descriptions théoriques. Ils ont permis en particulier l'observation au MIT par le groupe de W. Ketterle de l'émission stimulée d'atomes issus d'un

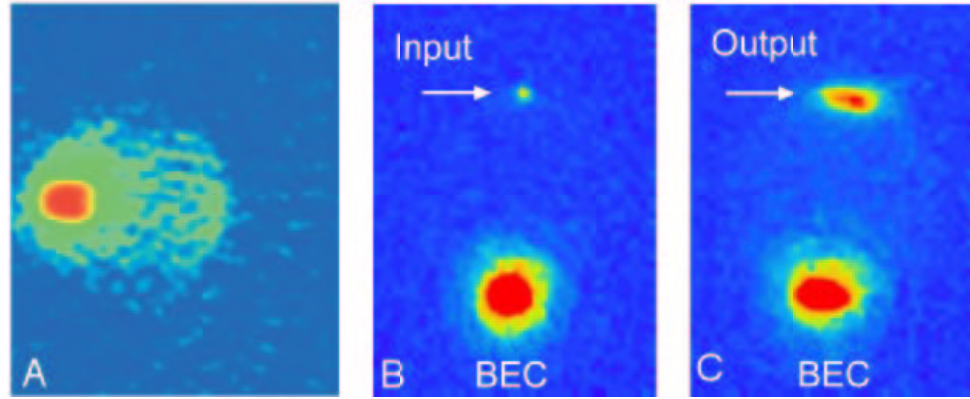


FIG. 1.6 – A : *Diffusion Rayleigh par un nuage d'atomes froids. Durant un processus d'absorption et d'émission d'un photon, les atomes reçoivent un recul aléatoire qui les disperse sur une sphère.* B, C : *Émission stimulée d'onde de matière. On place une petite fraction des atomes d'un condensat de Bose-Einstein immobile dans un état d'impulsion $\mathbf{P}$ (B, où l'on a attendu quelques millisecondes de façon à ce que les atomes d'impulsion $\mathbf{P}$ s'éloignent du condensat immobile). Après avoir généré le précurseur, on éclaire le condensat de façon à amorcer la diffusion Rayleigh (C). On constate que les atomes sont diffusés dans la même direction que le précurseur sous l'effet de l'émission stimulée bosonique.*

condensat de Bose-Einstein de Sodium⁷ (Fig. 1.5). Dans cette expérience, on éclaire un nuage d'atomes avec une lumière quasi-résonnante avec une transition atomique. Lorsqu'un photon est absorbé, un atome passe dans l'état excité, qu'il quitte rapidement par émission spontanée. La direction de cette émission est a priori aléatoire et le recul reçu par l'atome est donc lui aussi aléatoire. Ce processus dit de *diffusion Rayleigh* tend donc à disperser les atomes dans le halo que l'on observe sur la figure 1.6.A.

Imaginons à présent que juste avant d'allumer le faisceau Rayleigh on ait préparé une petite fraction des atomes dans un état d'impulsion $\mathbf{P}$ donnée (le nuage input de la figure 1.6.B). Lors d'une nouvelle diffusion Rayleigh, les atomes vont préférentiellement se désexciter vers l'état d'impulsion $\mathbf{P}$ déjà peuplé (nuage output de la figure 1.6.C).

⁷ S. Inouye, T. Pfau, S. Gupta, A. Chikkatur, A. Görlitz, D. Pritchard et W. Ketterle, Nature, **402**, 641 (1999).

1.3 Appendice technique : Théorie des perturbations et règle d'or de Fermi

1.3.1 Position du problème

Considérons un système quantique quelconque dont l'état est décrit en notation de Dirac par un vecteur (ou ket) $|\psi\rangle$ d'un espace de Hilbert $\mathcal{E}$. On sait que l'évolution du système est déterminée par l'opérateur⁸ hamiltonien $\hat{H}$ tel que

$$i\hbar \frac{d|\psi\rangle}{dt} = \hat{H}|\psi\rangle. \quad (1.13)$$

Le hamiltonien est un opérateur hermitien, c'est-à-dire qu'il est diagonalisable dans une base orthonormée avec des valeurs propres réelles⁹.

Supposons dans un premier temps $\hat{H}$ indépendant du temps et notons $|n\rangle$ et E_n ses vecteurs propres et ses valeurs propres. Les vecteurs $|n\rangle$ formant une base de l'espace des états, on peut écrire le vecteur d'onde $|\psi\rangle$ sous la forme

$$|\psi(t)\rangle = \sum_n c_n(t) |n\rangle, \quad (1.14)$$

où chaque c_n est par définition la projection de $|\psi\rangle$ sur un état $|n\rangle$. En utilisant la linéarité du hamiltonien, l'équation de Schrödinger s'écrit donc

$$i\hbar \sum_n \dot{c}_n |n\rangle = \sum_n c_n \hat{H} |n\rangle,$$

soit, puisque par définition des vecteurs propres et des valeurs propres, $\hat{H}|n\rangle = E_n |n\rangle$,

$$i\hbar \sum_n \dot{c}_n |n\rangle = \sum_n E_n c_n |n\rangle.$$

Si l'on projette cette relation sur un état $|n\rangle$ particulier, on voit sans difficulté que

8. Dans tout le cours, les observables quantiques sont notées avec un chapeau, pour les différencier des grandeurs classiques.

9. Cette condition d'hermiticité garantit que la norme de $|\psi\rangle$ reste constante durant l'évolution, ce qui garantit que la somme des probabilités de trouver l'atome dans chacun des états accessibles reste égale à 1.

$$i\hbar\dot{c}_n = E_n c_n,$$

et donc

$$c_n(t) = c_n(0)e^{-iE_n t/\hbar}. \quad (1.15)$$

Cette relation exprime que l'évolution des coefficients c_n se traduit par un simple déphasage dans le plan complexe, à une vitesse angulaire ω_n donnée par la relation de Planck-Einstein $E_n = \hbar\omega_n$. Connaissant l'état du système (i.e. les coefficients $c_n(0)$) à l'instant initial, l'équation (1.15) nous permet *a priori* de déterminer son état à tout instant ultérieur en calculant la somme (1.14). Cette méthode de résolution de l'équation de Shrödinger est certes exacte, mais se heurte à deux difficultés :

1. Même en dimension finie (supérieure à 5 tout du moins), la diagonalisation de $\hat{H}$ est en général impossible de façon exacte, ce qui empêche de connaître les valeurs propres et les vecteurs propres du hamiltonien.
2. La méthode décrite ci-dessus ne s'applique qu'à des hamiltoniens *indépendants du temps* et ne pourra donc pas être utilisée dans le cas où $\hat{H}$ dépend du temps. Ces situations sont malheureusement courantes lorsque l'on tente de décrire une expérience de spectroscopie où l'on doit décrire l'évolution d'un atome ou d'une molécule sous l'effet d'un champ électrique (par exemple en absorption optique - UV, visible ou infra-rouge) ou magnétique (c'est le cas de la RMN) variable.

1.3.2 Théorie des perturbations

La résolution de l'équation de Schrödinger étant impossible dans le cas général, on a été amené à développer des techniques plus ou moins raffinées de résolution approchée. Nous allons nous intéresser dans ce qui suit à la méthode de perturbations dans laquelle on suppose que le hamiltonien peut s'écrire sous la forme

$$\hat{H} = \hat{H}_0 + \hat{V},$$

où $\hat{H}_0$ est un hamiltonien indépendant du temps dont le spectre est supposé connu et où $\hat{V}$ est "petit". Le point important est que dans ce type de situation, le système évolue "pratiquement" comme sous l'action de $\hat{H}_0$ uniquement.

Écrivons le vecteur d'état sous la forme

$$|\psi(t)\rangle = \sum_n \tilde{c}_n(t) e^{-iE_n t/\hbar} |n\rangle,$$

où E_n et $|n\rangle$ désignent les énergies propres et les vecteurs propres du hamiltonien *non perturbé* $\hat{H}_0$. Si l'on se reporte à l'équation (1.15), on constate qu'avec les notations choisies ici les $\tilde{c}_n$ restent constants si $\hat{V} = 0$, et varient donc peu pour $\hat{V}$ petit non nul. Avec ces notations, l'équation de Schrödinger pour $|\psi\rangle$ s'écrit

$$i\hbar \sum_n e^{-iE_n t/\hbar} \dot{\tilde{c}}_n |n\rangle = \sum_n \tilde{c}_n e^{-iE_n t/\hbar} \hat{V} |n\rangle.$$

Si l'on projette cette relation sur un état $|m\rangle$, la condition d'orthonormalité $\langle n|m\rangle = \delta_{nm}$, où δ est le symbole de Kröneckner, nous donne

$$i\hbar \dot{\tilde{c}}_m = \sum_n e^{-i(E_n - E_m)t/\hbar} \langle m|\hat{V}|n\rangle \tilde{c}_n. \quad (1.16)$$

Faisons à présent l'hypothèse que le système était à $t = 0$ dans un état $|n_0\rangle$. Ceci signifie qu'à $t = 0$, on a $\tilde{c}_n(0) = \delta_{n,n_0}$. Pour $\hat{V} = 0$, les $\tilde{c}_n$ restent constants et le seul coefficient non nul est $\tilde{c}_{n_0}$ qui reste égal à 1. Si à présent $\hat{V}$ est non nul mais reste petit devant $\hat{H}_0$, les $\tilde{c}_{n \neq n_0}$ restent petits devant 1 alors que $\tilde{c}_{n_0}$ reste proche de 1. Dans la somme de l'équation (1.16), on voit que le terme dominant correspond à $n = n_0$ ce qui nous donne en première approximation

$$i\hbar \dot{\tilde{c}}_m \simeq e^{-i(E_{n_0} - E_m)t/\hbar} \langle m|\hat{V}|n_0\rangle,$$

Cette équation s'intègre formellement et on a pour $m \neq n_0$

$$\tilde{c}_m(t) \simeq \frac{1}{i\hbar} \int_0^t e^{-i(E_{n_0} - E_m)t'/\hbar} \langle m|\hat{V}(t')|n_0\rangle dt'. \quad (1.17)$$

1.3.3 Couplage stationnaire à un continuum, Règle d'or de Fermi

Un second cas particulièrement important est le cas où le potentiel perturbateur $\hat{V}$ est indépendant du temps et où l'état initial $|n_0\rangle$ est couplé à un continuum d'états finaux $|m\rangle$. Ce type de situation se rencontre dans de

nombreuses situations physiques parmi lesquelles on peut citer les exemples suivants :

1. *La diffusion d'une particule par un potentiel.* Dans ce cas l'état initial représente une particule de quantité de mouvement $\mathbf{p}_0$ qui est diffusée par un potentiel $V(\mathbf{r})$. Après la diffusion, la particule est en général déviée et se retrouve dans un état de quantité de mouvement $\mathbf{p}$ quelconque. La règle d'or de Fermi permet dans ce cas de calculer la *section efficace de diffusion*.
2. *Émission spontanée.* L'état initial est ici un atome placé dans un état excité. Après désexcitation radiative, l'atome se retrouve dans l'état fondamental en présence d'un photon de vecteur d'onde $\mathbf{k}$ quelconque. Nous reviendrons sur le problème de l'émission spontanée plus tard.
3. *Radioactivité.* Le cas de l'émission spontanée se généralise à celui des désintégrations radioactive, où le photon est remplacé par une particule émise par un noyau instable.

En faisant l'hypothèse d'un hamiltonien $\hat{V}$ indépendant du temps, l'équation (1.17) s'intègre immédiatement et on obtient

$$\tilde{c}_m(t) = \langle m | \hat{V} | n_0 \rangle \frac{e^{-i(E_{n_0} - E_m)t/\hbar} - 1}{E_m - E_{n_0}}.$$

La probabilité $P(t)$ d'avoir quitté l'état $|n_0\rangle$ à l'instant t vaut d'après l'interprétation probabiliste de la fonction d'onde

$$P(t) = \sum_{m \neq n_0} |\tilde{c}_m(t)|^2,$$

D'après la forme trouvée pour $\tilde{c}_m$, ceci peut se récrire de façon plus explicite

$$P(t) = t^2 \sum_{m \neq n_0} |\langle m | \hat{V} | n_0 \rangle|^2 |\text{sinc}((E_{n_0} - E_m)t/2\hbar)|^2, \quad (1.18)$$

où $\text{sinc}(x) = \sin(x)/x$ désigne la fonction sinus cardinal.

Lorsque t devient grand (on définira plus bas le domaine de validité de cette approximation), la fonction $u \mapsto \text{sinc}^2(tu)$ devient très piquée autour de $u = 0$. À la limite où $t \rightarrow \infty$, on peut donc l'approcher par une fonction δ de Dirac et on obtient

$$\text{sinc}^2(ut) \sim \frac{\pi}{t} \delta(u).$$

En utilisant cette expression dans l'expression de $P(t)$, on trouve donc

$$P(t) \sim t \frac{2\pi}{\hbar} \sum_m |\langle m | \hat{V} | n_0 \rangle|^2 \delta(E_m - E_{n_0}).$$

Posons $\langle m | \hat{V} | n_0 \rangle = \hbar \Omega(E_m)$. La probabilité de quitter l'état initial s'écrit alors

$$P(t) = \Gamma t, \quad (1.19)$$

où $\Gamma = 2\pi \hbar \sum_m \Omega^2(E_m) \delta(E_m - E_{n_0})$ désigne la *probabilité par unité de temps* de quitter l'état initial. Ce résultat est connu sous le nom de *Règle d'or de Fermi*. Cette relation appelle quelques commentaires :

1. La présence de la fonction $\delta(E_m - E_{n_0})$ garantit que seuls les états d'énergie $E_m = E_{n_0}$ contribuent au taux de probabilité : ceci traduit simplement la conservation de l'énergie, puisque partant d'un état d'énergie E_{n_0} , le système ne peut transiter que vers un état de même énergie $E_m = E_{n_0}$.
2. La probabilité de transition est proportionnelle à $|\langle E_m | \hat{V} | E_{n_0} \rangle|^2$. Si une règle de sélection aboutit à l'annulation de l'élément de matrice $\langle E_m | \hat{V} | E_{n_0} \rangle$, alors la transition est interdite.

Comportement aux temps longs

Pour que la technique de perturbation reste valable il faut que $\tilde{c}_{n_0}$ reste proche de 1. Or, par conservation des probabilité, on a $|\tilde{c}_{n_0}|^2 + P(t) = 1$ soit $|\tilde{c}_{n_0}|^2 = 1 - \Gamma t$. La dépendance linéaire de $P(t)$ que nous avons établie au paragraphe suivant n'est donc valable que pour $\Gamma t \ll 1$. Pour $\Gamma t \gtrsim 1$, on sort du domaine perturbatif et un formalisme plus raffiné doit être développé. Ces calculs montrent que $|c_0|^2$ décroît en fait exponentiellement avec un taux Γ , un résultat qui peut être retrouvé physiquement par le raisonnement suivant : imaginons que l'on observe N fois le système en répétant les mesures avec une période δt courte devant Γ et cherchons à trouver la probabilité d'avoir trouvé le système dans l'état $|n_0\rangle$ à chacune des N mesures.

La probabilité d'obtenir $|n_0\rangle$ à la p -ème mesure, sachant que l'on a obtenu $|n_0\rangle$ aux mesures précédentes, est exactement $1 - \Gamma \delta t$. En effet, en $t = (p -$

$1)\delta t^-$, juste avant la $p - 1$ -ème mesure, le système est dans un état $|\psi\rangle$ se décomposant *a priori* sur toute la base des $|m\rangle$. Comme par hypothèse j'ai obtenu le résultat $|n_0\rangle$ après la mesure, je sais d'après le postulat de réduction du paquet d'onde que le système se retrouve projeté dans l'état $|n_0\rangle$ à l'instant $t = (p-1)\delta t^+$. Entre les temps $(p-1)\delta t^+$ et $p\delta t^-$, l'évolution du système est décrite par le formalisme perturbatif, et on trouve bien le résultat annoncé.

Par récurrence, on voit donc que la probabilité d'obtenir $|n_0\rangle$ à *chacune* des $N = t/\delta t$ mesures vaut $(1 - \Gamma\delta t)^N$. Si l'on fait tendre N vers l'infini et δt vers 0 en gardant le temps $t = N\delta t$ constant, ce qui revient à observer continuellement le système, on trouve ainsi que la probabilité de trouver le système dans l'état $|n_0\rangle$ durant tout l'intervalle de temps $[0, t]$ vaut

$$P_0(t) = (1 - \Gamma\delta t)^N = e^{N \ln(1 - \Gamma\delta t)} \sim e^{-\Gamma t},$$

où l'on a utilisé le développement $\ln(1 + x) \sim x$ au voisinage de 0.

Aparté : l'effet Zénon quantique

La condition $\Gamma t \ll 1$ de validité de la méthode perturbative peut sembler en contradiction avec l'hypothèse t grand faite pour transformer le sinus cardinal en fonction δ de Dirac. Reprenons l'équation (1.18) que l'on écrit dans le cas d'un continuum de densité d'états ρ . Nous avons alors

$$P(t) = t^2 \sum_m \Omega^2(E_m) \text{sinc}^2((E_m - E_{n_0})t/2\hbar).$$

Pour pouvoir transformer le sinus cardinal par une fonction δ , il faut que sa largeur $\hbar/t$ devienne petite devant la largeur ΔE de la fonction $\Omega^2(E_m)$. La relation perturbative $P = \Gamma t$ n'est donc valable que pour

$$\frac{\hbar}{\Delta E} \ll t \ll \frac{1}{\Gamma}.$$

Pour les temps très courts ($t \ll \hbar\Delta E$), le sinus cardinal peut être remplacé par sa valeur en 0 et on trouve

$$P(t) = At^2,$$

avec $A = \sum_m \Omega^2(E_m)$.

Cette dépendance en t^2 aboutit à une conséquence assez surprenante lorsque l'on reprend l'argument de mesure continue qui nous avait permis

de retrouver le comportement aux temps longs. En effet, on voit que la probabilité d'être resté dans l'état $|n_0\rangle$ après un temps $t = N\delta t$ devient à présent

$$P_0(t) = (1 - A\delta t^2)^N = e^{N \ln(1 - A\delta t^2)} \sim e^{-At\delta t} \rightarrow 1.$$

Autrement dit, un système quantique n'évolue plus si on l'observe constamment ! Cet effet a été baptisé *Effet Zénon quantique*, en référence au paradoxe de Zénon (aussi connu sous le nom de paradoxe d'Achille et de la tortue) soulevé par le philosophe grec Zénon d'Elée ($\sim$ - 490 - -?) pour réfuter la possibilité du mouvement. Ce gel de l'évolution d'un système quantique sous l'effet d'observations répétées a été observé sur des transitions RMN que l'on empêche par des mesures optiques de plus en plus rapprochées¹⁰. Dans la plupart des cas cependant, la condition $\delta t \ll \hbar/\Delta E$ est trop difficile à remplir pour que l'Effet Zénon puisse se manifester. Ainsi, pour bloquer l'émission spontanée d'une atome, il faudrait être capable d'effectuer des mesures avec des taux de répétition de l'ordre des fréquences optiques, soit 10^{15} Hz !

10. W. Itano, Phys. Rev. A **41**, 2295 (1990).

Chapitre 2

Principe de fonctionnement d'un laser

2.1 Laser continu

Le laser, émettant une onde électromagnétique de fréquence bien déterminée, est un “oscillateur optique”. Le principe général de réalisation d'un oscillateur électronique est bien connu : il consiste simplement à boucler un amplificateur sur un filtre passe bande de fréquence centrale f_0 . Ce type d'oscillateur démarre sur du bruit : en effet, dans des situations réelles, le signal à l'entrée du filtre n'est jamais strictement nul. Il existe toujours un bruit (souvent faible) à n'importe quelle fréquence. Ce bruit est filtré par le passe-bande et à la sortie du filtre seul subsiste un bruit de fréquence f_0 . Ce bruit passe ensuite dans l'amplificateur et pour un gain suffisamment grand, il est réinjecté dans le filtre avec une amplitude plus grande qu'à la première étape. Après quelques tours, le bruit est complètement filtré et le signal oscille purement à la fréquence f_0 .

Le laser fonctionne sur le même principe. La rétroaction et le filtrage fréquentiel sont assurés par une cavité Fabry-Pérot et l'émission stimulée dans un milieu à “inversion de population” assure l'amplification optique. Lorsque le gain réalisé par émission stimulée est suffisamment grand, le laser se met à émettre spontanément une lumière quasi-monochromatique.

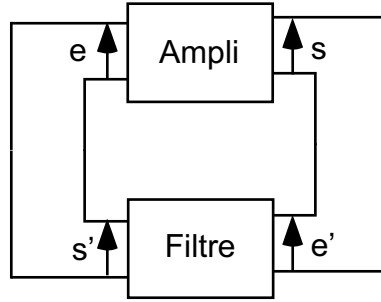


FIG. 2.1 – *Principe de réalisation d'un oscillateur électronique. On effectue une rétroaction entre un amplificateur et un filtre passe-bande de fréquence centrale f_0 . Lorsque le gain de l'amplificateur est suffisamment grand, le circuit se met à osciller spontanément à la fréquence f_0 .*

2.1.1 Cavité Fabry-Pérot

La cavité Fabry-Pérot (ou cavité résonnante) la plus simple est constituée de deux miroirs placés en vis-à-vis. Dans un laser, le rôle de la cavité Fabry-Pérot est double. D'une part il sert à fermer la "rétroaction" en maintenant les photons au sein de la cavité, ce qui permet d'entretenir l'émission stimulée. D'autre part, les interférences entre les ondes ayant fait plusieurs aller-retours dans la cavité ne laissent subsister que certaines longueurs d'onde bien particulières, ce qui confère à la cavité son rôle de filtre fréquentiel. La cavité Fabry-Pérot "standard" que nous venons de décrire n'est pas la seule stratégie envisageable. Nous verrons en particulier dans l'étude du gyrolaser qu'il est possible d'utiliser des cavités dites en anneaux (Fig. 2.2). Cette configuration présente l'avantage de ne pas créer de modulation d'intensité dans la cavité sous l'effet des interférences entre les ondes aller et retour : ces modulations sont en effet source d'une réduction de l'efficacité de l'émission stimulée, celle-ci étant en effet proportionnelle à l'intensité électromagnétique (*spatial hole burning*). Par souci de simplicité, nous ne considérerons donc ici que de telles cavités en anneau. Nous supposons que les trois miroirs sont sans perte et que seul l'un d'entre eux (appelé coupleur de sortie) possède une réflectivité différente de 1. On note r et t les coefficients de réflexion et de transmission en *amplitude* de ce miroir. On pose par ailleurs $R = |r|^2$ et $T = |t|^2$, avec $R + T = 1$ puisque le miroir est sans perte.

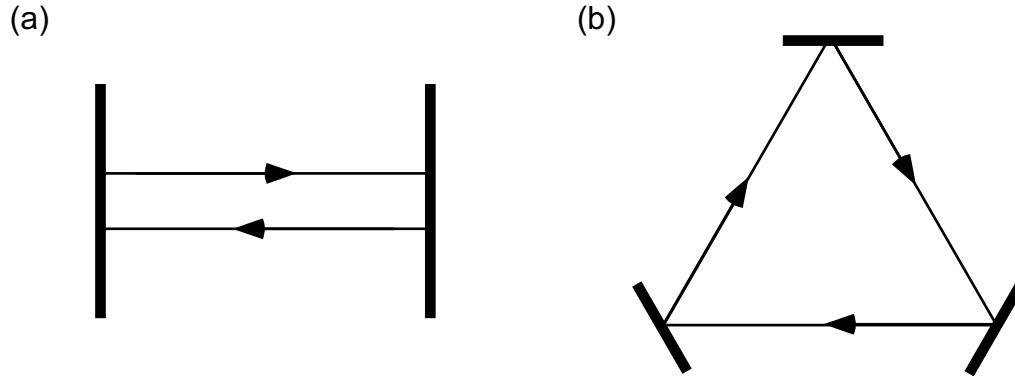


FIG. 2.2 – a) *Cavité Fabry-Pérot linéaire.* b) *Cavité Fabry-Pérot en anneaux.*

Recyclage et durée de vie d'un photon dans la cavité

Considérons pour commencer le cas d'un seul miroir que l'on éclaire avec une onde monochromatique de pulsation ω et de puissance P_0 . D'après la définition des coefficients R et T , les puissances réfléchies et transmises valent respectivement $P_r = RP_0$ et $P_t = TP_0$.

Tentons d'interpréter ce résultat en termes de photons. Sachant qu'un photon porte une énergie $\hbar\omega$, le flux de photons incidents Φ_0 est relié à la puissance incidente par la relation $P_0 = \hbar\omega\Phi_0$ et, de manière analogue, les flux de photons réfléchis et transmis valent simplement $\Phi_r = R\Phi_0$ et $\Phi_t = T\Phi_0$. Autrement dit, R et T s'interprètent comme la proportion de photons réfléchis et transmis. Si le faisceau ne contient qu'un seul photon, R et T peuvent donc s'interpréter en terme probabilistes comme les probabilités de réflexion et de transmission du photon à l'impact sur le miroir.

Revenons au cas de la cavité Fabry-Pérot en anneau. D'après ce qui précède, la probabilité qu'a le photon d'être resté dans la cavité à chaque passage sur le coupleur de ortie vaut R . Ce résultat se généralise à n tours après lesquels la probabilité d'avoir gardé le photon dans la cavité vaut R^n . Le photon se propageant à la vitesse de la lumière, au bout d'un temps t long il a fait $n \sim ct/L$ tours dans la cavité de longueur L . La probabilité qu'a le photon d'être resté piégé jusqu'au temps t s'écrit donc

$$P(t) = R^{ct/L} = \exp\left(\frac{ct}{L} \ln(R)\right).$$

Si à présent on suppose qu'un grand nombre $N_\gamma(0)$ de photons étaient

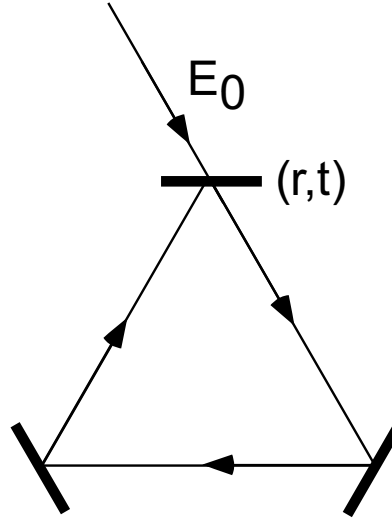


FIG. 2.3 – *Interférences multiples dans une cavité Fabry-Perot. Le champ électrique est la somme de toutes les ondes réfléchies à la surface des miroirs de la cavité.*

initialement présents dans la cavité, le nombre de photons N_γ contenus dans celle-ci suit une loi du type

$$N_\gamma(t) = N_\gamma(0) \exp(-t/\tau),$$

avec $\tau^{-1} = c|\ln(R)|/L$.

Autrement dit, l'évolution du nombre de photons présents dans la cavité sous l'effet des pertes peut se modéliser par une équation du type

$$\left(\frac{dN_\gamma}{dt} \right)_{\text{pertes}} = -\frac{N_\gamma}{\tau}.$$

Filtrage fréquentiel

L'analyse précédente reste très simpliste puisqu'elle néglige complètement la nature ondulatoire de la lumière et donc les éventuelles interférences pouvant se produire entre des rayons ayant fait un nombre différent d'aller-retours. Ces interférences aboutissent à ne conserver dans la cavité que les

modes dont la longueur d'onde interférant constructivement¹. Imaginons que l'on éclaire le miroir de couplage avec un faisceau monochromatique dont le champ électrique a pour amplitude E_0 , et cherchons à calculer le champ E_{int} à l'intérieur de la cavité. Si l'on somme les ondes réfléchies (Fig. 2.3), on trouve que E_{int} se met sous la forme d'une série infinie de termes correspondant aux réflexions multiples à la surface des miroirs et on obtient

$$E_{\text{int}} = E_0 t \sum_{n=0}^{\infty} r^n e^{inkL} = \frac{E_0 t}{1 - re^{ikL}},$$

la convergence de la série étant assurée par la condition $R < 1$. Si V désigne le volume de la cavité, l'énergie électromagnétique U stockée dans la cavité s'écrit après quelques manipulations

$$U = V \frac{\epsilon_0 |E_{\text{int}}|^2}{2} = V \frac{\epsilon_0 E_0^2}{2} \left(\frac{1+r}{1-r} \right) \frac{1}{1 + \mathcal{F} \sin^2(kL/2)}, \quad (2.1)$$

où $\mathcal{F} = 4r/(1-r)^2$ est baptisée *finesse* de la cavité. La raison de cette dénomination se comprend aisément en représentant graphiquement l'évolution de U en fonction de kL (Fig. 2.4). La courbe obtenue, dite courbe d'Airy, présente une succession de pics situés en $kL = n\pi$. Si l'on développe U au voisinage d'un maximum en posant $kL = n\pi + \epsilon$, l'équation (2.1) devient

$$U \sim V \frac{\epsilon_0 E_0^2}{2} \left(\frac{1+r}{1-r} \right) \frac{1}{1 + \mathcal{F} \epsilon^2/4},$$

c'est à dire une courbe lorentzienne de largeur à mi-hauteur $\Delta\epsilon \sim 1/\sqrt{\mathcal{F}}$. Autrement dit, plus $\mathcal{F}$ est grand, plus les pics sont étroits. La limite des grandes finesesses correspond à des miroirs très réfléchissants. Dans cette limite, on trouve à l'ordre dominant²

$$\Delta\epsilon \sim T \quad (2.2)$$

L'allure de la courbe d'Airy est compréhensible par des arguments relativement simples. Tout d'abord, la position des pics correspond à des interférences constructives entre les différentes ondes réfléchies. Après un tour

1. Cette discrétisation des modes est en tout point analogue à celle que nous avons obtenu au chapitre précédent lorsque nous avons introduits les conditions de Born-von Karman.

2. On utilise le développement $r = \sqrt{1-T} \sim 1 - T/2$ pour les petits T .

dans la cavité, le déphasage subi par une onde vaut kL . On aura donc interférences constructives pour $kL = 2n\pi$, soit $kL/2 = n\pi$.

