



## Rayonnement quantique

Alain Laverne

### ► To cite this version:

| Alain Laverne. Rayonnement quantique. 2006. cel-00092934

**HAL Id: cel-00092934**

**<https://cel.hal.science/cel-00092934>**

Submitted on 12 Sep 2006

**HAL** is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

# RAYONNEMENT QUANTIQUE

A\*\*\*L\*\*\*  
Université Paris 7

9 mars 1994



«Au total, cinquante années de spéculations conscientes ne m'ont pas rapproché de la réponse à la question “que sont les quanta de lumière?” » regrettait Einstein dans une lettre du 12 décembre 1951 à son vieil ami Besso. Qu'est-ce qu'un photon? La question fait toujours débat et, à défaut d'en éclairer la réponse, j'espère dans ce petit livre déjà trop long indiquer au moins ce que le photon n'est pas.

A vrai dire, jusqu'à ces dernières années la notion même de photon n'était pas réellement nécessaire pour expliquer le monde physique. Alors pourquoi depuis bientôt un siècle entend-on si souvent parler de ce photon dans les conversations des physicien(ne)s? Dans la plupart des cas les explications de processus physiques fondées sur des photons pourraient tout aussi bien se contenter de faire appel à des ondes électromagnétiques classiques, mais au prix de calculs généralement plus compliqués. De fait, la raison du succès populaire du photon tient à la simplicité et à la fécondité de la représentation qu'il propose à notre imagination; de l'impulsion et de l'énergie, quels ingrédients pourraient-être plus commodes pour notre intuition, ou nos calculs, généralement fondés sur des lois de conservation? Impulsion et énergie du champ électromagnétique classique existent certes, mais sont des quantités plus abstraites, variant continûment, plus complexes à manipuler qu'impulsion et énergie du photon qui, en se présentant par “quanta”, confèrent à celui-ci certains des attributs d'une particule.

Le photon, tellement pratique — qui ne l'a déjà invoqué? — trouve son fondement dans la théorie quantique du rayonnement, théorie qu'il est tout à fait, sinon parfaitement, possible d'exposer au niveau de la maîtrise de physique. Construire cette théorie n'exige rien de plus qu'un peu d'électrodynamique classique, un soupçon de relativité (généralement dissimulée dans les équations de Maxwell), et une connaissance de la mise en œuvre des principes de base de la théorie quantique qui ne dépasse guère les niveaux de l'oscillateur harmonique à une dimension. La lectrice endurente aura, je le souhaite, acquis à la fin de ce livre quelque idée de ce que l'on peut faire avec des photons et, ainsi, de ce que peut être une théorie quantique d'un champ, en l'occurrence le champ électromagnétique — ou encore champ vectoriel de masse nulle —, sans les complications et abstractions inhérentes à la théorie quantique des champs. Ici, tout va être quantique, mais seul le rayonnement sera représenté par un champ, la matière (les charges) non relativiste restant décrite par la fonction-d'onde solution d'une équation de Schrödinger usuelle. Bien entendu, ma maîtresse en physique a toute licence, si tel est son plaisir, de flâner au gré de ses tentations devant les titres aguicheurs du sommaire.

Une évolution relativement récente conduit à trouver de plus en plus d'ouvrages de physique qui exposent admirablement les réponses à des questions qui ne sont jamais explicitement posées. Le problème de la lectrice est alors la problématique! Je crois pour ma part m'être attaché ici, et plus par inclination que par devoir, à ces questions. En physique, lorsqu'elles celles-ci finissent par être clairement formulées, une bonne part de l'apprentissage est effectuée et leurs solutions, lorsqu'elles ne sont pas évidentes, ne réservent souvent que des difficultés techniques. Aussi, le

texte qui suit, indépendamment de mes défauts de style personnels (entre autres la verbosité et une invétérée propension au double-sens, voir à l’ambiguïté), est affligé d’une proportion de mots par rapport aux formules peut-être plus élevée que la moyenne. J’en demande pardon. Décidant néanmoins de m’adresser à un lectorat en maîtrise de physique et pas nécessairement virtuose des calculs, je me suis efforcé d’exposer ceux-ci de manière détaillée.

Le français n’autorise pas la neutralisation des genres. J’ai choisi de me confier à une hypothétique lectrice et d’évoquer plus volontiers observatrices et physiciennes. Vous pourrez voir là bien sûr une (à peine) subtile forme de machisme au second degré. Mais l’excès n’est pas encore à redouter dans ce registre, et lorsque ce reproche m’a parfois été formulé dans l’exercice de mon activité stipendiée, c’était invariablement par des éléments masculins. Et que n’entendrait pas une consœur qui pratiquerait le même choix? Pourquoi ne pas instaurer une nouvelle coutume éditoriale qui résoudrait ce dilemme le plus ancien du monde en s’inscrivant finalement dans notre traditionnelle, sinon archaïque, littérature courtoise: que l’auteur(e) s’adresse tout simplement à une personne du sexe opposé.

J’ai bien peur d’en avoir déjà fini avec l’exposé de ce qui pourrait être original dans ce que vous allez (j’imagine) lire. La théorie quantique du rayonnement compte heureusement d’excellents auteurs d’ouvrages parmi lesquels je citerai, des plus élémentaires aux plus complets: P.L. Knight et L. Allen [41], R. Loudon [52], C. Cohen-Tannoudji, J. Dupont-Roc et G. Grynberg [20, 21], et qui permettront à leur lectrice d’éclaircir mes considérations obscures et de corriger mes fautes. Ayant maintenant personnellement revendiqué celles-ci, je peux enfin brandir l’arme à double tranchant des remerciements aux nombreux collègues, ami(e)s et étudiant(e)s qui m’ont encouragé, assisté de leurs remarques, et même critiqué. En particulier, le tas de mots et de formules qui suit est le dernier avatar d’un cours que le responsable du D.E.A. de physique nucléaire de l’université de Grenoble n’avait pas hésité à me confier, il y a quelques années, dans un climat idéologique à vrai dire peu favorable.

Pour ma part, la conscience de mon impéritie reste mon seul espoir de me soustraire au jugement d’Einstein qui, après s’être apitoyé sur sa propre ignorance quant à la question primordiale de la nature du photon, retrouvait immédiatement sa causticité: «Il est vrai qu’aujourd’hui n’importe quel abruti croit connaître cette réponse, mais il se trompe.»

Bon courage.

# Sommaire

|           |                                                                 |           |
|-----------|-----------------------------------------------------------------|-----------|
| <b>I</b>  | <b>Prélude</b>                                                  | <b>1</b>  |
| I.1       | La théorie quantique . . . . .                                  | 1         |
| I.2       | Théorie et modèles . . . . .                                    | 3         |
| I.3       | Modèles quantiques . . . . .                                    | 4         |
| I.3.1     | Quanton à une dimension . . . . .                               | 5         |
| I.3.2     | Quanton à trois dimensions . . . . .                            | 6         |
| I.3.3     | Spin . . . . .                                                  | 7         |
| I.3.4     | Domaines de validité . . . . .                                  | 8         |
| I.3.5     | Autres modèles quantiques . . . . .                             | 9         |
| I.3.6     | Champ quantique . . . . .                                       | 10        |
| <b>II</b> | <b>A propos de l'équation de Schrödinger</b>                    | <b>11</b> |
| II.1      | L'invariance de jauge locale . . . . .                          | 13        |
| II.1.1    | Grandeurs physiques invariantes . . . . .                       | 17        |
| II.1.2    | L'interaction électromagnétique . . . . .                       | 18        |
| II.2      | Quanton de spin $1/2$ . . . . .                                 | 22        |
| II.2.1    | L'équation de Schrödinger libre réduite au premier ordre . . .  | 24        |
| II.2.2    | Hamiltonien de Pauli et facteur <b><math>g</math></b> . . . . . | 25        |
| II.3      | Matière quantique et rayonnement classique . . . . .            | 28        |
| II.4      | Le bon choix de jauge . . . . .                                 | 30        |
| II.4.1    | Equations des potentiels . . . . .                              | 31        |
| II.4.2    | La jauge de radiation . . . . .                                 | 31        |
| II.4.3    | Remarques sur le choix de jauge . . . . .                       | 33        |
| II.4.4    | Récapitulation . . . . .                                        | 33        |
| II.5      | Rayonnement classique, quelques caractéristiques . . . . .      | 35        |
| II.5.1    | Les ondes planes . . . . .                                      | 35        |
| II.5.2    | La boîte à modes . . . . .                                      | 36        |
| II.5.3    | Développement en modes plans . . . . .                          | 37        |
| II.5.4    | Sommes et intégrales . . . . .                                  | 39        |
| II.5.5    | Energie du champ électromagnétique . . . . .                    | 40        |

|                                                                              |           |
|------------------------------------------------------------------------------|-----------|
| II.6 La règle d'or de Fermi . . . . .                                        | 42        |
| II.7 Taux d'excitation . . . . .                                             | 48        |
| Exercices . . . . .                                                          | 49        |
| <b>III La théorie quantique du rayonnement</b>                               | <b>53</b> |
| III.1 Insuffisances et contradictions de la théorie semi-classique . . . . . | 54        |
| III.1.1 Le photon entre en scène . . . . .                                   | 54        |
| III.1.2 Deux sortes de rayonnement? . . . . .                                | 55        |
| III.1.3 Les avatars de $\hbar$ . . . . .                                     | 55        |
| III.2 L'espace de Fock des photons . . . . .                                 | 56        |
| III.2.1 La correspondance . . . . .                                          | 56        |
| III.2.2 L'espace de Fock . . . . .                                           | 57        |
| III.2.3 Excitation de la matière et absorption d'un photon . . . . .         | 57        |
| III.2.4 Opérateurs de création et d'annihilation . . . . .                   | 58        |
| III.3 L'opérateur de champ . . . . .                                         | 60        |
| III.4 L'opérateur énergie du rayonnement . . . . .                           | 61        |
| III.4.1 Quelques définitions . . . . .                                       | 62        |
| III.4.2 Nombre de photons . . . . .                                          | 64        |
| III.5 L'énergie du vide . . . . .                                            | 65        |
| III.5.1 La soustraction . . . . .                                            | 67        |
| III.5.2 Estimation de l'effet . . . . .                                      | 67        |
| III.5.3 Les modes d'une cavité résonante . . . . .                           | 69        |
| III.5.4 L'effet Casimir . . . . .                                            | 71        |
| III.5.5 Confirmation expérimentale . . . . .                                 | 75        |
| III.6 L'hamiltonien du rayonnement . . . . .                                 | 76        |
| III.7 L'impulsion du rayonnement . . . . .                                   | 80        |
| III.8 Le photon n'est pas une particule! . . . . .                           | 82        |
| Exercices . . . . .                                                          | 85        |
| <b>IV Propriétés des grandeurs physiques du rayonnement quantique</b>        | <b>87</b> |
| IV.1 Valeurs moyennes et dispersions . . . . .                               | 88        |
| IV.2 Commutateurs des opérateurs de champ . . . . .                          | 90        |
| IV.3 L'impulsion et le générateur des translations . . . . .                 | 95        |
| IV.3.1 Conservation de l'impulsion . . . . .                                 | 95        |
| IV.3.2 Invariance et générateur . . . . .                                    | 99        |
| IV.3.3 Translations et quanton . . . . .                                     | 101       |
| IV.3.4 Translations et champ . . . . .                                       | 103       |
| IV.3.5 Encore quelques considérations sur les translations . . . . .         | 105       |
| IV.4 Le moment angulaire et les rotations . . . . .                          | 108       |
| IV.4.1 Représentations des rotations . . . . .                               | 109       |
| IV.4.2 Polarisation circulaire . . . . .                                     | 112       |
| Exercices . . . . .                                                          | 115       |

|                                                                    |            |
|--------------------------------------------------------------------|------------|
| <b>V Quelques états du rayonnement quantique intéressants</b>      | <b>117</b> |
| V.1 Opérateurs amplitude et phase . . . . .                        | 119        |
| V.2 Les états de phase . . . . .                                   | 121        |
| V.3 Propriétés des états de nombre de photons . . . . .            | 123        |
| V.4 Propriétés des états de phase . . . . .                        | 125        |
| V.5 A la recherche d'états quasi classiques . . . . .              | 125        |
| V.5.1 Les états cohérents . . . . .                                | 127        |
| V.5.2 Propriétés algébriques . . . . .                             | 128        |
| V.5.3 Propriétés physiques . . . . .                               | 129        |
| V.5.4 La fabrication des états cohérents . . . . .                 | 134        |
| V.6 Pour les amateurs de classique . . . . .                       | 137        |
| V.7 Ondes et particules, le crépuscule des deux . . . . .          | 139        |
| V.8 Photons droits et onde tournante . . . . .                     | 146        |
| V.9 Les états comprimés . . . . .                                  | 148        |
| Exercices . . . . .                                                | 152        |
| <b>VI Emission et absorption des photons</b>                       | <b>159</b> |
| VI.1 Absorption d'un photon . . . . .                              | 160        |
| VI.2 Emission d'un photon . . . . .                                | 161        |
| VI.3 L'émission spontanée . . . . .                                | 163        |
| VI.4 L'émission dipolaire électrique . . . . .                     | 166        |
| VI.4.1 Pourquoi dipolaire électrique? . . . . .                    | 168        |
| VI.4.2 Règles de sélection . . . . .                               | 170        |
| VI.4.3 Taux de transition . . . . .                                | 172        |
| VI.4.4 Ordres de grandeur . . . . .                                | 175        |
| VI.5 Largeur naturelle . . . . .                                   | 176        |
| VI.5.1 Estimation . . . . .                                        | 177        |
| VI.5.2 Le modèle de Weisskopf et Wigner . . . . .                  | 180        |
| VI.5.3 Largeurs de raies, largeurs de niveaux . . . . .            | 188        |
| VI.6 Déplacement de niveau . . . . .                               | 191        |
| VI.6.1 Renormalisation de la masse de l'électron . . . . .         | 193        |
| VI.6.2 Déplacement Lamb . . . . .                                  | 195        |
| VI.7 Rayonnement multipolaire . . . . .                            | 200        |
| Exercices . . . . .                                                | 205        |
| <b>VII Histoires de photons</b>                                    | <b>207</b> |
| VII.1 La détection des photons . . . . .                           | 207        |
| VII.2 L'équilibre thermodynamique<br>matière-rayonnement . . . . . | 212        |
| VII.3 Diffusion d'un photon par un atome . . . . .                 | 218        |
| VII.3.1 La contribution diamagnétique . . . . .                    | 218        |
| VII.3.2 La contribution paramagnétique . . . . .                   | 219        |



|                                                       |            |
|-------------------------------------------------------|------------|
| VII.3.3 Graphes préhistoriques . . . . .              | 221        |
| VII.3.4 La formule de Kramers et Heisenberg . . . . . | 225        |
| VII.3.5 Diffusion élastique . . . . .                 | 228        |
| Exercices . . . . .                                   | 234        |
| <b>Bibliographie</b>                                  | <b>235</b> |

# Chapitre I

## Prélude

Les grandes généralités préliminaires sont rarement écoutées ou lues, encore plus rarement entendues. Si vous avez l'intention de transgresser cette règle statistique, alors vous découvrirez quelques réponses à des questions que vous vous êtes certainement, et plus ou moins consciemment, posées au cours de votre apprentissage de la théorie quantique, réponses qui ne nécessiteront rien de plus que ce que vous savez déjà de celle-ci. Mais il y a quelque chance que ces considérations finissent par vous paraître trop abstraites. Abandonnez sans remord et passez aux chapitres suivants. N'oubliez pourtant pas le but de ce prologue: qu'appelle-t-on système quantique et comment caractérise-t-on son espace des états? Et plus tard, lorsque ces questions vous paraîtront plus concrètes, à propos du rayonnement quantique, vous apprécierez peut-être un retour à ces quelques pages.

### I.1 La théorie quantique

Comme toute théorie physique, la quantique est très simple! Vous avez probablement suffisamment pratiquée cette dernière pour savoir qu'elles se résument à quelques règles de base:

- A un état du système quantique est associé un vecteur d'état (ou ket d'état), élément d'un espace des états, espace vectoriel qui présente la particularité d'être défini sur les nombres complexes.
- Une grandeur physique du système n'est plus seulement représentée par l'ensemble des valeurs qu'elle peut prendre. Elle n'a de valeur, disons  $a$ , que lorsque le système est dans un état spécifique. Notons en le ket  $|a\rangle$ .<sup>1</sup> Ainsi,

---

<sup>1</sup>Il faut bien sûr une valeur d'une autre grandeur pour identifier l'état si la valeur  $a$  est dégénérée. Mais ce n'est que brouille technique.

une grandeur physique est représentée par un ensemble de doublets (valeur, ket):  $\{(a, |a\rangle)\}$ .

- Une équation du mouvement, dite de Schrödinger,

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = H(t) |\psi(t)\rangle,$$

dans laquelle l'opérateur linéaire  $H$ , l'hamiltonien du système, génère l'évolution du vecteur d'état au cours des translations dans le paramètre temps.

- Enfin, une règle dite d'interprétation, aussi simple qu'incompréhensible: la probabilité de “trouver” l'état  $|\psi\rangle$  dans l'état  $|\phi\rangle$  est donnée par le module carré de leur produit scalaire hermitique,  $|\langle\psi|\phi\rangle|^2$ .