Plaçons nous à présent dans une situation légèrement non-résonnante pour laquelle $k = 2n\pi/L + \delta k$. À chaque tour, l'onde acquiert un déphasage δkL , soit après accumulation sur N tours, $N\delta kL$. Ce déphasage aboutira à une interférence destructive lorsque $N\delta kL \sim 1$.

L'effet de ce déphasage n'aura la possibilité de se manifester que si l'onde a le temps de faire N tours dans la cavité avant de la quitter. Or, nous avons vu au paragraphe précédent qu'un photon ne subsistait qu'un nombre de tours de l'ordre de $N \sim 1/|\ln(R)|$. Plaçons nous dans une situation de grande finesse, correspondant à des miroirs très réfléchissant. D'après la relation $R + T = 1$, on voit que $|\ln(R)| = |\ln(1 - T)| \sim T$, car la transmission des miroirs est faible, et on obtient $N \sim 1/T$.

En définitive, les interférences resteront constructives en dehors de la condition de résonance si l'on vérifie l'inégalité

$$\Delta\epsilon = \delta kL \lesssim T.$$

On retrouve donc la formule (2.2) donnant la largeur des pics d'Airy³.

Notons enfin qu'à résonance l'intensité lumineuse est supérieure à ce qu'elle serait en l'absence de cavité. En effet, si la lumière se propageait librement, l'énergie électromagnétique vaudrait

$$U_{\text{libre}} = V \frac{\epsilon_0 E_0^2}{2}.$$

Si l'on compare à l'énergie stockée dans une cavité résonnante ($\epsilon = 0$), on constate que

$$U = \frac{1+r}{1-r} U_{\text{libre}} > U_{\text{libre}}.$$

Cet effet d'accumulation (*build-up* en anglais) est l'équivalent ondulatoire du recyclage des photons dans la cavité et permet d'amplifier l'émission stimulée.

3. Cette largeur peut se retrouver d'une troisième manière, en invoquant une inégalité de type Heisenberg $\Delta\omega\tau \sim 1$, où $\Delta\omega$ représente l'incertitude sur la pulsation d'un photon et τ la durée de vie de ce photon dans la cavité.

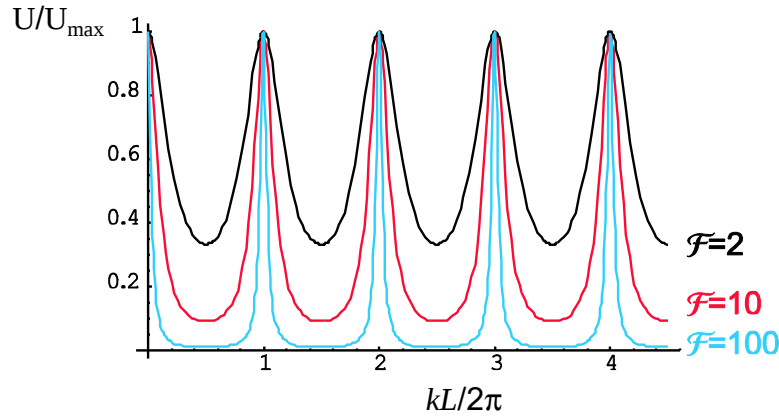


FIG. 2.4

2.1.2 Amplification optique et inversion de population

Comme noté précédemment, l'amplification nécessaire au démarrage d'un oscillateur est assurée par l'émission stimulée, qui impose à un photon d'être préférentiellement émis dans un mode déjà peuplé.

Considérons la cavité remplie d'un milieu constitué d'atomes dont une transition (de l'état fondamental $|f\rangle$ à l'état excité $|e\rangle$) est résonnante avec un mode de la cavité et cherchons à calculer la variation du nombre de photons N_γ dans ce mode sous l'effet des différents processus radiatifs. La variation de N_γ est due aux phénomènes d'émission stimulée et d'absorption et, en utilisant les coefficients d'Einstein introduits au chapitre précédent, on peut écrire que

$$\left(\frac{dN_\gamma}{dt} \right)_{\text{radia}} = B_{\text{ef}} u(\omega_{\text{ef}}) (N_e - N_f),$$

où $u(\omega_{\text{ef}})$ est la densité spectrale d'énergie électromagnétique à la pulsation ω_{ef} de la résonance atomique et N_e et N_f désignent le nombre d'atomes dans les états $|e\rangle$ et $|f\rangle$. D'après cette relation, l'amplification ne fonctionnera que pour $N_e > N_f$. Cette configuration nécessite ce que l'on appelle une "inversion de population" car, contrairement à la condition d'équilibre thermodynamique, l'état fondamental est moins peuplé que l'état excité. Dans de nombreux cas, l'inversion de population est réalisée par *pompage optique*⁴

4. Le pompage optique a été observé pour la première fois par Alfred Kastler dans les années 50, ce qui lui a valu de remporter le prix Nobel de physique 1966.

où l'on vide l'état fondamental optiquement.

Dans un système à deux niveaux atomiques seulement, le pompage optique ne permet pas d'effectuer l'inversion de population. En effet, supposons que l'on maintienne u constant dans la cavité à l'aide d'une source lumineuse externe. L'évolution du nombre d'atomes dans l'état excité suit l'équation déjà rencontrée dans l'étude du corps noir :

$$\frac{dN_e}{dt} = -AN_e + Bu(N_f - N_e).$$

La quantité intéressante est l'inversion de population $\Delta = N_e - N_f$ entre les populations de l'état fondamental et l'état excité. En notant que le nombre total d'atomes $N_0 = N_e + N_f$ est conservé, on trouve que Δ vérifie l'équation d'évolution

$$\dot{\Delta} = -AN_0 - (2Bu + A)\Delta.$$

Comme précédemment, on a $\dot{\Delta} = 0$ en régime stationnaire et on trouve donc

$$\Delta = -\frac{AN_0}{2Bu + A} < 0.$$

Autrement dit, l'inversion de population n'est jamais réalisée, même lorsque l'énergie lumineuse tend vers l'infini⁵. Afin de contourner ce problème, on a recours à une stratégie faisant intervenir plus de deux niveaux atomiques, en général trois ou quatre (Fig. 2.5).

Dans le mécanisme à trois niveaux (2.5.A), le pompage optique permet d'amener les atomes de l'état fondamental $|f\rangle$ à un état excité intermédiaire $|e'\rangle$. Cet état se désexcite rapidement (et en général non radiativement, c'est-à-dire sans émission de photon) vers un état de "longue" durée de vie $|e\rangle$. L'émission laser a alors lieu entre l'état $|f\rangle$ à présent vide, et l'état $|e\rangle$.

L'inconvénient de la stratégie utilisant trois niveaux est qu'il est nécessaire de presque totalement vider l'état fondamental afin de réaliser l'inversion de population. Afin de contourner ce problème, il est possible d'utiliser un schéma à quatre niveaux où l'on fait intervenir un nouveau niveau $|f'\rangle$, intermédiaire entre $|e\rangle$ et $|f\rangle$ et peu peuplé à l'équilibre thermodynamique

5. Pour $u \gg A/B$, on dit que la transition est saturée. Dans cette limite, les nombres d'atomes dans les deux niveaux deviennent égaux et on peut montrer que l'absorption du milieu tend vers 0.

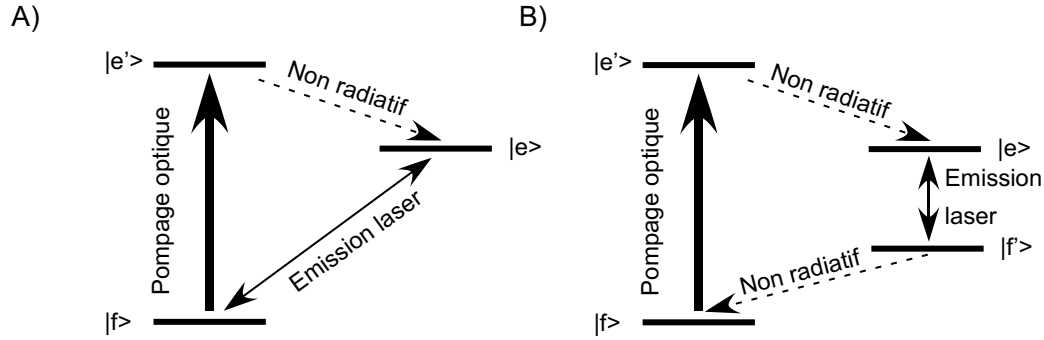


FIG. 2.5 – Principe du laser à 3 (A) et 4 (B) niveaux.

(2.5.B). La lumière laser est alors émise sur la transition $|e\rangle \leftrightarrow |f'\rangle$ sur laquelle il est très simple de réaliser une inversion de population.

2.1.3 Dynamique d'un laser à 3 niveaux pompé optiquement

Ce paragraphe est consacré à une discussion plus quantitative de la dynamique du démarrage d'un laser à 3 niveaux. Dans cette discussion, nous ferons les hypothèses suivantes :

- On note Q le taux de pompage de l'état fondamental vers l'état excité intermédiaire.
- On suppose que le taux de désexcitation non radiative γ de l'état $|e'\rangle$ vers l'état $|e\rangle$ est très rapide devant les autres temps caractéristiques du problème (en particulier le temps de désexcitation radiative vers l'état fondamental).
- Le nombre de photons N_γ dans la cavité est proportionnel à la densité spectrale d'énergie de la transition laser. On note $u = \rho N_\gamma$, où la constante de proportionnalité ρ dépend des paramètres de la cavité.

À l'aide des hypothèses précédentes, nous pouvons à présent écrire les équations satisfaites par les populations des différents niveaux atomiques, ainsi que le nombre de photons dans la cavité. On obtient le système non linéaire suivant

$$\dot{N}_f = -QN_f + AN_e + B\rho N_\gamma(N_e - N_f) \quad (2.3)$$

$$\dot{N}_{e'} = QN_f - \gamma N_{e'} \quad (2.4)$$

$$\dot{N}_e = \gamma N_{e'} - AN_e - B\rho N_\gamma (N_e - N_f) \quad (2.5)$$

$$\dot{N}_\gamma = -\frac{N_\gamma}{\tau} + B\rho N_\gamma (N_e - N_f) \quad (2.6)$$

Ce système d'équations à quatre inconnues ne possède pas de solution analytique dans le cas général. On peut cependant le simplifier en utilisant l'hypothèse de désexcitation rapide de l'état intermédiaire $|e'\rangle$. Cette hypothèse, analogue à l'approximation de l'état quasi-stationnaire en cinétique chimique permet de négliger $\dot{N}_{e'}$ devant $\gamma N_{e'}$, si le temps caractéristique d'évolution du système est très long devant γ^{-1} . D'après l'équation (2.4), ceci nous permet d'écrire que $\gamma N_{e'} \sim QN_f$ et ainsi d'éliminer $N_{e'}$ des équations.

En utilisant cette approximation et en introduisant le nombre total d'atomes N_0 , on aboutit à un système de deux équations ne portant que sur l'inversion de population $\Delta = N_e - N_f$ et le nombre de photon N_γ , soit

$$\dot{\Delta} = (Q - A)N_0 - (A + Q + B\rho N_\gamma)\Delta \quad (2.7)$$

$$\dot{N}_\gamma = -\frac{N_\gamma}{\tau} + B\rho N_\gamma \Delta \quad (2.8)$$

C'est le système d'équations (2.7) et (2.8) que nous allons à présent discuter.

Régime stationnaire

Intéressons nous tout d'abord aux solutions stationnaires pour lesquelles on a $\dot{N}_\gamma = \dot{\Delta} = 0$.

1. *Laser éteint.* D'après l'équation (2.8), on voit qu'une première solution est donnée par $N_\gamma = 0$, puis, selon (2.7),

$$\Delta = N_0 \frac{Q - A}{Q + A}. \quad (2.9)$$

Cette solution, où aucun photon stimulé n'est émis, correspond tout simplement à un laser *éteint*. Dans ce régime, et pour un taux de pompage très élevé ($Q \gg A$), on constate que l'inversion de population tend vers sa valeur maximale N_0 . En effet, l'émission stimulée étant bloquée,

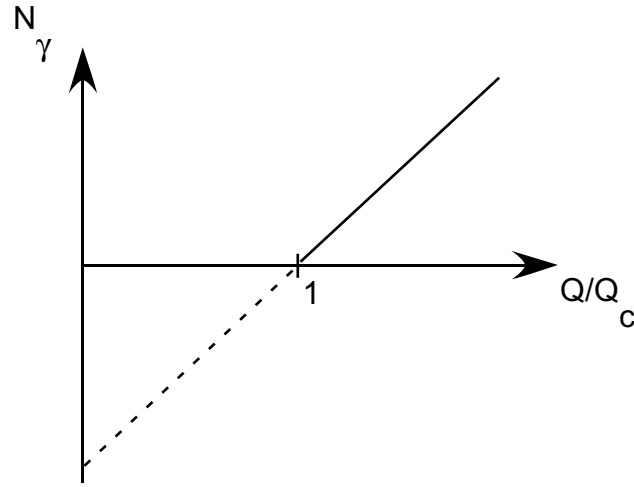


FIG. 2.6 – *Seuil de fonctionnement d'un laser. Le nombre de photon dans la cavité ne devient non nul que lorsque le taux de pompage Q dépasse une seuil critique Q_c .*

la seule voie de désexcitation de l'état $|e\rangle$ vers l'état $|f\rangle$ est par émission spontanée, relativement lente comparée à l'émission stimulée⁶.

2. *Laser allumé.* La deuxième solution stationnaire de (2.8) consiste à imposer $\Delta = 1/\tau B\rho$. En injectant cette relation dans (2.7), on obtient un nombre de photons non nul avec

$$N_\gamma = \frac{\tau}{2} \left(-A \left(N_0 + \frac{1}{B\rho\tau} \right) + Q \left(N_0 - \frac{1}{B\rho\tau} \right) \right).$$

Pour que cette solution ait un sens, il faut que le nombre de photons dans la cavité soit positif. Ceci impose d'une part d'avoir $N_0 > 1/B\rho\tau$, c'est-à-dire d'avoir suffisamment d'atomes dans la cavité. D'autre part, il faut disposer d'un taux de pompage Q plus grand qu'une valeur critique Q_c définie par

$$Q_c = A \frac{N_0 + \frac{1}{B\rho\tau}}{N_0 - \frac{1}{B\rho\tau}}.$$

6. Cette remarque, qui peut sembler pour le moment anodine, se révélera cruciale dans la réalisation de laser impulsionnels.

Stabilité de la solution éteinte et démarrage du laser

Alors qu'il n'existe qu'une seule solution acceptable lorsque $Q < Q_c$ (laser éteint), on ne peut savoir par l'étude du régime stationnaire quelle solution, éteinte ou allumée, sera observée en pratique lorsque $Q > Q_c$. Afin de lever cette ambiguïté, nous allons procéder à une étude de stabilité linéaire des solutions stationnaires. Plus précisément, nous allons partir d'une condition initiale très proche de la solution éteinte et nous allons chercher à savoir si le laser retourne vers cette solution éteinte, ou s'en éloigne.

Soit $\Delta_{e,0}$ la valeur stationnaire de l'inversion de population dans le régime du laser éteint (Eqn. (2.9)). Puisqu'en régime stationnaire, N_γ est strictement nul, l'équation (2.8) peut s'écrire à l'ordre 1 en perturbation

$$\dot{N}_\gamma = -\frac{N_\gamma}{\tau}(1 - B\rho\tau\Delta_{e,0}).$$

Cette équation se résout immédiatement et on constate que N_γ ne tendra vers 0 aux temps longs que si $B\rho\tau\Delta_{e,0}$ est plus grand que 1, soit, après quelques manipulations, si $Q < Q_c$.

Autrement dit, la solution éteinte n'est stable que pour $Q < Q_c$. Dès que Q devient plus grand que Q_c , la solution éteinte devient instable. La moindre perturbation - par exemple le premier photon spontané émis dans le mode de la cavité, suffit à ramener le système vers la solution allumée.

L'existence d'un seuil au dessus duquel le système change de régime de fonctionnement est appelée en mathématiques une *bifurcation*. Physiquement, c'est une caractéristique essentielle du fonctionnement du laser. Le seuil Q_c correspond à la valeur du taux de pompage au delà de laquelle le gain (c'est-à-dire l'émission stimulée) l'emporte sur les pertes (essentiellement l'émission spontanée et les pertes par les miroirs de sortie de cavité).

2.1.4 “Laserologie”

À titre d'illustration des principes que nous venons de développer, le paragraphe qui suit est consacré à une liste bien loin d'être exhaustive de quelques modèles courants de laser.

Le maser

Le maser n'est pas à strictement parler un laser, mais plutôt son ancêtre. Il a été inventé en 1954 par Townes. Son fonctionnement repose sur le

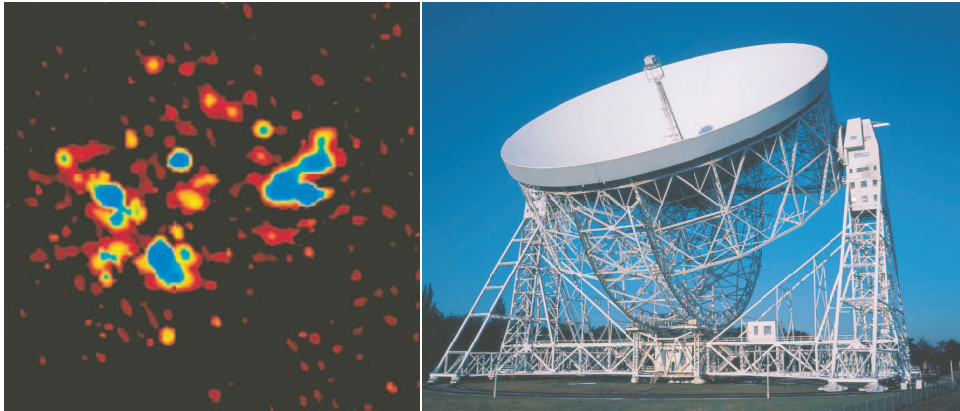


FIG. 2.7 – Maser à SiO observé à l’observatoire radioastronomique de Jodrell Bank (Grande Bretagne).

même principe que le laser, à ceci près qu’il émet un rayonnement micro-onde de longueur d’onde centimétrique (le “M” de maser signifie “Microwave”). Les masers ont été développés avant le laser en raison de la faiblesse de l’émission spontanée aux grandes longueurs d’ondes. L’émission stimulée est en fait tellement importante dans les micro-ondes que des masers naturels ont été observés par les astrophysiciens dans les régions de formation des étoiles. La transition du maser à eau, située à 22 GHz, est ainsi la raie la plus brillante de l’univers observée dans le domaine radio⁷, mais des masers impliquant d’autres transitions moléculaires (SiO ou méthanol par exemple) ont aussi été observés. Dans ces masers naturels, l’inversion de population est provoquée par des gradients de température entre les nuages de gaz et de poussière. En laboratoire, les masers à hydrogène, dont le milieu amplificateur est une vapeur d’hydrogène atomique, émettent un rayonnement à 21 cm et sont utilisés en métrologie comme horloge de grande précision.

Laser à gaz

- *Laser hélium-néon.* C’est le plus couramment utilisé dans les salles de classes. Le milieu amplificateur est constitué d’un mélange d’hélium et de néon gazeux. L’hélium est excité par une décharge électrique

7. Plus proche de nous cette raie à 22 GHz a été observée en 1994 lors de l’impact de la comète Shoemaker-Levy sur la surface de Jupiter.

dans un état situé à 20 eV au dessus du fondamental. Cet état excité coïncide presque parfaitement avec un des états excités du néon : les collisions hélium/néon permettent donc le transfert d'excitation d'une espèce chimique à l'autre. C'est ensuite depuis l'état excité du néon qu'a lieu l'émission laser. La longueur d'onde habituelle de l'hélium-néon est située dans le rouge à 633 nm. Cependant, il existe des modèles émettant à des longueurs d'ondes différentes, notamment dans le vert et l'infrarouge. La différence entre les différentes versions est simplement le type du miroir utilisé qui réfléchissent plus ou moins bien suivant la longueur d'onde et permettent de favoriser une raie d'émission par rapport à une autre.

- *Laser CO₂*. Le laser CO₂ utilise un principe similaire à l'hélium néon : on utilise une décharge électrique pour exciter de l'azote qui transfère son excitation au dioxyde de carbone. L'émission laser se fait dans l'infrarouge lointain à 10 μ m. Le principal intérêt de ce type de laser est son excellent rendement (quasiment 30 %) ainsi que la forte puissance accessible. Certains modèles peuvent ainsi fournir jusqu'à 10 kW de puissance continue. Cette caractéristique le rend particulièrement approprié dans l'usinage de matériaux ou la soudure laser.

Laser solide

- *Laser à rubis*. Son principal intérêt est historique, puisqu'il s'agit du premier laser jamais réalisé (Maiman, 1960). Son manque de fiabilité l'a fait tomber en désuétude... Le milieu amplificateur est un rubis, c'est-à-dire une matrice d'alumine Al₂O₃ dopée aux ions chrome Cr³⁺. Ce sont ces ions qui donnent sa couleur rouge au rubis et qui servent à l'émission laser. Le pompage optique est opéré par flash et le laser à rubis ne peut donc fonctionner qu'en impulsion (un tir à chaque flash de pompage).
- *Laser YAG*. Un autre laser de puissance, le laser YAG (Yttrium Aluminium Garnet, ou Grenat d'Aluminium et d'Yttrium, formule chimique Y₃Al₅O₁₂). La matrice de YAG est robuste mécaniquement, transparente et optiquement isotrope. Elle peut être dopée à l'aide de divers ions (néodyme, ytterbium..., suivant la longueur d'onde désirée) qui comme dans le cas du rubis vont servir à l'émission laser. Comme le laser CO₂, il est capable de fournir de fortes puissances lumineuses

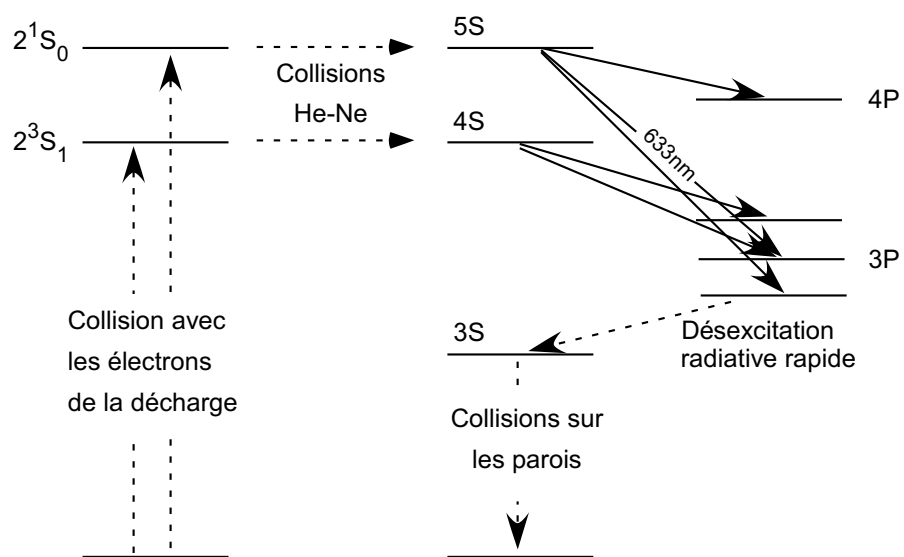


FIG. 2.8 – Schéma de niveau du laser hélium néon. La transition laser a lieu entre deux niveaux du néon. Plusieurs raies sont utilisables, à $3\ \mu\text{m}$, $1\ \mu\text{m}$, $633\ \text{nm}$ (dans le rouge, c'est la plus couramment utilisée) et dans le vert à $543\ \text{nm}$.

(quelques centaines de watts), à des longueurs d'ondes plus basses ($1\text{ }\mu\text{m}$ pour le Nd: YAG) où la matière absorbe plus que le CO_2 . Il est couramment utilisé dans les applications industrielles ou médicales, notamment en ophtalmologie.

- *Laser titane-saphir*. Comme les deux laser précédents, le milieu amplificateur est une matrice (ici de saphir) dopée aux ions titane. Sa principale qualité est de posséder une large bande d'émission (de 760 à 1100 nm), ce qui lui permet d'être accordable sur une grande plage de longueurs d'ondes et de pouvoir fonctionner en régime impulsif où il est capable de délivrer des impulsions de quelques femtosecondes seulement de durée.

Lasers à colorant

Les colorants ont une large bande d'émission et d'absorption optique ce qui les rend particulièrement adaptés à la réalisation de lasers accordables sur de larges plages de longueur d'onde. Afin de permettre la sélection de fréquence, on insère un réseau dans la cavité, de façon à ne conserver qu'une seule longueur d'onde, celle qui est diffractée normalement au miroir de fond de cavité. La rhodamine 6G, par exemple, permet de générer des longueurs d'onde allant de 565 à 595 nm. Les lasers à colorants sont utilisés dans les applications où une large bande de gain et/ou d'accordabilité sont nécessaires, notamment en spectroscopie, dans la génération d'impulsions courtes (cf. paragraphe sur le blocage de mode), la photothérapie, la dermatologie etc.

Diode laser

Ce sont très certainement les lasers les plus courants puisqu'on les trouve dans les lecteurs CD, les lecteurs de codes barres et sont massivement utilisés dans les télécommunications. Leur succès provient de l'immense développement de l'industrie optoélectronique qui a permis en particulier le développement de sources laser miniaturisées. Les propriétés optiques des diodes lasers proviennent directement de la structure de niveaux des semi-conducteurs la constituant. Comme dans tout cristal, les niveaux électroniques dans un semi-conducteur se présentent sous forme de bandes permises, séparées par des bandes interdites. Comme pour les atomes, chaque bande est remplie en accord avec

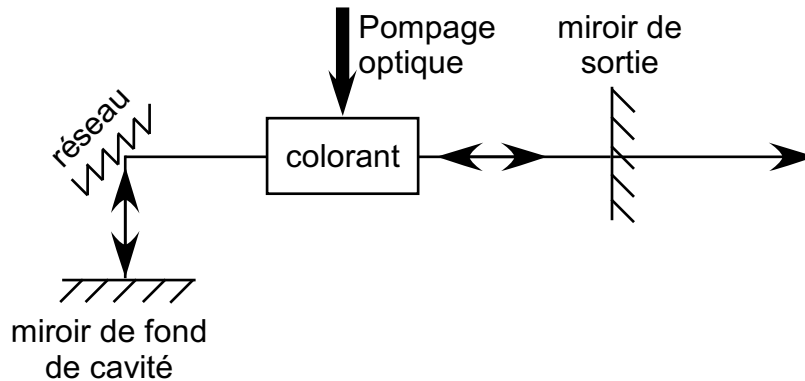


FIG. 2.9 – Principe de fonctionnement d'un laser à colorant. Le réseau est utilisé pour permettre la sélection de la longueur d'onde d'émission.

le principe de Pauli et, par définition, on appelle bande de valence la dernière bande complètement remplie et bande de conduction la bande suivante. Les isolants correspondent à des solides dans lesquels la bande de conduction est vide, au contraire des métaux pour lesquels cette bande est peuplée. Les semi-conducteurs sont à mi-chemin entre le métal et l'isolant. En effet, à température rigoureusement nulle, un semi-conducteur est un isolant. Cependant, la bande interdite séparant la bande de valence de la bande de conduction est relativement faible, ce qui fait que certains électrons peuvent passer thermiquement dans la bande de conduction, en laissant des trous dans la bande de valence. Ces trous peuvent être vus comme des particules à part entière, de charge opposée à celle de l'électron.

La conductivité d'un semi-conducteur peut être augmentée en le dopant, c'est-à-dire en ajoutant des impuretés pouvant soit capter des électrons de la bande de valence, soit libérer des électrons dans la bande de conduction de façon à augmenter respectivement le nombre de porteurs positifs (on parle alors de dopage P) ou négatifs (on parle de dopage N).

Les propriétés optiques des semi-conducteurs proviennent simplement de la possibilité qu'a un électron de conduction de remplir un trou de la bande de valence en émettant un photon (de façon imagée, on peut voir ce processus comme l'annihilation mutuelle d'un électron et d'un trou par émission d'un photon). La diode laser est réalisée en juxtapo-

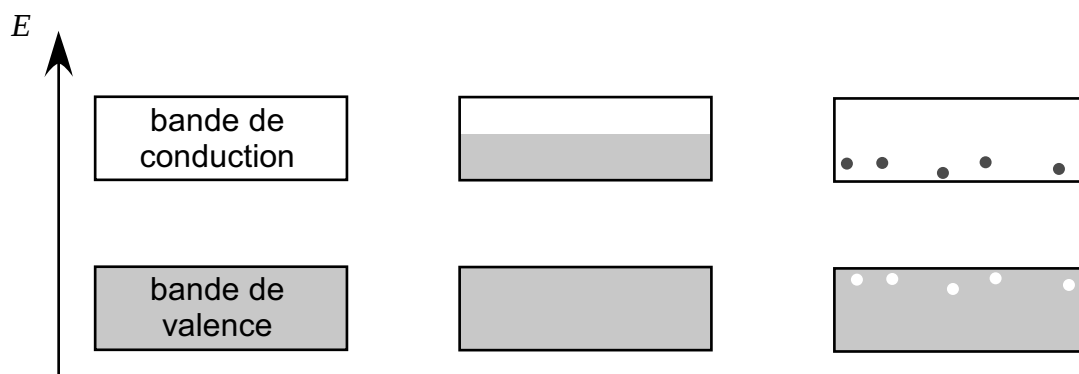


FIG. 2.10 – *Structure de bande dans les solides. De gauche à droite. Isolant : la bande de conduction est vide. Métal : la bande de conduction est partiellement remplie. Semi-conducteur : isolant à température nulle. La bande interdite est cependant étroite et à température non nulle, des électrons passent de la bande de valence à la bande de conduction. Les trous de la bande de valence se comportent comme des particules fictives de charge opposée à celle de l'électron.*

posant des semi-conducteurs dopés P et N⁸. En injectant un courant électrique dans la diode, on amène électrons et trous au niveau de la jonction, ce qui provoque l'émission de lumière. On a alors réalisé une diode électroluminescente. Afin de provoquer l'émission stimulée, on clive proprement les faces de la diode. Le gain dans la diode est tellement important que la réflexion (de l'ordre de 30%) de la lumière à l'interface semi-conducteur/air suffit à déclencher l'émission laser. Les diodes lasers sont développées dans trois gammes de longueurs d'ondes, suivant leurs applications.

- Infrarouge : dans les télécommunications optiques, la lumière est transportée par fibre optique. Afin de pouvoir transporter l'information sur des centaines de km sans pertes, il est nécessaire d'utiliser une longueur d'onde où l'absorption de la fibre est la plus faible. Les fibres optiques sont réalisées en verre (quartz, bromosilicate, silice...) et la principale source de pertes provient de l'absorption par des groupes hydroxyle dérivés de l'eau pouvant

8. Notons qu'une diode laser possède la même structure qu'une diode ou qu'une diode électroluminescente.

se former dans la fibre durant sa fabrication. Ces pertes sont minimales à 1550 nm, longueur d'onde utilisée par conséquent pour les télécommunications optiques.

- Domaine optique. On développe des diodes dans le domaine optique dans les applications où il est nécessaire de voir le faisceau, telles que les pointeurs lasers.
- Diodes bleues. Les capacités de stockage des CD ou des DVD, sont limitées par la longueur d'onde des lasers utilisés pour leur gravure et leur lecture. C'est pour cela qu'un intense effort est développé vers le développement de diodes lasers de courte longueur d'onde, en particulier dans le bleu. Les diodes laser utilisées actuellement pour la lecture de CD émettent dans le rouge. Elles pourraient donc bientôt être remplacées par des diodes bleues, ce qui porterait la capacité de stockage des DVDs à 27 Go (contre 5 Go actuellement) - projet Blue Ray Disk développé en collaboration par Hitashi, LG, Matsushita, Pioneer, Philips, Samsung, Sharp, Sony et Thomson.

Laser à un atome

- Une expérience récente réalisée à Caltech (J. McKeever et al., *Nature* **425**, 268 (2003)) a démontré la possibilité de réaliser un laser dont le milieu actif ne contient qu'un atome. L'atome utilisé est un atome de césium maintenu dans une cavité optique de haute finesse. Le laser fonctionne sur un schéma à 4 niveaux et délivre une puissance d'une fraction de nW.

2.2 Propriétés de cohérence

Au paragraphe précédent, l'approche du type "équations de robinets" que nous avons suivie pour décrire le démarrage de l'émission laser a passé sous silence une des propriétés fondamentales de la lumière laser, à savoir sa monochromaticité. La question assez naturelle se pose en particulier de savoir à quelle longueur d'onde le laser fonctionne (ce que l'on appelle le mode longitudinal de fonctionnement du laser). Le choix du mode résulte d'une compétition entre l'émission stimulée, qui ne peut avoir lieu qu'à la fréquence de résonance de l'atome, et la cavité Fabry-Pérot, qui impose des

contraintes sur les valeurs de longueurs d'ondes pouvant être stockées une durée non négligeable dans la cavité.

2.2.1 Modes longitudinaux d'un laser

Notion de courbe de gain : exemple de l'élargissement Doppler

Considérons un gaz d'atome éclairé par une lumière parfaitement monochromatique de pulsation ω_L . La densité spectrale d'énergie de l'onde est donc $u(\omega) = \varpi \delta(\omega - \omega_L)$, où ϖ est la densité d'énergie électromagnétique de l'onde (puisque par définition $\varpi = \int u(\omega) d\omega$). D'après la théorie des coefficients d'Einstein, le taux d'émission stimulée d'un photon par N_e atomes au repos est donné par la loi

$$\frac{dN_\gamma}{dt} = N_e B_{fe} \varpi \delta(\omega_a - \omega_L).$$

Autrement dit, l'émission stimulée n'aura lieu que si la fréquence du laser est exactement égale à celle de la transition.

Dans une situation réaliste, les atomes sont en mouvement sous l'effet de l'agitation thermique. Par effet Doppler, un atome animé d'une vitesse $\mathbf{v}$ voit l'onde électromagnétique avec la fréquence $\omega'_L = \omega_L - \mathbf{k} \cdot \mathbf{v}$, où $\mathbf{k}$ est le vecteur d'onde de l'onde. L'absorption n'aura plus lieu à la fréquence de résonance $\omega_L = \omega_a$, mais à $\omega'_L = \omega_a$, soit

$$\omega_L = \omega + \mathbf{k} \cdot \mathbf{v}. \quad (2.10)$$

Cette relation, purement ondulatoire, peut être retrouvée par un simple argument de conservation de l'énergie. En effet, avant l'absorption, le système atome+photon possède une énergie $E = \hbar\omega_L + mv^2/2$, où m désigne la masse de l'atome. Après absorption, l'atome dans l'état excité a absorbé l'impulsion $\hbar\mathbf{k}$ du photon et son énergie vaut donc $E' = m(\mathbf{v} + \hbar\mathbf{k}/m)^2/2 + \hbar\omega_a$. La conservation de l'énergie impose d'avoir $E = E'$, soit

$$\omega_L - \omega_a = \mathbf{k} \cdot \mathbf{v} + \frac{\hbar k^2}{2m}.$$

En pratique, la fréquence de recul $\hbar k^2/2m$ est très faible (typiquement 100 kHz) et peut donc être négligé devant le terme en $k\mathbf{v}$ de l'ordre du GHz à température ambiante, ce qui nous redonne bien la formule (2.10).

Supposons la vitesse des atomes répartie selon une distribution $f(\mathbf{v})$. Le nombre d^4N_γ de photons émis durant le temps dt par des atomes possédant une vitesse $\mathbf{v}$ à $d^3\mathbf{v}$ près est alors donné par

$$d^4N_\gamma = f(\mathbf{v})B_{\text{fe}}\delta(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})d^3\mathbf{v}dt.$$

Le nombre total de photons émis par unité de temps vaut donc

$$\frac{dN_\gamma}{dt} = \varpi \int f(\mathbf{v})B_{\text{fe}}\delta(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})d^3\mathbf{v}. \quad (2.11)$$

À l'équilibre thermique, la distribution de vitesse est donnée par la gaussienne de la distribution de Boltzmann, soit

$$f(\mathbf{v}) = N_e^0 \left(\frac{m\beta}{2\pi} \right)^{3/2} e^{-\beta m v^2/2},$$

où N_e^0 est le nombre total d'atomes dans l'état fondamental et $\beta = 1/k_B T$. Avec cette forme particulière, la relation (2.11) donne

$$\frac{dN_\gamma}{dt} = N_e^0 \varpi B_{\text{fe}} \left(\frac{m\beta}{2\pi} \right)^{1/2} e^{-\beta m (\omega_L - \omega_a)^2 / 2k^2}. \quad (2.12)$$

L'émission de la lumière par les atomes ne se fait donc plus uniquement à la fréquence de résonance $\omega_L = \omega_a$, mais sur toute une plage de largeur $\Delta\omega \sim k\sqrt{k_B T/m}$, 1 GHz dans le cas d'un laser hélium-néon à température ambiante. Le laser pouvant fonctionner lorsque le gain $G(\omega) = \dot{N}_\gamma$ est suffisamment grand, la relation (2.12) signifie donc que le laser est capable d'émettre sur tout la plage $\Delta\omega$.

En pratique, il existe un grand nombre de causes d'élargissement des raies atomiques, autres que l'effet Doppler. Citons ainsi les collisions entre atomes dans le cas d'un gaz ou encore l'interaction avec la matrice dans le cas de lasers solides. Dans une phase dense, les raies se transforment en bandes très larges, pouvant atteindre jusqu'à plusieurs dizaines de nanomètres dans un colorant ou même plusieurs centaines dans le cas des ions titane du titane-saphir.

Dans une cavité Fabry-Pérot typique, de quelques dizaines de cm de long, l'intervalle spectral libre, c'est-à-dire l'intervalle séparant de modes de la cavité vaut ~ 1 GHz, et la largeur de chaque pic est bien plus basse, moins d'un dixième de cette valeur. La courbe de gain est donc en général bien

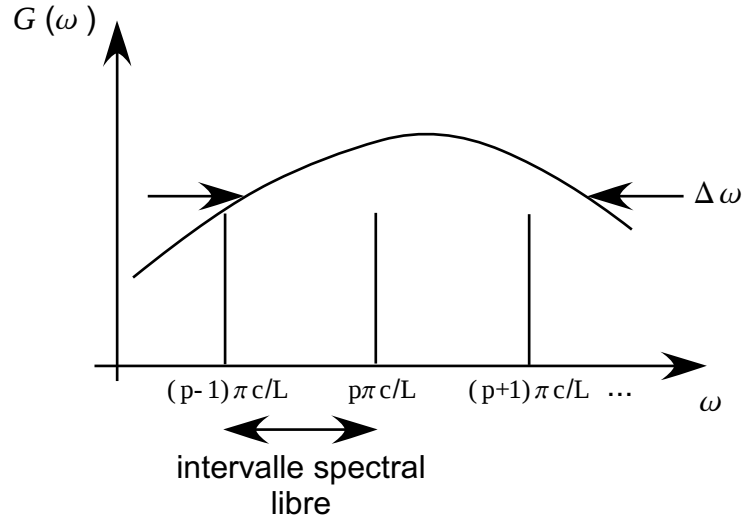


FIG. 2.11 – Comparaison de la courbe de gain d'émission stimulée et du peigne du Fabry-Pérot. La longueur d'onde d'émission du laser est fixée par les raies les plus étroites, et donc par le Fabry-Pérot.

plus large que la largeur d'un pic de la cavité. La sélection de mode est donc essentiellement le fait de la cavité qui impose des conditions plus restrictives quand à la fréquences des photons pouvant être utilisés dans le laser.

Compétition de modes

On dit que l'élargissement d'une raie est *inhomogène* si des atomes absorbant des photons à deux longueurs d'ondes différentes sont dans des états physiques différents. Dans le cas contraire l'élargissement est *homogène*. Par exemple, l'élargissement Doppler est inhomogène puisque l'absorption à deux longueurs d'ondes différentes correspondent à des vitesses différentes. L'élargissement dû à l'environnement est au contraire homogène, car dans une matrice, l'environnement est le même quel que soit l'endroit où est situé un ion.

Le caractère homogène ou inhomogène de l'élargissement va influencer sur le fonctionnement monomode ou multimode du laser. Si l'élargissement est homogène, les photons émis dans un mode de la cavité sont émis potentiellement par tous les atomes. Un atome ne pouvant émettre qu'un seul photon à la fois, les atomes ne pourront pas émettre dans un autre mode. Dans un

tel cas le laser fonctionne sur un seul mode, celui correspondant au gain le plus élevé.

Au contraire, dans le cas d'un élargissement inhomogène, la lumière émise dans chaque mode est issue d'atomes différents. Le laser peut donc fonctionner en multimode, 2 ou 3 modes dans le cas de l'hélium néon⁹.

Notons pour finir qu'un laser multimode peut être rendu monomode par ajout d'un filtre dans la cavité qui affine la courbe de gain de façon à la rendre plus étroite que l'intervalle spectral libre.

Que ce soit en fonctionnement monomode ou multimode, les raies d'émissions obtenues peuvent être très étroites. Elles sont en général limitées par les vibrations mécaniques et autres fluctuations extérieures. En travaillant avec soin et en utilisant diverses astuces que nous ne détaillerons pas ici, il est cependant possible de réduire la largeur de raies en dessous du kHz.

2.2.2 Quelques applications

Spectroscopie à haute résolution

Lorsque l'on étudie la structure des atomes en prenant en compte la relativité et les interactions magnétiques entre spins, on se rend compte que le spectre prédit par la théorie de Schrödinger doit être affiné. Ainsi, l'état fondamental des alcalins est dédoublé, avec un écart d'énergie de l'ordre d'1 GHz. Or nous avons vu que même en utilisant un laser infiniment fin, les raies d'absorptions atomiques étaient élargies par effet Doppler. Pour des atomes légers, cet élargissement devient plus grand que l'écart hyperfin et le dédoublement ne peut donc pas être résolu par simple spectroscopie d'absorption.

Afin de contourner cette difficulté, plusieurs stratégies ont été envisagées dont l'absorption saturée qui est l'objet de ce paragraphe. Nous allons dans un premier temps l'illustrer dans le cas d'une vapeur constituée d'atomes à deux niveaux dont les raies d'absorption sont élargies par effet Doppler. Le dispositif est le suivant : on envoie dans la cellule deux faisceaux contre-propageant. Ces deux faisceaux sont issus d'une même source laser mais l'un (la pompe) est supposée beaucoup plus puissante que l'autre (la sonde). On note $\varpi\phi(\omega - \omega_L)$ la densité spectrale de la pompe, où ϕ est une fonction d'intégrale unité, piquée autour de 0 et de largeur très faible notamment devant la largeur Doppler des raies atomiques. La densité spectrale de la

9. Ces modes peuvent être vus très facilement à l'aide d'un interféromètre de Michelson.

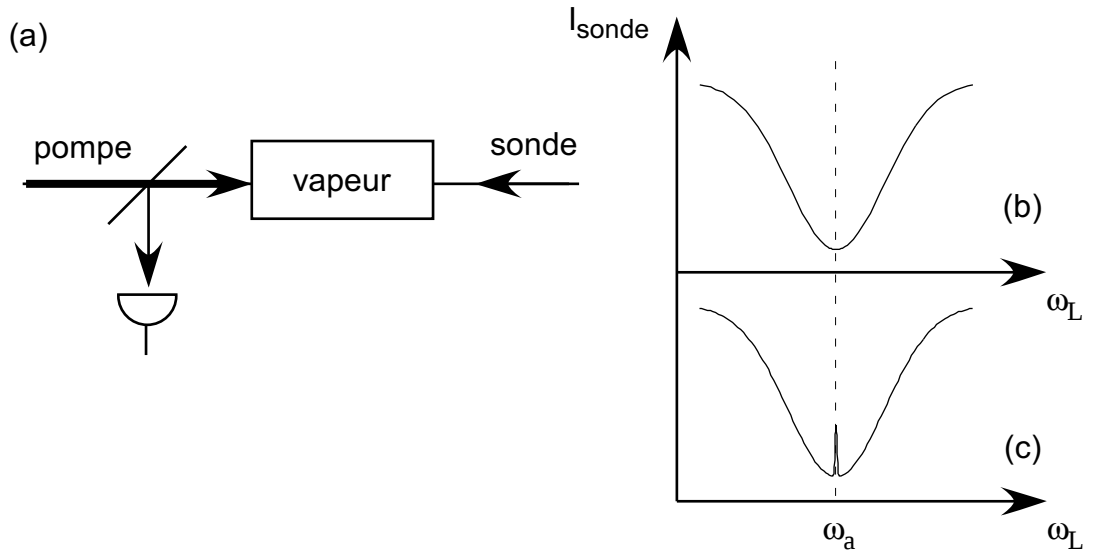


FIG. 2.12 – Principe de l'absorption saturée. a) On éclaire la vapeur d'atomes avec deux faisceaux de même longueur d'onde et contre propageants, le faisceau pompe étant plus puissant que la sonde et on mesure la puissance de la sonde en sortie de la cellule. b) Allure du signal du photodétecteur en l'absence de pompe. On retrouve la gaussienne caractéristique de l'élargissement Doppler. c) Allure du signal en présence de la pompe. Un pic étroit caractérisant précisément la position de la résonance apparaît au centre du signal d'absorption.

sonde est supposée de la forme $\epsilon \varpi \phi(\omega - \omega_L)$, avec $\epsilon \ll 1$. On s'intéresse ici à l'absorption en fonction de ω de la sonde par la vapeur atomique.

On commence par négliger l'influence de la sonde et on cherche à évaluer la population $d^3N_f(\mathbf{v})$ des atomes dans l'état fondamental et possédant une vitesse $\mathbf{v}$, à $d\mathbf{v}$ près. Pour simplifier les calculs, on pose $d^3N_f = x_f(\mathbf{v})f(\mathbf{v})d^3\mathbf{v}$ où $f(\mathbf{v})$ est la distribution de vitesse des atomes à l'équilibre thermodynamique. $x_f(\mathbf{v})$ désigne alors la proportion d'atomes de vitesse $\mathbf{v}$ dans l'état fondamental. De même, on définit $x_e(\mathbf{v})$, proportion d'atomes de vitesse $\mathbf{v}$ dans l'état excité.

De même que précédemment, les populations dans les états $|e\rangle$ et $|f\rangle$ vérifient l'équation de taux

$$\dot{x}_f = Ax_e + B\varpi\phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})(x_e - x_f).$$

Cette équation suppose que la vitesse des atomes change peu durant l'émission et l'absorption d'un photon. Cette hypothèse est justifiée dans la mesure où la vitesse de recul $v = \hbar k/m$ associée à ce type de processus vaut quelques centimètre par secondes pour les transitions optiques typiques et est petite devant la vitesse thermique des atomes (de l'ordre de la vitesse du son, soit quelques centaines de mètres par seconde).

En utilisant la relation $x_e + x_f = 1$, il vient ensuite

$$\dot{x}_f = Ax_e + B\varpi\phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})(1 - 2x_f).$$

Plaçons nous en régime stationnaire. On a alors $\dot{x}_f = 0$ et on obtient par conséquent

$$x_f = \frac{A + B\varpi\phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})}{A + 2B\varpi\phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})}.$$

On se place par ailleurs dans un régime de faible saturation, où l'émission spontanée l'emporte sur l'émission stimulée. On obtient donc à l'ordre 1 en B/A

$$x_f(\mathbf{v}) = 1 - \frac{B}{A}\varpi\phi(\omega_a - \omega_L - \mathbf{k} \cdot \mathbf{v}).$$

Si l'on prend $\mathbf{k}$ aligné selon z , on constate que $x_f(v_z)$ est pratiquement constamment égal à 1, sauf lorsque $v_z \sim (\omega_L - \omega_a)/k$, c'est-à-dire lorsque l'effet Doppler compense le désaccord entre le laser et la transition atomique. Dans ce cas, on a en effet un minimum de x_f , qui correspond au dépeuplement

de l'état fondamental par passage des atomes de l'état fondamental vers l'état excité.

Ajoutons à présent le faisceau sonde. Sa puissance étant faible, on peut supposer en première approximation qu'il perturbe peu les populations imposées par la pompe. Si l'on néglige comme précédemment l'émission stimulée, le taux d'absorption du faisceau sonde est proportionnel aux nombre d'atomes dans l'état fondamental. On a par conséquent

$$\left(\frac{dN_\gamma}{dt}\right)_{\text{sonde}} = - \int d^3N_f(\mathbf{v}) B \epsilon \varpi \phi(\omega_a - \omega_L - \mathbf{k} \cdot \mathbf{v}),$$

le signe de l'effet Doppler ayant changé de signe puisque le faisceau pompe et le faisceau sonde sont contre-propageants. Si l'on introduit la forme trouvée pour x_f , il vient

$$\left(\frac{dN_\gamma}{dt}\right)_{\text{sonde}} = -\epsilon B \varpi \int f(\mathbf{v}) \phi(\omega_a - \omega_L - \mathbf{k} \cdot \mathbf{v}) \left(1 - \frac{B}{A} \varpi \phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v})\right) d^3\mathbf{v}.$$

Cette intégrale est somme de deux termes. Le premier,

$$\left(\frac{dN_\gamma}{dt}\right)_1 = -\epsilon B \varpi \int f(\mathbf{v}) \phi(\omega_a - \omega_L - \mathbf{k} \cdot \mathbf{v}) d^3\mathbf{v},$$

correspond à l'absorption du faisceau sonde, s'il était seul. Puisque ϕ une fonction bien plus étroite que f , on peut en effet la remplacer par une fonction δ et on retrouve la courbe de gain gaussienne (2.11).

Le second terme s'écrit

$$\left(\frac{dN_\gamma}{dt}\right)_2 = \epsilon \frac{B^2}{A} \varpi \int f(\mathbf{v}) \phi(\omega_a - \omega_L - \mathbf{k} \cdot \mathbf{v}) \varpi \phi(\omega_a - \omega_L + \mathbf{k} \cdot \mathbf{v}) d^3\mathbf{v}.$$

Ce terme est toujours positif, mais négligeable pour presque toutes les valeurs de ω_L . En effet les fonctions $\phi(\omega_a - \omega_L \pm \mathbf{k} \cdot \mathbf{v})$ sont piqués autour des deux valeurs de vitesses opposées, $\pm(\omega_L - \omega_a)/k$. L'intégrand ne prend des valeurs non négligeables que lorsque $\omega_a - \omega_L = -(\omega_a - \omega_L)$, soit $\omega_L = \omega_a$. Si l'on trace l'absorption en fonction de $\omega_L - \omega_a$, on observe toujours le profil gaussien dû à l'effet Doppler, à ceci près que l'absorption est réduite au voisinage de la résonance exacte. Plus qualitativement, le phénomène d'absorption peut s'expliquer en notant que sous l'effet de la pompe, la population

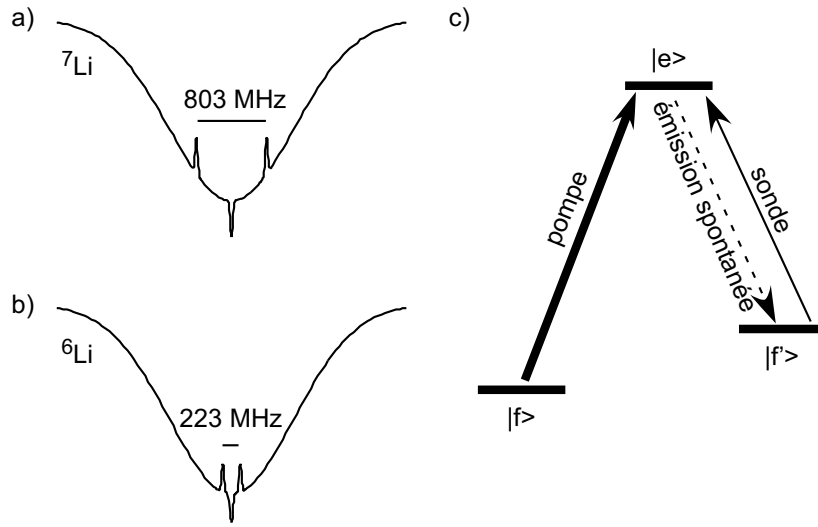


FIG. 2.13 – a) et b) Raies d'absorption saturée de la raie à 671 nm du ^7Li et du ^6Li . L'état fondamental du lithium est dédoublé par effet de l'interaction magnétique de l'électron de valence avec le noyau. Ces niveaux correspondent aux deux pics de réduction de l'absorption, séparés respectivement de 803 et 228 MHz. Le pic d'augmentation de l'absorption est un artefact dû au croisement de niveaux. c) Principe du croisement de niveaux. La pompe et la sonde sont résonantes sur deux transitions différentes. La pompe, par émission spontanée, augmente la population de l'état $|f'\rangle$, ce qui augmente l'absorption par la sonde.

de l'état fondamental n'est perturbée que pour des atomes de vitesse proches de $v_z = (\omega_L - \omega_a)/k$. La sonde, quant à elle, est absorbée par les atomes de vitesse $v_z = -(\omega_L - \omega_a)/k$. L'absorption de la sonde n'est donc perturbée par la pompe que lorsque $\omega_L = \omega_a$, puisque dans ce cas, ce sont les mêmes atomes qui absorbent à la fois la pompe et la sonde.

L'intérêt principal de cette technique est que la largeur du pic d'absorption saturée est donnée par la largeur de ϕ et donc par la résolution spectrale du laser. Ce qui permet de pouvoir localiser les raies atomiques à une précision limitée uniquement par la finesse spectrale de la source laser, comme on le constate sur la figure (2.13), où l'on distingue clairement la sous-structure de l'état fondamental du lithium.

Remarque complémentaire : raies de croisement de niveaux. Si l'état fondamental est constitué de plusieurs sous-niveaux $|f\rangle$ et $|f'\rangle$, des raies anormales peuvent apparaître sur le spectre d'absorption saturée (Fig. 2.13). Ces raies se produisent lorsque la pompe et la sonde sont sur chacune des transitions $|f\rangle \rightarrow |e\rangle$ et $|f'\rangle \rightarrow |e\rangle$. En effet, si l'on suppose que la pompe est résonnante avec la transition $|f\rangle \rightarrow |e\rangle$, l'émission spontanée de photons de $|e\rangle$ vers $|f'\rangle$ transfère des photons vers l'état fondamental $|f'\rangle$. Puisqu'il y a plus d'atomes dans cet état la sonde est plus absorbée lorsqu'elle est résonnante sur la transition $|f'\rangle \rightarrow |e\rangle$. Cette condition de double résonance est réalisée lorsque la pompe et la sonde sont résonnantes pour la même classe de vitesse, c'est-à-dire lorsque

$$\omega_L - \omega_a - \mathbf{k} \cdot \mathbf{v} = \omega_L - \omega'_a + \mathbf{k} \cdot \mathbf{v} = 0,$$

où ω'_a est la fréquence de transition $|f'\rangle \rightarrow |e\rangle$. On trouve alors en additionnant les deux premiers termes que cette condition est réalisée pour $\omega_L = (\omega_a + \omega'_a)/2$.

Exercice : Décrire le profil d'absorption saturée lorsque l'état fondamental est non dégénéré mais que l'état excité possède deux sous niveaux.

Effet Sagnac et Gyrolaser

L'Effet Sagnac est l'équivalent non galiléen de l'effet Doppler, puisqu'il décrit le décalage fréquentiel d'une onde dans un référentiel en rotation. Bien que la théorie rigoureuse de cet effet nécessite l'utilisation de la Relativité Générale et sort par conséquent du cadre de ce cours, il est possible d'en faire démonstration classique, donnant un résultat exact dans la limite des faibles vitesses de rotation. Ceci provient du fait que l'effet Sagnac, tout comme l'effet Doppler, est proportionnel à v/c et que l'on peut montrer que les phénomènes purement relativistes n'apparaissent qu'à l'ordre v^2/c^2 . Dans la suite, nous nous placerons donc dans un cadre classique, où l'on suppose notamment que la vitesse de la lumière est modifiée par un changement de référentiel.

Considérons pour simplifier une fibre optique d'indice $n = 1$ que l'on enroule sur un cercle de rayon R . La fibre est en rotation à une vitesse Ω et on cherche à calculer le déphasage subi par des ondes électromagnétiques se propageant dans le même sens et en sens opposé à la rotation. La réponse est relativement simple si l'on suppose que la vitesse de la lumière suit la relation de changement de référentiel classique. En effet, la vitesse de chacune

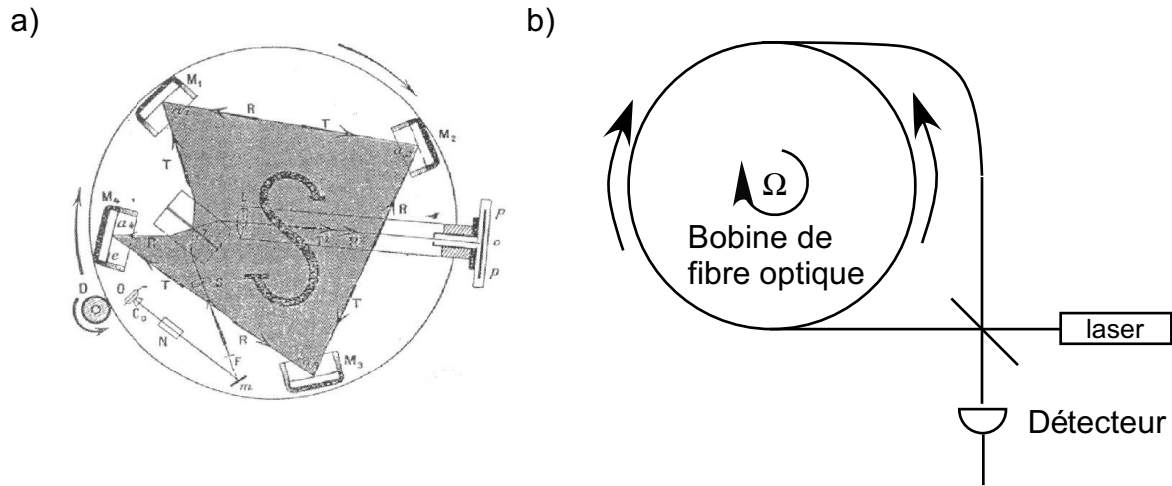


FIG. 2.14 – a) Schéma de principe de l'expérience de Sagnac présentée à l'Académie des Sciences le 22 décembre 1913. b) Dispositif plus moderne de gyrofibre. La longueur parcourue par la lumière est démultipliée par l'enroulement de la fibre sur elle même.

des ondes dans le référentiel tournant est $v_{\pm} = c \pm \Omega R$ et tout se passe donc comme si la lumière se propageait dans un milieu d'indice

$$n_{\pm} = c/v_{\pm} = (1 \pm \Omega R/c).$$

Le trajet optique $L_{\pm}$ sur un tour complet pour chacune des ondes vaut par définition

$$L_{\pm} = \oint n_{\pm} ds = \frac{2\pi R}{1 \pm \Omega R/c} \sim 2\pi R(1 \mp \Omega R/c).$$

où l'on s'est placé dans le domaine de l'approximation non relativiste de faible vitesse $\Omega R/c \ll 1$. Autrement dit, on obtient un écart de chemin optique entre les ondes se propageant dans chaque direction valant

$$\delta L = L_{-} - L_{+} = \frac{4\pi R^2 \Omega}{c}.$$

Cet écart de chemin optique est appelé effet Sagnac. Si l'on introduit le vecteur surface $\mathbf{S}$ associé au cercle sur lequel est enroulé la fibre optique, cette relation peut s'écrire

$$\delta L = \frac{4\mathbf{S} \cdot \boldsymbol{\Omega}}{c}, \quad (2.13)$$

relation qui est généralisable à des trajets de n'importe quelle forme.

Ce déphasage induit par la rotation a été démontré expérimentalement pour la première fois par Georges Sagnac en 1913, lors d'une publication à l'Académie des sciences. Le dispositif plus moderne présenté ici et utilisant une fibre optique est appelée gyro-fibre. Le rôle de la fibre optique est de multiplier le trajet parcouru autour de l'interféromètre en enroulant la fibre plusieurs fois sur elle même, le trajet parcouru par la lumière pouvant ainsi atteindre plusieurs km !