C'est tout! La règle d'interprétation implique la finitude de la norme des vecteurs d'état (il faut être certain de trouver un état dans lui-même). Autrement dit, à quelques exigences techniques supplémentaires près, l'espace des états a une structure dite d'espace de Hilbert. Deux valeurs différentes,  $a$  et  $a'$ , d'une grandeur physique s'excluant, la même règle conduit donc à l'orthogonalité des kets correspondants:  $\langle a|a'\rangle = 0$ . Cette propriété permet d'associer à toute grandeur physique du système, c'est-à-dire à l'ensemble des doublets  $(a, |a\rangle)$ , un opérateur linéaire dans l'espace des états:

$$A \stackrel{\text{df}}{=} \sum_a |a\rangle a \langle a|.$$

Ainsi défini, cet opérateur a pour valeurs propres, et kets propres associés, les  $a$  et  $|a\rangle$ . Au cas où il s'agit d'une grandeur physique à valeurs réelles (ce qui n'est nullement nécessaire, les réels ne sont pas moins imaginaires que les complexes), l'opérateur  $A$  est donc hermitique. (En cas de valeurs  $a$  complexes, l'opérateur  $A$  reste normal: il commute avec son adjoint.)

Mais est-ce bien tout? Il subsiste à vrai dire quelques “petites” difficultés liées à la règle d'interprétation, autrement dit au fameux problème de la mesure en théorie quantique. La question est non résolue, toujours objet de débat car elle obère profondément la cohérence logique de toute réalisation de la théorie, et d'autant plus surprenante qu'on ne lui connaît, en pratique, aucune manifestation évidente. Une expérience de mesure sur un système quantique nécessite un amplificateur pour obtenir des effets accessibles à nos sens classiques, donc un système à très grand nombre de degrés de liberté dont la description relève de la physique statistique, et est encore loin d'être claire. En d'autres termes, l'ambiguïté réside dans la notion de probabilité dont la définition, lorsqu'elle en a une, n'est assurément pas la même en mathématiques, en statistique, et en théorie quantique.

## I.2 Théorie et modèles

Voilà bouclé le tour de la théorie quantique y compris ses zones d'ombre. Mais une théorie physique seule est encore muette sur le réel. Intermédiaire obligé entre un objet et une théorie dont on attend des prédictions, il y a ce que l'on ne craignait pas d'appeler la "mise en équation du problème". Le terme plus séduisant de "modèle" est maintenant consacré par la mode. La notion est, la plupart du temps, d'autant moins expliquée que l'on reproche aux étudiant(e)s d'être incapable de sa mise en œuvre. Et pourtant il s'agit de l'activité principale des physicien(ne)s car bien peu ont eu la chance de créer effectivement une théorie physique.

Pour illustrer la chose, revenons à des théories plus familières, que vous avez le sentiment de mieux dominer. Convaincu(e) des mérites insignes de la gravitation newtonienne — ne serait-ce que la simplicité de ses règles —, vous désirez en calculer les prédictions pour le mouvement de la Terre. Et commence la modélisation: vous n'envisagez d'abord qu'un univers constitué du Soleil et de la Terre seuls, puis vous assimilez ceux-ci à des masses ponctuelles. Ne reste qu'à intégrer les équations de la théorie, devenues les équations du mouvement du modèle, moyennant les valeurs empiriques...

- des paramètres de la théorie, à savoir la constante de la gravitation;
- des paramètres du modèle, en l'occurrence les masses du Soleil et de la Terre;
- des paramètres de la solution, autrement dit les conditions initiales.

Vos prédictions sont brillamment confirmées par les observations ultérieures du Soleil, mais les techniques expérimentales progressant, quelques légers désaccords se font jour. Mettez-vous alors en doute la théorie newtonienne de la gravitation? Non, bien sûr. Plus modestement vous modifiez votre modèle en meublant son univers d'autres masses ponctuelles représentant la Lune, Vénus, Mars, Saturne, Jupiter *etc.* Ça complique certes un peu l'intégration des équations du mouvement mais donne d'excellents résultats, et la théorie de Newton s'en sort intacte. Enquêteur sceptique sur l'astrologie, vous en analysez les vieilles prétentions et réalisez que les positions du Soleil n'ont maintenant plus rien à voir avec celles correspondant aux périodes convenues il y a 2200 ans pour les différents signes du zodiac. Reprenant votre modèle, vous représentez la Terre non plus par un point mais par une sphère dont la densité ne dépend que du rayon. Demi-échec: les mêmes prédictions, ni plus ni moins, et une Terre qui peut tourner sur elle-même, certes, mais autour d'un axe qui garde une direction fixe. Opiniâtre, vous corrigez encore votre modèle en représentant la Terre par un sphéroïde légèrement aplati; la distribution de masse possède maintenant un moment quadrupolaire. Victoire: le Soleil et la Lune ont pour effet une rotation de l'axe polaire de la Terre, vous expliquez la précession des équinoxes. La théorie newtonienne de la gravitation est encore juste. Elle le reste tant que vous ne cherchez pas à expliquer l'avance du périhélie de Mercure, là où tous les modèles échouent; une théorie plus profonde, celle d'Einstein par exemple, en fait désormais une théorie fausse et, partant, dont on connaît précisément le domaine de

validité. Quant à vos divers modèles, sont-ils justes ou faux? Vous voyez à présent que la question ne se pose pas en ces termes. Chacun de ces modèles a son domaine d'application.

La dualité théorie-modèle est inévitable. On ne “vérifie” pas directement une théorie. Elle n'est capable de prédiction qu'à travers des modèles construits pour les besoins de la cause, ou plutôt de l'effet. Elle ne peut donc tout nous dire sur l'essence des objets de la nature. La théorie de Newton est elle-même sans opinion sur le modèle de la Terre: un point, une sphère, un ellipsoïde? Alors, tout est permis? Non, car le modèle est une vision idéalisée qui doit quand même être compatible avec les grands principes de la théorie. Pour prendre exemple d'une autre théorie — l'électrodynamique classique —, ce ne sont pas les équations de Maxwell-Lorentz qui peuvent décréter que les particules “élémentaires” sont des charges ponctuelles ou réparties, d'une forme ou d'une autre. Les particules élémentaires sont des modèles, mais contraints par la théorie, par exemple à satisfaire la conservation de la charge,  $\partial\rho/\partial t + \nabla \cdot \mathbf{j} = 0$ , conséquence des équations du Maxwell. A un niveau plus technique, la théorie permet aussi de simplifier les modèles. Par exemple, la gravitation newtonienne nous dit qu'il est inutile d'envisager une forme de masse compliquée si l'on se borne à en étudier le mouvement du centre de gravité: ce mouvement est le même que si toute la masse y était concentrée.

Ma lourde insistance sur cette dualité devrait dissiper l'incrédulité quasi obligée de la lectrice lorsqu'elle s'entendait affirmer qu'une théorie physique est simple. Etait-ce acceptable à propos de l'électrodynamique classique, par exemple, alors qu'une année n'était pas de trop pour l'étudier? Vous réalisez maintenant, je l'espère, que les complications résidaient, et résident toujours, dans la description et l'utilisation des charges ponctuelles ou réparties, des milieux isolants, conducteurs, diélectriques, magnétiques, linéaires, isotropes, homogènes... ou non, tous modèles adaptés à des situations spécifiques, et accessoires optionnels mais essentiels de la théorie, en aval de celle-ci. Physicien(ne)s ordinaires, notre art s'exerce dans les complications de l'élaboration et de la résolution de ces modèles. Nous n'avons ni la chance de vivre au moment où la remise en cause de la théorie devient impérieuse (lorsqu'échouent tous les modèles), ni le génie qu'exige l'édification de la monumentale simplicité d'une nouvelle théorie physique.

### I.3 Modèles quantiques

Qu'en est-il de la théorie quantique? J'en ai rappelé les règles de base et il est clair que celles-ci n'ont *a priori* pas grand chose à nous dire sur la nature des particules élémentaires par exemple. Où se situe donc la notion de modèle en théorie quantique?

Au système quantique modélisé correspond un espace des états, mais déclarer que c'est un espace de Hilbert, même si cela impressionne, est encore trop général.

Ce sont les grandeurs physiques indépendantes (comme  $\mathbf{r}(t)$  et  $\mathbf{p}(t)$  pour une particule classique, ou  $\mathbf{E}(\mathbf{r}, t)$  et  $\mathbf{B}(\mathbf{r}, t)$  pour un champ électromagnétique) qui caractérisent notre système, c'est-à-dire les opérateurs associés agissant dans l'espace des états. Mais deux opérateurs qui commutent ont des propriétés très particulières:

- soit ils sont fonction l'un de l'autre, proportionnels par exemple, auquel cas ils ne peuvent décemment être qualifiés d'indépendants;
- soit ce sont les prolongements, à un espace produit, de deux opérateurs agissant dans des espaces différents.

On peut donc définir un modèle, autrement dit la structure de son espace des états, par la donnée:

- d'un *ensemble irréductible d'opérateurs*, à savoir une liste d'opérateurs telle que tout autre opérateur agissant dans cet espace et qui commuterait avec *tous* les opérateurs de l'ensemble irréductible, serait trivial, autrement dit un multiple de l'identité;
- et des relations de commutation des opérateurs de cet ensemble.

Tout opérateur est ainsi une fonction des opérateurs de l'ensemble irréductible et on peut exprimer son commutateur avec chacun de ceux-ci.

Nous avons notre système. Reste à en étudier l'évolution ultérieure, c'est-à-dire à se donner un hamiltonien. Là encore, la théorie quantique nous laisse le choix, mais pour nous cet hamiltonien  $H$  est maintenant une fonction des opérateurs de l'ensemble irréductible (et éventuellement du temps si le système est soumis à un environnement variable), à construire en respectant quelques grands principes d'invariance, si l'on y croit, et guidés par l'espoir d'en tirer des prédictions valides. (Poursuivant le parallèle avec la particule classique, nous avons les grandeurs physiques indépendantes  $\mathbf{r}$  et  $\mathbf{p}$ , reste à se donner la force  $\mathbf{F}(\mathbf{r}, \mathbf{p}, t)$ .)

### I.3.1 Quanton à une dimension

A tout cela qui vous semble terriblement formel, vous connaissez de multiples exemples. Commençons par le modèle le plus simple, défini par l'ensemble irréductible  $\{X, P\}$  et la relation de commutation  $[X, P] = i\hbar$ , ce que vous appeliez probablement particule à une dimension. Mais la notion de particule en théorie quantique ne recouvre pas exactement celle de la mécanique classique. (Sinon qu'y aurait-il de nouveau?) La controverse sur la prétendue dualité onde-corpuscule n'a pas peu contribué à brouiller les esprits. Aussi, par désir de neutralité et plutôt que de m'exposer, par le recours à des avatars tels que "particule" ou "ondicule", au soupçon ou au ridicule, j'utiliserai le nom, guère euphonique, de *quanton* [14, § 6.1]. Dans notre modèle du quanton à une dimension, toute grandeur physique, à commencer par l'hamiltonien, s'exprime donc par définition en fonction de  $X$  et  $P$ . (Sinon, il s'agit d'un nouvel opérateur, à ajouter à l'ensemble irréductible, et d'un nouveau modèle.) Vous n'avez que l'embarras du choix pour en proposer des exem-

ples, à commencer par  $H = P^2/2m$ , et plus généralement la forme

$$H = \frac{1}{2m}P^2 + V(X, t).$$

Vous avez déjà pratiqué de multiples réalisations du modèle du quanton à une dimension, en d'autres termes différentes fonctions potentiel  $V$  rendant plus ou moins bien compte de toutes sortes propriétés physiques: puits infinis, créneaux, merlons, oscillateur  $V(X) = m\omega^2 X^2/2$ , potentiel de Morse  $V(X) = V_0(e^{-2\lambda X} - 2e^{-\lambda X})$ , champ électrique oscillant  $V(X, t) = -qE_0 X \cos \omega t$ , ou encore  $V(X) = -V_0/\cosh^2 \lambda X$ , *etc.*

Le système quantique est défini par son ensemble irréductible d'opérateurs et l'algèbre (les commutateurs) de ceux-ci, mais l'ensemble n'est pas unique. On peut trouver d'autres opérateurs, fonctions de ceux-là, qui définiront tout aussi bien le même espace des états, et en fonction desquels les grandeurs physiques s'exprimeront. Introduisons par exemple dans notre quanton à une dimension,  $\{X, P\}$ ,  $[X, P] = i\hbar$ , un paramètre  $l$  ayant la dimension d'une longueur, et définissons l'opérateur

$$a \stackrel{\text{df}}{=} \frac{1}{\sqrt{2}}\left(\frac{1}{l}X + i\frac{l}{\hbar}P\right).$$

On a alors  $[a, a^+] = 1$ , et le même espace des états, le même modèle, est tout aussi bien défini par l'ensemble irréductible d'opérateurs  $\{a, a^+\}$ ,  $[a, a^+] = 1$ , en fonction desquels s'exprimeront toutes les grandeurs physiques de ce système:  $X = (\frac{l}{\sqrt{2}})(a + a^+)$ ,  $P = (\frac{\hbar}{i\sqrt{2}})(a - a^+)$ ,  $P^2/2m = \dots$ ,  $H = \dots$ , *etc.* En particulier, un quanton de masse  $m$  dans un environnement d'oscillateur de pulsation  $\omega$ , comporte une longueur caractéristique  $l = \sqrt{\hbar/m\omega}$ , et vous avez déjà eu l'occasion d'apprécier tout le parti que l'on peut tirer d'une description équivalente en termes de  $a \stackrel{\text{df}}{=} \sqrt{m\omega/2\hbar}X + (i/\sqrt{2m\hbar\omega})P$  et de son conjugué.

### I.3.2 Quanton à trois dimensions

La fécondité du modèle du quanton ne se borne pas à un environnement unidimensionnel. On peut maintenant construire, par produit de trois espaces de quanton à une dimension, un nouveau modèle dont l'ensemble irréductible est constitué de six opérateurs,  $\{R_x, R_y, R_z, P_x, P_y, P_z\}$  ayant pour commutateurs:

$$\begin{aligned} [R_i, R_j] &= i\hbar\delta_{ij}, \\ [R_i, R_j] &= [P_i, P_j] = 0, \end{aligned}$$

puisque par définition deux opérateurs prolongements d'opérateurs agissant dans des espaces différents commutent. L'invariance par rotation à laquelle on croit *a priori* pour un quanton, sinon pour son environnement, oblige alors à prendre le

même paramètre  $m$  dans les trois directions, et vous reconstruisez ainsi votre modèle du quanton favori:

$$H = \frac{1}{2m} \mathbf{P}^2 + V(R_x, R_y, R_z).$$

Ce n'est déjà pas si mal, mais l'adhésion bien comprise au principe d'invariance par rotation réserve quelques richesses supplémentaires. Par cette invariance on entend que des observatrices dont les points de vue ne diffèrent que par une rotation, vont fournir des descriptions équivalentes d'une expérience sur le système, en particulier les mêmes probabilités (mêmes taux de comptage). Le comportement du vecteur d'état par rapport à cette rotation de point de vue est ainsi caractérisé par l'opérateur  $\mathbf{J}$ , le générateur des rotations du système, le moment angulaire en langue vernaculaire, dont les composantes doivent satisfaire les relations de commutation

$$[J_i, J_j] = i\hbar \varepsilon_{ijk} J_k.$$

Grandeur physique du système,  $\mathbf{J}$  doit s'exprimer en fonction des opérateurs de l'ensemble irréductible,  $\mathbf{R}$  et  $\mathbf{P}$ , qui définissent ce système. De fait vous connaissez l'opérateur moment orbital,  $\mathbf{L} \stackrel{\text{df}}{=} \mathbf{R} \wedge \mathbf{P}$ , dont les relations de commutation sont précisément celles-ci. Le moment orbital est donc un candidat *bona fide* au rôle de générateur des rotations de notre quanton.

### I.3.3 Spin

Imaginons maintenant un nouveau modèle, avatar du quanton, défini à partir d'un ensemble irréductible constitué de  $\mathbf{R}$ ,  $\mathbf{P}$  et de  $\mathbf{J}$ . En d'autres termes, le générateur des rotations, qui joue un rôle fondamental, est indépendant de  $\mathbf{R}$  et  $\mathbf{P}$ . Analysons les conséquences de cette idée saugrenue en reprenant le moment orbital  $\mathbf{L} \stackrel{\text{df}}{=} \mathbf{R} \wedge \mathbf{P}$ , et en définissant l'opérateur  $\mathbf{S} \stackrel{\text{df}}{=} \mathbf{J} - \mathbf{L}$ , traditionnellement appelé "spin", que l'on peut, sans rien changer à la nature du modèle, substituer à  $\mathbf{J}$  dans la liste des opérateurs de l'ensemble irréductible.

Les commutateurs se calculent facilement et vous vous retrouvez ainsi en possession d'un modèle fondé sur un ensemble de neuf opérateurs,

$$\{R_x, R_y, R_z, P_x, P_y, P_z, S_x, S_y, S_z\},$$

avec les commutateurs:

$$\begin{aligned} [R_i, P_j] &= i\hbar \delta_{ij}, \\ [S_i, S_j] &= i\hbar \varepsilon_{ijk} S_k, \\ [R_i, R_j] &= [P_i, P_j] = [S_i, R_j] = [S_i, P_j] = 0. \end{aligned}$$

L'opérateur  $\mathbf{S}$  (surprise) commute avec les autres opérateurs,  $\mathbf{R}$  et  $\mathbf{P}$ , de l'ensemble irréductible! L'espace des états de ce modèle a donc la structure, inattendue, d'un



produit d'espace des états d'un quanton et de l'espace des états du spin, lui même défini...

- par l'ensemble irréductible  $\{S_x, S_y, S_z\}$ ,
- et par l'algèbre  $[S_i, S_j] = i\hbar\epsilon_{ijk}S_k$ .