L'équation (2.13) suggère une méthode particulièrement élégante de mesure de vitesse de rotation (Fig. 2.15). En effet, considérons un laser dont la cavité possède une géométrie en anneaux. Si ce laser est mis en rotation, on constate que les conditions de résonance (chemin optique égal à un nombre entier de fois la longueur d'onde) sont différentes pour les modes co- et contre propagatifs. Les deux modes n'oscillent donc pas à la même fréquence et si on les mélange sur une photodiode, le signal obtenu est de la forme

$$|E_0 e^{-i\omega_+ t} + E_0 e^{-i\omega_- t}|^2 = 4E_0^2 \cos((\omega_+ - \omega_-)t/2).$$

On obtient donc un battement à une fréquence proportionnelle à $\omega_+ - \omega_-$, c'est-à-dire proportionnelle à Ω .

Les gyrolasers (et des gyrofibres qui en sont une alternative) trouvent leur application principale dans les "centrales de navigation inertielle" où combinés à des accéléromètres ils permettent de repérer la position du système sur lequel ils sont embarqués, sans aucune référence extérieure (telle que les étoiles ou les satellites GPS).

Dans ce type d'application, le gyrolaser est venu progressivement remplacer les gyroscopes mécaniques dont l'usure inhérente posaient des problèmes de vieillissement. Gyro-lasers et gyrofibres sont aujourd'hui amplement utilisés dans des applications hautement stratégiques telles que la navigation de sous-marins ou le guidage de missiles.

Holographie

Le principe de Huygens stipule que pour connaître le champ électromagnétique en tout point de l'espace, il suffit de le connaître sur un plan (ou une surface), le champ en un point extérieur se déduisant par la propagation d'ondes

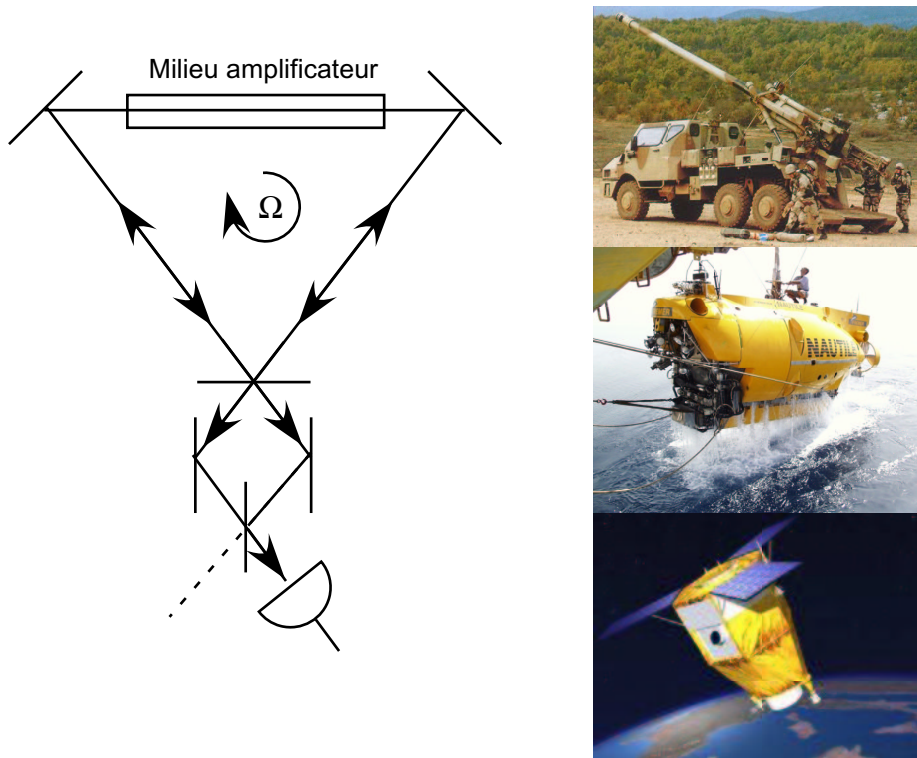


FIG. 2.15 – À droite, principe du gyrolaser : la rotation du laser est mesurée par le battement des ondes se propageant dans le sens positif et négatif d'une cavité en anneau. À gauche, quelques utilisations des gyrolasers : navigation et pointé du Canon Caesar 155mm de Giat ; Repérage du sous-marin Nautilus ; pointage des satellites Pléiades, successeurs des satellites d'observation SPOT.

sphériques secondaires. Dans le principe, ceci permet donc de stocker sur une surface bidimensionnelle toute l'information nécessaire à la reconstruction de l'image d'un objet tridimensionnel. La photographie classique ne stocke que l'information sur l'intensité lumineuse et donc le module du champ électrique, ce qui explique qu'elle ne puisse pas rendre la profondeur d'une image réelle. Au contraire, l'holographie est un processus interférométrique permettant de reproduire à la fois le module et la phase d'un champ électrique, ce qui permet ainsi de créer des images tridimensionnelles.

Au delà de son simple côté ludique, les techniques holographiques ont d'importantes applications pratiques dont nous en développerons certaines à la fin de cette section. Les techniques interférentielles permettent par exemple de mesurer des déplacements égaux à une fraction de longueur d'onde (soit typiquement des fractions de micron) ce qui permet d'utiliser l'holographie dans l'étude des propriétés mécaniques des matériaux.

En pratique, on enregistre l'hologramme d'un objet tridimensionnel en divisant en deux le faisceau issu d'un laser. Un des faisceaux est utilisé comme référence alors que le second est utilisé pour éclairer l'objet dont on veut reproduire l'image. Les faisceaux de références et diffusés sont combinés sur un film photographique qui enregistre leur interférogramme (Fig. 2.16.A). Afin de reconstruire l'image de l'objet, on éclaire le film à l'aide d'une lumière cohérente : nous allons alors montrer que la lumière diffractée par le film est identique (dans une certaine limite que nous préciserons) à celle issue de l'objet initial (Fig. 2.16.B).

Pour comprendre pourquoi ces deux étapes permettent bien de reconstruire l'image tridimensionnelle de l'objet initial, on note $E_{\text{dif}}(x,y)$ et $E_{\text{ref}}(x,y)$, les champs électriques des ondes diffusées et de référence dans le plan (x,y) du film photographique. On pose alors

$$\begin{aligned} E_{\text{dif}} &= E_{\text{dif},0} \cos(\omega t - \psi(x,y)) \\ E_{\text{ref}} &= E_{\text{ref},0}(x,y) \cos(\omega t - \phi(x,y)), \end{aligned}$$

les pulsations des deux ondes étant identiques car les faisceaux sont issus de la même source. Notons que l'onde de référence est une onde plane et que son amplitude est prise indépendante de x et y . Si $\mathbf{k}$ est son vecteur d'onde $\phi(x,y) = k_x x + k_y y$. La phase ϕ est donc reliée à l'inclinaison de l'onde de référence par rapport au film.

Le film photographique est sensible à l'intensité lumineuse totale avec

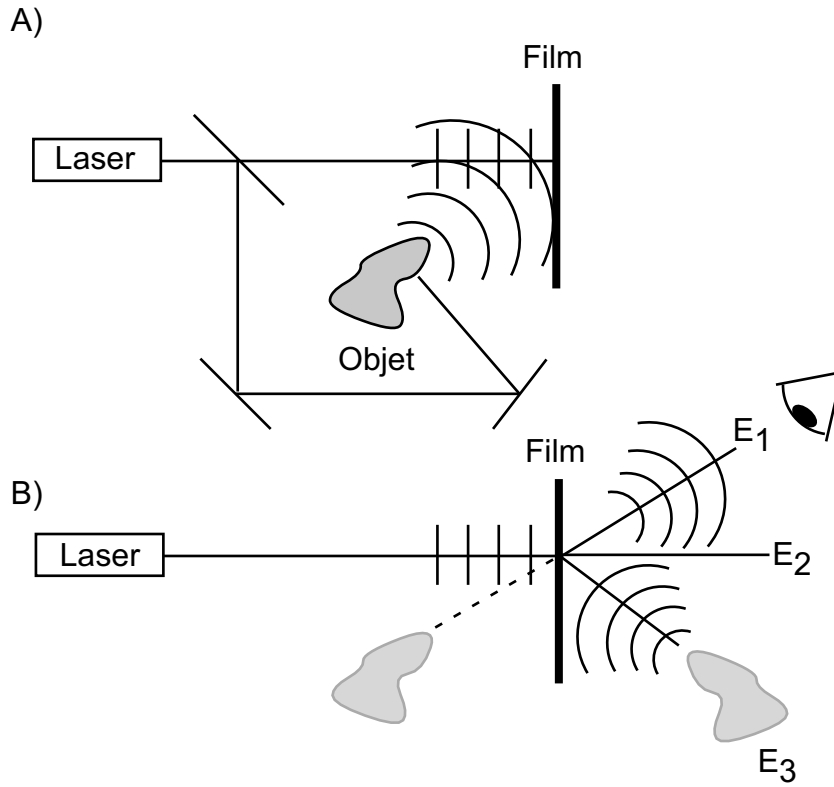


FIG. 2.16 – Principe de l'holographie. A) Enregistrement : on sépare en deux le faisceau issu d'une source laser et un des faisceaux est diffusé par l'objet dont on souhaite obtenir l'hologramme. Les deux faisceaux sont ensuite recombinaés sur un film photographique qui stocke leur figure d'interférence. B) Lecture. On éclaire le film avec le faisceau de référence seul. La figure de diffraction due à la traversée du film fait apparaître trois rayons distincts : E_1 reconstruit l'image virtuelle de l'objet initial. C'est cette image qui constitue l'hologramme proprement dit. E_2 est le rayon transmis directement par le film. E_3 enfin reconstruit une image réelle de l'objet initial.



FIG. 2.17 – Train and Bird, le premier hologramme de l'histoire, réalisé en 1964 par Emmeth Leith et Juris Upatnieks à l'université du Michigan.

laquelle il est éclairé. Après impression, sa transmission T vaut en un point (x,y)

$$T(x,y) = (1 - 2\beta I_0(x,y)),$$

où β est un paramètre caractérisant la sensibilité du film et I_0 est l'intensité lumineuse avec laquelle le film est éclairé au point (x,y) . Cette intensité vaut ici

$$I_0(x,y) = \langle |E_{\text{dif}} + E_{\text{ref}}|^2 \rangle = \frac{1}{2} (E_{\text{dif},0}^2 + E_{\text{ref},0}^2 + 2E_{\text{ref},0}E_{\text{dif},0} \cos(\phi - \psi)).$$

Passons à la phase de reconstruction de l'image. On éclaire le film avec l'onde référence seule. Le champ transmis E_{rec} vaut TE_{ref} , soit

$$\begin{aligned} E_{\text{rec}}(x,y) &= E_{\text{ref}}(1 - \beta(E_{\text{dif},0}^2 + E_{\text{ref},0}^2)) \cos(\omega t - \phi) \\ &\quad - \beta E_{\text{ref}}^2 E_{\text{dif}}(x,y) (\cos(\omega t - \psi) + \cos(\omega t + \psi - 2\phi)). \end{aligned}$$

Le champ électrique se met donc sous la forme d'une somme de trois termes $E_{i=1,2,3}$ respectivement égaux à

$$\begin{aligned} E_1 &= -\beta E_{\text{ref}}^2 E_{\text{dif}}(x,y) \cos(\omega t - \psi) \\ E_2 &= E_{\text{ref}}(1 - \beta(E_{\text{dif},0}^2 + E_{\text{ref},0}^2)) \cos(\omega t - \phi) \\ E_3 &= -\beta E_{\text{ref}}^2 E_{\text{dif}}(x,y) \cos(\omega t + \psi - 2\phi) \end{aligned}$$

et que nous allons chercher à interpréter

- Le champ E_1 est celui qui produit l'hologramme, puisqu'à un facteur près, il est identique au champ électrique diffusé par l'objet au niveau du film.
- Le champ E_2 possède la même phase ϕ que le champ de référence E_{ref} . Comme nous l'avons vu, cette phase est reliée aux coordonnées du vecteur d'onde et indique donc la direction de propagation de l'onde. Le champ E_2 donne donc lieu à une onde se propageant dans la même direction que l'onde de référence, avec un profil d'amplitude spatiale modifié par la traversée du film.

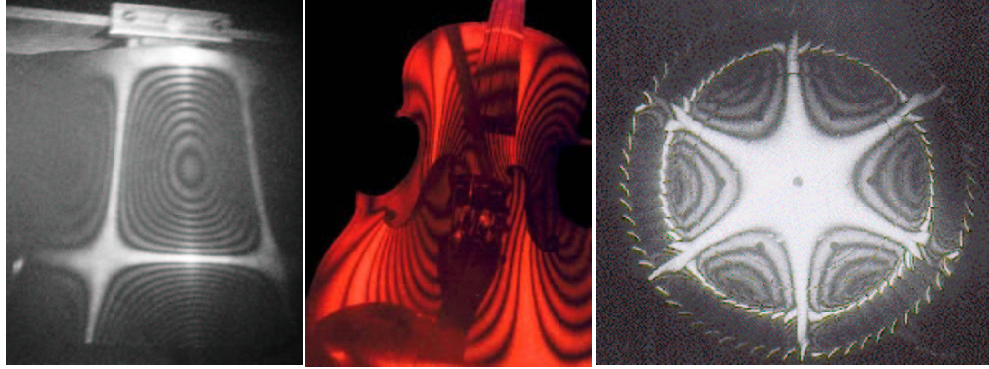


FIG. 2.18 – Structures spatiales des modes de vibrations d'une cloche, d'un violon et d'un rotor de turbine, obtenues par holographie.

- La structure de l'onde engendrée par le champ E_3 est un peu plus complexe que dans les deux premiers cas. Formellement, la phase $\cos(\omega t + \psi - 2\phi)$ est fonction de ψ et contient donc effectivement l'information de phase nécessaire à la reconstruction de l'objet, mais “dans le désordre”. En effet, considérons pour commencer le rôle de la phase ϕ . Comme nous l'avons déjà vu, cette phase est associée à une onde se propageant selon la direction du faisceau de référence. Le champ E_3 va donc créer un faisceau incliné par rapport à l'image reconstruite par E_1 . Par ailleurs, si l'on omet ϕ , la phase n'est toujours pas la bonne, puisque le cosinus s'écrit $\cos(\omega t + \psi)$, au lieu de $\cos(\omega t - \psi)$. On peut cependant se ramener à cette deuxième forme en notant que

$$\cos(\omega t + \psi) = \cos(-\omega t - \psi) = \cos(\omega(-t) - \psi).$$

Autrement dit, E_3 possède la même phase que l'onde diffusée, à condition de changer t en $-t$, autrement dit de renverser le sens du temps. En terme moins “ésotérique”, ceci signifie que le champ E_3 va générer une image réelle de l'objet diffusant.

Applications: La principale application de l'holographie est la détection de petits mouvements, permettant ainsi l'étude de la réponse de matériaux à des vibrations acoustiques: l'atout de l'holographie provient de sa nature interférométrique, qui lui permet de détecter des mouvements de l'ordre d'une fraction de longueur d'onde, soit moins d'une centaine de nanomètres! Une technique utilisée est de procéder par moyenne temporelle: on enregistre

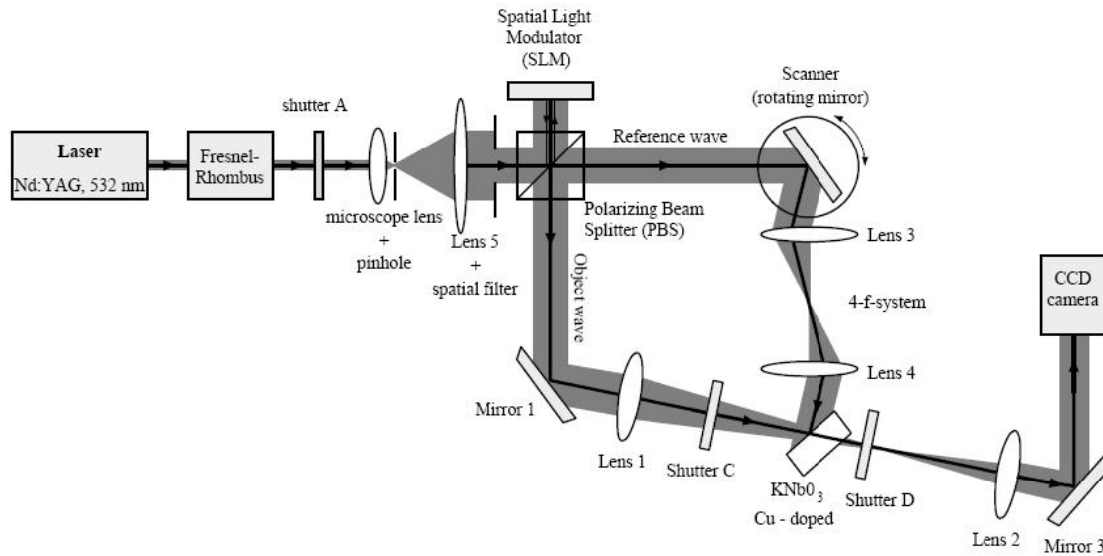


FIG. 2.19 – *Exemple de réalisation expérimentale d'une mémoire holographique. La série des métamorphoses de Escher est affichée sur l'écran à cristaux liquides (SLM). À chaque image, l'angle du miroir tournant miroir est modifié. La série des images est ensuite reconstruite en n'éclairant le cube qu'à l'aide de la référence et en tournant le cube pour faire apparaître les images l'une après l'autre sur la caméra (<http://www.crypto.ethz.ch/przydatek/hms.html>). On trouvera la série des images initiales et reconstruites dans les deux pages qui suivent.*

l'hologramme alors que l'objet est en train de vibrer. Au points où l'objet est immobile l'hologramme n'est pas perturbé et l'objet initial est reconstruit fidèlement. En revanche, aux points où l'objet vibre, les interférences entre la référence et la lumière diffusée se brouillent et on obtient des franges sombres lors de la reconstruction de l'hologramme (Fig. 2.18). Une deuxième technique, plus adaptée à une étude en temps réel, consiste à superposer l'objet vibrant à son hologramme figé. De nouveau, le mouvement de l'objet se manifeste par l'apparition de franges d'interférences lorsque l'on observe à travers le film.

Une seconde application, plus futuriste, est la réalisation de mémoires holographiques qui, à terme, pourraient remplacer les CD ou les DVD. Les cubes holographiques actuellement en développement permettent en effet de

stocker de l'information en volume, et non plus en surface, comme cela est fait dans les supports usuels de l'information. Ce nouveau type de mémoire devrait permettre de stocker jusqu'à 1000 Go dans un cube de la taille d'un sucre en cube, ces données pouvant être récupérées à très grande vitesse (un DVD lu en 30 s)! Dans ce dispositif, l'information à stocker est tout d'abord transformée en une série de pixels noirs et blancs (des 0 et des 1) d'un écran à cristaux liquides. Comme dans l'holographie classique, on effectue ensuite l'interférence d'un faisceau de référence avec la lumière diffusée par l'écran à cristaux liquide, à ceci près que l'on réalise cette interférence non sur un film photographique, mais au sein d'un cube photoréfractif (KNbO_3 par exemple), dont l'indice est modifié de façon permanente par la lumière. Comme précédemment, la lecture de l'hologramme se fait en illuminant le cube à l'aide de la référence seule. Le point central est ici que l'angle d'incidence du faisceau de référence doit être strictement le même dans le cas de la lecture et de l'écriture. En effet, on peut imaginer le cube comme une superposition de films photographiques, créant chacun un hologramme. Si l'angle d'incidence du faisceau de reconstruction n'est pas le bon, ces hologrammes vont être créés dans des directions différentes et interférer destructivement¹⁰. Plusieurs hologrammes différents peuvent par conséquent être stockés, en choisissant pour chacun un angle d'incidence différent pour le faisceau de référence.

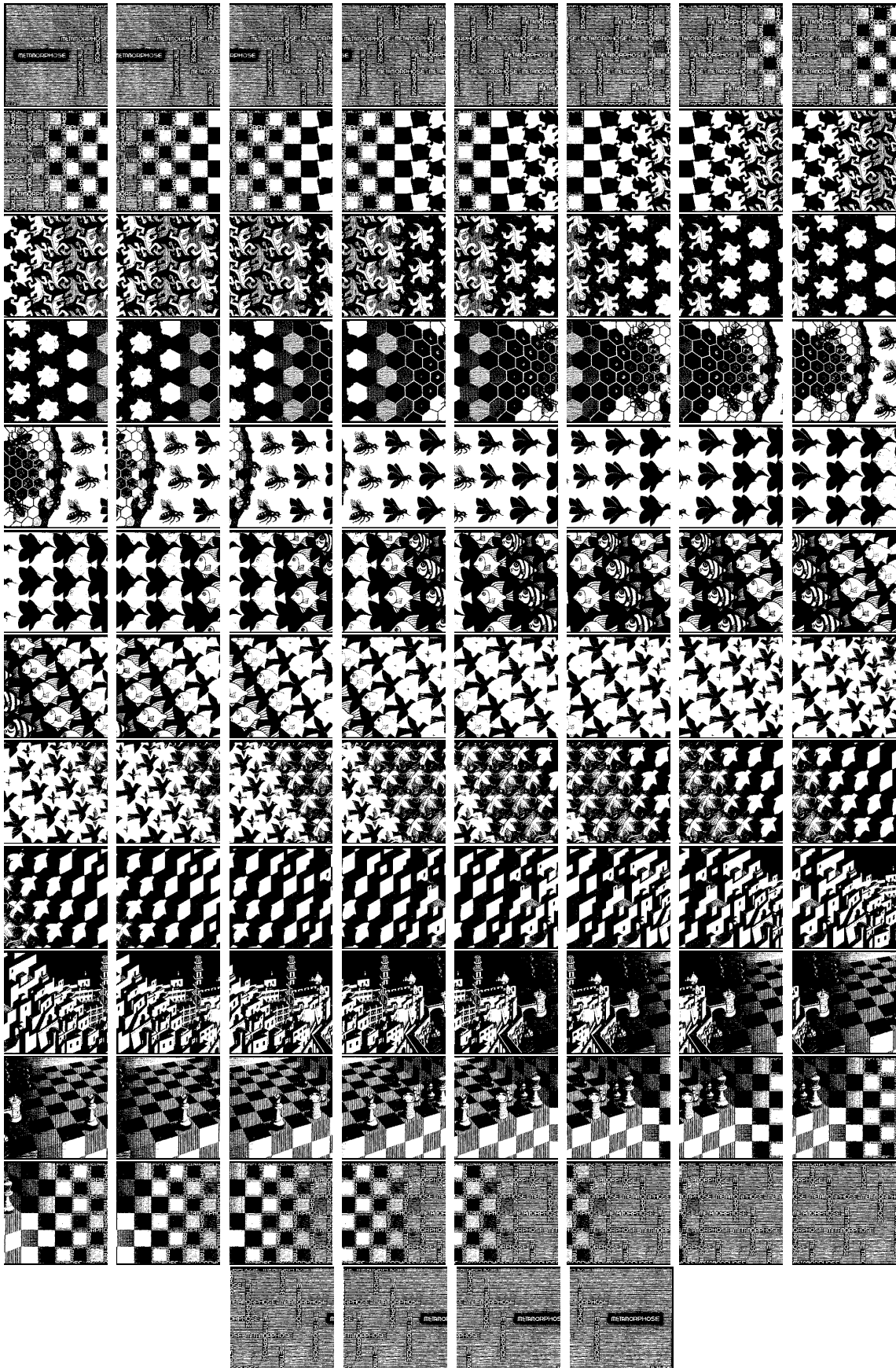
2.3 Lasers à impulsion

L'étude que nous venons de faire se limite au cas des lasers continus. D'importantes applications aussi bien pratique que fondamentales reposent cependant sur un fonctionnement *pulsé*, dans lequel le laser délivre l'énergie électromagnétique sous forme d'impulsions pouvant atteindre une fraction de femtoseconde ($1 \text{ fs} = 10^{-15} \text{ s}$).

Parmi les exemples d'utilisation des lasers pulsés on peut citer :

- *L'usinage laser*. L'énergie électromagnétique est en effet déposée tellement vite dans le matériau à usiner que la chaleur n'a pas le temps de diffuser, ce qui permet de réaliser une découpe très fine, particulièrement adaptée au micro-usinage (2.20).

10. Notons au passage que c'est un procédé similaire qui est utilisé pour les hologrammes en "lumière blanche" déposés sur les cartes de crédit.

Figure 4.5: The original images (*Metamorphosis II* by M. C. Escher, scanned from [5])

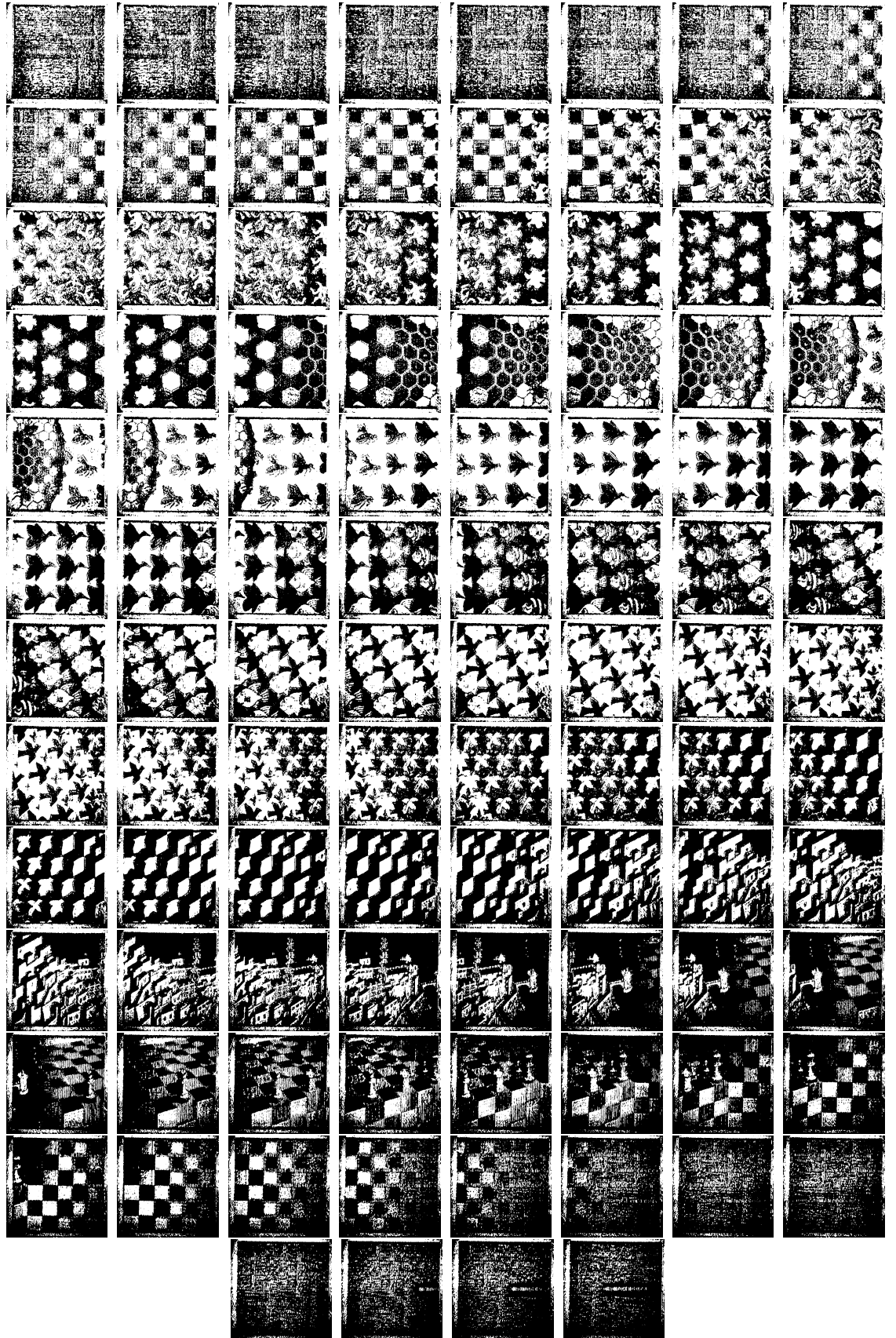


Figure 4.6: The read-out images

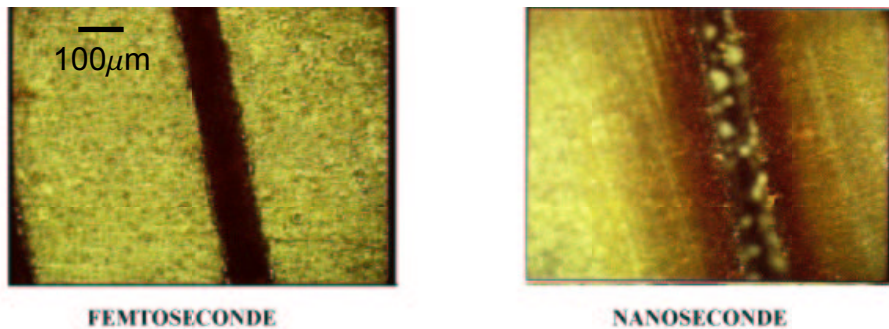


FIG. 2.20 – Usinage sur fer à l'aide d'impulsion femtosecondes ou nanosecondes. Dans le cas d'impulsion ultra brève, l'ablation de matière se fait sans diffusion thermique, rendant ainsi plus propre le travail du matériau.

- *La fusion inertielle.* Le laser Mégajoule (ou son concurrent américain le NIF - National Ignition Facility) actuellement en construction devrait pouvoir délivrer des impulsions d'un mégajoule et d'une durée de l'ordre de la nanoseconde, soit une puissance crête de quelques 10^{15} W ! Son but est essentiellement militaire, puisqu'il vise à étudier les mécanismes de fusion deux noyaux d'hydrogène, en vue de développer de nouvelles armes nucléaires. Les retombées pour les sciences fondamentales ne sont cependant pas négligeables puisque ce laser permet d'atteindre des conditions de température et de pression analogues à celles du cœur des étoiles.
- *Femtochimie.* Les lasers femtosecondes ont ouvert la porte à un nouveau type de spectroscopie optique résolue en temps et non plus en fréquence uniquement. Les applications sont particulièrement fascinantes en chimie puisque l'on peut à présent observer et orienter certaines réactions chimiques simples pratiquement en temps réel (cf. les travaux d'Ahmed Zewail, prix Nobel de chimie 1999) !