Nul doute qu'après avoir étudié le moment angulaire quantique, les propriétés spectrales qu'impliquent cette algèbre vous soient bien connues. L'opérateur  $\mathbf{S}^2$  commute avec chacun des opérateurs  $S_x$ ,  $S_y$  et  $S_z$ . Il existe donc une base d'états propres communs à  $\mathbf{S}^2$  et  $S_z$ , par exemple. Les valeurs propres de  $\mathbf{S}^2$  sont de la forme  $\hbar^2 s(s+1)$ , avec  $2s$  entier naturel. Le sous espace des vecteurs propres de  $\mathbf{S}^2$  correspondant à la valeur  $s$  a la dimension  $2s+1$ . Ce sous espace est sous-tendu, par exemple, par les vecteurs propres de  $S_z$  dont les  $2s+1$  valeurs propres vont de  $-\hbar s$  à  $\hbar s$  par pas de  $\hbar$  (base standard). Chaque sous espace,  $\mathcal{E}_s$ , a la propriété d'être invariant par rotation: une rotation de point de vue change un vecteur du sous espace  $\mathcal{E}_s$  en un vecteur du *même* sous espace. Outre son invariance, ce sous espace est irréductible: il ne contient pas de sous espace invariant.

Notre idée de particule est celle d'un objet doué de certaines caractéristiques, une masse inerte, une charge électrique (constante de couplage avec le champ électromagnétique) et, pour ce qui nous intéresse ici, son comportement par rapport aux rotations (sa "forme"). Et si cette particule est élémentaire, cela signifie que justement on ne veut pas y trouver de constituants eux-mêmes élémentaires. Nous avons donc maintenant un modèle candidat tout indiqué pour représenter une particule élémentaire dans le cadre de la théorie quantique: un espace des états  $\mathcal{E} = \mathcal{E}_0 \otimes \mathcal{E}_s$ , produit de l'espace des états de notre premier modèle du quanton à trois dimensions (dans lequel agissent les opérateurs  $\mathbf{R}$  et  $\mathbf{P}$ ) et du sous espace des états propres de  $\mathbf{S}^2$  correspondant à la valeur  $s$  (et dans lequel agit l'opérateur  $\mathbf{S}$ ). Le sous espace  $\mathcal{E}_s$  étant irréductible, c'est bien l'espace des états d'un modèle "élémentaire" que je nous autoriserai à appeler quanton de spin  $s$ .

### I.3.4 Domaines de validité

Comme tout modèle, ceux que je viens de construire ont, s'ils correspondent à une quelconque situation expérimentale, un domaine probablement limité. Vous ne pouvez ignorer les immenses succès du modèle du quanton de spin  $\frac{1}{2}$  en ce qui concerne l'électron tout au moins non relativiste ou, plus exactement, relativiste galiléen, c'est-à-dire pour des énergies inférieures à  $m_e c^2 = 0,5 \text{ Mev}$ . Justement, si dans une expérience on fait subir à cet électron une interaction mettant en jeu des échanges d'énergie supérieures à  $0,5 \text{ Mev}$ , autrement dit si cette expérience est sensible à des détails inférieurs à la longueur d'onde Compton  $\lambda_e = \hbar/m_e c = 4 \times 10^2 \text{ fm}$ , on sort évidemment du domaine galiléen et, de plus, se manifestent d'étranges phénomènes tels que des créations de paires électron-positron. Est-ce à dire que la théorie quantique est invalidée et que l'électron n'est pas une particule élémentaire? Restons modestes et contentons nous d'en conclure que notre modèle du quanton de spin  $\frac{1}{2}$

n'est opérant pour l'électron qu'à des énergies inférieures à 0,5 Mev ou, autrement dit, que dans le cadre de ce modèle, l'électron est "ponctuel", mais qu'en tout état de cause les expériences correctement décrites par ce modèle permettront au mieux de conclure que la taille de l'électron est inférieure à 400 fm.

Posons nous la même question pour le proton. Les succès de la physique nucléaire dans les prédictions des propriétés des noyaux comme agrégats (produits d'espaces) de nucléons attestent encore la pertinence du modèle du quanton. Mais ce modèle n'est plus en accord avec les résultats expérimentaux lorsqu'on examine le proton avec une résolution inférieure au fermi, autrement dit lorsqu'on a une énergie d'interaction de l'ordre de 200 Mev. (Souvenez vous:  $\hbar c = 200 \text{ Mev fm.}$ ) Le modèle atteint là sa limite, limite qui n'est même pas celle de la relativité galiléenne (la masse du proton est de l'ordre du Gev). Conclusion: le modèle du quanton convient pour le proton jusqu'à 1 fm. Au delà, force est de songer à d'autres modèles, comportant d'autres degrés de liberté que **R**, **P**, et **S**, par exemple un espace des états produit de trois espaces des états de quantons de spin  $\frac{1}{2}$ : deux quarks *u* et un quark *d*.

On peut bien sûr en dire autant pour le neutron, parfaitement décrit jusqu'à 1 fm par le modèle du quanton de spin  $\frac{1}{2}$ , encore que... le neutron soit instable! Après une attente d'une attente d'une dizaine de minutes, le neutron que l'on surveillait à de bonnes chances de s'être désintégré, ceci en complet désaccord avec notre modèle dont l'espace des états même est associé au quanton (et un vecteur de cet espace à un état du quanton). Ce quanton est donc aussi éternel que son espace des états (l'espace des états n'évolue pas!). Dans un style plus orthodoxe: l'état du neutron est représenté par un vecteur d'état de norme constante, dont l'évolution est générée par un hamiltonien hermitique. Toujours est-il que le modèle du quanton de spin  $\frac{1}{2}$  représente bien le neutron sur des distances supérieures à 1 fm et surtout pendant une durée inférieure à 10 minutes. Au delà, nul besoin encore de remettre en cause la théorie quantique, mais il faut se chercher un nouveau modèle.

### I.3.5 Autres modèles quantiques

Les modèles que vous connaissiez permettent pas mal de jongleries. A partir de l'espace défini par l'ensemble irréductible  $\{S_x, S_y, S_z\}$ ,  $[S_i, S_j] = i\hbar\epsilon_{ijk}S_k$ , nous avons déjà envisagé le sous espace  $\mathcal{E}_s$  des vecteurs propres de  $\mathbf{S}^2$  associés à la valeur *s*. Ce modèle, disons du spin *s*, peut très bien s'appliquer à la description d'un noyau dans l'environnement d'un analyseur de Stern et Gerlach. Les effets quantiques sur les autres degrés de liberté du noyau s'avèrent alors négligeables. La particule ne joue ici qu'un rôle de porteur classique d'un spin quantique qui voit son environnement magnétique évoluer à l'aune du mouvement de la particule. L'évolution du vecteur d'état du spin est régie, dans cette approximation, par l'hamiltonien

$$H(t) = -g\frac{q}{2m}\mathbf{B}(t) \cdot \mathbf{S}.$$

Mais, contre exemple, ce modèle ne convient pas du tout pour un électron: les comportements quantiques de  $\mathbf{R}$  et  $\mathbf{P}$  ne sont pas négligeables par rapport à celui de  $\mathbf{S}$ . Cela n'empêche pas, avec le modèle du spin, de construire par produit d'espaces un modèle pour un ensemble de spins portés par des sites, et des hamiltoniens (Ising ou Heisenberg) d'un intérêt avéré pour l'étude du magnétisme de la matière.

### I.3.6 Champ quantique

Si nous voulons décrire la disparition ou la création de particules observées dans la nature, il nous faut abandonner le modèle du quanton, et remettre en cause notre notion même de particule. (Après tout, classiquement, les particules sont éternelles, elles ne peuvent être créées ou détruites mais seulement agrégées ou désagrégées.) Les grandes lignes du modèle quantique à mettre au point pour cela apparaissent maintenant. Positions  $\mathbf{R}$  et impulsions  $\mathbf{P}$ , spécifiques d'un quanton, ne peuvent plus être des variables dynamiques. Notre modèle doit avoir un espace des états plutôt structuré par un ensemble irréductible d'opérateurs de création et d'annihilation dotés de relations de commutation compatibles avec leur rôle. Historiquement, cette opération s'est appelée la "seconde quantification" (l'existence ou non de la particule devenant une variable dynamique quantique). Ayant reconnu que la position a perdu son statut de variable dynamique pour être réduite à l'état de simple paramètre, au même titre que le temps, il est préférable d'adopter la terminologie plus moderne de "champ quantique". C'est à la construction de l'espace des états de ce genre de modèle que nous allons nous consacrer, dans le cas particulier du champ électromagnétique. Plus généralement, on parle de "théorie quantique des champs", mais il peut être rassurant de se souvenir qu'il ne s'agit toujours en fait que de modèles dans le cadre de la théorie quantique et d'une théorie de relativité, toutes théories qui vous sont déjà plus ou moins familières.

## Chapitre II

# A propos de l'équation de Schrödinger

La moindre des contraintes phénoménologiques que doit satisfaire une éventuelle théorie quantique du rayonnement est de reproduire les résultats des théories antérieures lorsqu'ils sont satisfaisants. Je vais donc commencer par rappeler l'expression de l'hamiltonien d'un quanton chargé, aussi vieux que la théorie quantique elle-même, puisque ses premiers succès remontent à la description des niveaux de l'atome d'hydrogène.

Cet hamiltonien s'obtient de façon particulièrement économique à partir d'un principe d'invariance de jauge. Un tel choix n'est pas seulement dicté par la mode récente dont jouit une propriété déjà bien ancienne. Le procédé orthodoxe, peut-être le seul connu de la lectrice, reposait sur l'analogie formelle entre hamiltoniens quantique et classique, érigée en “principe de correspondance”. Or ce principe est autant affecté d'insuffisances techniques (quelle grandeur quantique va correspondre par exemple à la grandeur classique  $xp$ : l'opérateur  $XP$ , ou  $PX$ , ou une quelconque moyenne? quelle grandeur classique peut revendiquer la paternité directe du spin quantique?), que de faiblesse logique (la nature est, pour l'instant, quantique et il serait vain d'espérer en bâtir une théorie cohérente en se contentant de s'appuyer sur des manifestations classiques qui, pour abondantes qu'elles soient, ne constituent qu'une vue partielle des choses).

Faites l'effort — fondamental dans le processus d'apprentissage — d'éliminer les schémas acquis et imaginez Sophie, brillante physicienne qui connaît bien son équation de Schrödinger pour un quanton,

$$i\hbar \frac{\partial}{\partial t} \psi(\mathbf{r}, t) = \left\{ \frac{\mathbf{P}^2}{2m} + V(\mathbf{r}, t) \right\} \psi(\mathbf{r}, t), \quad (\text{II.1})$$

encore qu'elle soit assez ingénue pour tout ignorer de l'interaction électromagnétique. L'opérateur du membre de droite, l'*hamiltonien*, a pour rôle, comme l'indique le premier membre, de générer l'évolution temporelle de la fonction-d'onde. L'expression de l'hamiltonien dépend:

- de l'environnement du quanton par la fonction potentiel  $V(\mathbf{r}, t)$ ;
- de l'opérateur impulsion,  $\mathbf{P} = (\hbar/i)\nabla$  en représentation de position qui, de la même façon que  $i\hbar(\partial/\partial t)$ , génère l'évolution spatiale de la fonction-d'onde lors d'une translation.

Depuis ses premiers pas dans l'univers quantique, Sophie a noté que rien n'y dépend d'un changement de phase global de la fonction-d'onde. Ainsi, non seulement la fonction

$$\psi'(\mathbf{r}, t) \stackrel{\text{df}}{=} e^{i\omega} \psi(\mathbf{r}, t) \quad (\text{II.2})$$

satisfait-elle la même équation d'évolution (II.1), mais de plus les grandeurs physiques du quanton ont les mêmes distributions (valeurs moyennes et divers moments), que sa fonction-d'onde soit  $\psi$  ou  $\psi'$ . Les fonctions-d'onde de la classe (II.2) sont physiquement équivalentes; elles représentent le même état. Une transformation telle que (II.2) est qualifiée d'*interne* car elle n'implique aucun déplacement spatio-temporel. On parle ici de *transformation de jauge globale* du groupe  $U(1)$ , le groupe des matrices unitaires de rang 1, qui ne désigne rien d'autre — dans un style ampoulé — que les nombres complexes de module 1. Un collègue de Sophie qui utiliserait la fonction  $\psi'$  pour décrire la même situation n'y verrait aucune différence, et l'une comme l'autre n'ont nul besoin de se mettre d'accord sur un choix de phase absolu; fort heureusement d'ailleurs puisque aucun n'est privilégié par la nature.

Vous avez déjà rencontré la même ambiguïté à propos de deux observateurs en mouvement relatif uniforme; les équations du mouvement d'un phénomène revêtent pour eux la même forme, même si les valeurs qu'ils attribuent respectivement aux grandeurs physiques peuvent être différentes. Le fait, érigé en principe, que rien dans les équations d'une théorie ne doive, jusqu'à présent, permettre de distinguer l'un d'entre eux est connu sous le nom d'invariance galiléenne; les équations du mouvement d'un pendule dans le T.G.V. sont les mêmes pour la voyageuse et pour le chef de gare.

Bien sûr, selon les conditions initiales (qui n'ont rien à voir avec les équations du mouvement), un point de vue peut se révéler techniquement plus commode que l'autre; la trajectoire du pendule a une expression plus simple pour la voyageuse. Il en ira de même pour le choix de phase — le choix de la jauge — beaucoup plus tard, selon que l'on souhaitera mettre plus particulièrement en évidence l'invariance de jauge ou l'invariance de Lorentz de l'électrodynamique quantique, discuter du mécanisme de Higgs ou s'interroger sur la renormalisabilité d'une théorie de jauge.

## II.1 L'invariance de jauge locale

Enhardie par la prise de conscience de cette invariance de la théorie quantique, Sophie essaye un autre type de transformation interne du même groupe,

$$\psi'(\mathbf{r}, t) \stackrel{\text{df}}{=} e^{i\frac{q}{\hbar}\omega(\mathbf{r}, t)} \psi(\mathbf{r}, t). \quad (\text{II.3})$$

Le paramètre de la transformation dépend maintenant du lieu et de l'instant; on parle de *transformation de jauge locale*  $U(1)$ . La licence d'extraire de la définition de  $\omega$  les facteurs  $q$  (une constante réelle) et  $\hbar$ , n'a d'autre but que le confort ultérieur.

Se pourrait-il que la théorie quantique affiche encore la même indifférence à ce type de transformation? Pour cela, il n'est que de déterminer l'équation d'évolution satisfaite par  $\psi'$ , sachant que  $\psi$  est régie par (II.1). De la relation inverse

$$\psi(\mathbf{r}, t) = e^{-i\frac{q}{\hbar}\omega(\mathbf{r}, t)} \psi'(\mathbf{r}, t),$$

on déduit immédiatement (Exercice 1)

$$\begin{aligned} i\hbar \frac{\partial}{\partial t} \psi &= e^{-i\frac{q}{\hbar}\omega} \left( i\hbar \frac{\partial}{\partial t} + q \left( \frac{\partial \omega}{\partial t} \right) \right) \psi', \\ \mathbf{P} \psi &= e^{-i\frac{q}{\hbar}\omega} \left( \mathbf{P} - q(\nabla \omega) \right) \psi'. \end{aligned}$$

L'action de  $\mathbf{P}^2$  sur  $\psi$  s'obtient en calculant d'abord la divergence du produit d'un champ scalaire exponentiel par un champ vectoriel  $\mathbf{u}(\mathbf{r}, t)$  (Exercice 2):

$$\mathbf{P} \cdot \left( e^{-i\frac{q}{\hbar}\omega} \mathbf{u} \right) = e^{-i\frac{q}{\hbar}\omega} \left( \mathbf{P} - q(\nabla \omega) \right) \cdot \mathbf{u}.$$

Appliquant cette relation à  $\mathbf{P}\psi$ , on trouve finalement que  $\psi'$  a pour équation d'évolution

$$i\hbar \frac{\partial}{\partial t} \psi' = \left\{ \frac{1}{2m} \left( \mathbf{P} - q(\nabla \omega) \right)^2 + V - q \frac{\partial \omega}{\partial t} \right\} \psi',$$

dont la forme semble irrémédiablement différente de l'équation (II.1). Sophie doit-elle alors abandonner tout espoir de voir  $\psi$  et  $\psi'$  offrir des descriptions du quanton équivalentes?

Les malheurs de Sophie viennent apparemment des opérations de dérivation à travers lesquelles l'exponentielle passe, certes, mais en laissant quelques plumes,  $\partial \omega / \partial t$  ou  $\nabla \omega$ , lorsque  $\omega$  est variable, plumes qui se retrouvent dans l'équation d'évolution de  $\psi'$ . Surmontant sa déception, Sophie décide d'examiner ces opérations et commence par adopter les notations plus commodes, et tout aussi populaires que modernes:

$$\partial_t \stackrel{\text{df}}{=} \frac{\partial}{\partial t},$$

et, de même, pour les composantes du nabla,

$$\partial_x \stackrel{\text{df}}{=} \frac{\partial}{\partial x}, \quad \partial_y \stackrel{\text{df}}{=} \frac{\partial}{\partial y}, \quad \partial_z \stackrel{\text{df}}{=} \frac{\partial}{\partial z}.$$

Ces dérivées sont dites *covariantes*. Mais covariantes par rapport à quoi au juste? Les coordonnées  $x, y, z$  d'un point représentent les composantes du vecteur position de ce point, par rapport à un point origine, sur trois vecteurs indépendants, dits de base:

$$\mathbf{r} = x\mathbf{e}_x + y\mathbf{e}_y + z\mathbf{e}_z,$$

tandis que sur une autre base,

$$\mathbf{r} = x'\mathbf{e}'_x + y'\mathbf{e}'_y + z'\mathbf{e}'_z.$$

Prenons par exemple le plus simple des changements de base:  $\mathbf{e}'_x \stackrel{\text{df}}{=} 2\mathbf{e}_x$ ,  $\mathbf{e}_y$  et  $\mathbf{e}_z$  restant inébranlables. Alors,  $x' = x/2$ , et

$$\partial_{x'} = \frac{\partial x}{\partial x'} \partial_x = 2\partial_x.$$

La dérivée  $\partial_x$  se transforme comme le vecteur de base correspondant; elle co-varie avec  $\mathbf{e}_x$ . De façon générale, les dérivées  $\partial_x, \partial_y, \partial_z$  sont covariantes avec les vecteurs de base  $\mathbf{e}_x, \mathbf{e}_y, \mathbf{e}_z$  dans toute transformation linéaire de ceux-ci.