On peut distinguer deux méthodes permettant de réaliser un fonctionnement pulsé. La première consiste à fonctionner en mode déclenché : le démarrage de l'impulsion est provoqué par une brusque modification d'un des paramètres de fonctionnement du laser, comme le taux de pompage ou le facteur de qualité de la cavité. La deuxième méthode est le blocage de mode, qui permet de placer le laser dans une configuration où il émet spontanément des trains d'impulsions émises à intervalles réguliers. Cette seconde technique

est permise par l'utilisation d'un milieu amplificateur présentant une large bande passante dont on verrouille la phase relative des différents modes.

2.3.1 Impulsions déclenchées

Pour faire fonctionner un laser en mode déclenché, on peut essentiellement faire varier deux paramètres, le gain ou les pertes.

Pour faire varier le gain, on module le taux de pompage. Ainsi, dans le cas d'une diode laser on module le courant électrique amenant les porteurs à la jonction. Dans le cas d'un laser pompé optiquement, on réalise un pompage par lampe flash – c'est cette méthode qui est utilisée par exemple dans le laser Mégajoule.

Par analogie avec les oscillateurs électriques ou mécanique, on peut définir le facteur de qualité Q d'une cavité. Plus le pics du Fabry-Pérot sont étroits et plus ce facteur de qualité est grand. En jouant sur les pertes de la cavité, on modifie le facteur de qualité, d'où la dénomination de "Q-switch" donné à ce type de fonctionnement. Le point clef du fonctionnement en Q-Switch, dont une implémentation est proposée Fig. 2.21.b), est l'amplification de l'inversion de population Δ lorsque le gain du laser est en dessous du seuil d'amorçage (cf. paragraphe (2.1.3)). Dans une première phase, on maintient un niveau de pertes suffisamment grand dans la cavité de façon à empêcher le démarrage du laser et ainsi obtenir une forte inversion de population Δ . Lorsque le régime stationnaire est atteint, on diminue brutalement les pertes de façon à repasser au dessus du seuil de démarrage. Le taux d'émission stimulée de photons, proportionnel à Δ , est alors bien supérieur à sa valeur stationnaire. On observe alors une brusque impulsion lumineuse en sortie du laser, qui s'atténue ensuite progressivement au fur et à mesure que l'inversion de population rejoint sa valeur stationnaire.

2.3.2 Blocage de mode

L'autre méthode couramment utilisée pour réaliser un laser pulsé est le mode bloqué. Pour comprendre ce mécanisme, considérons un lase pouvant émettre sur N modes de la cavité Fabry-Pérot, ces modes émettant tous avec une même phase et une même amplitude. Le champ électrique de l'onde émise par le laser est alors

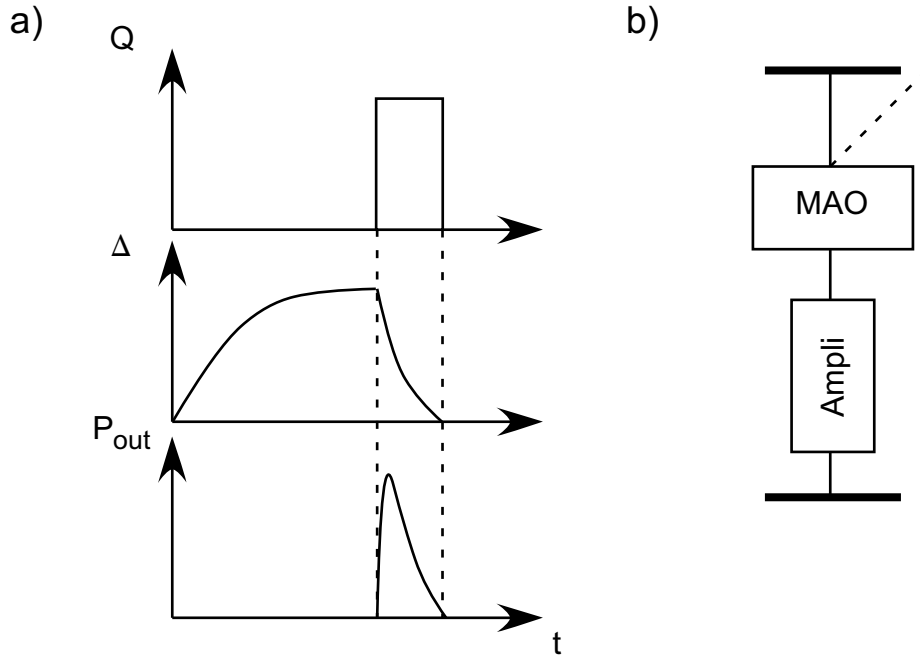


FIG. 2.21 – a) Évolution du facteur de qualité, de l'inversion de population et de la puissance de sortie d'un laser en Q-Switch. On maintient un faible facteur de qualité dans la cavité de façon à empêcher l'émission laser et à favoriser une importante inversion de population dans le milieu amplificateur. On déclenche l'émission laser en augmentant le facteur de qualité de la cavité au delà du seuil ce qui provoque une impulsion laser puissante. b) Exemple de réalisation d'un dispositif à impulsion déclenchée. On introduit dans la cavité un modulateur acousto-optique (MAO). Il s'agit d'un cristal transparent dans lequel on peut faire propager une onde acoustique. Si cette onde est présente, elle crée une modulation périodique de densité et donc d'indice optique qui diffracte la lumière et accroît donc les pertes - ligne pointillée. L'impulsion est alors déclenchée en supprimant l'onde acoustique.

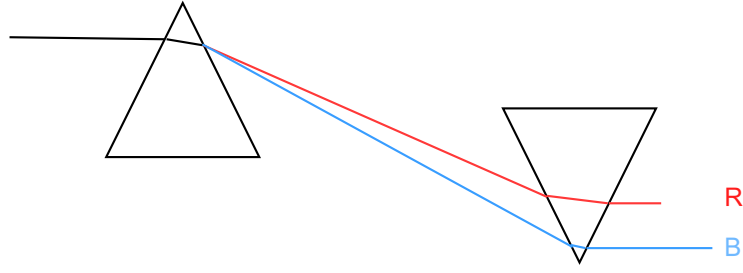


FIG. 2.22 – Dispositif de compensation de la dispersion dans un laser à impulsion. Considérons un faisceau “bicolore” dont le spectre possède essentiellement deux composantes dans le rouge et le bleu. L’indice du bleu étant dans la matière plus grand que le rouge, la composante bleue du spectre prend du retard durant la propagation dans la cavité. Ce retard est compensé en faisant passer le faisceau dans un premier prisme qui sépare spatialement les deux composantes du spectre (faisceaux R et B). Ceci permet de leur faire traverser le second prisme à différents endroits. On s’arrange alors pour que le bleu traverse une plus faible épaisseur de verre que le rouge, ce qui lui permet de rattraper son retard.

$$E(t) = \sum_{n=n_0}^{n_0+N-1} E_0 e^{in\pi ct/L},$$

soit en resommant cette série géométrique

$$E(t) = E_0 e^{i2n_0\pi ct/2L} \frac{1 - e^{iN\pi ct/L}}{1 - e^{i\pi ct/L}}.$$

L’intensité lumineuse $I \propto |E|^2$ varie dans le temps et vaut

$$I(t) \propto \frac{\sin^2(N\pi ct/2L)}{\sin^2(\pi ct/2L)}.$$

Le graphe de $I(t)$ (analogue à la fonction de transfert du Fabry-Pérot) est constitué d’une série périodique de pics de largeur $\delta t \sim 2L/Nc$. L’intervalle de temps séparant deux pics est $T = 2L/c$, c’est-à-dire le temps mis par la lumière pour faire un aller-retour dans la cavité.

Bien que simplifiée à l’extrême, l’analyse que nous venons de réaliser permet de bien comprendre les conditions essentielles à remplir pour obtenir

un régime d'impulsions courtes.

- Utiliser un milieu amplificateur large bande de façon à obtenir des impulsions les plus courtes possibles (la durée des impulsions est inversement proportionnelle aux nombres de modes émis).
- Minimiser la dispersion dans la cavité. En effet, si le milieu amplificateur est dispersif, avec un indice $n(\nu)$ dépendant de la fréquence ν la condition de résonance dans la cavité s'écrit $2n(\nu)L = p\lambda$, soit pour la fréquence de l'onde

$$\nu = p2L/cn(\nu).$$

Autrement dit, les fréquences ne sont plus dans des rapports entiers. Pour pallier à ce défaut, on utilise par exemple des couples de prismes permettant de compenser la dispersion. Le premier prisme sépare spatialement les différentes longueurs d'ondes, ce qui permet de leur faire parcourir des chemins différents dans le second prisme, et ainsi de compenser les effets dispersifs (Fig. 2.22). Ce dispositif permet par ailleurs d'accorder le laser sur une longueur d'onde bien définie en plaçant une fente en sortie du second prisme de façon à ne laisser passer qu'une partie du spectre.

- Enfin, il faut favoriser le fonctionnement pulsé par rapport au fonctionnement continu, toujours possible. Pour cela, on réalise une cavité dans laquelle les pertes sont plus importantes pour des ondes basses puissance que haute puissance. Le laser régime impulsionnel possédant une puissance crête élevée sera par conséquent le seul à pouvoir démarrer. Les deux méthodes les plus couramment utilisées consistent soit à introduire un absorbant saturable dans la cavité, soit à jouer sur l'effet d'autofocalisation des faisceaux intenses¹¹. Dans ce second cas, le faisceau en régime continu est moins intense, donc moins focalisé. Son diamètre est donc plus grand que celui du mode pulsé et on le bloque en ajoutant dans la cavité un diaphragme ne laissant passer que le faisceau pulsé.

Ces concepts sont illustrés dans le cas du laser titane-saphir (Fig. 2.23). Ce laser commercial permet de délivrer des impulsions d'une durée égale à une

11. Cet effet sera étudié au chapitre suivant et résulte de la dépendance de l'indice avec l'intensité lumineuse.

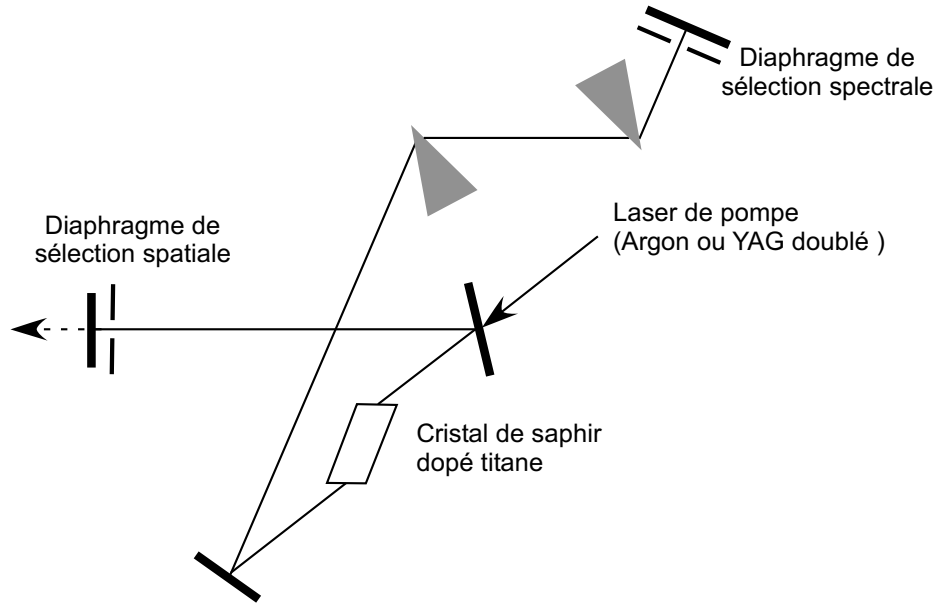


FIG. 2.23 – Schéma de principe d'un laser saphir dopé titane. On reconnaît les deux prismes de compensation de la dispersion, ainsi que les deux diaphragmes de sélection du mode spectral et spatial.

dizaine de femtoseconde et portant une énergie de 10 nJ. Les durées d'impulsions les plus courtes obtenues à ce jour atteignent la centaine d'attoseconde (une fraction de femtoseconde). Celles-ci sont obtenues par génération d'harmoniques d'ordre élevé en éclairant des matériaux non-linéaires à l'aide de laser de puissance. Si la fréquence de l'onde incidente est ν_0 on obtient les harmoniques $n\nu_0$, avec n grand. Si ν_0 est une fréquence optique, on obtient une impulsion sub-femtoseconde centrée dans le domaine X.

2.3.3 Applications à la spectroscopie résolue en temps

Parmi les nombreuses possibilités ouvertes par la possibilité de générer des impulsions courtes, nous allons détailler dans le présent chapitre les applications à la spectroscopie pompe-sonde résolue en temps¹². Dans cette technique on excite très brièvement le système étudié par une première impulsion, dite de pompe. On laisse ensuite le système évoluer pendant un temps

12. L'utilisation d'impulsion courte dans les matériaux optiquement non-linéaires et la fusion inertielle seront aussi abordés dans les chapitres suivants.

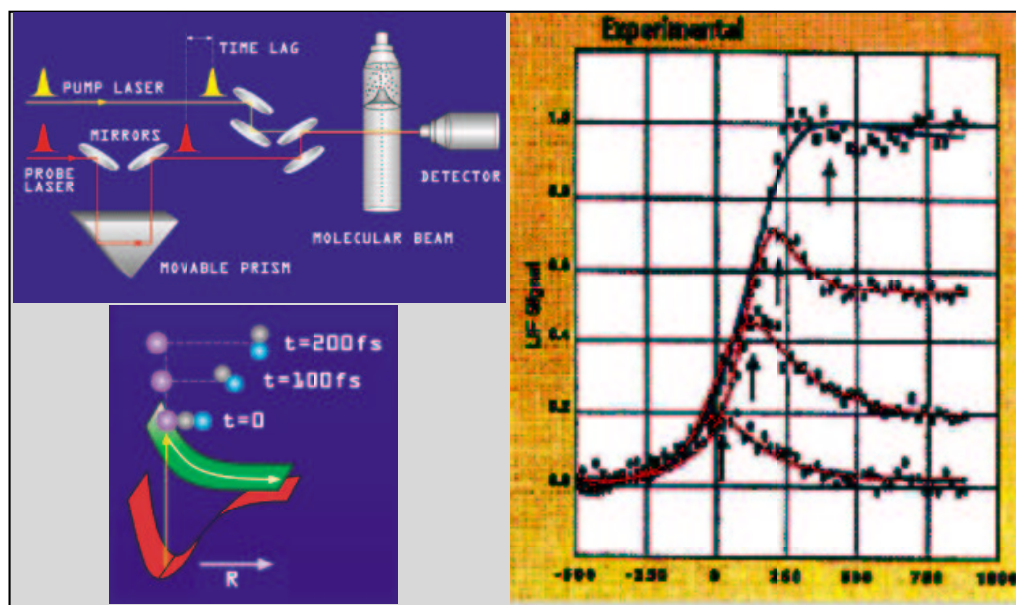


FIG. 2.24 – Principe de la première expérience de femtochimie. L'étude portait sur la dissociation de ICN^* excitée par une impulsion de pompe. On note la "ligne à retard" constituée d'un prisme et de deux miroirs qui permet de décaler les temps d'arrivée des impulsions de pompe et de sonde (un décalage du prisme de $3\text{ }\mu\text{m}$ correspond à un délai de 10 fs). La figure de droite représente le nombre de composés détectés par fluorescence en fonction du temps et de la longueur d'onde de la sonde. Les différentes courbes sont associées à des longueurs d'ondes différentes et sont donc sensibles à des intermédiaires réactionnels différents.

τ , après lequel on sonde ses propriétés optiques par une deuxième impulsion, par exemple en mesurant l'absorption du milieu.

Une des applications les plus remarquable des lasers femtosecondes est le développement de la femtochimie, une discipline née au cours des années 80 notamment grâce aux travaux d'Ahmed Zewail (Prix Nobel de chimie 1999¹³). L'objet de ce domaine très récent est l'étude "en temps réel" des réactions chimiques, et en particulier l'observation directe d'intermédiaires

13. Au sujet de l'histoire du développement de la femtochimie, on pourra se référer au discours prononcé par A. Zewail devant l'académie Nobel, accessible en ligne à l'adresse <http://nobelprize.org/chemistry/laureates/1999/zewail-lecture.html>

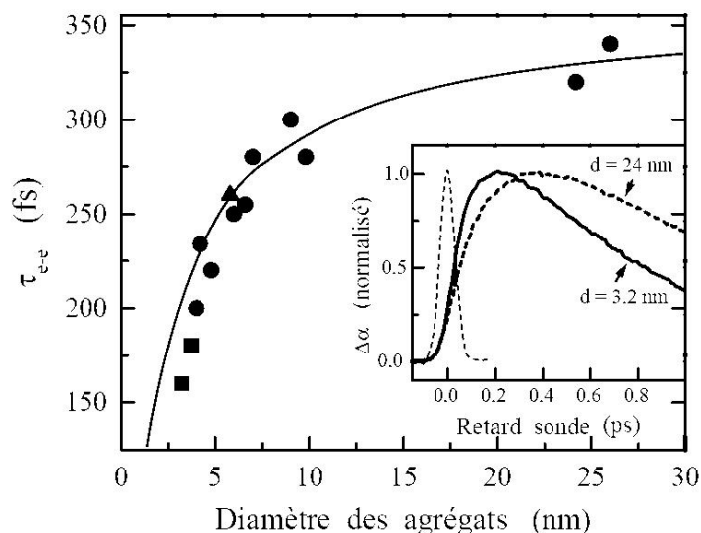


FIG. 2.25 – Évolution du temps de rethermalisation τ_{e-e} du nuage électronique de nanoparticules d'argent de rayon diamètre d variable. En insert, on a représenté l'absorption $\Delta\alpha$ de nanoparticules de diamètres $d = 3.2$ nm et $d = 24$ nm e, fonction du temps (la courbe pointillée représente l'impulsion de pompe).

réactionnels, souvent trop éphémères pour être isolés. Cet objectif ambitieux nécessite des résolutions temporelles très grandes. En effet, on sait qu'un atome ou une molécule en phase gazeuse se déplace à une vitesse de l'ordre du kilomètre par seconde. La réaction chimique se produisant lorsque les différents réactifs sont distants de quelques rayons atomiques soit quelques angström, on voit que la durée d'une réaction élémentaire est de l'ordre d'une centaine de femtosecondes. La première expérience de femtochimie a concerné l'étude de la rupture de la liaison I-C dans ICN*. Dans cette expérience, l'impulsion de pompe amène les électrons dans une orbitale antiliante, alors que l'impulsion de sonde permet la détection des CN par fluorescence.

Une seconde illustration de cette méthode est l'étude de la dynamique électronique au sein de nano-particules métalliques¹⁴. Ces études visent à comprendre les différences de comportements des électrons dans des métaux massifs et des agrégats. Ces questions deviennent d'actualités avec la minia-

14. On pourra trouver plus de détails dans N. del Fatti et F. Vallée, Spectroscopie femtoseconde de nano-objets, Images de la physique 2002, p 50.

turisation de plus en plus poussée des circuits électroniques, pour lesquels on approche la limite où les transistors d'un microprocesseur ne contiennent plus que quelques électrons chacun. Dans l'expérience décrite ici, la première impulsion excite les électrons et modifie leur distribution d'énergie et l'amène en particulier hors de l'équilibre thermique. La mesure de l'absorption de la sonde par les nanoparticules permet d'étudier la dynamique du retour à l'équilibre, comme on le constate sur la figure 2.25. Les temps caractéristiques mis en jeu sont ici de l'ordre de quelques centaines de femtosecondes, ce qui nécessite donc bien l'utilisation de laser à impulsions courtes. À temps plus long, cette technique a permis l'étude du transfert d'énergie du gaz d'électron au réseau ionique.

2.4 Complément : sources non-classiques

2.4.1 Sources à photons uniques

Une question qui n'a pas jusqu'ici été abordée est la nature quantique du champ électromagnétique issu d'un laser.

Classiquement, on sait que le laser est une source cohérente, bien décrite par une onde monochromatique de la forme $\mathbf{E} = \mathbf{E}_0 e^{i(\mathbf{k}\mathbf{r} - \omega t)}$. Quantiquement, on s'attend donc à ce que l'état $|\psi_0\rangle$ du champ électromagnétique vérifie

$$\langle \psi_0 | \hat{\mathbf{E}} | \psi_0 \rangle = \mathbf{E}_0 e^{i(\mathbf{k}\mathbf{r} - \omega t)}.$$

Essayons de chercher $|\psi_0\rangle$ sous la forme d'un état de Fock contenant N_0 photons dans le mode $\mathbf{k}$. Le champ électrique dans cet état se calcule en utilisant les opérateurs créations et annihilation

$$\langle \psi_0 | \hat{\mathbf{E}} | \psi_0 \rangle = \sqrt{\frac{\hbar\omega_k}{2\epsilon_0 V}} \langle N_0 | (\hat{a} e^{i\mathbf{k}\mathbf{r}} + \hat{a}^\dagger e^{-i\mathbf{k}\mathbf{r}}) | N_0 \rangle,$$

soit en utilisant les propriétés des $\hat{a}$ et $\hat{a}^\dagger$

$$\langle \psi_0 | \hat{\mathbf{E}} | \psi_0 \rangle = \sqrt{\frac{\hbar\omega_k}{2\epsilon_0 V}} (\sqrt{N_0} e^{i\mathbf{k}\mathbf{r}} \langle N_0 | N_0 - 1 \rangle + \sqrt{N_0 + 1} e^{-i\mathbf{k}\mathbf{r}} \langle N_0 | N_0 + 1 \rangle) = 0,$$

car les produits scalaires $\langle N_0 | N_0 \pm 1 \rangle$ sont nuls. Autrement dit, les états à nombre fixé de photons possèdent un champ électrique nul et ne peuvent

donc pas décrire le rayonnement cohérent issu d'un laser. Cette propriété s'interprète par l'existence d'une relation d'incertitude de type Heisenberg entre le nombre de photons et la phase φ du champ électromagnétique, soit $\Delta N \Delta \varphi \sim 1$. Lorsque le nombre de photon est parfaitement déterminé, la phase du champ électrique est complètement aléatoire, ce qui conduit à une annulation de la valeur du champ lors d'une série de mesures.

Nous ne chercherons pas ici à déterminer la nature exacte du champ émis par le laser, mais nous allons retourner la question initialement posée et nous demander si il est possible de générer des états à nombre de photons bien déterminé. Pour répondre à cette question, il est intéressant de noter qu'un atome unique ne peut qu'émettre qu'un seul photon à la fois. Les fluctuations de nombre de photons dans une source classique comportant un grand nombre d'atomes proviennent donc du fait que les photons sont émis à des instants différents par des atomes différents. Pour réaliser une source de photon unique, il est donc nécessaire d'isoler une source (atome, molécule etc.) unique.

- *Molécules fluorescentes.* L'utilisation de molécules fluorescente est la première option pouvant venir à l'esprit, par exemple une molécule de rhodamine isolée. Si celle-ci est pompée dans un état excité, elle retombe dans l'état fondamental en émettant un photon unique. L'inconvénient des molécules fluorescentes est le phénomène de "photoblanchiment" qui correspond à la disparition des propriétés de fluorescence après une irradiation trop longues. C'est pourquoi des recherches intensives sont menés pour développer des composés fluorescents à longue durée de vie, tels que les centres colorés ou les boîtes quantiques.
- *Centres colorés.* Une voie possible est l'utilisation de centres colorés (i.e. de défauts) dans les cristaux, en particulier les centres N-V du diamant, correspondant à l'association d'un atome d'azote avec une lacune de la matrice de carbone. Chaque centre N-V se comporte comme une molécule fluorescente et peut donc être utilisé comme source de photon unique.
- *Boîtes quantiques.* Une seconde alternative est l'utilisation de boîtes quantiques, des agrégats d'InAs inclus dans une matrice de GaAs. Ces agrégats sont obtenus en déposant InAs sur un substrat de GaAs. Du fait de la différence de maille entre les deux structures cristallines, il existe une énergie de surface associée à l'interface InAs/GaAs. Pour minimiser cette interface, InAs forme des "gouttes" à la sur-

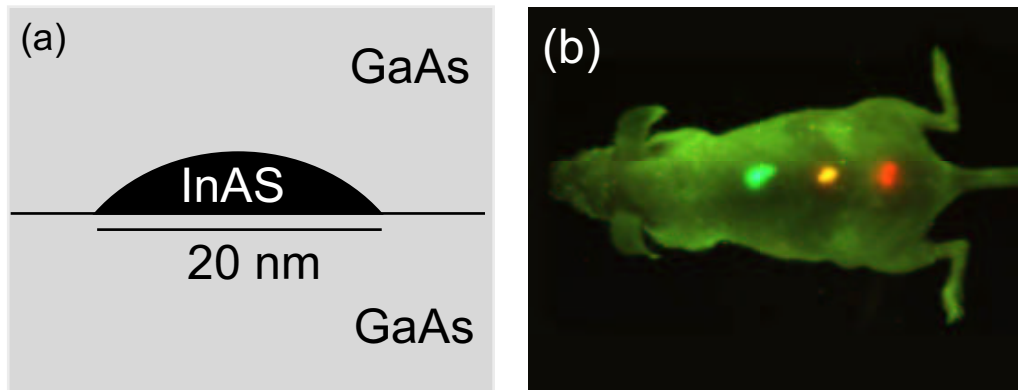


FIG. 2.26 – a) a) Principe de réalisation d'une boîte quantique semi-conductrice. b) Observation des organes d'une souris vivante par fluorescence de boîtes quantiques semi-conductrices.

face du substrat de GaAs. Ces gouttes, de quelques dizaines de nanomètres de diamètre et quelques nanomètres d'épaisseur, sont ensuite recouvertes par un dépôt de GaAs. Bien que stables uniquement à basse température, les techniques aujourd'hui parfaitement maîtrisées de croissance cristallines des semi-conducteurs permettent de les inclure dans des édifices microscopiques plus vastes telles que des micro-cavités qui permettent d'améliorer la directionnalité de la source.

2.4.2 Application à la cryptographie quantique

Le commerce sur internet ou la téléphonie mobile ont étendu du domaine militaire au domaine civil la nécessité de communications cryptés. Le protocole de cryptographie le plus sûr est le cryptage par clef secrète. Il consiste à écrire le message à transmettre en une série m_i de 0 et de 1 et de le coder en la sommant modulo 2 à une suite aléatoire de c_i 0 et de 1 (la clef) de même longueur que le message initial. La suite $m'_i = m_i + c_i$ est transmise d'un correspondant à l'autre sur un canal public puis décodée en effectuant l'opération $m'_i + c_i = m_i$.

Ce protocole bien qu'en principe parfaitement sûr possède deux faiblesses. D'une part il est nécessaire que les deux partenaires (Alice et Bob conventionnellement) connaissent la clef de cryptage c_i . Il se pose donc un problème de distribution de cette clef sur un canal secret. Deuxièmement, une clef ne

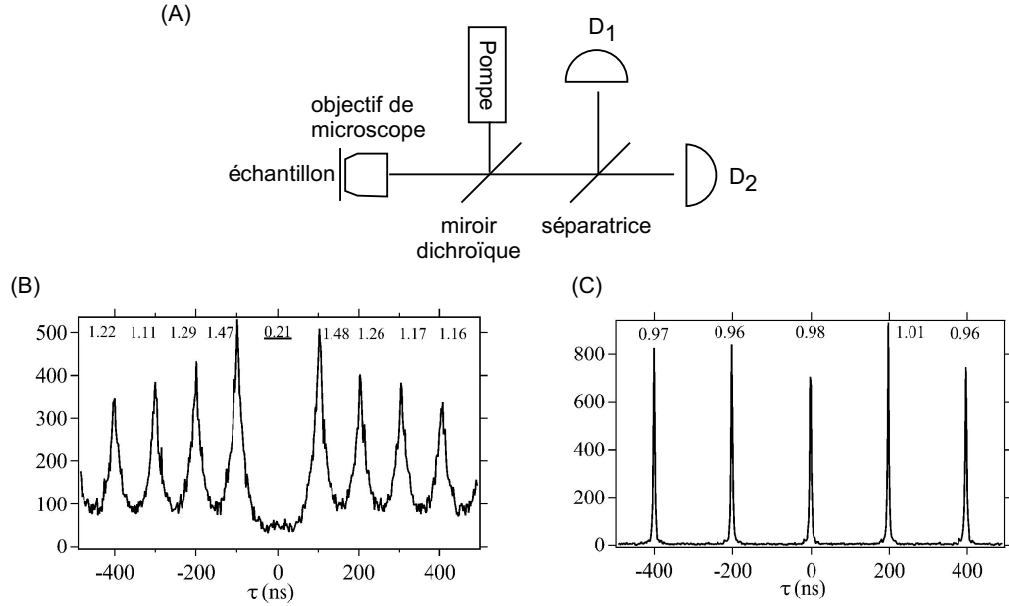


FIG. 2.27 – Caractérisation d'une source à un photon par corrélation d'intensité (extrait de Beveratos A. et al., *Eur. Phys. J. D*, **18** 191 (2002)). A) Principe de l'expérience. La source à un photon, ici un centre N-V, est placée sur le porte échantillon d'un objectif de microscope. Le pompage est assuré par un laser pulsé émettant des impulsions de 0.8 ns avec un taux de répétition de 5.3 MHz. Le miroir dichroïque réfléchit la longueur d'onde du laser de pompe (532 nm) et transmet la longueur d'onde d'émission du centre coloré (690 nm). Le nombre de photons émis par la source est mesuré par corrélation temporelle. On mesure le nombre d'événements où un photon est détecté par D_1 à l'instant t et un photon est détecté par D_2 à $t + \tau$. Si la source n'émet qu'un photon à la fois, ce nombre d'événements doit être nul pour $\tau = 0$. B) et C) Résultat de la mesure pour une source à photon unique (B) et une source classique (C). Les pics latéraux correspondent à des photons arrivés dans deux impulsions successives.

peut servir qu'une seule fois. On peut en effet montrer que l'on peut reconstruire la clef de cryptage à partir de deux messages codés avec une même clef.