Dans le cas des transformations de jauge globales, Sophie n'avait aucun problème parce que les dérivées partielles de  $\psi$  se transformaient exactement comme  $\psi$  elle-même, par un simple facteur, c'est-à-dire de manière covariante. Imaginative, elle décide de chercher une nouvelle manière de mesurer les variations de  $\psi$  d'un instant à l'autre, d'un point à un autre, qui soit covariante par rapport aux transformations de jauge locales. Il lui faut des dérivées covariantes  $D_t, D_x, D_y, D_z$ , telles qu'au cours de la transformation (II.3) elles deviennent:

$$\begin{aligned} D_t\psi &\rightarrow D'_t\psi' = e^{i\frac{q}{\hbar}\omega} D_t\psi, \\ \mathbf{D}\psi &\rightarrow \mathbf{D}'\psi' = e^{i\frac{q}{\hbar}\omega} \mathbf{D}\psi. \end{aligned}$$

Prenons d'abord la dérivée par rapport au temps. Sophie sait déjà que

$$\partial_t\psi' = e^{i\frac{q}{\hbar}\omega} (\partial_t + i\frac{q}{\hbar}(\partial_t\omega))\psi,$$

qu'elle peut réécrire:

$$(\partial_t - i\frac{q}{\hbar}(\partial_t\omega))\psi' = e^{i\frac{q}{\hbar}\omega} i\frac{q}{\hbar}\partial_t\psi.$$

Cette expression a presque la forme cherchée. Elle suggère d'introduire un "potentiel"  $\phi(\mathbf{r}, t)$  et d'ajouter membre à membre l'identité

$$i\frac{q}{\hbar}\phi\psi' \equiv e^{i\frac{q}{\hbar}\omega} i\frac{q}{\hbar}\phi\psi,$$

pour donner

$$\left(\partial_t + i\frac{q}{\hbar}(\phi - (\partial_t\omega))\right)\psi' = e^{i\frac{q}{\hbar}\omega}\left(\partial_t + i\frac{q}{\hbar}\phi\right)\psi.$$

C'est gagné! La "dérivée"

$$D_t \stackrel{\text{df}}{=} \partial_t + i\frac{q}{\hbar}\phi(\mathbf{r}, t)$$

est bien covariante par rapport aux transformations de jauge locales (II.3) de la fonction-d'onde dans la mesure où conjointement le potentiel se transforme en:

$$\phi(\mathbf{r}, t) \rightarrow \phi'(\mathbf{r}, t) \stackrel{\text{df}}{=} \phi(\mathbf{r}, t) - \partial_t\omega(\mathbf{r}, t).$$

Sophie n'a plus maintenant qu'à procéder de manière analogue pour les dérivées spatiales. Elle a d'abord

$$(\nabla - i\frac{q}{\hbar}(\nabla\omega))\psi' = e^{i\frac{q}{\hbar}\omega}\nabla\psi,$$

relation qu'elle écrit plus généralement en introduisant un autre "potentiel",  $\mathbf{A}(\mathbf{r}, t)$ , et en retranchant<sup>1</sup> membre à membre l'identité

$$i\frac{q}{\hbar}\mathbf{A}\psi' = e^{i\frac{q}{\hbar}\omega}i\frac{q}{\hbar}\mathbf{A}\psi,$$

pour obtenir:

$$\left(\nabla - i\frac{q}{\hbar}(\mathbf{A} + (\nabla\omega))\right)\psi' = e^{i\frac{q}{\hbar}\omega}\left(\nabla - i\frac{q}{\hbar}\mathbf{A}\right)\psi.$$

Sophie trouve ainsi sa "dérivée"

$$\mathbf{D} \stackrel{\text{df}}{=} \nabla - i\frac{q}{\hbar}\mathbf{A}(\mathbf{r}, t),$$

covariante à condition que le potentiel  $\mathbf{A}$  selon la prescription:

$$\mathbf{A}(\mathbf{r}, t) \rightarrow \mathbf{A}'(\mathbf{r}, t) \stackrel{\text{df}}{=} \mathbf{A}(\mathbf{r}, t) + \nabla\omega(\mathbf{r}, t).$$

Partant de l'équation de Schrödinger libre, elle écrit celle-ci en termes de dérivées covariantes, soit:

$$i\hbar D_t\psi = \frac{1}{2m}\left(\frac{\hbar}{i}\mathbf{D}\right)^2\psi,$$

---

<sup>1</sup>Pourquoi, cette fois, retrancher plutôt qu'ajouter? On a le choix, le facteur  $iq/\hbar$  étant par contre essentiel. Disons, pour la lectrice curieuse, qu'avec ce choix Sophie obtient des dérivées  $D_t$ ,  $D_x$ ,  $D_y$ ,  $D_z$  qui sont aussi covariantes par rapport aux transformations de Lorentz, dans la mesure où les potentiels  $\phi$  et  $\mathbf{A}$  se transforment de manière contravariantes, c'est-à-dire comme  $t$ ,  $x$ ,  $y$  et  $z$ . Sophie se prépare ainsi des équations de Maxwell (§II.1.2) dont la forme sera invariante par rapport aux transformations de Lorentz. Le principe d'invariance de jauge  $U(1)$  ne permet quand même pas, à lui seul, de trouver la relativité!



ou, explicitement,<sup>2</sup>

$$i\hbar \frac{\partial}{\partial t} \psi(\mathbf{r}, t) = \left\{ \frac{1}{2m} \left( \mathbf{P} - q\mathbf{A}(\mathbf{r}, t) \right)^2 + q\phi(\mathbf{r}, t) \right\} \psi(\mathbf{r}, t). \quad (\text{II.4})$$

Sophie est alors assurée que cette nouvelle équation de Schrödinger à une forme invariante, à savoir

$$i\hbar D'_t \psi' = \frac{1}{2m} \left( \frac{\hbar}{i} \mathbf{D}' \right)^2 \psi',$$

ou, explicitement,

$$i\hbar \frac{\partial}{\partial t} \psi'(\mathbf{r}, t) = \left\{ \frac{1}{2m} \left( \mathbf{P} - q\mathbf{A}'(\mathbf{r}, t) \right)^2 + q\phi'(\mathbf{r}, t) \right\} \psi'(\mathbf{r}, t), \quad (\text{II.5})$$

par rapport aux transformations

$$\begin{aligned} \psi &\xrightarrow{\omega} \psi' \stackrel{\text{df}}{=} e^{i\frac{q}{\hbar}\omega} \psi, \\ \phi &\xrightarrow{\omega} \phi' \stackrel{\text{df}}{=} \phi - \partial_t \omega, \end{aligned} \quad (\text{II.6})$$

$$\mathbf{A} \xrightarrow{\omega} \mathbf{A}' \stackrel{\text{df}}{=} \mathbf{A} + \nabla \omega. \quad (\text{II.7})$$

Bien entendu, l'équation de Schrödinger (II.4), ou (II.5), peut comporter en outre n'importe quel potentiel  $V(\mathbf{r}, t)$  invariant. Toujours est-il que cette équation peut encore décrire un quanton libre, si  $\phi$  et  $\mathbf{A}$  sont nuls par exemple, ou  $\phi' = -\partial_t \omega$  et  $\mathbf{A}' = \nabla \omega$ . Mais en général le quanton évolue maintenant en interaction avec un environnement caractérisé par le potentiel scalaire  $\phi$  et le potentiel vecteur  $\mathbf{A}$ , collectivement appelés *champs de jauge*. Remarquez que les caractères scalaire ou vecteur des potentiels sont inhérents à la méthode: ces potentiels ont été introduits pour corriger les variations d'opérateurs  $\partial_t$  et  $\nabla$  eux-mêmes scalaire et vectoriel respectivement.<sup>3</sup>

Sophie semble donc en bonne voie d'établir un modèle invariant par rapport aux transformations de jauge locales, puisque les équations d'évolution (II.4) et (II.5) des deux fonctions-d'onde transformées (II.3) ont absolument (ou relativement?) la même forme. Mais ce n'est qu'une partie de l'histoire, car reste à examiner la correspondance entre valeurs des grandeurs physiques dans les deux descriptions.

---

<sup>2</sup>Curieusement cet hamiltonien, indépendamment de toute propriété de transformation des potentiels  $\phi$  et  $\mathbf{A}$  est aussi la forme la plus générale compatible avec l'invariance par rapport aux transformations de Galilée. Je ne connais guère d'autre référence sur cette étonnante cïncidence que [38, 45] et [51].

<sup>3</sup>La nature nécessairement vectorielle de l'un des champs de jauge se traduit plus tard, dans la version quantique de ces champs, par des particules de jauge qui sont vectorielles par rapport aux rotations, autrement dit de spin 1. Atoute théorie de jauge correspond des *bosons de jauge* (les photons dans le cas de la théorie de jauge  $U(1)$ ).

### II.1.1 Grandeurs physiques invariantes

L'hypothétique environnement du quanton doit, pour sa part et s'il existe, être totalement indifférent à ces manipulations de fonctions-d'onde, ce qui signifie que l'on doit pouvoir en faire une description complète en termes indépendants des choix  $\phi, \mathbf{A}$  ou  $\phi', \mathbf{A}'$ , si l'on veut préserver la toute fraîche invariance du modèle. Or la transformation (II.6–II.7) correspond à une “translation interne”, un décalage des valeurs des potentiels (ce n'est pas une translation dans l'espace ou le temps!), et ce sont donc plutôt les différences de ces valeurs qui ont quelque chance de rester indifférentes.

On a une situation analogue en mécanique classique, lorsqu'un environnement donné peut être décrit par une infinité de potentiels équivalents

$$U'(\mathbf{r}) = U(\mathbf{r}) + \text{cte.} \quad (\text{II.8})$$

Ce sont les différences de potentiel

$$U'(\mathbf{r}_2) - U'(\mathbf{r}_1) = U(\mathbf{r}_2) - U(\mathbf{r}_1) \quad \forall \mathbf{r}_1, \mathbf{r}_2$$

qui, évidemment, sont indépendantes de ce choix et se trouvent avoir une signification physique. Entre points voisins,  $\mathbf{r}$  et  $\mathbf{r} + d\mathbf{r}$ , un développement traduit cette égalité par

$$\nabla U'(\mathbf{r}) \cdot d\mathbf{r} = \nabla U(\mathbf{r}) \cdot d\mathbf{r}, \quad \forall \mathbf{r}, d\mathbf{r},$$

soit  $\nabla U' \equiv \nabla U$ , qui est par ailleurs conséquence directe de (II.8). On a ainsi retrouvé que c'est le gradient du potentiel (la force) qui rend compte de l'environnement d'une manière invariante par rapport aux transformations (II.8).

Mais notre cas local (II.6–II.7) est un peu plus subtil. Les dérivées  $\partial_t \omega$  et  $\nabla \omega$  sont des fonctions de  $\mathbf{r}$  et  $t$ ; les “translations” ne sont pas les mêmes partout, tout le temps:

$$\begin{aligned} \phi' &= \phi - \partial_t \omega, \\ A'_i &= A_i + \partial_i \omega, \end{aligned}$$

et l'on en déduit, pour les diverses dérivées des potentiels,

$$\begin{aligned} \partial_t \phi' &= \partial_t \phi - \partial_t \partial_t \omega \\ \partial_i \phi' &= \partial_i \phi - \partial_i \partial_t \omega \\ \partial_t A'_i &= \partial_t A_i + \partial_t \partial_i \omega \\ \partial_j A'_i &= \partial_j A_i + \partial_j \partial_i \omega. \end{aligned}$$

Reste à éliminer ces intempestives dérivées secondes de  $\omega$ . Il n'y a pas trente-six façons mais seulement six. Ajoutant la deuxième et la troisième équations membres à membres, on obtient

$$\partial_i \phi' + \partial_t A'_i = \partial_i \phi + \partial_t A_i, \quad (\text{II.9})$$

soit trois quantités invariantes, tandis qu'avec la quatrième équation il vient

$$\partial_j A'_i - \partial_i A'_j = \partial_j A_i - \partial_i A_j, \quad (\text{II.10})$$

soit six invariants non identiquement nuls, mais dont trois seulement sont indépendants (pour cause d'antisymétrie). Reste la première équation dont on ne peut rien faire car on ne dispose d'aucune autre relation pour éliminer  $\partial_t \partial_t \omega$ .

Les trois grandeurs (II.9) sont les sommes de composantes d'un gradient et de la dérivée temporelle d'un vecteur. Elles sont donc elles-mêmes composantes d'un vecteur. Dans les quantités (II.10) on retrouve les trois composantes d'un rotationnel; encore un vecteur. Les six grandeurs (les "forces") caractérisant de façon invariante l'hypothétique environnement qui se manifeste dans l'équation (II.4), sont donc les deux champs vectoriels

$$\mathbf{E}(\mathbf{r}, t) \stackrel{\text{df}}{=} -\nabla\phi(\mathbf{r}, t) - \partial_t \mathbf{A}(\mathbf{r}, t) \quad (\text{II.11})$$

$$\mathbf{B}(\mathbf{r}, t) \stackrel{\text{df}}{=} \nabla \wedge \mathbf{A}(\mathbf{r}, t). \quad (\text{II.12})$$

En ce qui concerne les grandeurs physiques du quanton lui-même, la question n'est pas absolument triviale, car certaines des quantités qui lui sont ordinairement associées ne sont manifestement pas invariantes, au sens en particulier où leurs valeurs moyennes varient selon le choix  $\psi, \phi, \mathbf{A}$  ou  $\psi', \phi', \mathbf{A}'$ . C'est le cas, par exemple, de  $\mathbf{P}$ ,  $\mathbf{A} \cdot \mathbf{P}$ , etc., contrairement à d'autres grandeurs qui, elles, sont invariantes comme  $\mathbf{r}$ ,  $\mathbf{P} - q\mathbf{A}$ ,  $\mathbf{E} \cdot \mathbf{r}$ , etc. (Exercice 4).

## II.1.2 L'interaction électromagnétique

En résumé, enfourcher un principe d'invariance par rapport à la transformation locale

$$\psi(\mathbf{r}, t) \rightarrow \psi'(\mathbf{r}, t) = e^{i\frac{q}{\hbar}\omega(\mathbf{r}, t)} \psi(\mathbf{r}, t)$$

implique (avec plus ou moins de rigueur, nous allons le voir) l'existence d'une interaction médiée par les potentiels  $\phi$  et  $\mathbf{A}$  (les champs de jauge) dans l'équation de Schrödinger

$$i\hbar\partial_t\psi = \left\{ \frac{1}{2m} (\mathbf{P} - q\mathbf{A})^2 + q\phi \right\} \psi,$$

ainsi que la loi de transformation du champ de jauge

$$\begin{aligned} \phi &\rightarrow \phi' = \phi - \partial_t \omega \\ \mathbf{A} &\rightarrow \mathbf{A}' = \mathbf{A} + \nabla \omega. \end{aligned}$$

Cet environnement dans lequel se meut le quanton doit, s'il est effectivement réalisé, être entièrement caractérisé par les champs invariants de jauge

$$\begin{aligned} \mathbf{E} &= -\nabla\phi - \partial_t \mathbf{A} \\ \mathbf{B} &= \nabla \wedge \mathbf{A}. \end{aligned}$$

La nature, bonne fille, nous gratifie en fait de ce type d'environnement, à savoir, comme s'en doutait la lectrice pas si ingénue, le champ électromagnétique  $\mathbf{E}, \mathbf{B}$  lorsqu'y évolue un quanton de charge électrique  $q$ .

Notre simple principe d'invariance nous a donc apporté directement l'hamiltonien d'un quanton en interaction électromagnétique, en faisant l'économie du discutable principe de correspondance. Remarquons que, conséquence immédiate de leurs définitions (II.11) et (II.12), les champs  $\mathbf{E}$  et  $\mathbf{B}$  satisfont pour leur part les équations

$$\begin{aligned}\nabla \cdot \mathbf{B} &= 0 \\ \nabla \wedge \mathbf{E} + \partial_t \mathbf{B} &= 0,\end{aligned}$$

dans lesquelles on reconnaît les deux *équations de Maxwell homogènes*.

Notre principe, étant appliqué à l'équation de Schrödinger du quanton, devrait être muet en ce qui concerne la dynamique de l'environnement de celui-ci, c'est-à-dire le procédé de fabrication des champs  $\mathbf{E}$  et  $\mathbf{B}$ . Et pourtant... on peut pousser le bouchon plus loin. Dans notre recherche des équations du mouvement possibles pour  $\mathbf{E}$  et  $\mathbf{B}$  — c'est-à-dire des équations contraignant les dérivées spatiales et temporelles de ces champs — nous avons déjà usé des dérivées spatiales de  $\mathbf{E}$  et temporelle de  $\mathbf{B}$  sous une forme vectorielle, et des dérivées spatiales de  $\mathbf{B}$  sous forme scalaire. Restent disponibles...