Ces deux difficultés rendent la cryptographie par clef aléatoire inopérante en pratique et on lui préfère plutôt les algorithmes dits à clef public actuellement en usage de nos jours sur internet. La sécurité de ces protocoles n'est cependant pas absolue : la difficulté du décryptage d'un message codé se fonde sur la difficulté à décomposer un grand nombre en facteurs premiers et est donc tributaire de la puissance de calcul des ordinateurs : on peut notamment montrer que le développement d'un éventuel ordinateur quantique rendrait caduque tous ces algorithmes.

La cryptographie quantique est née en vue de résoudre le problème la communication d'une clef secrète entre Alice et Bob en se fonde sur un des fondements de la mécanique quantique, qui est que la mesure perturbe irrémédiablement l'état d'un système. Dans le protocole le plus simple de cryptographie quantique, les 0 et 1 de la clef secrète sont représentés par les états de polarisation d'un photon : un photon polarisé selon x correspond à 0 et selon y elle correspond à 1.

Imaginons donc qu'Alice génère une succession de photons, dont les polarisations sont choisies aléatoirement dans deux bases de polarisations $\mathcal{B}$ et $\mathcal{B}'$ inclinée à 45° l'une de l'autre. Chaque photon émis est donc caractérisé par deux grandeurs aléatoires, respectivement la base choisie et la polarisation choisie dans cette base, ces deux informations étant conservées par Alice.

Bob, lorsqu'il reçoit les photons émis par Alice, ne connaît pas la base choisie pour coder les photons et va donc mesurer la polarisation de chacun d'eux aléatoirement dans $\mathcal{B}$ et $\mathcal{B}'$. Pour un photon donné, si Alice et Bob ont fait un même choix de base, l'état de polarisation mesuré par Bob est identique à celui choisi par Alice. En revanche, si la base est mal choisie, Bob a une chance sur deux d'obtenir un résultat différent d'Alice. Une fois les photons tous lus par Bob, Alice communique sur un canal public le choix de base pour chaque photon envoyé : Alice et Bob peuvent donc construire une clef commune en ne conservant que les polarisations de photons mesurés dans une même base.

Pour vérifier la sécurité de ce protocole, imaginons qu'Eve cherche à intercepter la communication entre Alice et Bob. Eve ne connaît pas non plus le choix de polarisation d'Alice et doit donc choisir à chaque photon émis une base de lecture parmi $\mathcal{B}$ et $\mathcal{B}'$. Statistiquement, elle choisit la même base qu'Alice une fois sur deux. Dans ce cas, elle mesure le bon état de polarisation



FIG. 2.28 – *Système commercial de cryptographie quantique développé par la firme genevoise ID-Quantique (<http://www.idquantique.com/>).*

du photon et celui-ci est ensuite réémis dans le même état de polarisation à Bob. En revanche, lorsque Eve choisit la mauvaise base ($\mathcal{B}'$ pour fixer les idées), va se retrouver aléatoirement polarisé selon x' ou y' lorsqu'il sera émis vers Bob et la mesure de Bob va s'en trouver affecté. En effet, si Bob fait la mesure de l'état de polarisation du photon dans la même base qu'Alice, il ne trouvera pas le même résultat que celle-ci avec une probabilité de 100 %.

Pour détecter la présence d'Eve, Alice et Bob vont donc sacrifier une partie de la clef qu'ils s'étaient transmis en communiquant publiquement une partie de celle-ci. Comme nous venons de le voir, les polarisations mesurées par Alice et Bob doivent être identiques en l'absence de tout espion, puisque les états des photons conservés par les deux interlocuteurs ont été mesurés dans des bases identiques. Un rapide calcul de probabilité montre que lorsque N photons sont sacrifiés, la probabilité qu'a Eve de passer inaperçu vaut $(3/4)^N$ et tend donc vers 0 rapidement.

Pour être efficace, le protocole de cryptographie quantique nécessite l'utilisation de sources à photons uniques. En effet, si les impulsions envoyées par Alice contenaient plusieurs photons, Eve pourrait en prélever un seul et faire la mesure sur celui-ci, sans perturber le reste des photons : ceci justifie donc la nécessité de recourir à des sources de photons non classiques telles que les centres colorés ou les boîtes quantiques que nous venons de décrire.

Chapitre 3

Propagation lumineuse

3.1 Optique gaussienne

3.1.1 Aspects qualitatifs de la propagation lumineuse

Une des particularités remarquables de la lumière laser est sa grande directionnalité. Au contraire de la lumière issue de sources lumineuses classiques qui est émise dans toutes les directions, le faisceau issu d'un laser est pratiquement unidirectionnel. Cette conséquence de l'émission stimulée n'est en réalité vraie qu'à "courte distance". En effet, on constate que loin du laser le rayon diverge légèrement. Cette propriété est simplement due à la diffraction, conséquence du diamètre fini d du faisceau laser. Plus précisément, on sait que pour un faisceau de longueur d'onde λ , l'angle de diffraction vaut $\alpha \sim \lambda/d$ et le rayon du laser a une distance L de la source vaut

$$d(z) \sim L\alpha \sim L\lambda/d. \quad (3.1)$$

Cette relation n'est valable qu'à longue distance. En effet, lorsque L tend vers 0, le diamètre $d(z)$ doit tendre vers d , ce qui contredit l'équation (3.1). Cette relation n'est donc valable que pour $d(z) \gg d$ soit $L \gg d^2/\lambda$. Pour L plus petit, on peut supposer que le diamètre du rayon n'a pas été sensiblement modifié (Fig.).

Application au laser Terre-Lune

Dans le cas d'un laser typique, le diamètre du faisceau est de l'ordre du millimètre et de longueur d'onde d'un micron. La distance en dessous de

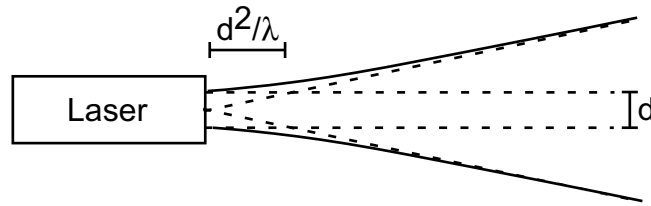


FIG. 3.1 – Évolution du diamètre d’un faisceau laser. Pour un faisceau de diamètre d , la divergence due à la diffraction ne se manifeste qu’à partir d’une distance $\sim d^2/\lambda$ de la source.

laquelle on pourra supposer le faisceau comme collimaté est donc de l’ordre du mètre. Dans les applications courantes, on peut donc considérer le faisceau issu du laser comme parfaitement parallèle.

Une application où cette divergence devient dramatique est la mesure de la distance Terre-Lune. Actuellement, cette mesure est réalisée par télémétrie laser : on envoie des impulsions lumineuses depuis la Terre vers la Lune. Grâce à des miroirs déposés par les missions Appolo XI, XIV et XV et Lunakhod 17 et 21, la lumière est renvoyée vers les télescopes terrestres. La mesure du temps mis par la lumière pour réaliser l’aller-retour fournit alors une mesure précise de la distance Terre-Lune. Du fait de la divergence du faisceau, l’intensité de l’onde est considérablement diminuée au retour. Pour la maximiser, on agrandit le faisceau à l’aide d’un miroir de télescope de 1,5 m (celui du plateau de Calern en France, ainsi que d’autres situés au Japon, aux États-Unis et en Australie), ce qui donne un rayon du faisceau au niveau de la Lune de “seulement” 100 m. Malgré la réduction de la divergence, l’intensité lumineuse collectée au retour est extrêmement faible : à peine un photon tous les cents tirs ! Malgré cela, la précision des mesures est bonne : la distance Terre-Lune est connue aujourd’hui au millimètre près, soit une précision de l’ordre de $1/10^{11}$.

3.1.2 Approximation de l’enveloppe lentement variable

Après l’introduction qualitative présentée au paragraphe précédent, nous allons essayer de préciser plus quantitativement le profil d’intensité d’un faisceau laser se propageant dans le vide.

Considérons donc une onde électromagnétique, supposée polarisée rectilignement. Si l’on note $E(\mathbf{r}, t)$ l’amplitude de son champ électrique, celui-ci

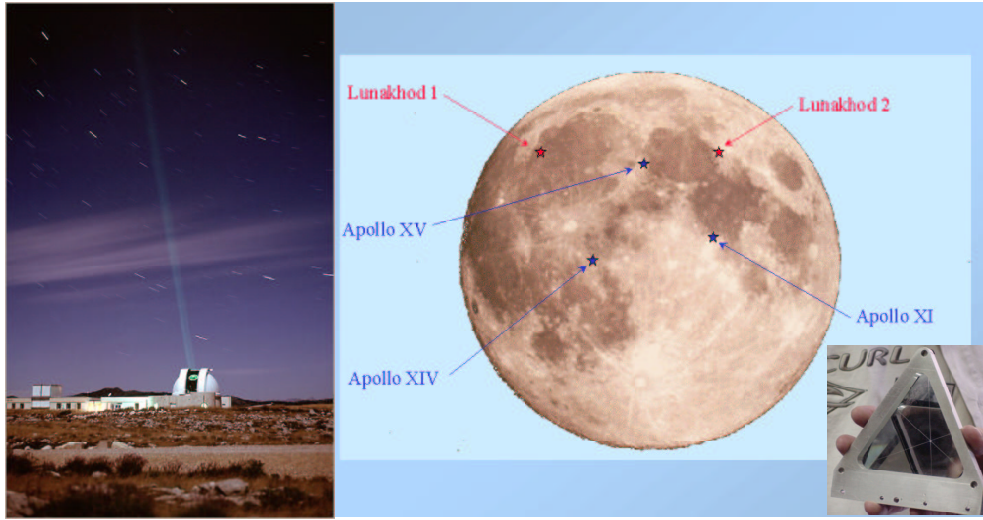


FIG. 3.2 – À gauche, coupole du télescope abritant le laser Terre-Lune (Projet MéO – Métrologie Optique – de l’Observatoire de Côte d’Azur). À droite carte de la Lune indiquant l’emplacement des miroirs déposés sur la Lune. Insert : réplique d’un de ces miroirs.

sera solution de l’équation de propagation

$$\square E = 0,$$

où $\square$ désigne l’opérateur d’Alembertien. Supposons par ailleurs l’onde monochromatique, de pulsation ω . Si l’on cherche E sous la forme $\mathcal{E}e^{i(kz-\omega t)}$, avec $k = \omega/c$, alors l’enveloppe $\mathcal{E}$ est solution de

$$2ik\partial_z\mathcal{E} + \partial_z^2\mathcal{E} + \Delta_\perp\mathcal{E} = 0,$$

où $\Delta_\perp = \partial_x^2 + \partial_y^2$ est le laplacien transverse.

En général, l’intensité lumineuse, proportionnelle à $|E|^2$, et donc $|\mathcal{E}|^2$, varie sur des échelles spatiales $L_\perp$ et L_z dans les directions transverses et longitudinales. En ordre de grandeur on peut écrire que $\Delta_\perp\mathcal{E} \sim \mathcal{E}/L_\perp^2$ et $\partial_z\mathcal{E} \sim \mathcal{E}/L_z$. Dans la limite où L_z est grand devant la longueur d’onde, les termes en $k\partial_z\mathcal{E}$ et $\partial_z^2\mathcal{E}$ sont dans un rapport $kL_z \gg 1$. Il est donc raisonnable de négliger le terme en ∂_z^2 ce qui nous donne l’équation dite de l’enveloppe lentement variable :

$$2ik\partial_z\mathcal{E} + \Delta_\perp\mathcal{E} = 0. \quad (3.2)$$

Notons qu'en interprétant la coordonnée z comme un temps, on obtient une analogie formelle entre cette équation et l'équation de Schrödinger d'une particule libre.

3.1.3 Faisceau gaussien

On sait qu'une solution particulière de l'équation d'onde est une onde sphérique de la forme Ae^{ikr}/r . En développant cette solution au voisinage de l'axe z , on en déduit une solution de (3.2) de la forme

$$\mathcal{E} = \frac{A}{z} e^{ik\rho^2/2z},$$

où $\rho^2 = x^2 + y^2$. Dans cette solution, on a pris le centre de l'onde en $z = 0$. On peut de manière plus général placer le centre de l'onde en ξ , de sorte que si l'on pose $q(z) = z - \xi$, on obtienne une nouvelle solution de la forme

$$\mathcal{E} = \frac{A}{q(z)} e^{ik\rho^2/2q(z)}.$$

Bien que nous ayons pris ξ réel pour l'interpréter comme un décalage du centre de l'onde, rien ne nous empêche formellement de le supposer complexe. La solution que nous venons d'obtenir reste toujours valable, mais son comportement physique est profondément modifié. Le cas du *faisceau gaussien* correspond à prendre ξ imaginaire, avec $z_R = \text{Im}(\xi)$. On obtient alors

$$\mathcal{E} = \frac{A}{q(z)} e^{-kz_R\rho^2/2|q(z)|^2} e^{ikz\rho^2/2|q(z)|^2},$$

soit, pour l'intensité lumineuse :

$$I(r, z) = c\epsilon_0 \frac{|\mathcal{E}|^2}{2} = \frac{c\epsilon_0 |A|^2}{kz_R W^2(z)} e^{-2\rho^2/W^2(z)},$$

avec $W^2(z) = 2(z_R^2 + z^2)/kz_R$.

$W(z)$ s'interprète comme le "rayon" du faisceau. On peut l'écrire sous la forme

$$W(z) = W_0 \sqrt{1 + z^2/z_R^2},$$

où W_0 , appelé waist, ou col, est le rayon minimum du faisceau. La relation liant le waist à z_R (appelé longueur de Rayleigh) s'écrit

$$z_R = \frac{\pi W_0^2}{\lambda}.$$

On constate que l'évolution de W avec z reproduit bien les caractéristiques qualitatives concernant la propagation lumineuse, que nous avons établi précédemment. En effet, en identifiant d et W_0 on constate que

1. À grand z le rayon du faisceau vaut $W(z) \sim z\lambda/\pi W_0 \sim z\lambda/\pi d$, ce qui est l'équivalent de la formule de la diffraction par un objet de taille $d = W_0$.
2. Pour $z \ll z_R$ le rayon du faisceau est inchangé.

Les modes gaussiens que nous venons de construire ont leur foyer en $z = 0$. On généralise aisément au cas d'un faisceau focalisé en un point z_0 quelconque en considérant $\xi = z_0 + iz_R$.

Les faisceaux gaussiens que nous venons d'étudier peuvent se généraliser à des bases de fonctions solutions de l'équation (3.2). Les deux principales familles de solutions sont constituées respectivement par les modes de Hermite-Gauss et Laguerre-Gauss.

Les modes de Hermite-Gauss, notés $HG_{m,n}$ correspondent à un profil transverse de la forme

$$\mathcal{E}(x,y) = P_n(x)P_m(y) \exp(-(x^2 + y^2)/W^2),$$

où P_n est le polynôme de Hermite de degré n . Le polynôme P_n possède n racines, ce qui implique l'existence de n (resp. m) annulations de l'intensité dans la direction x (resp. y).

Exprimé en coordonnées polaire (ρ, θ) , le profil transverse d'un mode de Laguerre-Gauss LG_{mn} s'écrit

$$\mathcal{E}(x,y) = P_{m,n}(\rho) e^{-\rho^2/W^2} e^{im\theta},$$

où les $P_{m,n}$ sont des polynômes de Laguerre.

Des profils expérimentaux de ces deux classes de solutions sont représentés sur la figure (3.3). On note en particulier la mise en évidence du profil de phase des modes de Laguerre-Gauss. La structure en spirale est associée à une propagation "en hélice" de l'énergie le long du faisceau. En effet, les

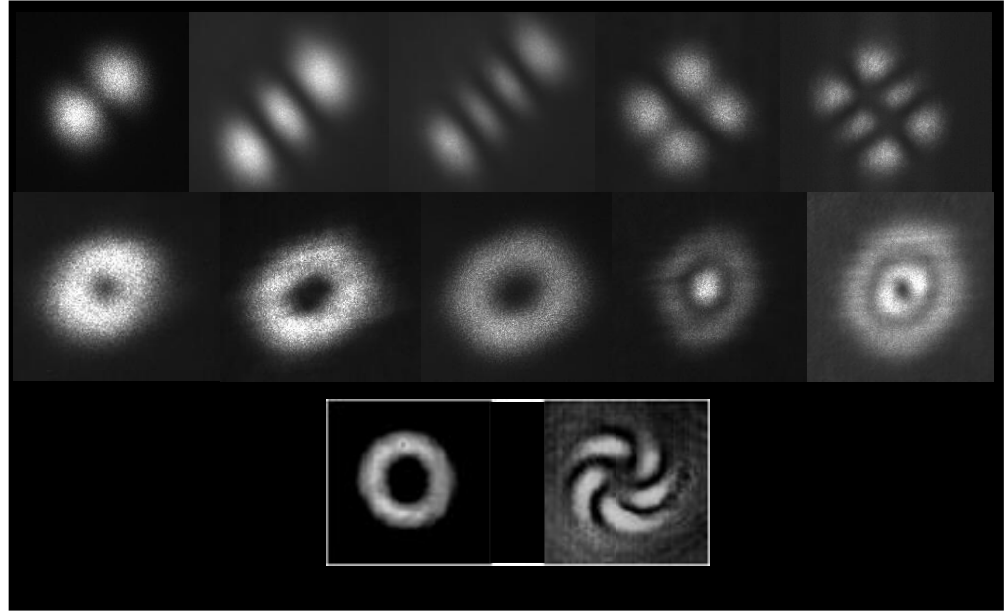


FIG. 3.3 – *Ligne du haut: visualisation de modes de Hermite-Gauss. De gauche à droite, HG01, HG02, HG03, HG11, HG21. Ligne du milieu, modes de Laguerre-Gauss LG10, LG20, LG30, LG11, LG21. Ligne du bas, visualisation du profil de phase d'un mode de Laguerre-Gauss par interférence avec un faisceau gaussien.*

surfaces équiphasées, dont la normale définit le vecteur de Poynting, sont des hélicoïdes données par $kz + m\theta = \text{cte}$.

3.1.4 Optique gaussienne

Relation de conjugaison pour les faisceaux gaussiens

La traversée d'une lentille mince se traduit essentiellement par un déphasage des faisceaux lumineux dû à la variation de l'épaisseur de verre à traverser en fonction de la distance à l'axe optique. Dans la limite paraxiale, et en utilisant la symétrie par rapport à l'axe optique, on peut écrire qu'à l'ordre non trivial en ρ le déphasage se met sous la forme

$$\delta\phi(x,y) = a - b(x^2 + y^2)/2,$$

où a et b sont des paramètres dépendant de la géométrie de la lentille ainsi que de l'indice du verre la constituant.

Après traversée de la lentille, supposée localisée en $z = 0$, un faisceau d'amplitude $\mathcal{E}$ avant la traversée aura pour amplitude

$$\mathcal{E}'(x,y) = e^{i\delta\phi} \mathcal{E}(x,y)$$

après la traversée. Si l'on pose $q_0 = q(0)$, où $q(z)$ est le paramètre du faisceau incident on trouve

$$\mathcal{E}'(x,y) = \frac{Ae^{ia}}{q_0} e^{ik\rho^2/2q'_0},$$

avec $q_0'^{-1} = q_0^{-1} - b/k$. Autrement dit, le faisceau en sortie de lentille est toujours gaussien, avec un paramètre en 0 changé de q_0 à q'_0 .

La relation liant q_0 à q'_0 est une généralisation de la formule de conjugaison pour les lentilles minces. En effet, dans l'approximation de l'optique géométrique, on suppose que la focalisation d'un faisceau peut-être parfaite, ce qui revient à supposer la longueur de Rayleigh faible. Si l'on note z_0 et z'_0 la position des waists des faisceaux incidents et transmis, on a dans cette approximation $z_0 \gg z_R$ et $z'_0 \gg z_R$. Ceci nous amène donc à la relation

$$\frac{1}{z'_0} = \frac{1}{z_0} + \frac{b}{k},$$

c'est-à-dire la relation de conjugaison pour une lentille de focale $f = k/b$. La relation de passage pour un faisceau gaussien s'écrit donc simplement comme

$$\frac{1}{q'_0} = \frac{1}{q_0} - \frac{1}{f}, \quad (3.3)$$

équation (3.3) qui généralise la relation de conjugaison des lentilles minces au cas de faisceaux gaussiens.

Faisceaux gaussiens et cavité résonnante

L'équation (3.3) permet de montrer que les modes propres transverses d'une cavité linéaire constituée de deux miroirs sphériques sont les modes gaussiens.

Considérons pour simplifier une cavité linéaire constituée de deux miroirs identiques de focale f et placés à une distance L l'un de l'autre : un aller-retour dans la cavité sera donc équivalent à la traversée de deux lentilles de

même focale f , séparées d'une distance L . Par définition, un mode stationnaire de la cavité doit rester inchangé par cette série de transformation.

Comme la traversée d'une lentille n'altère pas le caractère gaussien du faisceau, il suffit de s'intéresser à la transformation du paramètre q lors de la propagation lumineuse. Considérons donc un faisceau caractérisé par un paramètre q_0 à l'entrée de la première lentille. La succession des étapes pendant la propagation est la suivante :

1. *Traversée de la première lentille.* D'après la relation de conjugaison (3.3), q_0 est transformé en $q'_0 = \mathcal{L}_f(q_0)$ tel que

$$\frac{1}{\mathcal{L}_f(q_0)} = \frac{1}{q_0} - \frac{1}{f}.$$

2. *Propagation jusqu'à la seconde lentille.* Par définition, $q'_0 = q'(z = 0)$, l'origine des z étant choisi au niveau de la première lentille. Au niveau de la seconde lentille, le paramètre du faisceau vaut

$$q''_0 = \mathcal{P}_L(q'_0) = q'(z = L) = L + q'_0.$$

3. *Suite de la propagation.* Les fonctions de transfert pour la traversée de la seconde lentille et la propagation sur une distance L sont données par les mêmes formules que précédemment

D'après ce qui précède, un mode stable doit donc satisfaire après un aller-retour entre les miroirs

$$q_0 = \mathcal{P}_L(\mathcal{L}_f(\mathcal{P}_L(\mathcal{L}_f(q_0)))),$$

soit (en s'aidant éventuellement du formalisme matriciel développé au paragraphe suivant)

$$q_0 = \frac{Aq_0 + B}{Cq_0 + D},$$

avec

$$\begin{pmatrix} A & B \\ C & D \end{pmatrix} = \begin{pmatrix} (1 - L/f)^2 - L/f & L(2 - L/f) \\ (-2 + L/f)/f & 1 - L/f \end{pmatrix}$$

Formellement, l'équation sur q_0 se ramène à la résolution de l'équation du second degré

$$Cq_0^2 + (D - A)q_0 - B = 0.$$

Rappelons que $q_0 = -z_0 - iz_R$, avec $z_R > 0$. Autrement dit q_0 est toujours imaginaire. Puisque les coefficients du polynôme sont réels, la seule façon d'obtenir des racines complexes est d'imposer au discriminant de devenir négatif, c'est-à-dire

$$(D - A)^2 + 4BC < 0.$$

En utilisant le fait que les matrices ABCD sont toutes de déterminant 1, on peut récrire cette inégalité comme $\text{Tr}(ABCD)^2 < 4$. En utilisant l'expression des coefficients ABCD, on en déduit l'inégalité :

$$\frac{L}{f} < 4.$$

Autrement dit, la cavité admet des modes gaussiens lorsque sa longueur est suffisamment courte comparé au rayon de courbure des miroirs. Il est au passage tout à fait notable de retrouver un critère de stabilité des modes gaussiens identique à celui de l'optique géométrique.

Lien avec les matrices ABCD

De manière générale, le paramètre q_0 d'un faisceau gaussien est de la forme $-z_0 - iz_R$ et appartient donc au demi-plan Π_- constitué des q_0 de partie réelle négative. Le passage à travers un dispositif optique représente une application de Π_- vers Π_- , analytique (on est en physique!) et inversible, puisqu'il est toujours possible de traverser le dispositif en sens inverse. Un théorème d'analyse complexe¹ stipule que de telles transformations sont nécessairement des homographies de la forme

$$q' = f_M(q) = \frac{Aq + B}{Cq + D},$$

avec $AD - BC \neq 0$ et où M désigne la matrice ABCD

$$M = \begin{pmatrix} A & B \\ C & D \end{pmatrix}$$

1. Voir par exemple H. Cartan, Théorie élémentaire des fonctions analytiques (Hermann), Chapitre VI.

L'intérêt de cette notation matricielle est double. D'une part, il existe un isomorphisme de groupe existant entre les homographies et les matrices. On voit en effet que l'on a pour tout couple (M, M')

$$f_M \circ f_{M'} = f_{M \cdot M'}.$$

Ce qui, d'un point de vue pratique, ramène le calcul de q'_0 à celui d'un produit de matrices.

D'autre part, les matrices ABCD de l'optique gaussienne sont identiques aux matrices ABCD de l'optique géométrique, comme on le constate sur les exemples de la propagation dans l'espace libre et au travers d'une lentille mince. Dans ces deux cas, les matrices ABCD notées respectivement M_p et M_l sont de la forme

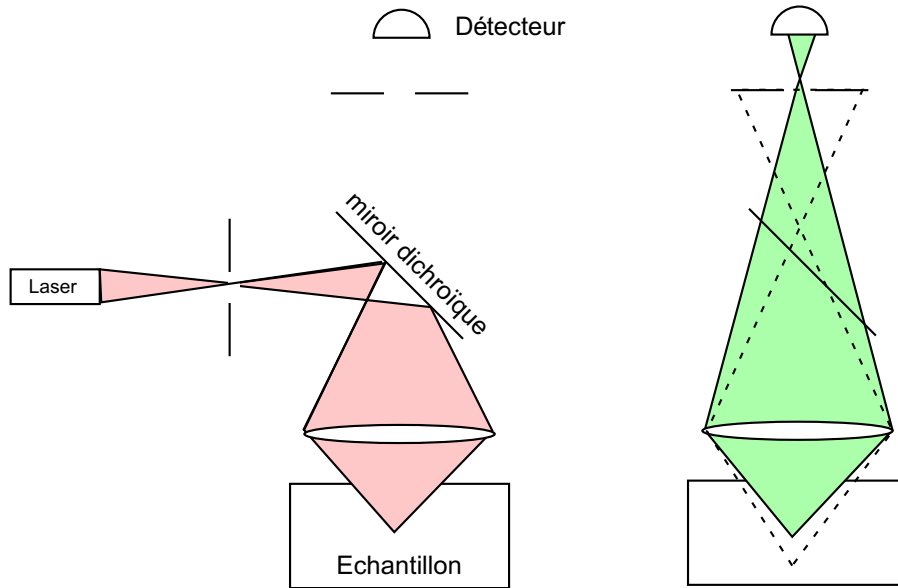
$$M_p = \begin{pmatrix} 1 & L \\ 0 & 1 \end{pmatrix} \quad M_l = \begin{pmatrix} 1 & 0 \\ -1/f & 1 \end{pmatrix}$$

3.1.5 Microscopie à haute résolution

Comparée à l'optique géométrique, la description d'un faisceau laser en terme de faisceau gaussien offre la possibilité de décrire de façon simple à la fois les propriétés de propagation et de diffraction de la lumière. Elle permet en particulier de prédire la résolution de tel ou tel dispositif optique, en calculant le waist du faisceau (et donc la focalisation) au point image. Les lasers sont principalement utilisés dans les dispositifs de microscopie par sonde fluorescentes, tels que la microscopie confocale et la microscopie par onde évanescente.

Microscopie confocale

La microscopie confocale est une technique très puissante permettant de reconstruire une image tridimensionnelle de l'échantillon étudié. Son principe est décrit sur la fig. 3.4. On commence par focaliser la lumière laser sur l'échantillon après passage dans d'un diaphragme de quelques dizaines de microns de diamètre (appelé *pinhole* en anglais). Cette lumière de pompe sert à exciter des molécules fluorescentes dans le milieu. La lumière de fluorescence émise depuis le point de focalisation de la pompe est ensuite focalisée sur un deuxième diaphragme. Cette technique permet de ne récolter que la lumière provenant du foyer image du premier diaphragme. En effet, la lumière venant

FIG. 3.4 – *Principe de la microscopie confocale*

d'un autre point est bloquée par le deuxième pinhole et ne parvient donc pas au détecteur. Par rapport à la microscopie classique, les dispositifs confocaux permettent donc la réalisation de "coupes" dans l'échantillon. En pratique, la résolution transverse est donnée par le waist du faisceau d'excitation et la résolution longitudinale par sa longueur de Rayleigh. Le volume observé vaut donc $\sim w_0^4/\lambda$.

Microscopie par onde évanescente

Un second dispositif permettant d'obtenir une résolution axiale est la microscopie par onde évanescente. Cette technique tire partie du phénomène de réflexion totale pouvant se produire lorsque l'angle d'incidence d'un faisceau est trop important lors du passage d'un milieu d'indice fort à un milieu d'indice faible. En effet, d'après les lois de Descartes, lors du passage d'un milieu d'indice n_1 à deuxième d'indice n_2 on a la relation

$$n_1 \sin(i) = n_2 \sin(r)$$

où i et r sont les angles d'incidence et de réfraction. On a par conséquent

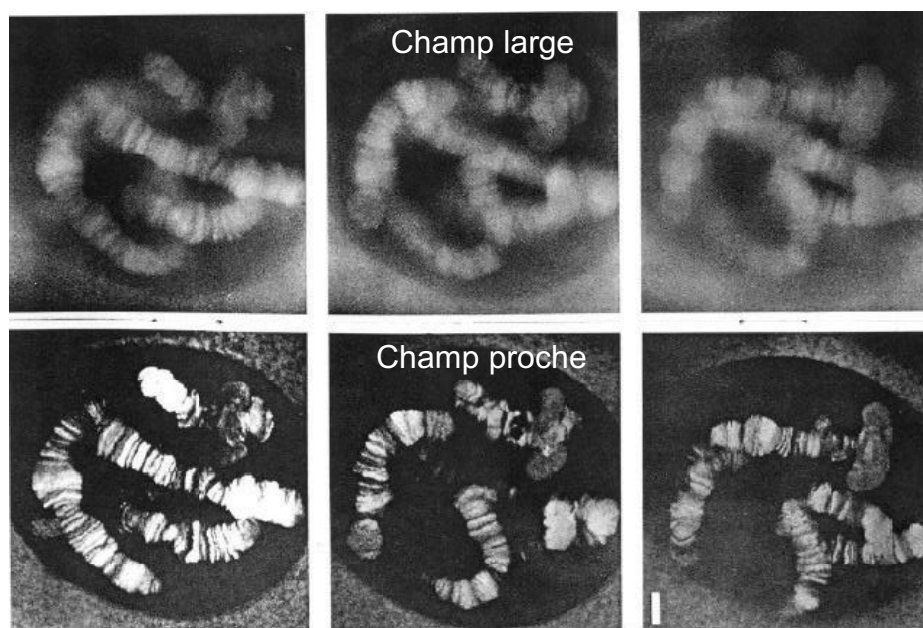


FIG. 3.5 – Exemple d'application de la microscopie confocale : visualisation des chromosomes dans le noyau de glandes salivaires de *Chironomus Thummi* – un insecte – en microscopie classique (haut) et en champ proche (bas). Les coupes sont séparées de $4\ \mu\text{m}$ et la barre blanche correspond à $10\ \mu\text{m}$.

$$\sin(r) = \frac{n_1}{n_2} \sin(i).$$

Pour que cette relation ait un sens, il faut que $\sin(r)$ reste inférieur à 1. Si elle ne l'est pas, il n'existe pas de faisceau transmis, situation effectivement observée si $\sin(i) > n_2/n_1$.

Même la transmission de l'onde est nulle, une partie de l'énergie électromagnétique pénètre quand même dans le milieu sous forme d'onde évanescente. La profondeur de pénétration est de l'ordre de la longueur d'onde, et l'énergie électromagnétique reste donc localisée près de l'interface. Dans ces conditions, l'onde évanescente peut être utilisée pour n'exciter que des molécules proches de la transition entre les deux milieux.

En pratique, l'onde évanescente est réalisée grâce à un prisme ou une fibre optique étirée. Outre la possibilité de réaliser des coupes, l'intérêt de la microscopie en champ proche est la réduction du bruit de fond.

3.2 Optique non-linéaire

3.2.1 Retour sur la théorie des perturbations

Dans le chapitre 1, nous avons montré que si l'on traitait l'interaction entre un atome et un photon au premier ordre de la théorie de perturbation, le seul processus envisageable était l'absorption d'un seul photon résonnant avec la transition atomique. Dans ce paragraphe nous allons montrer qu'à forte intensité lumineuse, l'absorption de deux photons, voire plus, est possible.

Pour cela, repartons d'un hamiltonien $\hat{H} = \hat{H}_0 + \hat{V}$, où $\hat{H}_0$ décrira le système atome + champ sans interaction et $\hat{V} = -\hat{\mathbf{d}} \cdot \hat{\mathbf{E}}$ décrit le couplage entre atomes et photons. Nous avons vu que si l'on décrivait l'état du système par

$$|\psi(t)\rangle = \sum_n \tilde{c}_n(t) e^{-iE_n t/\hbar} |n\rangle,$$

alors les coefficients $\tilde{c}_n$ étaient donnés par une équation différentielle du type

$$i\hbar \dot{\tilde{c}}_m = \sum_n \langle m | \hat{V} | n \rangle e^{-i(E_n - E_m)t/\hbar} \tilde{c}_n. \quad (3.4)$$

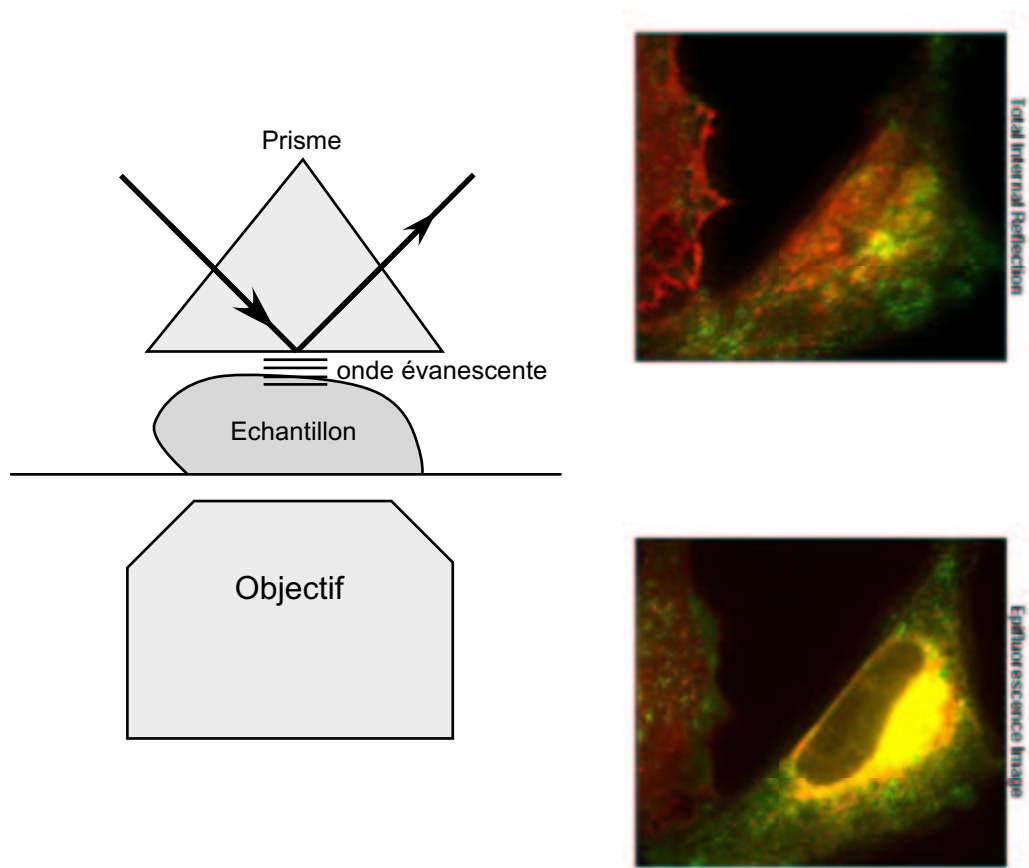


FIG. 3.6 – *Gauche : principe de la microscopie de Fluorescence en champ proche. Droite : comparaison de résultats de microscopie en champs proche (haut) et en champ large (bas). Les couleurs rouge et vertes correspondent aux deux sondes fluorescentes utilisées, On constate que la lumière parasite est considérablement réduite lorsque la technique en champ proche est utilisée.*

Le terme de droite de cette expression est analogue au produit d'une matrice et d'un vecteur. Si l'on introduit formellement le ket $|\tilde{\psi}\rangle = \sum_m \tilde{c}_m |m\rangle$, l'équation (3.4) peut s'écrire comme une équation de Schrödinger

$$i\hbar \frac{d}{dt} |\tilde{\psi}\rangle = \tilde{V}(t) |\tilde{\psi}\rangle, \quad (3.5)$$

où l'opérateur $\tilde{V}(t)$ a pour éléments de matrice

$$\langle m | \tilde{V}(t) | n \rangle = \langle m | \hat{V} | n \rangle e^{-i(E_n - E_m)t/\hbar}.$$

Pour le moment, aucune approximation n'a été faite et l'équation que nous venons d'établir n'est qu'une réécriture de l'équation de Schrödinger. Comme au premier chapitre, l'approche perturbative consiste à remarquer qu'en l'absence du couplage $\hat{V}$, le vecteur $|\tilde{\psi}\rangle$ est constant. Dans le cas où $\hat{V}$ est non nul, mais reste petit, on peut par conséquent supposer que $|\tilde{\psi}(t)\rangle$ reste proche de sa valeur initiale $|\tilde{\psi}_0\rangle$.

Les considérations précédentes suggèrent de poser $|\tilde{\psi}\rangle = |\tilde{\psi}_0\rangle + |\tilde{\psi}_1\rangle$, où $|\tilde{\psi}_1\rangle$ est une correction supposée faible devant $|\tilde{\psi}_0\rangle$. l'équation (3.5) s'écrit alors à l'ordre 1 en perturbation comme

$$i\hbar \frac{d}{dt} |\tilde{\psi}_1\rangle = \tilde{V}(t) |\tilde{\psi}_0\rangle,$$

soit,

$$|\tilde{\psi}_1\rangle = \frac{1}{i\hbar} \int_0^t \tilde{V}(t') |\tilde{\psi}_0\rangle dt',$$

expression qui redonne les $\tilde{c}_n$ obtenus dans la partie sur les perturbations à l'ordre 1.

Si l'on suppose que l'onde électromagnétique n'est résonnante avec aucune transition atomique, ce terme d'ordre 1 de la théorie de perturbation n'induit aucune échange d'énergie entre l'atome et le champ. L'interaction est donc décrite par les termes suivants du développement perturbatif. On doit donc prendre en compte le terme d'ordre 2 et écrire $|\tilde{\psi}\rangle$ sous la forme

$$|\tilde{\psi}\rangle = |\tilde{\psi}_0\rangle + |\tilde{\psi}_1\rangle + |\tilde{\psi}_2\rangle + \dots$$

ce qui nous donne à l'ordre 2

$$i\hbar \frac{d}{dt} |\tilde{\psi}_2\rangle = \tilde{V}(t) |\tilde{\psi}_1\rangle,$$

soit, en utilisant l'expression explicite de $|\tilde{\psi}_1\rangle$:

$$|\tilde{\psi}_2\rangle = \frac{1}{i\hbar} \int_0^t dt_1 \int_0^{t_1} dt_2 \tilde{V}(t_1) \tilde{V}(t_2) |\tilde{\psi}_0\rangle. \quad (3.6)$$

Grâce à la relation $\tilde{c}_n = \langle n | \tilde{\psi}_2 \rangle$, la relation (3.6) permet en principe de calculer la probabilité $P_n = |\tilde{c}_n|^2$ de transition de l'état initial $|\tilde{\psi}_0\rangle$ vers un état final $|n\rangle$. On remarque que la forme particulière de $|\tilde{\psi}_2\rangle$ fait intervenir deux fois le hamiltonien dipolaire, et nécessite donc l'absorption ou l'émission de *deux* photons. Si l'état $|\tilde{\psi}_0\rangle$ correspond à l'atome dans son état fondamental en présence de N photons ($N \gg 1$) de fréquence ω_L , la conservation de l'énergie impose donc d'avoir $2\omega_L = \omega_a$ pour pouvoir transiter vers un état d'énergie $\hbar\omega_a$. Par ailleurs, le calcul de l'élément de matrice $\langle n | \tilde{V} \tilde{V} | \tilde{\psi}_0 \rangle$ fait intervenir deux fois l'opérateur d'annihilation. Il apparaît donc un facteur $\sqrt{N(N-1)} \sim N$ nous donnant une probabilité de transition $|\tilde{c}_n|^2$ proportionnelle à N^2 , et non N comme dans le cas des processus à un photon.

Par analogie avec les transitions à un photon, on peut donc s'attendre sans faire de calcul à ce que le taux d'absorption à deux photons s'écrivent dans le cas d'une radiation de densité spectrale $u(\omega)$:

$$P_{f \rightarrow e}^{(2)} = B_{fe}^{(2)} u^2(\omega_a/2),$$

où $B_{ef}^{(2)}$ est le coefficient d'absorption à deux photons ne dépendant que de l'atome considéré. Exactement comme pour les transitions à 1 photon, on peut définir des coefficients $B_{ef}^{(2)}$ et $A_{ef}^{(2)}$ d'émissions stimulés et spontanée à deux photons, avec $B_{ef}^{(2)} = B_{fe}^{(2)}$.

Exercice : généraliser ce développement à l'ordre n quelconque de la théorie de perturbation et montrer que l'ordre n correspond à des absorptions simultanées à n photons.

La première application de ces effets non linéaire est la possibilité de modifier la longueur d'onde d'un laser de manière cohérente. En effet, en éclairant un milieu avec une pulsation ω_1 , il est possible d'émettre de la lumière à la pulsation $2\omega_1$. C'est par exemple ce procédé qui est utilisé pour réaliser des lasers verts par doublage de fréquence d'un laser YAG (on passe ainsi de 1 μm à 500 nm).

Plus généralement, si plusieurs radiations éclairent le milieu non linéaire, il est possible de réaliser des mélanges à plusieurs fréquences. Si on éclaire avec deux fréquences ω_1 et ω_2 , on peut créer la fréquence $\omega_1 + \omega_2$.

3.2.2 Applications

Spectroscopie non-linéaire

Les techniques usuelles de fluorescence font appel à des transitions à un photon. Avec le développement de lasers de plus en plus puissants, de nouvelles sondes fluorescentes, utilisant des transitions à deux photons, sont aujourd'hui disponibles. Deux propriétés fondamentales les rendent particulièrement intéressantes.

1. Sélectivité spatiale. L'absorption à deux photons est proportionnelle à l'intensité lumineuse et la fluorescence ne se manifeste qu'aux points de forte focalisation.
2. Sélectivité spectrale. La lumière d'excitation est en général bien plus puissante que la lumière de fluorescence et il peut être très délicat de séparer l'une de l'autre sur les détecteurs (en particulier si ceux-ci sont déjà saturés par la lumière de pompe). Or, dans le cas de la fluorescence à deux photons, l'absorption (à deux photons) et la désexcitation (à un seul photon) se font à des longueurs d'ondes différentes. En utilisant des miroirs dichroïques, il est ainsi possible de séparer la lumière de pompage et la lumière de fluorescence sur le photodétecteur, facilitant ainsi la mesure de l'intensité de fluorescence.

Oscillateur paramétrique optique

Le phénomène d'émission stimulée non linéaire permet la réalisation d'oscillateurs paramétriques optiques (OPO), qui constituent une alternative au laser en tant que source lumineuse cohérente. Tout comme le laser, un OPO est constitué d'une cavité résonnante dans laquelle est placé un milieu non linéaire pompé par une onde de pulsation ω_1 . Par "fluorescence" stimulée, un photon de pompe est converti en deux photons de pulsations ω_2 et ω_3 . Ces deux fréquences sont conditionnées par les contraintes suivantes

- *Résonance avec un mode de la cavité.* Il faut, comme pour le laser, que les modes émis par l'OPO soient résonnants avec la cavité, i.e. que les chemins optiques parcourus dans la cavité par les faisceaux 2 et 3 soient des multiples entiers de leur longueur d'onde.
- *Condition d'accord de phase.* Cette seconde condition résulte de la conservation de l'impulsion lors de la conversion du photon de pompe.

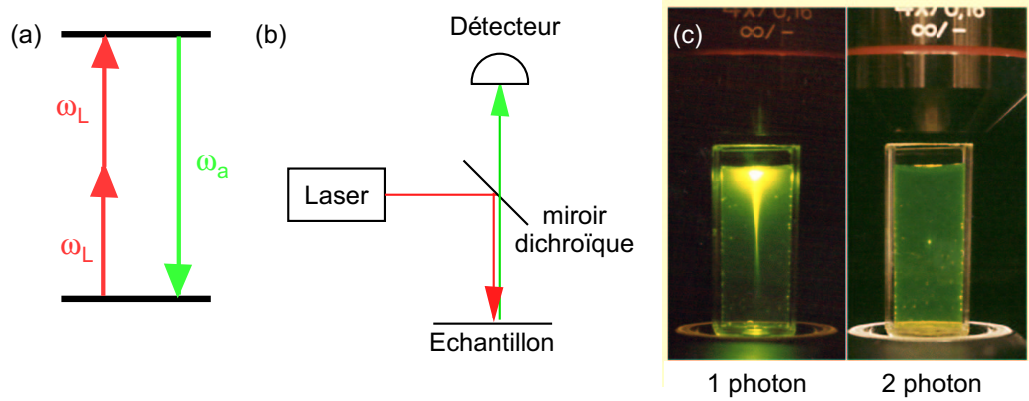


FIG. 3.7 – a) Principe de la fluorescence à deux photons. On excite une transition à l'aide d'une radiation de pulsation $\omega_L \sim \omega_a/2$. Si l'intensité lumineuse est suffisante, une transition à deux photons peut se produire vers l'état excité. L'atome se désexcite ensuite en émettant un photon de pulsation ω_a . b) Visualisation de la fluorescence à deux photons. La lumière émise par le laser (en général un laser à impulsion de façon maximiser l'intensité lumineuse crête) est réfléchié sur un miroir dichroïque puis focalisée sur l'échantillon fluorescent. La lumière de fluorescence, de longueur d'onde différente de celle d'excitation, est transmise par le miroir et focalisée sur le détecteur. c) Comparaison de la sélectivité spatiale des sondes fluorescentes à un et deux photons.

Si l'on note $\mathbf{k}_i$ le vecteur d'onde du faisceau de pulsation ω_i , on doit avoir

$$\mathbf{k}_1 = \mathbf{k}_2 + \mathbf{k}_3.$$

Puisque les trois modes se propagent dans la même cavité, les vecteurs d'ondes k_i sont parallèles. En introduisant l'indice $n_i = n(\omega_i)$, on obtient

$$n_1\omega_1 = n_2\omega_2 + n_3\omega_3.$$

Ces trois conditions (deux conditions de résonance avec la cavité et une condition d'accord de phase) fixent les trois valeurs ω_i . Autrement dit, une fois la géométrie de cavité et le milieu non-linéaire choisis, la fréquence de pompe est déterminée de façon univoque! Cette particularité rend les OPO très sensibles aux vibrations, qui tendent à modifier les conditions de résonances.

Leur intérêt pratique est de fournir des sources cohérentes accordables sur de larges plages de longueur d'ondes, la plage d'accordabilité dépendant essentiellement de l'importance du couplage non linéaire dans le milieu. Un second intérêt est apparu plus récemment. On a en effet montré que le rayonnement émis par un OPO avait la particularité de générer un faisceau comprimé, c'est-à-dire un mode dont les fluctuations de phase ou de nombres de photons sont réduites par rapport à un rayonnement laser standard.

Effet Kerr

En électromagnétisme classique, l'absorption d'un milieu peut s'interpréter comme la partie imaginaire de l'indice optique. Une dépendance de l'absorption avec l'intensité lumineuse I peut donc généraliser à une variation $n(I) \sim n_0 + n_2 I$ de l'indice avec I , ce que l'on appelle *l'effet Kerr*.

- *Automodulation de phase.* La première conséquence de l'effet Kerr est l'automodulation de phase. Dans le cas d'un laser à impulsion l'indice n vu par l'onde électromagnétique varie avec le temps, du fait de la variation temporelle de l'intensité lumineuse. Ceci provoque un déphasage $\delta\phi = n_2 I(t) k_0 z$. Ce déphasage temporel correspond à l'élargissement du spectre de l'impulsion, par l'ajout de fréquence $\delta\omega = \dot{\delta\phi}$.

Cette technique est couramment utilisée pour la génération de continuum "cohérents" de lumière, comme on l'observe sur la fig. 3.8.

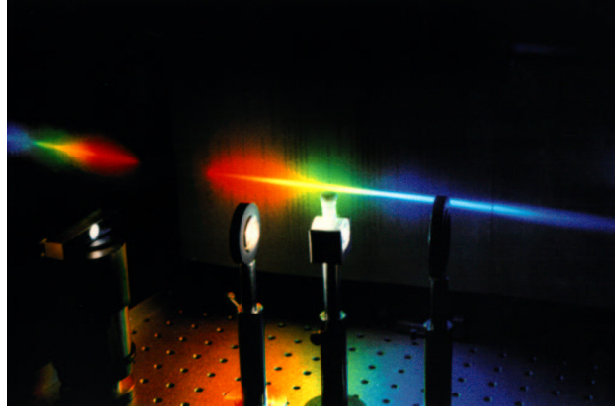


FIG. 3.8 – Génération de continuum à l'aide d'impulsion femtosecondes.

- *Soliton optique*: On remarque que dans l'effet d'automodulation de phase, les nouvelles fréquences apparaissant dans l'impulsion ne sont pas localisées spatialement de la même façon : les hautes fréquences sont en avant, et les basses fréquences en arrière. On parle d'impulsion *chirpée* (de l'anglais *chirp* qui signifie gazouiller). En combinant cet effet avec la dispersion naturelle du milieu de propagation, il est possible de réaliser des impulsions se propageant sans déformation, ce que l'on appelle des *solitons*. En effet, la dispersion tendant à retarder les hautes fréquences par rapport aux basses, il existe une intensité pour laquelle dispersion et automodulation se compensent. Cette propriété est particulièrement intéressante dans ses applications aux télécommunications numériques optiques. En effet, dans les systèmes habituelles, chaque bit est codé par une impulsion optique. Lors du transport du signal lumineux, la dispersion du milieu constituant la fibre optique provoque un étalement du paquet d'onde qui constitue une limite au débit des lignes par fibres optiques : si l'on fait se propager sur de grandes distances des impulsions émises à un taux très élevé, deux impulsions successives vont se chevaucher, ce qui va aboutir à un brouillage du message (Fig. 3.9.a).

Les solitons constituent un phénomène caractéristique de la propagation dans des milieux non-linéaire. Historiquement, Lord Russel fut le premier à étudier leur physique, suite à une course à cheval à la poursuite d'une telle onde solitaire remontant l'Union Canal près d'Edinburgh (Fig. 3.9.b,c).

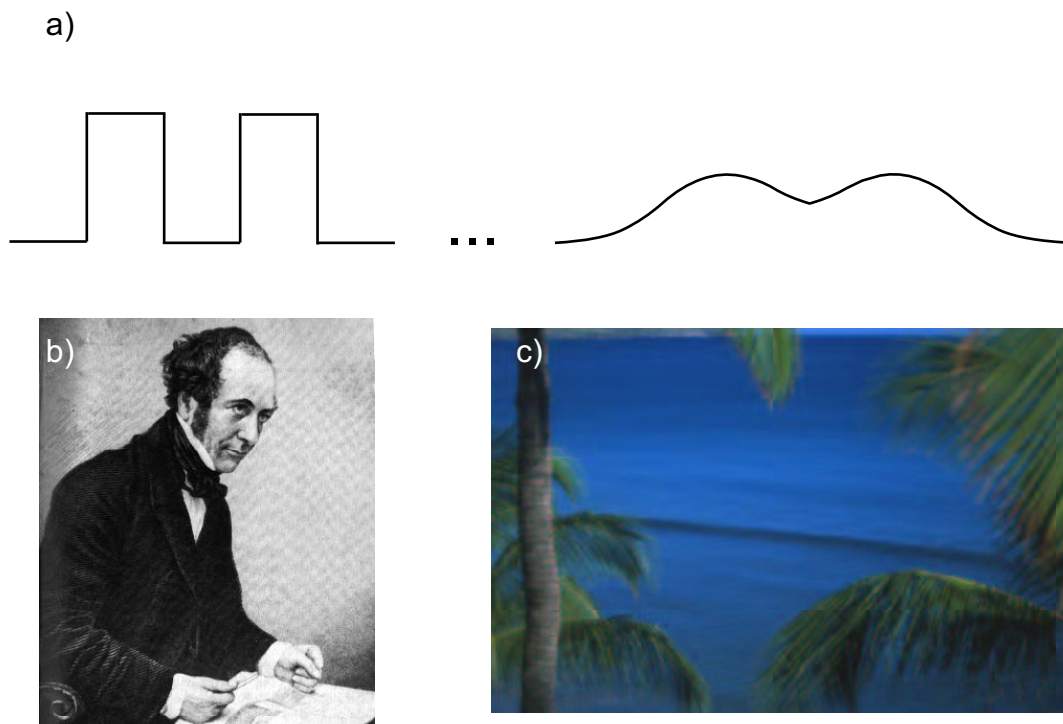


FIG. 3.9 – a) *Etalement du paquet dans les télécommunications numériques.* b) *Lord Russel, premier à avoir étudié la propagation des solitons.* c) *Exemple de soliton naturel : vague sur la plage de Maui.*

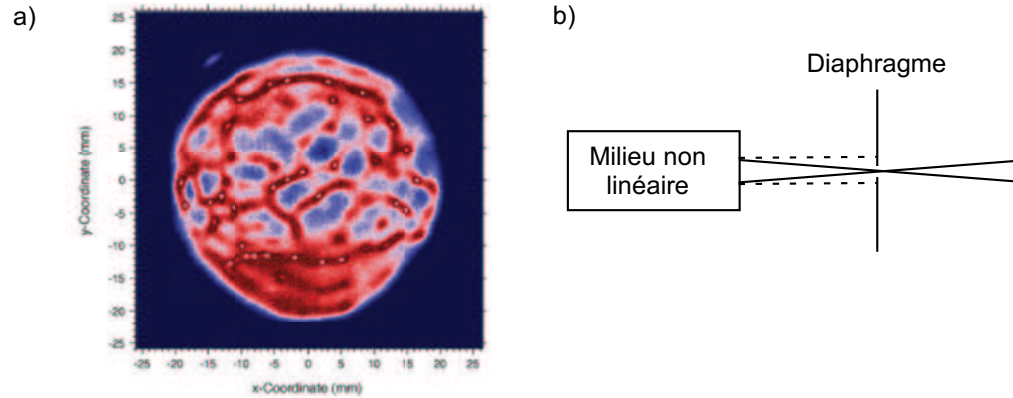


FIG. 3.10 – a) Profil transverse d'un faisceau après propagation dans l'air. L'intensité crête de l'impulsion est de 3.56 TW et la distance de propagation de 10 m . b) Blocage de mode par lentille Kerr. Après propagation dans un milieu non linéaire, les modes intenses subissent l'effet d'auto-focalisation (trait plein). En plaçant un diaphragme dans une cavité laser, on bloque le mode de fonctionnement continu et peu focalisés (trait pointillé), de façon à permettre le démarrage du fonctionnement pulsé.

- *Lentille Kerr*: Un troisième effet est l'auto-focalisation. Lorsque la lumière se propage l'indice vu par les bords du faisceaux est plus faible qu'au centre. Ce gradient spatial d'indice est en tout point équivalent au passage au travers d'une lentille, où le bord du faisceau traverse moins de verre que le centre. Le faisceau est donc focalisé par la traversée d'un milieu homogène !

Dans de nombreux cas, ce phénomène peut être néfaste. L'intensité au foyer d'autofocalisation peut en effet devenir trop intense et endommager le matériau. Par ailleurs, pour de fortes non linéarités, le profil spatial du faisceau devient instable et se filamente comme on l'observe sur la figure (3.10.a).

L'effet d'autofocalisation est utilisé dans les cavités des lasers Titane-Saphir pour amorcer le fonctionnement impulsionnel. En effet, en ajoutant un diaphragme au foyer d'autofocalisation, on ne s'arrange pour ne laisser passer que le mode de plus grande intensité lumineuse crête, c'est à dire le mode impulsionnel (Fig. 3.10.b).

Peignes de fréquences

Un des problèmes de la métrologie est de pouvoir mesurer des fréquences optiques ($\sim 10^{15}$ Hz), à partir de l'étalon de temps défini par la transition à 9.2 GHz du césium. Jusqu'à récemment, cette opération nécessitait l'utilisation de chaînes de fréquences impliquant plusieurs étapes de doublement pour, partant d'une horloge à césium, pouvoir comparer un signal de fréquence contrôlée avec la fréquence optique à mesurer (fig. 3.12). Afin de simplifier la procédure, T. W. Hänsch, physicien de Munich lauréat du prix Nobel 2005, a développé une stratégie originale permettant, à l'aide d'un laser à impulsions, de franchir en une seule étape les six ordres de grandeur séparant les domaines optiques et micro-ondes. L'idée directrice se fonde sur la particularité du spectre des lasers à mode bloqués de se présenter sous la forme de peignes de fréquences. On a pu montrer expérimentalement que la séparation, donnée par le taux de répétition des impulsions, était très stable, ce qui en fait d'excellentes "règles" à fréquence. Le taux de répétition de l'ordre de la centaine de MHz peut être facilement mesurée à l'aide d'une horloge atomique (le standard de temps étant défini par la transition hyperfine à 10 GHz du césium) et stabilisé par une boucle de rétroaction sur la longueur de la cavité laser. Cette propriété avait déjà été utilisée dès 1978 par T.W. Hänsch pour la mesure fine d'écarts de fréquences (en l'occurrence la largeur du doublet sodium, J.N. Eckstein *et al.*, Phys. Rev. Lett., **13**, 847 (1978)). Malgré cette première réussite, la mesure de fréquences absolues se heurtait à l'impossibilité de repérer le zéro du peigne de fréquence (par exemple, le laser titane saphir ne couvre que l'intervalle 700-1000 nm). L'idée permettant de contourner cet obstacle est de mesurer l'écart de fréquence entre f_0 (la fréquence à mesurer) et sa première harmonique $2f_0$, obtenue après passage dans un cristal doubleur. Dans ce cas, il suffit d'utiliser une règle de fréquence s'étendant sur une octave seulement. Pour réaliser cette règle, on part d'impulsions issues d'un titane saphir, dont on élargit le spectre par génération de continuum dans une fibre optique choisie de façon à conserver la stabilité fréquentielle du peigne. Le peigne et les faisceaux à f_0 et $2f_0$ sont ensuite mélangés puis dispersés par un réseau. L'écart de fréquence $2f_0 - f_0$ est ensuite mesuré en comptant le nombre de dents du peigne les séparant (Fig. 3.11). Couplé aux progrès de la précision des horloges optiques, ce dispositif de comparaison de

fréquence permet d'envisager une nouvelle définition du temps, utilisant à présent des transitions optiques plutôt que micro-onde.