- les dérivées spatiales de  $\mathbf{E}$  sous forme de combinaison scalaire, qui permettent de définir une *densité de charge électrique*,

$$\rho_s \stackrel{\text{df}}{=} \varepsilon_0 \nabla \cdot \mathbf{E},$$

- où la constante  $\varepsilon_0$  définit l'unité adoptée pour mesurer la charge électrique;
- la dérivée temporelle de  $\mathbf{E}$ , permettant de définir une *densité de courant*,  $\partial_t \mathbf{E}$ , à un facteur près, et même mieux, exactement  $-\varepsilon_0 \partial_t \mathbf{E}$ , si l'on tient à la *conservation locale de la charge* (qui doit se traduire par une *équation de continuité*, du type  $\partial_t \rho_s + \nabla \cdot \mathbf{j}_s = 0$ );
- les dérivées spatiales de  $\mathbf{B}$  sous forme de combinaison vectorielle, avec lesquelles on peut définir une autre densité de courant, proportionnelle à  $\nabla \wedge \mathbf{B}$ , toujours conservée car la divergence d'un rotationnel est nulle.

On peut donc construire une théorie dans laquelle les contributions électrique et magnétique au courant, conservé, sont regroupées, dans la proportion  $1/\mu_0$ ,

$$\mathbf{j}_s \stackrel{\text{df}}{=} -\varepsilon_0 \frac{\partial \mathbf{E}}{\partial t} + \frac{1}{\mu_0} \nabla \wedge \mathbf{B},$$

où  $\mu_0$  est une constante de plus de la théorie. C'est même précisément ce qu'il faut faire si l'on croit par ailleurs que cette théorie doit être relativiste<sup>4</sup>, c'est-à-dire avoir

---

<sup>4</sup> *Credo* qui ne dispense quand même pas d'effectuer les expériences de magnétostatique permettant de mesurer  $\mu_0$ .

une forme qui soit invariante par transformations de Lorentz: dans le vide (hors des sources)  $\mathbf{E}$  et  $\mathbf{B}$  sont alors solutions d'équations d'onde, de forme invariante à condition que

$$\mu_0 = \frac{1}{\varepsilon_0 c^2},$$

où  $c$  est la constante fondamentale de la théorie de la relativité einsteinienne, dite aussi — pour de malencontreuses raisons historiques — vitesse de la lumière.<sup>5</sup> En résumé,  $\mathbf{E}$  et  $\mathbf{B}$  satisfont donc aussi deux *équations de Maxwell inhomogènes*, régissant leur fabrication à partir des sources  $\rho_s(\mathbf{r}, t)$  et  $\mathbf{j}_s(\mathbf{r}, t)$  (généralement vendues par E.D.F.):

$$\nabla \cdot \mathbf{E} = \frac{\rho_s}{\varepsilon_0} \quad (\text{II.13})$$

$$\nabla \wedge \mathbf{B} - \frac{1}{c^2} \partial_t \mathbf{E} = \frac{\mathbf{j}_s}{\varepsilon_0 c^2}. \quad (\text{II.14})$$

Ces succès ne sauraient estomper les quelques zones d'ombre qui font apprécier le relief de toute théorie physique. Si nous sommes maintenant habitués à l'indifférence de la théorie quantique vis à vis d'un choix de phase global, imposer cette invariance en tous lieux, tous temps, manque de motivation évidente. Sur ce point, on n'a guère réalisé de progrès par rapport à l'argument ancien de cohérence avec l'idée d'interaction locale. L'équation de Schrödinger conduit l'évolution de la fonction-d'onde par pilotage à vue, en se contentant d'informations puisées sur place (les valeurs de la fonction-d'onde et des potentiels) ou au voisinage immédiat (par le jeu des opérateurs dérivatifs). S'il en est ainsi pour tous les phénomènes, pourquoi une vénusienne et un martien devraient-ils utiliser la même convention de phase relative en chaque point d'espace-temps pour les fonctions-d'onde qu'ils assignent respectivement au quanton qu'ils étudient (sur la Terre par exemple)?

L'impossibilité de spécifier instantanément un changement de phase universel est une autre justification pour la localité de la transformation de jauge (II.3). A ce stade, invoquer la relativité einsteinienne plutôt que galiléenne semble tout aussi arbitraire, mais la plus grande généralité de celle-là nous permet par ce rapprochement d'espérer tenir avec l'invariance de jauge locale un nouveau principe universel. Bien que l'opérationnalisme commode de l'énoncé de ce genre de principe en reste au stade formel, puisque les transformations dont il s'agit sont pour la plupart potentielles (changement de repère inertiel par exemple) sinon abstraites (changement de phase), les concepts de transformation et d'invariance sont devenus une marque distinctive de la physique de ce siècle : *«...l'idée est que le changement n'est pas dans les choses, il est simplement dans le changement de repère de l'observateur. Or l'observateur change constamment, ne serait-ce que parce qu'il vieillit. C'est une*

---

<sup>5</sup>C'est la lumière (le champ électromagnétique, une interaction parmi d'autres) qui se propage à la vitesse fondamentale  $c$ , et non pas  $c$  qui est égale à la vitesse de la lumière!

problématique très moderne. La physique fondamentale est fondée sur les règles qui permettent d'obtenir le consensus intersubjectif à partir des visions d'observateurs différents.»[72]

D'autre part, la forme de l'hamiltonien adopté dans l'équation (II.4) n'est certainement pas la seule qui soit invariante par transformation de jauge locale. Il n'y a aucune difficulté à imaginer des termes tout aussi invariants et que l'on peut donc ajouter sans vergogne à l'hamiltonien, par exemple (pas tout à fait au hasard)

$$q(\nabla\phi + \partial_t\mathbf{A}) \cdot \mathbf{r} \quad \text{ou} \quad -\frac{q}{2m}(\nabla \wedge \mathbf{A}) \cdot (\mathbf{r} \wedge \mathbf{P}).$$

Pour notre soulagement, il faut nous contenter de noter que parmi l'*infinité* de modèles invariants de jauge ainsi imaginables, la nature choisit finalement de n'en réaliser qu'un, celui où comme par hasard ne figure aucun de ces termes additionnels, c'est à dire le modèle dont l'hamiltonien est obtenu en administrant à (II.1) la simple prescription

$$\begin{aligned} \mathbf{P} &\rightarrow \mathbf{P} - q\mathbf{A} \\ i\hbar\partial_t &\rightarrow i\hbar\partial_t - q\phi, \end{aligned}$$

dite de *couplage minimal*.

A propos de couplage justement, le principe d'invariance de jauge n'apporte aucune restriction aux valeurs de la constante  $q$ , la *charge électrique* du quanton, qu'il faut donc se résigner à déterminer empiriquement, selon les cas d'espèce. Notre modèle de l'interaction électromagnétique ne rend donc compte ni de la stricte contenance de la nature qui restreint ses charges électriques à des valeurs multiples d'une charge  $e$ , pour le coup fondamentale<sup>6</sup>, ni *a fortiori* de la valeur de cette charge, ou de la valeur de la *constante de structure fine*

$$\alpha \stackrel{\text{df}}{=} \frac{e^2/4\pi\epsilon_0}{\hbar c} = \frac{1}{137,035\,989\,5(61)}.$$

Qu'elle soit sans dimension n'est pas son moindre intérêt pratique de cette dernière et confirme son caractère fondamental... qui n'a rien d'éternel, puisqu'à la merci d'une éventuelle théorie qui viendrait prédire sa valeur (n'est fondamental que ce dont on ignore l'origine!).<sup>7</sup>

Dans la jactance des physiciens, le procédé que nous avons utilisé pour parvenir à l'hamiltonien (II.4) constitue une *théorie de jauge de l'interaction électromagnétique*. A dessein, j'ai pour un temps qualifié celle-ci de modèle, selon la hiérarchie discutée dans le préluce. Ce modèle/théorie s'inscrit en effet strictement dans le

<sup>6</sup>C'est plutôt  $e/3$  qui est maintenant, par l'intervention des quarks, fondamentale. Mais cela ne change rien à la question posée.

<sup>7</sup>C'est encore moins une constante, puisque sa valeur fluctue au gré des déterminations empiriques comme nous le rappelle l'incertitude qui lui est attachée; voir [48].

cadre de la théorie quantique, et si ses succès sont nombreux, à commencer par la prédiction des niveaux de l'atome d'hydrogène, les situations qu'il se révèle incapable de décrire ne manquent pas, soit qu'elles échappent au modèle du quanton (la création de rayonnement lors d'une transition entre niveaux d'un atome par exemple, ou l'annihilation d'un électron et d'un positron et leur réincarnation en deux photons), soit que l'interaction apparaisse manifestement d'une autre nature (comme c'est le cas entre les nucléons, dans un noyau). Le champ d'application du modèle reste plus restreint que celui de la théorie quantique et les limites que connaît celui-là n'exigent aucune remise en cause de celle-ci.

## II.2 Quanton de spin 1/2

La discussion précédente, dans laquelle je n'ai pas fait intervenir le spin ne concerne donc que le quanton de spin nul ou les situations (elles sont nombreuses) où les effets de spin sont énergétiquement négligeables. Bien que ce ne soit pas une nécessité dans ce hors d'œuvre à la théorie quantique du rayonnement, je ne peux résister au plaisir d'aborder la question de l'invariance de jauge pour un quanton de spin 1/2.

Rappelons tout d'abord que, contrairement à une opinion encore répandue<sup>8</sup>, il n'est nul besoin de faire appel à la relativité einsteinienne et au modèle de Dirac pour établir la généalogie de la grandeur physique spin. Le consensus souhaitable entre observateurs dont les repères diffèrent par une rotation dicte le caractère vectoriel — ainsi que l'algèbre des composantes — de l'opérateur qui génère, au cours de cette rotation, la transformation du vecteur d'état d'un système quantique<sup>9</sup>, caractère vectoriel qui s'exprime par les commutateurs des composantes:

$$[J_i, J_j] = i\hbar \varepsilon_{ijk} J_k. \quad (\text{II.15})$$

L'opérateur  $\mathbf{J}$ , *générateur* des rotations, est plus traditionnellement appelé *moment angulaire* du système. Chacun des indices de composantes,  $i$ ,  $j$  ou  $k$ , dans les relations (II.15), peut prendre les trois valeurs 1, 2 ou 3 (ou  $x$ ,  $y$  ou  $z$ ). Les 27 nombres  $\varepsilon_{ijk}$  (les composantes du *tenseur de Levi-Civita*) sont définis par:

- leur totale antisymétrie ( $\varepsilon_{ijk}$  change de signe dans toute transposition de deux indices);
- la valeur  $\varepsilon_{123} \stackrel{\text{df}}{=} 1$ .

On a ainsi, par exemple,  $\varepsilon_{321} = -1$ ,  $\varepsilon_{112} = 0$ , *etc.* Enfin, nous utilisons dorénavant la *convention d'Einstein*: toute répétition d'indice — une fois seulement, sinon il y a une erreur! — dans un monôme implique une sommation sur les trois valeurs accessibles.

<sup>8</sup>Prenez le florilège offert dans [45, app. A].

<sup>9</sup>Ces considérations sur l'invariance par rotation sont à la base de la théorie quantique et devraient vous être connues; nous y reviendrons par la suite (page 108). Quoi qu'il en soit, vous pouvez consulter, avec profit, les références [53, § XIII-13] et/ou [51, vol. II].

Dans le cas du quanton, on peut, à l'aide de sa position  $\mathbf{R}$  et de son impulsion  $\mathbf{P}$  définir une nouvelle grandeur physique

$$\mathbf{L} \stackrel{\text{df}}{=} \mathbf{R} \wedge \mathbf{P}$$

appelée *moment orbital*. A cause des relations de commutation entre composantes de  $\mathbf{R}$  et de  $\mathbf{P}$  (conséquences de l'invariance par translation),  $[R_i, P_j] = i\hbar\delta_{ij}$ , les composantes de  $\mathbf{L}$  se trouvent obéir aux relations (II.15), et la lectrice peut vérifier qu'il en va de même pour la grandeur physique associée (Exercice 7)

$$\mathbf{S} \stackrel{\text{df}}{=} \mathbf{J} - \mathbf{L} \quad (\text{II.16})$$

appelée *spin* du quanton

Mais le caractère vectoriel de  $\mathbf{R}$  dicte le comportement de ses composantes lors des rotations, et donc les valeurs des commutateurs de celles-là avec les composantes du générateur de celles-ci,

$$[J_i, R_j] = i\hbar\varepsilon_{ijk}R_k,$$

et de même pour les composantes de  $\mathbf{P}$ . Dans ces conditions, les composantes du spin  $\mathbf{S}$  (définition (II.16)) commutent avec celles de  $\mathbf{R}$  et de  $\mathbf{P}$  et ne sont donc fonctions ni des unes ni des autres. Nous nous retrouvons ainsi avec une nouvelle grandeur physique indépendante associée au quanton, et la possibilité, entre autres, de construire un espace des états comme produit de l'espace des états de  $\mathbf{R}$  et  $\mathbf{P}$  et du sous-espace des états de  $\mathbf{S}$  correspondant à une valeur propre de  $\mathbf{S}^2$  déterminée. Que cette valeur propre, de la forme  $\hbar^2 s(s+1)$  avec  $2s$  entier naturel (du fait que les composantes du spin satisfont (II.15)), soit fixée *ad æternam* est donc une caractéristique de construction de ce nouveau modèle de quanton, justifiée par sa pertinence pour toutes sortes de situations. Envisager une variation de  $s$  au cours du temps signerait l'arrêt de mort du quanton, éventualité qui sort complètement du cadre du modèle, et en vue de laquelle il faudra justement élaborer un modèle plus performant, la *théorie quantique des champs*.

Dans le cas d'un quanton de spin  $s = 1/2$ , il est commode d'adopter la représentation en position, dans laquelle la fonction-d'onde  $\psi(\mathbf{r}, t)$  a deux composantes; à chaque valeur de  $\mathbf{r}$  et  $t$ , elle associe un vecteur de l'espace des états de spin.<sup>10</sup>

L'hamiltonien du quanton libre, avec ou sans spin, est  $\mathbf{P}^2/2m$ . On ne voit donc pas *a priori* comment le spin pourrait éventuellement jouer un rôle dans l'interaction à naître de l'application du principe d'invariance de jauge locale. Un moyen sûr pour que le spin intervienne dans l'hamiltonien final serait de le faire figurer dans l'hamiltonien libre de départ. A cette fin, adoptons pour les états de spin une base

---

<sup>10</sup>Autrement dit,  $\psi(\mathbf{r}, t)$  est un champ de spineurs, comme il existe des champs scalaires ou vectoriels, voir [53, § XIII-20].



standard ( $S_z$  diagonale). Les matrices représentatives des composantes de  $\mathbf{S}$  dans le sous-espace  $s = 1/2$  nous permettent de définir les trois *matrices de Pauli*

$$\sigma_i \stackrel{\text{df}}{=} \frac{2}{\hbar} S_i,$$

soit

$$\sigma_x = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma_y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad \sigma_z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

L'intérêt pratique de ces matrices apparaît de lui-même lorsqu'on réalise les merveilleuses propriétés algébriques dont elles sont douées:

$$\sigma_i \sigma_j \equiv i \varepsilon_{ijk} \sigma_k + \delta_{ij} I. \quad (\text{II.17})$$

Ainsi, les matrices de Pauli anticommulent, elles ont pour carré la matrice unité  $I$ , et surtout engendrent une remarquable identité,

$$(\boldsymbol{\sigma} \cdot \mathbf{A})(\boldsymbol{\sigma} \cdot \mathbf{B}) \equiv \mathbf{A} \cdot \mathbf{B} + i \boldsymbol{\sigma} \cdot (\mathbf{A} \wedge \mathbf{B}), \quad (\text{II.18})$$

pour deux opérateurs vectoriels  $\mathbf{A}$  et  $\mathbf{B}$  par ailleurs quelconques.