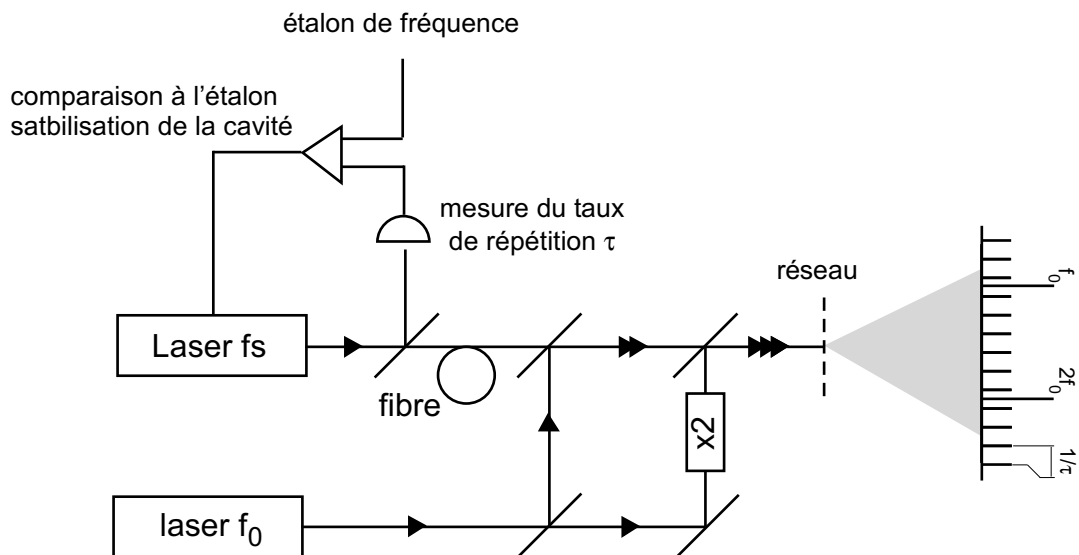


FIG. 3.11 – Principe de la mesure de fréquence optique à l'aide d'un peigne de fréquence.

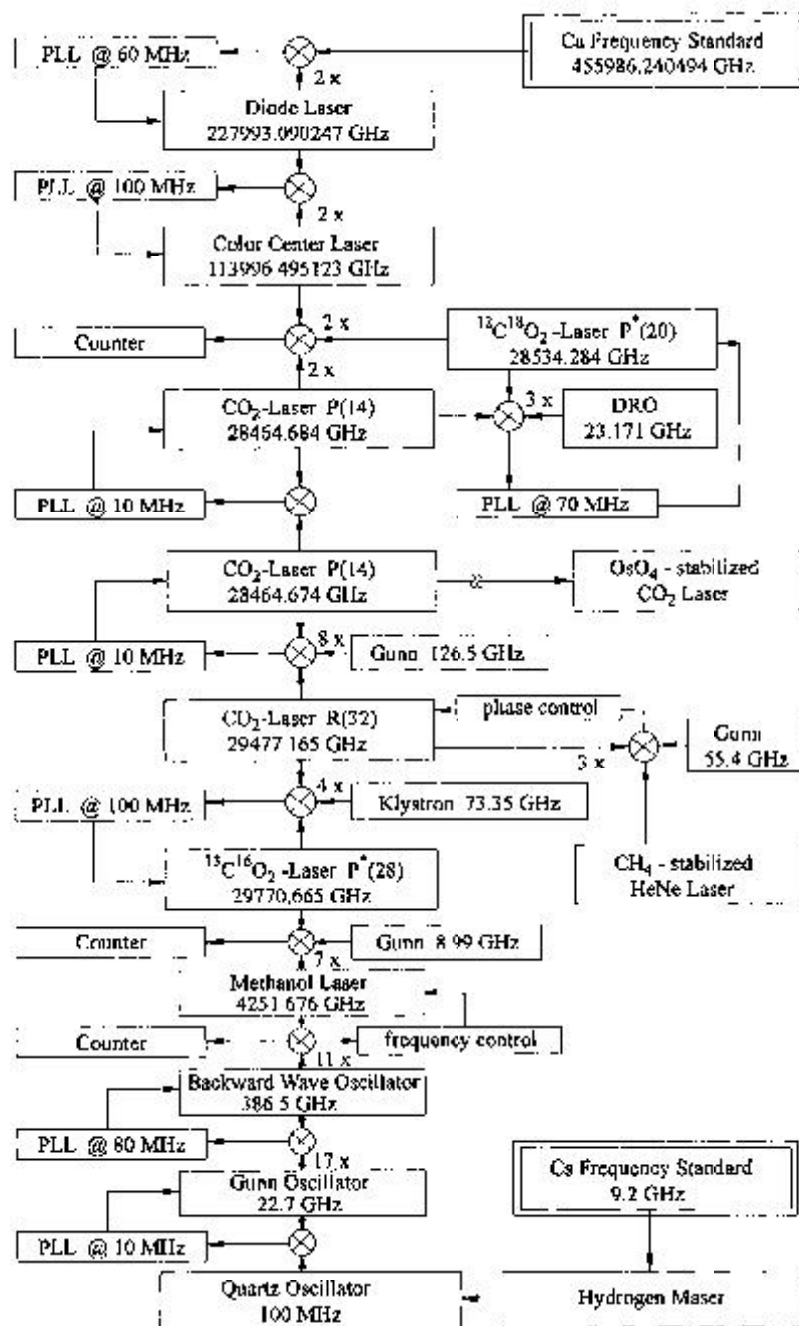


FIG. 1. PTB's frequency chain to the Ca intercombination line (PLL = phase locked loop, details are given in the text).

FIG. 3.12 – Exemple de chaîne de fréquence utilisée pour la comparaison d'une horloge à césium et des raies du calcium (figure extraite de H. Schnatz et al. *Phys. Rev. Lett.* **76**, 18 (1996)).

Chapitre 4

Forces radiatives

4.1 Force dipolaire et pression de radiation

La manifestation la plus spectaculaire de l'effet mécanique de la lumière sur la matière est l'effet du rayonnement solaire qui repousse la queue des comètes comme on l'observe sur la figure 4.1).

En pratique, on distingue deux types de forces. Les forces dipolaires sont des forces conservatives, correspondant à des transferts cohérents de photons d'un mode peuplé à un autre par absorption et émission stimulée. Les forces de pressions de radiation, qui sont à l'œuvre dans le cas de la queue des comètes, sont associées à des cycles absorption-émission spontanée et donnent lieu à une force dissipative. Ce chapitre est consacré à l'étude de ces deux phénomènes dont nous discuterons quelques applications.

4.1.1 Force dipolaire

Cas des atomes

Retournons au modèle quantique de l'interaction d'un atome avec un unique mode du champ électromagnétique polarisé selon z et peuplé par N photons. Si l'on néglige les phénomènes spontanés couplant l'atome aux modes vides du champ, le hamiltonien du système s'écrit à l'approximation dipolaire

$$\hat{H} = \hbar\omega_a|e\rangle\langle e| + \hbar\omega_L\hat{a}^\dagger\hat{a} - \sqrt{\frac{\hbar\omega_L}{2\epsilon_0 V}}d_z(\hat{a}^\dagger|f\rangle\langle e| + \hat{a}|e\rangle\langle f|),$$



FIG. 4.1 – *Effet de la lumière sur la queue d’une comète. On observe en réalité deux queues. La queue du bas, jaunâtre et légèrement incurvée correspond à des poussières repoussées par la pression de radiation. La queue bleuâtre est constituée interagissant préférentiellement avec le vent solaire, i.e. un plasma d’ions et d’électrons éjecté par le Soleil. Sa couleur provient de la fluorescence des ions CO^+ .*

où l'on a repris les notations du premier chapitre. Dans ce hamiltonien, les espaces $\mathcal{E}_N$ engendrés par les états $|f, N\rangle$ et $|e, N-1\rangle$ sont découplés, ce qui correspond à une dynamique “conservative”, dans laquelle le nombre de photons dans le mode du laser reste constant.

La restriction $\hat{H}_N$ de $\hat{H}$ à $\mathcal{E}_N$ a alors pour expression matricielle

$$\hat{H}_N = \begin{pmatrix} N\hbar\omega_L & -\hbar\Omega\sqrt{N} \\ -\hbar\Omega\sqrt{N} & (N-1)\hbar\omega_L + \hbar\omega_a \end{pmatrix}$$

avec $\hbar\Omega = d_z \sqrt{\hbar\omega_L/2\epsilon_0 V}$.

Les valeurs propres de $\hat{H}_N$ peuvent être calculées analytiquement et l'on obtient

$$E_{\pm} = N\hbar\omega_L - \frac{\hbar\Delta}{2} \pm \sqrt{N\hbar^2\Omega^2 + \hbar^2\Delta^2/4},$$

où $\Delta = \omega_L - \omega_a$ désigne le désaccord à résonance.

L'allure du spectre de $\hat{H}_N$ en fonction de Δ est représenté sur la figure 4.2. On constate que pour des désaccords importants, les états sont relativement peu perturbés, ce qui correspond à un couplage faible entre l'atome et le rayonnement.

À l'ordre dominant en perturbation, l'énergie du niveau $|f, N\rangle$ est modifiée et vaut

$$E(\Delta) \sim N\hbar \frac{\Omega^2}{\Delta}.$$

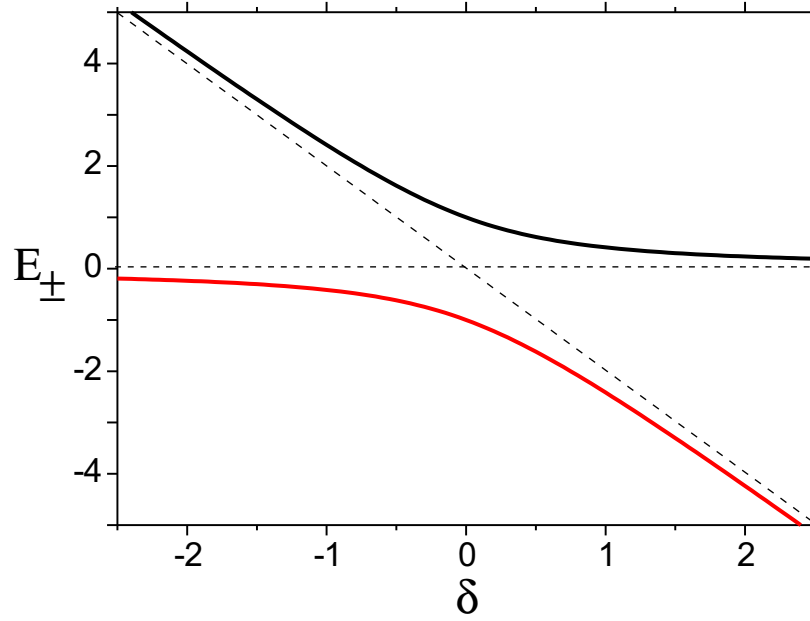
En utilisant l'expression de Ω , on voit apparaître la quantité $N\hbar\omega_L/V$, qui est en fait la densité d'énergie électromagnétique ϖ . En introduisant l'intensité lumineuse $I_0 = \varpi c$, on obtient finalement

$$E(\Delta) \sim \frac{d_z^2}{2\epsilon_0 \hbar c} \frac{I_0}{\Delta}.$$

Conventionnellement, on introduit l'intensité de saturation I_{sat} telle que

$$E(\Delta) \sim \hbar A \frac{A}{\Delta} \frac{I}{8I_{\text{sat}}}.$$

Dans le cas où l'intensité lumineuse est inhomogène (par exemple dans le cas d'un faisceau gaussien) et dépend de la position de l'atome, cette relation signifie que l'atome évolue dans un potentiel extérieur $U(\mathbf{r})$ donné par

FIG. 4.2 – Spectre de $\hat{H}_N$

$$U(\mathbf{r}) = \hbar A \frac{A I(\mathbf{r})}{\Delta I_{\text{sat}}},$$

ce qui aboutit à une force

$$\mathbf{F} = -\mathbf{grad}(U) = -\hbar A \frac{A \mathbf{grad}(I(\mathbf{r}))}{\Delta I_{\text{sat}}}.$$

Pour $\Delta > 0$, cette force attire l'atome vers les régions de faible intensité lumineuse. Au contraire, pour $\Delta < 0$, l'énergie d'interaction est négative et l'atome tend à rejoindre les régions de forte intensité.

À l'échelle microscopique, l'existence de cette force *dipolaire*¹ fournit un outil particulièrement flexible de manipulation des atomes à l'aide de la lumière. Deux applications spectaculaires en sont le miroir à atomes, réalisé grâce à une onde évanescente sur laquelle rebondit les atomes, et les réseaux optiques, dans lesquels on piège des atomes aux nœuds de la figure d'interférence de plusieurs faisceaux 4.3.

1. Cette force est ainsi nommée car en physique classique, l'énergie U que nous venons d'écrire correspond à l'énergie d'interaction $-\mathbf{d} \cdot \mathbf{E}$ entre un dipôle et un champ électrique.

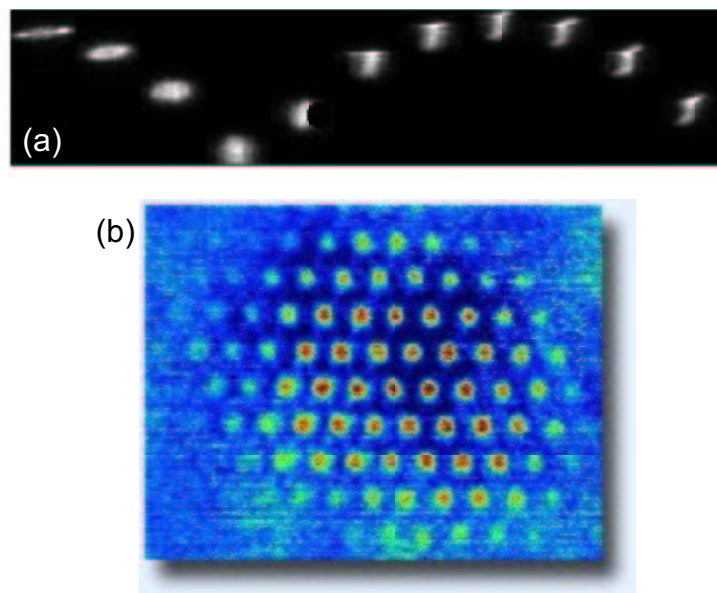


FIG. 4.3 – Exemple d'action de la lumière sur les atomes. a) Condensat de Bos-Einstein rebondissant sur un miroir à onde évanescente. b) Atomes piégés au nœuds de la figure d'interférence de plusieurs faisceaux.

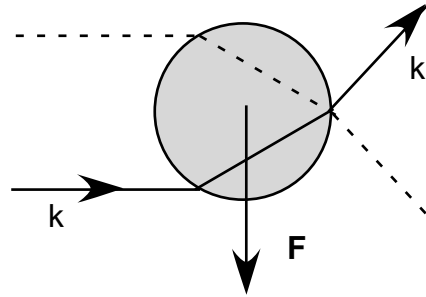


FIG. 4.4 – *Interprétation macroscopique de la force dipolaire : du fait de la deflection de la lumière par la bille, la quantité de mouvement des photons est modifiée de $\hbar\mathbf{k}$ à $\hbar\mathbf{k}'$. Par conservation de la quantité de mouvement, il en résulte donc une force sur la matière. Dans le cas d'un éclairage uniforme (faisceau en trait plein égal en intensité au faisceau pointillé), les forces exercées par les différents rayons se compensent. Dans le cas d'un gradient d'intensité (trait plein plus intense que le trait pointillé), la force dipolaire tend à attirer la bille vers les régions de forte intensité.*

Cas des objets “macroscopiques”

Bien que les résultats soient qualitativement semblables, le calcul de la force dipolaire sur un objet méso- ou macroscopique est compliqué par la modification du champ électrique diffusé par les autres atomes.

Dans le cas d'objets dont la taille est de l'ordre de la longueur d'onde λ , la théorie complète est particulièrement complexe, mais elle se simplifie considérablement dans le cas où la taille devient très grande devant λ et où l'approximation de l'optique géométrique devient valable. En effet, dans ce cas le transfert d'impulsion du rayonnement à la matière provient de la déflexion des photons à la traversée des interfaces entre le milieu environnant et l'objet (figure 4.4).

Cet effet est illustré sur la figure 4.5, dans laquelle la force dipolaire exercée par un faisceau laser focalisé au voisinage de la frontière entre deux fluides provoque une déformation de l'interface².

2. Figure extraite de A. Casner et J.-P. Delville, J. Phys. IV France **12**, 85 (2002), et Ulm PC 2004.

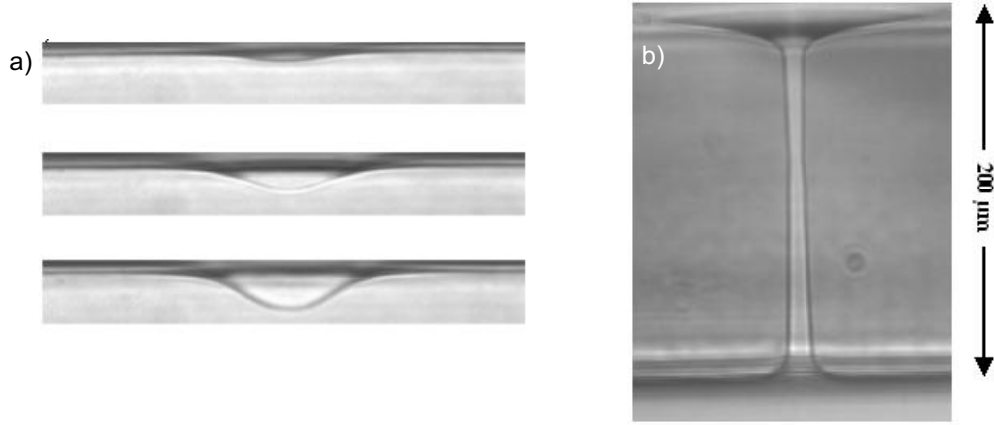


FIG. 4.5 – a) Déformation d’une interface sous l’effet de la force dipolaire. De haut en bas, la puissance lumineuse a pour valeurs 270, 540 et 810 mW respectivement. b) Formation d’un pont liquide : on focalise le faisceau sur deux interfaces. Lorsque la déformation est suffisamment longue, les protubérances fusionnent de façon à former un pont liquide.

4.1.2 Pression de radiation

La pression de radiation est due aux cycles d’absorption et d’émission spontanée. C’est un processus non conservatif (on dit aussi dissipatif), car le nombre de photons dans les modes initialement peuplés n’est pas conservé.

Considérons un atome de quantité de mouvement $\mathbf{P}$ éclairé par un faisceau lumineux constitué de photons placé dans un mode de vecteur d’onde $\mathbf{k}_L$. Lorsque l’atome absorbe un photon, il reçoit à la fois son énergie et sont impulsions. Après absorption, la quantité de mouvement de l’atome vaut

$$\mathbf{P}' = \mathbf{P} + \hbar \mathbf{k}_L.$$

Si l’intensité du faisceau incident est suffisamment faible, la désexcitation de l’atome se fait par émission spontanée. Le photon est alors émis avec un vecteur d’onde $\mathbf{k}'_L$ de direction aléatoire. Après un cycle complet absorption-émission, l’atome a pour quantité de mouvement

$$\mathbf{P}'' = \mathbf{P} + \hbar \mathbf{k}_L + \hbar \mathbf{k}'_L.$$

Ces cycles se produisent avec un taux $\Gamma = Bu(\omega_a)$. En moyennant sur un

temps long devant Γ^{-1} , on trouve par conséquent que la variation de quantité de mouvement de l'atome est donnée par

$$\frac{d\mathbf{P}}{dt} \sim \Gamma(\mathbf{P}'' - \mathbf{P}) = \Gamma\langle\hbar\mathbf{k}_L + \hbar\mathbf{k}'_L\rangle,$$

où $\langle\cdot\rangle$ désigne la moyenne temporelle. Or, la direction de $\mathbf{k}'_L$ est aléatoire, on a $\langle\mathbf{k}'_L\rangle = 0$ et $\mathbf{k}_L$ est constant. On en déduit que puisque l'émission stimulée est par hypothèse négligée

$$\frac{d\mathbf{P}}{dt} = \Gamma\hbar\mathbf{k}_L = Bu(\omega_a)\hbar\mathbf{k}_L.$$

Sous l'effet de cette force, l'atome est par conséquent poussé par le faisceau : c'est ce qui explique la forme de la queue des comètes. Plus futuriste, il avait un temps (au début des années 90) été imaginer de réaliser une course à la voile de la Terre à la Lune. Les voiles réfléchissantes en questions auraient alors utilisé la pression de radiation solaire comme force motrice.

4.2 Applications

4.2.1 Pincettes optiques

Les forces radiatives sont en générale trop faibles (de l'ordre du pN) pour avoir un effet sur des objets de taille macroscopique. On peut en revanche les utiliser pour manipuler des particules microscopiques à l'aide de ce que l'on appelle des pincettes optiques (*optical tweezers* en anglais).

Ces pincettes sont utilisées pour déplacer des billes sub-micrométriques qui peuvent servir à la réalisation de micromachines. Un exemple de telle réalisation est présentée sur la figure 4.6 où un micro-rotor usiné par réaction-photochimique est mis en rotation par la force dipolaire exercée par la lumière, et entraîne les deux micro-engrenages avec lesquels il est lié³. Ces dispositifs peuvent notamment être utilisés dans des applications de type microfluidiques pour réaliser des valves ou des pompes dans des microcanaux.

En biophysique, ces mêmes pincettes sont utilisées afin d'étudier les propriétés mécaniques de l'ADN, l'action d'enzymes ou le fonctionnement de moteurs moléculaires naturels (tels que les filaments d'actine/myosine dans les muscles).

3. Voir D.G. Grier, Nature **424**, 811 (2003).

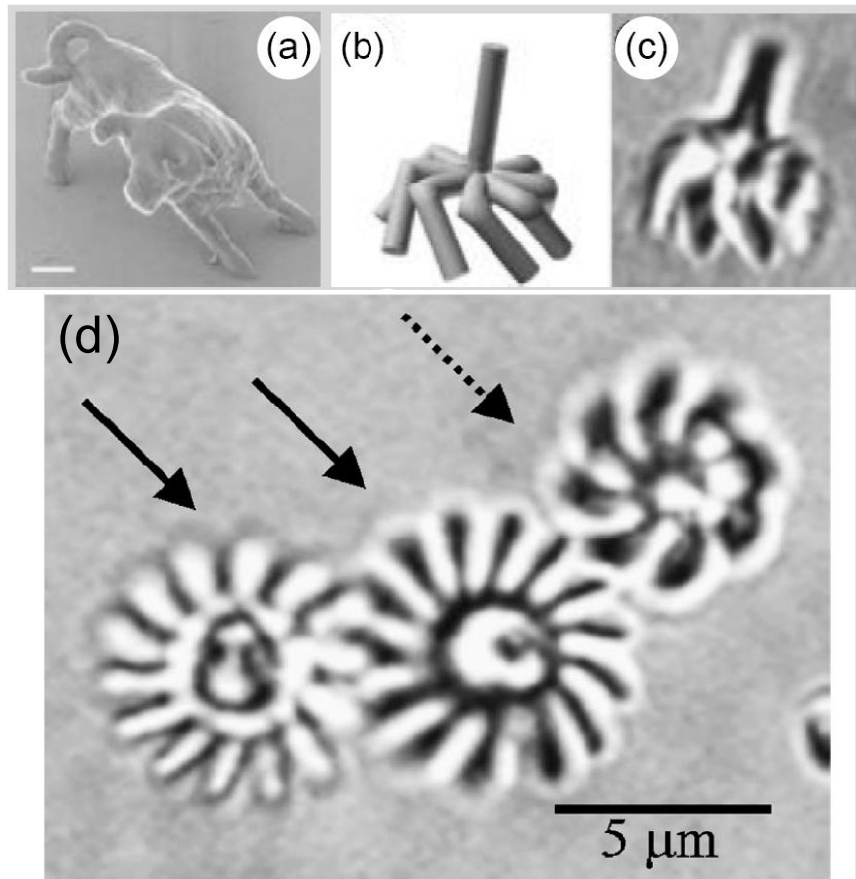


FIG. 4.6 – *Exemple de micromachine actionnée par laser. a-c) En combinant des techniques d'holographie et de photopolymérisation à deux photons, il est possible de réaliser des structures de taille micrométrique, ici un chien (a), un rotor (b : théorie, c : image du rotor réel). Ce rotor peut être combiné avec deux engrenages (d) de façon à réaliser un micromoteur actionné par force dipolaire.*



FIG. 4.7 – Piège magnéto-optique. La boule rouge au centre de l'image correspond à un nuage de 10^9 atomes de lithium refroidis à des températures de l'ordre du mK.

4.2.2 Refroidissement d'atomes par laser

À une échelle encore plus réduite, les accélérations subies par un atome sous l'effet de la pression de radiation peuvent être considérables. En effet, l'impulsion de recul $\hbar k$ associée à l'absorption d'un photon correspond à une variation de vitesse de l'ordre du cm/s (variable d'un atome à l'autre). À résonance, le taux d'émission de photons est limité par l'émission spontanée et l'accélération subie par l'atome vaut par conséquent $a \sim A\hbar k/m \sim 10^4 g$! Cette accélération colossale peut être mise à profit pour ralentir et refroidir des atomes, technique développée dans les années 80 et qui valut à ses découvreurs (C. Cohen-Tannoudji, W. Philips et S. Chu) le prix Nobel de physique 1997. Cette méthode utilise la pression de radiation pour piéger et refroidir des atomes à des températures de l'ordre du microkelvin.

La mélasse optique, qui utilise la combinaison de la pression de radiation et de l'effet Doppler est un outil à présent courant dans les laboratoires de physique atomique. Son principe consiste à éclairer les atomes avec trois paires de faisceaux contre-propageant. Ces faisceaux ont même fréquence ω_L , choisie inférieure à la fréquence de transition atomique.

Considérons dans un premier temps le cas d'une seule paire de faisceaux. Si l'atome est immobile, les forces de pression de radiations exercées par les deux faisceaux se compensent par symétrie et la force totale est nulle. Si l'atome se déplace à une vitesse $\mathbf{v}$, l'effet Doppler décale la fréquence du faisceau contre-propageant vers la résonance, alors que le faisceau co-propageant est repoussé de résonance. En conséquence, la force de pression de radiation exercée par le faisceau contre-propageant est plus forte que celle du faisceau propageant, ce qui provoque une force nette opposée au mouvement de l'atome.

Cet effet s'interprète comme une force de friction s'opposant au mouvement de l'atome et tendant à le ralentir, d'où le nom de "mélasse optique" donné à ce dispositif. La vitesse la plus basse, et donc la température limite, accessible dans une mélasse est limitée par l'émission spontanée inhérente au processus de pression de radiation. La quantité de mouvement des atomes ne peut en effet pas être diminuée en dessous du recul associé à un photon et on voit donc que la température limite du refroidissement par laser est donnée par

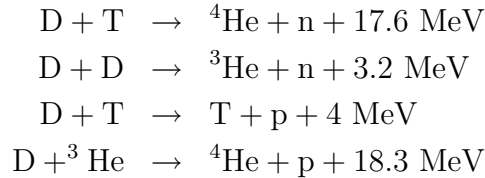
$$k_B T_R \sim \frac{\hbar^2 k_L^2}{2m},$$

ce qui, dans le cas d'un atome tel que le césium donne une température de 100 nK.

4.2.3 Fusion inertielle

En comparaison des réactions chimiques dégageant des énergies de l'ordre de l'eV, les réactions nucléaires, dont l'échelle est de l'ordre du MeV, fournissent de bien meilleures sources de puissance par kilogramme. La filière nucléaire est actuellement fondée sur la technique de fission, qui consiste à briser en deux les noyaux d'atomes lourds, tels que l'uranium. Cette technique qui se veut une alternative sérieuse aux combustibles fossiles issus du pétrole se heurte cependant à deux problèmes. Le premier, d'ordre économique, est que les réserves d'uranium, comme celles de pétrole, ne sont pas inépuisables. Le second, écologique, est que les produits de la fission contiennent des nucléides radioactifs de très longue durée de vie (tels que ^{90}Sr ou ^{137}Cs dont les durées de vies peuvent atteindre jusqu'à plusieurs millions d'années) dont le stockage (ou le retraitement éventuel) pose un problème majeur.

Une alternative aux réacteurs à fission est le développement de la fusion contrôlée. Les réactions en jeu se fondent sur la fusion de noyaux d'hydrogène (deutérium D ou tritium T) suivant les réactions suivantes



Dans cette voie, qui est aussi celle fournissant l'énergie au Soleil, les réactifs sont issus de l'hydrogène et donc virtuellement inépuisables⁴. Par ailleurs, cette série de réactions ne produit aucune substance radioactive, ce qui résout le problème de stockage des déchets.

Au contraire de la fission, la fusion ne se réalise pas spontanément. En effet, si l'on désire approcher deux noyaux à des distances de l'ordre du femtomètre (la portée des interactions fortes assurant la cohésion nucléaire), il est d'abord nécessaire de franchir la barrière de potentielle engendrée par

4. Avec une petite restriction pour le tritium, radioactif, qui ne peut être obtenu qu'à partir de la désintégration du lithium, qui reste cependant 15 fois plus abondant que l'uranium dans l'écorce terrestre.

la répulsion coulombienne entre deux noyaux chargés positivement. Cette contrainte nécessite de porter les réactifs à des températures de plusieurs centaines de millions de degrés (la température au centre d'une étoile). Ces températures extrêmement élevées posent des problèmes de confinement, puisqu'aucune paroi matérielle ne résiste à une telle température. Une première solution à cet obstacle est le confinement magnétique dans les Tokamak. C'est cette solution qui est au cœur du projet ITER (International Thermonuclear Experimental Reactor) devant être développé en commun par l'Union Européenne, les États-Unis, la Russie la Chine et le Japon. La seconde solution est le confinement inertiel.

Dans la stratégie de fusion inertielle, les réactifs sont stockés dans des capsules (de l'ordre du millimètre de diamètre) que l'on illumine à l'aide de plusieurs faisceaux de haute puissance. Sous l'effet de la pression de radiation, les parois de la capsule se désintègre et génèrent une onde de choc comprimant et chauffant l'hydrogène jusqu'au seuil d'amorçage des réactions thermonucléaires. Les deux projets actuellement en construction en France et aux États Unis nécessitent des puissances laser considérables : l'énergie de chaque impulsion est de 1 MJ, qui est déposée en une nanoseconde, soit une puissance instantanée de 10^{15} W, c'est-à-dire la puissance moyenne consommée par les États Unis ! Le laser MégaJoule à Bordeaux devrait entrer en service courant 2010, mais les retombées civiles, mais d'éventuelles retombées civiles (i.e. la réalisation d'une centrale nucléaire civile) devront encore attendre plusieurs dizaines d'années !

Bibliographie

- [1] G. Grynberg, A. Aspect, C. Fabre, *Introduction aux lasers et à l'optique quantique*, Ellipses.
- [2] B. Cagnac et J.P.Faroux, *Lasers, interaction lumière-atomes*, EDP Sciences.
- [3] A. E. Siegman, *Laser*, University Science Books.