### II.2.1 L'équation de Schrödinger libre réduite au premier ordre

En vertu de l'identité (II.18), on a  $(\boldsymbol{\sigma} \cdot \mathbf{P})^2 = \mathbf{P}^2$ , et l'équation de Schrödinger

$$i\hbar \partial_t \psi(\mathbf{r}, t) = \frac{1}{2m} \{a \mathbf{P}^2 + (1-a)(\boldsymbol{\sigma} \cdot \mathbf{P})^2\} \psi(\mathbf{r}, t) \quad (\text{II.19})$$

est tout aussi acceptable pour un quanton libre ( $a$  est une constante réelle, et la matrice unité  $I$  est sous-entendue partout où elle devrait apparaître en facteur!), mais conduit à des théories invariantes de jauge différentes. Nous disposons maintenant, selon les valeurs de  $a$ , d'une infinité de formes de l'hamiltonien libre équivalentes. Parmi celles-ci il en est une, celle qui correspond à  $a = 0$ , avec laquelle on peut tout aussi bien commencer, et qui va même s'avérer meilleure que les autres:

$$i\hbar \partial_t \psi = \frac{(\boldsymbol{\sigma} \cdot \mathbf{P})^2}{2m} \psi. \quad (\text{II.20})$$

En effet, en définissant la fonction-d'onde à deux composantes

$$\chi(\mathbf{r}, t) \stackrel{\text{df}}{=} -\frac{\boldsymbol{\sigma} \cdot \mathbf{P}}{2m} \psi(\mathbf{r}, t)$$

la lectrice vérifiera que l'équation de Schrödinger (II.20) est réductible au système du premier ordre de rang quatre

$$\begin{pmatrix} \boldsymbol{\sigma} \cdot \mathbf{P} & 2m \\ i\hbar \partial_t & \boldsymbol{\sigma} \cdot \mathbf{P} \end{pmatrix} \begin{pmatrix} \psi \\ \chi \end{pmatrix} = 0,$$

chose impossible avec l'équation (II.19) en général. La raison d'un tel miracle est dans la possibilité de décomposer  $(\boldsymbol{\sigma} \cdot \mathbf{P})^2 \psi$  en  $(\boldsymbol{\sigma} \cdot \mathbf{P})(\boldsymbol{\sigma} \cdot \mathbf{P})\psi$ , soit l'action successive de deux opérateurs scalaires (invariants par rotation), tandis que  $\mathbf{P}^2$  ne peut (qu'il s'agisse d'un quanton de spin 1/2 ou zéro) se décomposer qu'en  $\mathbf{P} \cdot \mathbf{P} \psi$ , où n'apparaît aucun opérateur différentiel du premier ordre *scalaire*. Qu'advient-il de l'équation (II.20) lors d'une transformation de jauge locale du type (II.3)? L'action de  $\mathbf{P}$  sur  $\psi$  a déjà été calculée (page 13); on en déduit

$$\boldsymbol{\sigma} \cdot \mathbf{P} \psi = e^{-i\frac{q}{\hbar}\omega} \boldsymbol{\sigma} \cdot (\mathbf{P} - q\nabla\omega) \psi'.$$

Ensuite de quoi — miracle des spineurs — le calcul est en fait plus simple que dans le cas du spin nul car une simple réitération de cette expression donne

$$(\boldsymbol{\sigma} \cdot \mathbf{P})^2 = e^{-i\frac{q}{\hbar}\omega} \left( \boldsymbol{\sigma} \cdot (\mathbf{P} - q(\nabla\omega)) \right)^2 \psi',$$

sans plus faire appel à la divergence du produit d'un scalaire et d'un vecteur. La fonction transformée  $\psi'$  évolue donc selon

$$i\hbar\partial_t\psi' = \left\{ \frac{1}{2m} \left( \boldsymbol{\sigma} \cdot (\mathbf{P} - q(\nabla\omega)) \right)^2 - q(\partial_t\omega) \right\} \psi',$$

et l'invariance de jauge locale se trouve ainsi finalement assurée pour l'équation de Schrödinger

$$i\hbar\partial_t\psi = \left\{ \frac{1}{2m} (\boldsymbol{\sigma} \cdot (\mathbf{P} - q\mathbf{A}))^2 + q\phi \right\} \psi.$$

## II.2.2 Hamiltonien de Pauli et facteur $g$

Pour étudier le rôle de la nouvelle grandeur de spin et déterminer si toutes ces spéculations ont un quelconque rapport avec l'interaction électromagnétique, essayons d'exprimer l'hamiltonien en fonction des champs  $\mathbf{E}$  et  $\mathbf{B}$  (définitions (II.11–II.12)), *a priori* plus significatifs que  $\mathbf{A}$  et  $\phi$ , puisque invariants. Grâce à l'identité (II.18), on a:

$$\begin{aligned} (\boldsymbol{\sigma} \cdot (\mathbf{P} - q\mathbf{A}))^2 &= (\mathbf{P} - q\mathbf{A})^2 - iq\boldsymbol{\sigma} \cdot (\mathbf{P} \wedge \mathbf{A} + \mathbf{A} \wedge \mathbf{P}) \\ &= (\mathbf{P} - q\mathbf{A})^2 - q\hbar\boldsymbol{\sigma} \cdot (\nabla \wedge \mathbf{A} + \mathbf{A} \wedge \nabla). \end{aligned}$$

Reste à développer l'action de l'opérateur  $\nabla$  sur le produit  $\mathbf{A}\psi$ . On a (Exercice 1):

$$\nabla \wedge \mathbf{A}\psi = (\nabla \wedge \mathbf{A})\psi - \mathbf{A} \wedge \nabla\psi,$$

soit, en fonction du spin lui-même — plutôt que  $\boldsymbol{\sigma}$  — et du champ  $\mathbf{B}$ ,

$$i\hbar\partial_t\psi = \left\{ \frac{1}{2m} (\mathbf{P} - q\mathbf{A})^2 - \frac{q}{m} \mathbf{S} \cdot \mathbf{B} \right\} \psi.$$

Au second membre apparaît l'*hamiltonien de Pauli*, introduit dans les débuts de la théorie quantique par analogie formelle avec l'hamiltonien classique d'une particule chargée porteuse d'un moment magnétique.

En particulier, dans le cas

$$\phi(\mathbf{r}, t) = 0, \quad \mathbf{A}(\mathbf{r}, t) = \begin{cases} -By/2 \\ Bx/2 \\ 0 \end{cases},$$

pas (totalement) académique puisqu'il représente un champ magnétique uniforme et constant

$$\mathbf{E}(\mathbf{r}, t) = 0, \quad \mathbf{B}(\mathbf{r}, t) = \begin{cases} 0 \\ 0 \\ B \end{cases},$$

on a

$$\nabla \cdot \mathbf{A} = 0,$$

et

$$\mathbf{A} \cdot \mathbf{P} = B(xP_y - yP_x)/2 = \frac{1}{2}\mathbf{B} \cdot \mathbf{L}.$$

L'équation d'évolution prend alors la forme

$$i\hbar\partial_t\psi = \left\{ \frac{\mathbf{P}^2}{2m} - \frac{q}{2m}\mathbf{B} \cdot (\mathbf{L} + 2\mathbf{S}) + \frac{q^2B^2}{8m}(x^2 + y^2) \right\} \psi \quad (\text{II.21})$$

où les deuxième et troisième termes du second membre nous font signe que leurs contributions sont respectivement de types *paramagnétique* et *diamagnétique*.<sup>11</sup>

Le terme paramagnétique peut s'écrire sous une forme plus canonique

$$-\frac{q\hbar}{2m}\mathbf{B} \cdot \left( \frac{\mathbf{L}}{\hbar} + 2\frac{\mathbf{S}}{\hbar} \right),$$

où apparaît le *magnéton de Bohr*, caractéristique du quanton de charge  $q$  et de masse  $m$ :

$$\mu_B \stackrel{\text{df}}{=} \frac{q\hbar}{2m}.$$

Historiquement, l'étonnant facteur 2 devant  $\mathbf{S}$ , appelé *facteur  $g$*  du quanton, a d'abord été déterminé empiriquement, dans le cas de l'électron, pour rendre compte de la déviation à travers un aimant de Stern et Gerlach des faisceaux atomiques d'éléments comportant un électron de valence sur une couche  $s$ : H, Li, Na, K et Ag<sup>12</sup>. Le calcul des *facteurs de Landé* et l'explication de l'*effet Zeeman anomal* ont ensuite confirmé cette anomalie du facteur  $g$  de l'électron.

<sup>11</sup>Voir par exemple [19, complément D<sub>VII</sub>].

<sup>12</sup>Dans ces éléments, les contributions des autres électrons, ainsi que celle du moment orbital de l'électron de valence, sont nulles. Comme pour la plupart des atomes, on peut de plus négliger, en première approximation, la contribution du spin du noyau puisqu'elle est inversement proportionnelle à sa masse; voir l'équation (II.21).

Notre principe d'invariance de jauge locale nous a donc finalement conduits à une interaction tout à fait réelle, avec en particulier ce facteur  $g = 2$  encore souvent attribué à la relativité einsteinienne.<sup>13</sup> Mais nous venons de voir (si l'on peut dire!) qu'il n'y a à cela aucune raison, comme on pouvait s'y attendre d'ailleurs pour une quantité sans dimension, indépendante de la vitesse fondamentale  $c$ . Dans le passé, cette évidence était cachée, car c'est grâce au choix d'unités du *Système International* pour la charge électrique et le champ magnétique que la constante  $c$  ne figure pas dans l'équation d'évolution (II.21). Il ne faudrait évidemment pas en conclure hâtivement que  $c$  n'intervient pas dans la théorie de l'électromagnétisme; le système légal permet juste de repousser cette constante dans les équations de Maxwell inhomogènes (II.13–II.14) qui engendrent les champs, et d'où elle ne peut être éliminée que par l'adoption d'une unité commune pour longueurs et temps.<sup>14</sup>

C'est Dirac qui le premier a justifié théoriquement l'existence du spin  $s = 1/2$  et la valeur du facteur  $g = 2$  en établissant une équation d'évolution quantique relativiste pour l'électron; d'où la légende persistante sur l'origine relativiste de ces deux grandeurs. Il n'en est finalement rien comme j'espère en avoir convaincu la lectrice. Le succès de Dirac sur ce point est à mettre à l'actif d'une utilisation cohérente de l'invariance par rotation et du “truc” heuristique de linéarisabilité de l'équation d'évolution, tout comme nous l'avons fait pour justifier le choix  $a = 0$  dans l'équation (II.19). Abandonnant cette restriction, la lectrice curieuse<sup>15</sup> trouvera un facteur  $g = 2(1 - a)$ .

La relativité einsteinienne trouve quand même son mot à dire dans des corrections au facteur  $g$ , de l'ordre de  $10^{-3}$  (dans le cas de l'électron ou du muon), dont la petitesse témoigne d'un exceptionnel succès de notre modèle. L'existence de particules dont le facteur  $g$  diffère de 2, par exemple le proton avec  $g_p = 5,6$ , ou le neutron avec  $g_n = \infty$  (puisque bien que de charge nulle, le neutron voit son spin interagir avec le champ magnétique; en fait  $g_n = -3,8$  rapporté à la charge du proton), nous signifie clairement qu'il y a des limites au domaine de validité du modèle du quanton. L'invariance par rotation impose toujours la forme scalaire  $\mathbf{B} \cdot \mathbf{S}$  pour l'interaction du spin et du champ magnétique. Mais ce terme, en dépit de sa simplicité, constitue déjà une sonde sensible à la structure interne des particules; le proton (comme l'atome ou le noyau) n'est pas élémentaire, c'est une particule composite dont la description exige, à partir d'un certain degré de précision, d'autres grandeurs physiques que les simples  $\mathbf{R}$ ,  $\mathbf{P}$  et  $\mathbf{S}$  constitutives du modèle du quanton.

---

<sup>13</sup>Voir encore [45].

<sup>14</sup>Légaliser cette possibilité, c'est justement accepter la relativité einsteinienne [5].

<sup>15</sup>Elle pourra lire à ce sujet la référence [39], ou même traiter l'exercice 9.

### II.3 Matière quantique et rayonnement classique

Pour établir le modèle semi-classique du rayonnement, partons de l'hamiltonien en représentation des positions

$$H = \sum_{i=1}^Z \left\{ \frac{1}{2m} \left( \mathbf{P}_i - q\mathbf{A}(\mathbf{r}_i, t) \right)^2 + q\phi(\mathbf{r}_i, t) \right\} + V(\mathbf{r}_1, \dots, \mathbf{r}_Z, t) \quad (\text{II.22})$$

qui permet de décrire un ensemble de  $Z$  quantons (ici de mêmes masse et charge)...

- dans le potentiel  $V$ , représentant aussi bien un potentiel extérieur qu'une interaction mutuelle;
- et dans le potentiel électromagnétique  $\phi, \mathbf{A}$ .

Il peut s'agir par exemple des  $Z$  électrons d'un atome (pour lesquels  $V$  est alors l'interaction coulombienne avec le noyau et entre les électrons eux-mêmes), dans la mesure où la masse du noyau oppose une inertie suffisante à toute tentative d'entraînement et où les contributions des spins des électrons ne sont pas importantes.

La distinction, au sein de la seule interaction électromagnétique, entre interaction avec le rayonnement et interaction coulombienne peut sembler aussi arbitraire que dépendante du choix de jauge; j'y reviendrai lors de la discussion de ce choix. Ce modèle, modifié le cas échéant pour tenir compte des effets négligés, est connu pour fournir d'excellents résultats concernant quantités de processus tels que

- la désexcitation d'un niveau atomique (stimulée) ou son excitation (et donc l'effet photo-électrique) pour ce qui est de l'interaction avec le rayonnement, de l'interaction électron-électron et électron-noyau;
- l'effet Zeeman pour ce qui est des interactions statiques.

Remarquons tout d'abord que les potentiels  $\mathbf{A}(\mathbf{r}, t)$  et  $\phi(\mathbf{r}, t)$  peuvent être des fonctions compliquées de la position  $\mathbf{r}$ , ce qui ne va pas arranger les choses dans l'hamiltonien (II.22) où ces potentiels deviennent des fonctions des opérateurs positions des quantons. En fait, de simples raisons techniques lors de l'établissement d'un modèle quantique du rayonnement justifieront amplement l'écriture de cet hamiltonien sous une forme différente que je vais maintenant présenter.

En développant l'expression (II.22), on peut mettre en évidence deux contributions:

$$H = H_{\text{mat}} + H_{\text{int}} \quad (\text{II.23})$$

dont l'une, l'hamiltonien de la matière (l'atome isolé, ou plutôt ses électrons, dans notre exemple)

$$H_{\text{mat}} \stackrel{\text{df}}{=} \sum_i \frac{\mathbf{P}_i^2}{2m} + V(\mathbf{r}_1, \dots, \mathbf{r}_Z, t), \quad (\text{II.24})$$

est indépendante des potentiels du rayonnement, tandis que l'autre,

$$H_{\text{int}} \stackrel{\text{df}}{=} \sum_i \left\{ -\frac{q}{2m} (\mathbf{P}_i \cdot \mathbf{A}(\mathbf{r}_i, t) + \mathbf{A}(\mathbf{r}_i, t) \cdot \mathbf{P}_i) + \frac{q^2}{2m} \mathbf{A}^2(\mathbf{r}_i, t) + q\phi(\mathbf{r}_i, t) \right\}, \quad (\text{II.25})$$

vient s'ajouter lorsqu'un rayonnement extérieur interagit avec cette matière.

Pour simplifier la dépendance fonctionnelle par rapport aux *variables dynamiques* que sont ici les opérateurs position des quantons  $\mathbf{R}_i$  qui se sont substitués au *paramètre*  $\mathbf{r}$  position du champ, introduisons l'opérateur *densité de matière* en représentation de position,

$$\rho(\mathbf{r}, \mathbf{r}_1, \dots, \mathbf{r}_Z) \stackrel{\text{df}}{=} \sum_{i=1}^Z \delta^3(\mathbf{r} - \mathbf{r}_i), \quad (\text{II.26})$$

où  $\mathbf{r}$  est un paramètre, tandis que les  $\mathbf{r}_i$  sont des valeurs propres des variables dynamiques  $\mathbf{R}_i$ . On obtient ainsi des représentations intégrales des termes d'interaction,

$$\begin{aligned} \sum_i \phi(\mathbf{r}_i, t) &= \int d^3\mathbf{r} \rho(\mathbf{r}, \mathbf{r}_1, \dots, \mathbf{r}_Z) \phi(\mathbf{r}, t) \\ \sum_i \mathbf{A}^2(\mathbf{r}_i, t) &= \int d^3\mathbf{r} \rho(\mathbf{r}, \mathbf{r}_1, \dots, \mathbf{r}_Z) \mathbf{A}^2(\mathbf{r}, t), \end{aligned}$$

où les fonctions compliquées que peuvent être  $\phi$  et  $\mathbf{A}$  ne dépendent plus que du paramètre  $\mathbf{r}$ , toute la dépendance dans les variables dynamiques  $\mathbf{r}_i$  étant reportée dans la densité, sous la forme canonique de “fonctions” de Dirac. La même manipulation est faisable sur les termes d'interaction dépendant de l'impulsion, à l'aide de la *densité de courant de matière*<sup>16</sup>

$$\mathbf{j}(\mathbf{r}, \mathbf{r}_1, \dots, \mathbf{r}_Z) \stackrel{\text{df}}{=} \frac{1}{2m} \sum_{i=1}^Z \left( \mathbf{P}_i \delta^3(\mathbf{r} - \mathbf{r}_i) + \delta^3(\mathbf{r} - \mathbf{r}_i) \mathbf{P}_i \right). \quad (\text{II.27})$$

Remarquons que ce courant n'est pas invariant de jauge. L'opérateur  $\mathbf{P}_i/m$  (ici, on a  $\mathbf{P}_i = (\hbar/i)\nabla_{\mathbf{r}_i}$ ) n'est pas l'opérateur *vitesse quantique*  $\mathbf{V}_i$ , plutôt défini par l'ensemble de ses valeurs moyennes

$$\langle \psi(t) | \mathbf{V}_i | \psi(t) \rangle \stackrel{\text{df}}{=} \frac{d}{dt} \langle \psi(t) | \mathbf{R}_i | \psi(t) \rangle.$$

<sup>16</sup>Attention, ce n'est pas une densité de courant électrique; nous n'avons pas encore donné de sens quantique à celle-ci, et il lui manque une charge électrique en facteur pour en avoir la dimension. La même remarque vaut évidemment pour la densité de charge.

On en déduit (Exercice 10) l'expression

$$\mathbf{V}_i = \frac{1}{m} (\mathbf{P}_i - q\mathbf{A}(\mathbf{r}_i, t)),$$

qui est bien invariante de jauge, ainsi que le courant invariant

$$\mathbf{j}_{\text{inv}} \stackrel{\text{df}}{=} \mathbf{j} - \frac{1}{2m} \rho q \mathbf{A}(\mathbf{r}, t),$$

dont les deux termes, de par les signes de leurs contributions dans l'interaction, reçoivent les noms de *courants paramagnétique* et *diamagnétique* respectivement.

A l'aide de ces quantités, l'hamiltonien d'interaction (II.25) peut finalement s'écrire sous forme d'intégrale de volume

$$H_{\text{int}} = \int d^3\mathbf{r} \mathcal{H}_{\text{int}}(\mathbf{r}, t, \mathbf{r}_1, \dots, \mathbf{r}_Z, \mathbf{P}_1, \dots, \mathbf{P}_Z) \quad (\text{II.28})$$

d'une densité d'hamiltonien

$$\mathcal{H}_{\text{int}} \stackrel{\text{df}}{=} -q \mathbf{j} \cdot \mathbf{A}(\mathbf{r}, t) + \frac{q^2}{2m} \rho \mathbf{A}^2(\mathbf{r}, t) + q \rho \phi(\mathbf{r}, t). \quad (\text{II.29})$$

La factorisation

$$\left( \begin{array}{c} \text{partie} \\ \text{dynamique} \end{array} \right) \times \left( \begin{array}{c} \text{fonction de} \\ \text{l'environnement} \end{array} \right)$$

ainsi réalisée dans chacun des termes de l'interaction va s'avérer particulièrement pratique lors des calculs perturbatifs.

## II.4 Le bon choix de jauge

Au point où nous en sommes, notre modèle est, de façon plus ou moins évidente (le mode de séparation de l'interaction laisse encore quelques doutes), invariant de jauge. C'est même sur cette prescription que nous l'avons érigé.

Si toutes les jauges se valent, il en est qui valent plus que les autres, selon les processus que l'on veut décrire ou les propriétés que l'on veut mettre en évidence. La latitude de jauge est en cela l'exacte analogue de l'arbitraire de la constante du potentiel en mécanique. Par exemple, selon la situation analysée, vous n'hésitez pas à choisir l'origine du potentiel gravitationnel créé par la Terre, à l'infini, au niveau de la mer (lequel?) ou en tout autre endroit. Ainsi, il ne viendrait à l'idée de personne (hormis un enseignant sadique à la veille d'un examen) de décrire les effets électromagnétiques d'une charge ponctuelle immobile classique autrement que par le potentiel  $\phi = (q/4\pi\epsilon_0)(1/r)$ , avec (cela va littéralement sans dire)  $\mathbf{A} = 0$ , alors qu'il existe une infinité d'autres potentiels qui rendront compte, moins brièvement, de la même situation.

### II.4.1 Equations des potentiels

Pour étudier les termes du bon choix de  $\mathbf{A}$  et  $\phi$ , rappelons-nous que ceux-ci sont déterminés par la soumission de leurs dérivées, sous forme des champs électrique (II.11) et magnétique (II.12), aux densités de charge et courant sources  $\rho_s$  et  $\mathbf{j}_s$  (rien à voir avec les densités de matière de la section précédente), par le jeu des équations de Maxwell inhomogènes (II.13–II.14). En substituant les relations (II.11–II.12) dans ces équations, et en se souvenant que, pour un champ vectoriel (Exercice 5),

$$\nabla \wedge (\nabla \wedge \mathbf{A}) = \nabla(\nabla \cdot \mathbf{A}) - \Delta \mathbf{A},$$

on obtient les *équations des potentiels*

$$\begin{aligned} \Delta \phi + \partial_t \nabla \cdot \mathbf{A} &= -\frac{\rho_s}{\varepsilon_0} \\ \Delta \mathbf{A} - \frac{1}{c^2} \partial_t \partial_t \mathbf{A} - \nabla \left( \nabla \cdot \mathbf{A} + \frac{1}{c^2} \partial_t \phi \right) &= -\frac{\mathbf{j}_s}{\varepsilon_0 c^2}. \end{aligned}$$

Ainsi, l'ensemble des quatre équations de Maxwell qui déterminent le champ électromagnétique en fonction de ses sources est réduit, en écrivant les équations des potentiels sous une forme plus raffinée (inutile d'ailleurs pour l'usage que nous en ferons), au système

$$\left( \Delta - \frac{1}{c^2} \partial_t \partial_t \right) \phi + \partial_t \left( \nabla \cdot \mathbf{A} + \frac{1}{c^2} \partial_t \phi \right) = -\frac{\rho_s}{\varepsilon_0} \quad (\text{II.30})$$

$$\left( \Delta - \frac{1}{c^2} \partial_t \partial_t \right) \mathbf{A} - \nabla \left( \nabla \cdot \mathbf{A} + \frac{1}{c^2} \partial_t \phi \right) = -\frac{\mathbf{j}_s}{\varepsilon_0 c^2} \quad (\text{II.31})$$

$$\mathbf{E} = -\nabla \phi - \partial_t \mathbf{A} \quad (\text{II.32})$$

$$\mathbf{B} = \nabla \wedge \mathbf{A}. \quad (\text{II.33})$$

La recherche de  $\mathbf{A}$  et  $\phi$  est considérablement allégée par rapport à celle de  $\mathbf{E}$  et  $\mathbf{B}$ : quatre équations aux dérivées partielles (une pour  $\phi$  et trois pour les composantes de  $\mathbf{A}$ ) au lieu de huit (deux équations vectorielles et deux scalaires). Les champs  $\mathbf{E}$  et  $\mathbf{B}$  s'obtiennent ensuite sans difficulté — tout au moins formelle — par simple dérivation.

### II.4.2 La jauge de radiation

Il reste que les potentiels  $\phi$  et  $\mathbf{A}$ , solutions des équations (II.30) et (II.31), sont loin d'être uniques ne serait-ce qu'en raison du principe d'invariance de jauge auquel ils doivent leur paternité. Le choix de  $\phi$  et  $\mathbf{A}$  particuliers parmi la multitude de ceux



qui décrivent la même situation physique est affaire de convenance, selon le genre de question que nous nous posons.

En ce qui nous concerne, imaginons une situation décrite par des potentiels  $\phi'$  et  $\mathbf{A}'$  donnés, et considérons les potentiels équivalents

$$\begin{aligned}\phi &\stackrel{\text{df}}{=} \phi' - \partial_t \omega, \\ \mathbf{A} &\stackrel{\text{df}}{=} \mathbf{A}' + \nabla \omega.\end{aligned}$$

Quelle que soit la fonction de jauge  $\omega(\mathbf{r}, t)$ , ces potentiels conviennent tout aussi bien, les champs  $\mathbf{E}$  et  $\mathbf{B}$  sont les mêmes. Choisissons pour cete transformation de jauge la fonction

$$\omega(\mathbf{r}, t) \stackrel{\text{df}}{=} \int_0^t dt_1 \phi'(\mathbf{r}, t_1) + f(\mathbf{r}).$$

Alors, quelle que soit la fonction  $f(\mathbf{r})$ , on a  $\phi(\mathbf{r}, t) = 0$ , tandis que

$$\mathbf{A}(\mathbf{r}, t) = \mathbf{A}'(\mathbf{r}, t) + \int_0^t dt_1 \nabla \phi'(\mathbf{r}, t_1) + \nabla f(\mathbf{r}).$$

Mais nous pouvons encore tenter de choisir  $f(\mathbf{r})$  au mieux de nos intérêts ultérieurs. Remarquons d'abord que  $\nabla \phi' = -\mathbf{E} - \partial_t \mathbf{A}'$ , et calculons la divergence de  $\mathbf{A}$ . On a:

$$\begin{aligned}\nabla \cdot \mathbf{A}(\mathbf{r}, t) &= \nabla \cdot \mathbf{A}'(\mathbf{r}, t) - \int_0^t dt_1 \left( \nabla \cdot \mathbf{E}(\mathbf{r}, t_1) + \frac{\partial \nabla \cdot \mathbf{A}'}{\partial t}(\mathbf{r}, t_1) \right) + \Delta f(\mathbf{r}) \\ &= \nabla \cdot \mathbf{A}'(\mathbf{r}, 0) - \frac{1}{\varepsilon_0} \int_0^t dt_1 \rho_s(\mathbf{r}, t_1) + \Delta f(\mathbf{r}).\end{aligned}$$

Chose merveilleuse, la divergence de  $\mathbf{A}$  est indépendante du temps dans les régions où la densité de charge source  $\rho_s$  reste nulle. Dans une telle région, il suffit en particulier de trouver une fonction  $f$  solution de l'équation  $\Delta f(\mathbf{r}) = -\nabla \cdot \mathbf{A}'(\mathbf{r}, 0)$ , pour que la divergence de  $\mathbf{A}$  reste nulle. Il n'y a aucune raison pour que cette *équation de Poisson* nous arrête; il suffit d'en pêcher une solution *ad hoc*,

$$f(\mathbf{r}) = \frac{1}{4\pi} \int d^3 \mathbf{r}' \frac{\nabla \cdot \mathbf{A}'|_{\mathbf{r}', 0}}{|\mathbf{r} - \mathbf{r}'|}, \quad (\text{II.34})$$

obtenue grâce à l'analogie formelle avec le potentiel électrostatique ordinaire d'une distribution de charge.

En résumé, il est toujours possible de décrire un champ électromagnétique par des potentiels dans la *jauge de radiation*, c'est-à-dire qui, dans le vide (hors des sources du champ), ont les propriétés

$$\begin{aligned}\phi(\mathbf{r}, t) &= 0 \\ \nabla \cdot \mathbf{A}(\mathbf{r}, t) &= 0.\end{aligned} \quad (\text{II.35})$$

La condition (II.35) seule définit la *jauge de Coulomb*. Mais notre choix, plus spécifique, satisfait aussi

$$\nabla \cdot \mathbf{A} + \frac{1}{c^2} \partial_t \phi = 0,$$

c'est-à-dire la condition de *jauge de Lorenz*.

### II.4.3 Remarques sur le choix de jauge

Notre application, la théorie semi quantique du rayonnement, envisagera toujours un système de charges quantiques soumis à un rayonnement électromagnétique dont les sources sont ailleurs et ne jouent pas d'autre rôle. Aussi, dans la région du système, il sera commode d'adopter la jauge de radiation qui permet, en particulier, de dériver champs électrique et magnétique d'un seul potentiel (vecteur, il est vrai):

$$\mathbf{E} = -\partial_t \mathbf{A}, \quad (\text{II.36})$$

$$\mathbf{B} = \nabla \wedge \mathbf{A}. \quad (\text{II.37})$$

Ce choix simplifie au maximum l'exposé. Il me faut quand même reconnaître que le potentiel  $\mathbf{A}$  ne va représenter l'action d'un rayonnement sur le système que si ce dernier est loin des sources. En effet, ce potentiel contient encore des contributions du type convection qui, fort heureusement pour nous, ont une décroissance plus rapide que les contributions du type rayonnement lorsqu'on s'éloigne des sources. La chose peut surprendre, alors que notre potentiel en jauge de radiation satisfait, hors des sources et en vertu de (II.31) l'équation d'onde homogène

$$\left( \Delta - \frac{1}{c^2} \partial_t^2 \right) \mathbf{A}(\mathbf{r}, t) = 0,$$

qui semblerait impliquer, à tout coup, une solution du type rayonnement (qui se propage). Mais c'est une idée fausse. Par exemple, une charge statique à l'origine crée les champs  $\mathbf{E} = (q_s/4\pi\epsilon_0) \mathbf{r}/r^3$ , et  $\mathbf{B}(\mathbf{r}) = 0$ , fort bien décrits en jauge de radiation (!) par le potentiel  $\mathbf{A}(\mathbf{r}, t) = -(q_s/4\pi\epsilon_0) \mathbf{r} t/r^3$ , évidemment solution de l'équation d'onde, bien que n'ayant rien à voir avec le rayonnement.

### II.4.4 Récapitulation

Revenons à notre "atome", autrement dit notre système de  $Z$  quantons sans spin, chargés, dans l'environnement d'un noyau inerte, régi, comme la lectrice le savait, par l'hamiltonien:

$$H_{\text{at}} = \sum_{i=1}^Z \frac{\mathbf{P}_i^2}{2m} - \sum_{i=1}^Z \frac{Zq^2}{4\pi\epsilon_0 r_i} + \sum_{i>j=1}^Z \frac{q^2}{4\pi\epsilon_0 |\mathbf{r}_i - \mathbf{r}_j|},$$

Mais cette lectrice a maintenant appris qu'elle faisait ainsi de la théorie quantique de l'interaction coulombienne sans le savoir! Vis-à-vis de cette interaction, les quantons jouent aussi bien le rôle de charge source que de charge d'épreuve. La distinction a disparu, mais pas complètement puisque j'ai pris soin de ne pas inclure dans le terme d'interaction mutuelle les contributions d'auto-interaction  $i = j$  : un quanton ne peut éprouver sa propre source. Cette *soustraction* se justifie par la remarque que l'hamiltonien ne joue son rôle de générateur de l'évolution temporelle de la fonction-d'onde qu'à une constante près (ici, la somme des énergies propres de chaque quanton) qui ne vient changer que la phase globale de la fonction-d'onde et l'origine des énergies. Ce n'est probablement pas la première fois que vous vous autorisez, sur de simples arguments heuristiques, à négliger des quantités infinies: ce problème se rencontre en électrodynamique classique lorsqu'il s'agit de calculer l'énergie du champ d'une charge ponctuelle.<sup>17</sup> Ce ne sera pas la dernière fois; aucune théorie physique n'est totalement cohérente!

Le même atome mis en présence de charges et courants sources classiques extérieurs va avoir, d'après (II.22), l'hamiltonien en jauge de radiation:

$$H = \frac{1}{2m} \sum_i \left( \mathbf{P}_i - q\mathbf{A}(\mathbf{r}_i, t) \right)^2 - \sum_i \frac{Zq^2}{4\pi\epsilon_0 r_i} + \sum_{i < j} \frac{q^2}{4\pi\epsilon_0 |\mathbf{r}_i - \mathbf{r}_j|} .$$

Si le système est proche des sources, le potentiel  $\mathbf{A}$  peut encore rendre compte de contributions du type champ électrique statique ou champ magnétique statique, causes d'effets Stark ou Zeeman par exemple. Mais ces champs statiques décroissent vite et les régions lointaines restent l'apanage du voyageur au long cours, le rayonnement qui nous permet de voir Sirius par delà les espaces désolés à travers lesquels s'est tarie la partie statique de son champ électromagnétique.

Nous connaissons le succès des quantons dans leur rôle de sources du "champ" longitudinal (avec des guillemets parce que ce n'est qu'un intermédiaire non nécessaire pour rendre compte de l'interaction coulombienne) que nous avons dissimulé dans le potentiel  $V$  des hamiltoniens (II.22) et (II.24). Dans le cadre du modèle que je viens de rétablir, il nous est par contre strictement interdit d'envisager ces quantons comme sources de rayonnement classique. N'oubliez pas que toute la "mécanique" quantique, à partir du modèle de Bohr, a été construite précisément dans le but d'interdire aux quantons de rayonner, afin d'expliquer la stabilité de l'atome. Ces quantons chargés ne peuvent que mettre le rayonnement à l'épreuve, ils sont incapables de briller. Si, comme nous allons le voir, ce modèle explique fort bien les transitions atomiques sous l'action d'un rayonnement extérieur, il lui est impossible de rendre compte (autrement que par un principe de correspondance surajouté de manière *ad hoc*) de l'émission ou de l'absorption de rayonnement au cours du processus et, *a fortiori*, de l'émission spontanée en absence de tout ray-

---

<sup>17</sup>Voir par exemple [29, I.II § 28-1], psalmodié: Feynman livre II, chapitre XXVIII, verset 1.

onnement extérieur. A vouloir stabiliser l'atome isolé, on l'a figé dans des états irrémédiablement stationnaires.

Cette description dans laquelle la matière quantique (régie par l'équation de Schrödinger) subit un rayonnement classique (créé par les équations de Maxwell) porte habituellement le nom de *théorie semi-classique du rayonnement*. Encore une fois, il s'agit plutôt d'un modèle, et qui pourrait être qualifié, tout aussi bien, de semi-quantique!

## II.5 Rayonnement classique, quelques caractéristiques

Je vais juste indiquer les quelques propriétés du rayonnement qui nous seront utiles à propos du modèle semi-classique et de l'extension ultérieure de celui-ci à un modèle quantique du rayonnement.

### II.5.1 Les ondes planes

Bornons nous pour l'instant à l'étude du rayonnement loin de ses sources. L'équation du potentiel (II.31) devient, dans la jauge de radiation, une *équation d'onde* homogène:

$$\left(\Delta - \frac{1}{c^2}\partial_t^2\right)\mathbf{A} = 0. \quad (\text{II.38})$$

Cette équation admet en particulier comme solution toute fonction vectorielle dont la dépendance spatio-temporelle est de la forme  $e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)}$ , à condition que le *vecteur d'onde*  $\mathbf{k}$  et la *pulsation*  $\omega$  satisfassent la *relation de dispersion*

$$\omega = kc, \quad (\text{II.39})$$

où  $k \stackrel{\text{df}}{=} |\mathbf{k}|$ . La pulsation  $\omega$  est donc une fonction de  $\mathbf{k}$ .

Ainsi, le champ vectoriel

$$\boldsymbol{\varepsilon} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \quad (\text{II.40})$$

— où  $\boldsymbol{\varepsilon}$  est un vecteur constant (indépendant de  $\mathbf{r}$  et  $t$ ) — est une solution. Celle-ci est qualifiée d'*onde*, probablement en raison de lointaines analogies formelles avec la houle océanique, *plane* parce qu'à un instant  $t$  donné l'ensemble des points où sa phase prend une valeur donnée est une famille de plans, *monochromatique* car en un point donné sa variation temporelle est purement harmonique, et *polarisée rectiligne* puisque, en ce même point, elle est immuablement dirigée spatialement selon le vecteur  $\boldsymbol{\varepsilon}$ . Pour solution qu'elle soit, elle doit encore satisfaire deux conditions pour représenter un rayonnement.

D'une part, elle doit obéir à la condition de jauge (II.35). Or la lectrice trouvera sans peine — soit au moyen du “truc” de Feynman (Exercice 1), soit par calcul direct — que, de manière générale,  $\nabla \cdot (\boldsymbol{\varepsilon} e^{i\mathbf{k} \cdot \mathbf{r}}) = \boldsymbol{\varepsilon} \cdot \nabla e^{i\mathbf{k} \cdot \mathbf{r}}$ , et  $\nabla e^{i\mathbf{k} \cdot \mathbf{r}} = i\mathbf{k} e^{i\mathbf{k} \cdot \mathbf{r}}$ . Etant donné un vecteur d'onde  $\mathbf{k}$ , le vecteur  $\boldsymbol{\varepsilon}$  ne peut donc être quelconque, et la condition de jauge de Coulomb lui impose

$$\boldsymbol{\varepsilon} \cdot \mathbf{k} = 0. \quad (\text{II.41})$$

La direction de propagation de la solution (II.40) étant donnée par  $\mathbf{k}$ , on comprend alors la raison du qualificatif de transverse attribué au champ vectoriel de divergence nulle.

D'autre part, le potentiel vecteur est réel. Pourquoi alors envisager d'entrée des solutions complexes (II.40)? Parce que, comparées aux sinus et cosinus, les exponentielles offrent bien des facilités formelles en particulier pour la dérivation et la multiplication. Une solution physiquement significative sera donc plutôt constituée, par exemple, par la partie réelle de l'onde-plane-monochromatique-polarisée-rectiligne-transverse (!).

Pourquoi, parmi la multitude de solutions de l'équation d'onde libre (II.38), singulariser les ondes planes? On peut en fait, avec des combinaisons d'exponentielles  $e^{i\mathbf{k} \cdot \mathbf{r}}$  — flanquées de la dépendance temporelle idoine  $e^{i\omega t}$  — représenter commodément toutes les solutions de l'équation d'onde libre qui nous intéresseront. Cette représentation par une *intégrale de Fourier*, pour commode qu'elle soit, est affectée d'un défaut, inhérent à sa condition même d'intégrale, lorsqu'on essaye d'y recourir pour un champ vectoriel uniforme et constant. Dans ce cas, aussi banal que d'apparence inoffensive, l'amplitude de la décomposition en ondes planes du champ (sa *transformée de Fourier*) est donnée par une intégrale divergente!

À cette pathologie, on peut trouver plusieurs remèdes. De par son origine, ce genre d'intégrale ne peut être totalement dénué de sens et, après tout, ce n'est pas tant cette amplitude de Fourier qui nous intéresse, que son intégrale, c'est à dire le champ lui-même. On conjure le sort en donnant à cette intégrale malade, destinée en dernier ressort à n'être utilisée elle-même qu'intégrée, le nom de “fonction” de Dirac (ou distribution du même si l'on veut faire sérieux).<sup>18</sup>

## II.5.2 La boîte à modes

Une autre façon d'éviter le problème consiste à reconnaître qu'un champ uniforme dans *tout* l'espace, ça n'existe pas! On ne sait même pas très bien jusqu'où s'étend l'espace. Aussi peut-on se cantonner aux solutions de (II.38) dans une boîte<sup>19</sup> — le plus souvent parallélépipédique de volume  $\mathcal{V} = L_x L_y L_z$  — qui satisfont une condition donnée sur les parois.

<sup>18</sup>Voir par exemple [53, app. A].

<sup>19</sup>... ou une “caisse” dans la langue du même bois; voir [30, p. 11].

Cette limitation à un espace confiné, ainsi que l'adoption d'une condition limite arbitraire, peuvent être angoissantes pour une néophyte, mais la désinvolture avec laquelle — dans leur pratique — les experts traitent généralement la question, dénote le peu d'importance de celle-ci. Notre ignorance de ce qui se passe dans les contrées lointaines nous permet toutes les suppositions. En termes plus concrets, la plupart de nos expériences, effectuées en laboratoire et limitées dans le temps, sont totalement indifférentes aux valeurs simultanées du rayonnement aux confins d'un univers dont nous ne savons à vrai dire même pas s'il est fini. Mais nous pouvons tabler sans risque sur cette indifférence, trop heureux de perdre notre pari si jamais quelqu'un parvenait à monter une expérience sensible aux valeurs limites du rayonnement. Pour reprendre l'argument présenté par Peierls<sup>20</sup>, lorsque la boîte est assez grande l'expérience est terminée bien avant que le rayonnement ait fait le voyage aller et retour jusqu'aux limites pour nous dire ce qu'il y valait.

Le choix des valeurs limites, tout comme en mécanique quantique<sup>21</sup>, n'est plus alors qu'affaire de commodité, sans réelle conséquence physique. S'imposer des conditions aux limites a pour avantage de ramener l'infinité continue de solutions de base  $e^{i\mathbf{k}\cdot\mathbf{r}}$  à une infinité dénombrable. Nous adopterons généralement la condition de périodicité sur les parois opposées, qui implique donc que les valeurs du vecteur d'onde soient restreintes à

$$k_x = \frac{2\pi}{L_x}n_x, \quad k_y = \frac{2\pi}{L_y}n_y, \quad k_z = \frac{2\pi}{L_z}n_z, \quad (\text{II.42})$$

où  $n_x, n_y$  et  $n_z$  sont des entiers algébriques. C'est le choix le plus commode lorsqu'on a à faire à un rayonnement qui se propage.

L'autre cas pratique important est celui d'une cavité résonante, bien réelle, ou d'un guide d'onde. Ce sont alors les modes stationnaires qui nous intéressent plutôt, obtenus par des conditions aux limites qui dépendent de la nature physique du champ considéré (  $\mathbf{A}$ ,  $\mathbf{E}$  ou  $\mathbf{B}$  en électrodynamique) et des parois (constituées de conducteur plus ou moins parfait).

### II.5.3 Développement en modes plans

Quoi qu'il en soit, écrivons notre solution physique élémentaire sous la forme

$$\mathbf{A}(\mathbf{r}, t) = A_{\mathbf{k}} \hat{\mathbf{e}}_{\mathbf{k}} \frac{e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)}}{\sqrt{\mathcal{V}}} + A_{\mathbf{k}}^* \hat{\mathbf{e}}_{\mathbf{k}}^* \frac{e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)}}{\sqrt{\mathcal{V}}} . \quad (\text{II.43})$$

Les paramètres de cette solution sont le vecteur d'onde  $\mathbf{k}$ , pris dans la famille (II.42), le vecteur polarisation  $\hat{\mathbf{e}}_{\mathbf{k}}$ , et l'amplitude complexe  $A_{\mathbf{k}}$ .

<sup>20</sup>Dans un délicieux petit livre, à lire et à relire, [60, p. 64].

<sup>21</sup>Voir par exemple [53, § V-11].

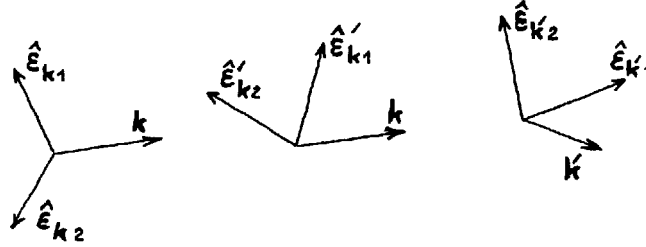


Figure II.1: Quelques choix possibles pour des vecteurs polarisation rectiligne de base.

Le vecteur polarisation a, par définition, une norme unité, ce que sa notation avec un accent circonflexe nous rappellera. Il dépend de  $\mathbf{k}$  dans une certaine mesure, puisqu'il est astreint à respecter la condition de transversalité (II.41). Plus exactement, il ne dépend que de la direction  $\hat{\mathbf{k}}$ , mais n'abusons pas des notations! Comme il intervient en facteur d'une exponentielle complexe et de l'amplitude complexe, je prévois la possibilité (inutilisée dans le cas de la polarisation rectiligne) de lui inclure un facteur complexe, aussi sa condition de normalisation s'écrit

$$\hat{\mathbf{e}}_{\mathbf{k}} \cdot \hat{\mathbf{e}}_{\mathbf{k}}^* = 1. \quad (\text{II.44})$$

Un couple  $(\mathbf{k}, \hat{\mathbf{e}}_{\mathbf{k}})$  constitue un *mode* de la boîte. A vecteur d'onde  $\mathbf{k}$  donné, on n'a que l'embarras du choix pour sélectionner deux vecteurs unitaires indépendants, orthogonaux à  $\mathbf{k}$ , que l'on adoptera comme *vecteurs de base de polarisation*. Il serait particulièrement pervers de ne pas les choisir orthogonaux; nous les noterons  $\hat{\mathbf{e}}_{\mathbf{k}1}$  et  $\hat{\mathbf{e}}_{\mathbf{k}2}$ , et ils sont par ailleurs arbitraires (voir fig. II.1).

De la définition de l'amplitude complexe  $A_{\mathbf{k}}$ , extrayons le dénominateur  $\sqrt{\mathcal{V}}$  pour que nos exponentielles de base jouissent d'une normalisation indépendante du volume de la boîte. En effet,

$$\int_0^{L_x} dx e^{ik_x x} = \begin{cases} L_x, & \text{si } k_x = 0, \\ 0, & \text{si } k_x \neq 0, \end{cases}$$

et donc

$$\int_{\mathcal{V}} d^3\mathbf{r} \frac{e^{i\mathbf{k} \cdot \mathbf{r}}}{\sqrt{\mathcal{V}}} \frac{e^{-i\mathbf{k}' \cdot \mathbf{r}}}{\sqrt{\mathcal{V}}} = \delta_{\mathbf{k}\mathbf{k}'}. \quad (\text{II.45})$$

Finalement, nous écrirons donc une solution physique générale, périodique sur les parois, sous forme d'un développement sur les modes  $(\mathbf{k}, T)$  de la boîte:

$$\mathbf{A}(\mathbf{r}, t) = \sum_{\mathbf{k}} \sum_{T=1}^2 \left( A_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \frac{e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)}}{\sqrt{\mathcal{V}}} + A_{\mathbf{k}T}^* \hat{\mathbf{e}}_{\mathbf{k}T}^* \frac{e^{-i(\mathbf{k} \cdot \mathbf{r} - \omega t)}}{\sqrt{\mathcal{V}}} \right). \quad (\text{II.46})$$

### II.5.4 Sommes et intégrales

En acceptant de nous confiner à un volume  $\mathcal{V}$ , nous sommes passés de modes distribués de façon continue à des modes discrets. Du fait de la relation de dispersion (II.39) et de la condition de périodicité (II.42) il n'existe alors aucun mode de pulsation inférieure à

$$\omega_0 \stackrel{\text{df}}{=} \frac{2\pi c}{\sup(L_x, L_y, L_z)} . \quad (\text{II.47})$$

Il importera donc généralement de se donner une boîte de dimensions minimales compatibles avec le rayonnement que l'on veut décrire.

Les grandeurs physiques calculées, susceptibles de comparaison avec des valeurs empiriques, s'expriment en général sous forme de sommes sur tous les modes, du type

$$\sum_{\mathbf{k}} f(\mathbf{k}) \stackrel{\text{df}}{=} \sum_{k_x} \sum_{k_y} \sum_{k_z} f(\mathbf{k}), \quad (\text{II.48})$$

qu'il est parfois plus commode de ramener à une intégrale. Pour ce faire, on note que dans une somme partielle sur  $k_x$ , étendue, rappelons le, aux valeurs  $k_x = (2\pi/L_x)n_x$ , avec  $n_x$  entier algébrique, ce  $k_x$  varie par pas constants

$$\Delta k_x \stackrel{\text{df}}{=} \frac{2\pi}{L_x} . \quad (\text{II.49})$$

On peut donc écrire, trivialement,

$$\frac{2\pi}{L_x} \sum_{k_x} g(k_x) = \Delta k_x \sum_{k_x} g(k_x),$$

où, moins trivialement, on reconnaît dans le deuxième membre la formule des trapèzes utilisée pour l'évaluation numérique de l'intégrale

$$\int_{-\infty}^{\infty} dk_x g(k_x).$$

Lorsque la longueur  $L_x$  est assez grande pour que les valeurs  $g(k_x + \Delta k_x)$  et  $g(k_x)$  soient voisines, on a une bonne approximation de la valeur de l'intégrale par la formule des trapèzes, et donc

$$\frac{2\pi}{L_x} \sum_{k_x} g(k_x) \underset{L_x \rightarrow \infty}{\sim} \int_{-\infty}^{\infty} dk_x g(k_x).$$

Dans notre cas à trois dimensions, lorsqu'on travaille dans une grosse boîte, on a donc finalement

$$\frac{1}{\mathcal{V}} \sum_{\mathbf{k}} f(\mathbf{k}) \underset{\mathcal{V} \rightarrow \infty}{\sim} \int \frac{d^3\mathbf{k}}{(2\pi)^3} f(\mathbf{k}). \quad (\text{II.50})$$



C'est la souplesse dont nous disposons pour les dimensions de la boîte qui autorise ce genre de manipulation dont le succès repose sur le caractère volumique des grandeurs physiques qui nous occuperont. Les effets de surface jouent alors un rôle négligeable et le résultat est finalement indépendant de la forme même de la boîte.<sup>22</sup> Lorsque la fonction  $f(\mathbf{k})$  est assez douce (ou les dimensions de la boîte suffisantes), le résultat de la somme (II.48) dépend plus de la densité des modes que de leurs valeurs exactes, c'est à dire plutôt de la grandeur du volume que de sa forme.

La physique statistique recourt souvent à ce “truc” de calcul, mais présente aussi un spectaculaire contre-exemple dans le cas d'un gaz parfait de bosons dont le nombre d'occupation moyen du niveau  $\mathbf{k}$ , énergie  $\varepsilon_{\mathbf{k}}$ , est

$$\bar{n}(\mathbf{k}) = \frac{1}{e^{\beta(\varepsilon_{\mathbf{k}} - \mu)} - 1}.$$

Lorsque la température décroît, le potentiel chimique  $\mu$  négatif croît. Quand  $\mu$  devient nul,  $\bar{n}(\mathbf{k})$  varie trop brutalement au voisinage de  $\mathbf{k} = 0$  pour que l'approximation précédente soit justifiée;  $\bar{n}(0)$  est infini, ou de façon plus réaliste prend une macro(scopique)-valeur correspondant à une condensation des bosons dans l'état  $\mathbf{k} = 0$ . Le remède consiste alors à extraire le terme  $\mathbf{k} = 0$  de la sommation dont on évalue le reste par l'approximation intégrale.

### II.5.5 Energie du champ électromagnétique

Après ces considérations techniques, revenons au champ électromagnétique dont l'énergie  $U$  est obtenue par intégration de la densité

$$u(\mathbf{r}, t) \stackrel{\text{df}}{=} \frac{\varepsilon_0}{2} \left( \mathbf{E}^2(\mathbf{r}, t) + c^2 \mathbf{B}^2(\mathbf{r}, t) \right).$$

Un cas particulier, aussi théorique que pratique, est celui du *faisceau monomode* (II.43) dans une boîte  $\mathcal{V}$ . On trouve alors:

$$\begin{aligned} \mathbf{E} &= -\partial_t \mathbf{A} \\ &= i\omega A_{\mathbf{k}} \hat{\mathbf{e}}_{\mathbf{k}T} \frac{e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)}}{\sqrt{\mathcal{V}}} + \text{complexe conjugué}, \\ \mathbf{E}^2 &= 2\omega^2 \frac{|A_{\mathbf{k}}|^2}{\mathcal{V}} - 2\omega^2 \Re \left( A_{\mathbf{k}}^2 \frac{e^{2i(\mathbf{k} \cdot \mathbf{r} - \omega t)}}{\mathcal{V}} \right), \\ \mathbf{B} &= \nabla \wedge \mathbf{A} \\ &= iA_{\mathbf{k}} \mathbf{k} \wedge \hat{\mathbf{e}}_{\mathbf{k}T} \frac{e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)}}{\sqrt{\mathcal{V}}} + \text{complexe conjugué}, \end{aligned}$$

---

<sup>22</sup>Cette question est étudiée dans [43] (sous l'angle didactique), et dans [6] (de manière encyclopédique). Son actualité est rappelée, tambour battant, dans [68].

$$\mathbf{B}^2 = 2\frac{\omega^2}{c^2}\frac{|A_{\mathbf{k}}|^2}{\mathcal{V}} - 2\frac{\omega^2}{c^2}\Re\left(A_{\mathbf{k}}^2\frac{e^{2i(\mathbf{k}\cdot\mathbf{r}-\omega t)}}{\mathcal{V}}\right),$$

où j'ai utilisé l'unitarité du vecteur polarisation, éq. (II.44), ainsi que la propriété

$$(\hat{\mathbf{e}}_{\mathbf{k}T} \wedge \mathbf{k}) \cdot (\hat{\mathbf{e}}_{\mathbf{k}T}^* \wedge \mathbf{k}) = \mathbf{k}^2 = \frac{\omega^2}{c^2}$$

qui découle de sa transversalité ( $\hat{\mathbf{e}}_{\mathbf{k}T}$  est orthogonal à  $\mathbf{k}$ ). La densité d'énergie du faisceau,

$$u = \frac{\varepsilon_0}{2}(\mathbf{E}^2 + \mathbf{B}^2),$$

a donc dans ce cas pour moyenne temporelle,

$$\bar{u} = 2\varepsilon_0\omega^2\frac{|A_{\mathbf{k}}|^2}{\mathcal{V}}. \quad (\text{II.51})$$

La moyenne temporelle du *vecteur de Poynting*

$$\mathbf{S} \stackrel{\text{df}}{=} \varepsilon_0 c^2 \mathbf{E} \wedge \mathbf{B} \quad (\text{II.52})$$

n'est autre que l'intensité  $I$  du faisceau (l'énergie propagée par unité de section et par unité de temps<sup>23</sup>) et un calcul analogue permettra à la lectrice consciencieuse de retrouver qu'elle a bien  $I = \bar{u}c$ , c'est à dire

$$I = \frac{\bar{u}L_x L_y c \Delta t}{L_x L_y \Delta t}$$

dans la géométrie de la fig. II.2.

On obtient ainsi l'expression du module de l'amplitude complexe, qui va nous servir dans l'analyse théorique, en fonction de l'intensité du faisceau (des watts par mètre carré)

$$|A_{\mathbf{k}}|^2 = \frac{\mathcal{V}}{2\varepsilon_0\omega^2 c} I.$$

De par son origine, le vecteur de Poynting n'a la signification d'un débit d'énergie qu'intégré sur une surface fermée.<sup>24</sup> Il en existe bien d'autres définitions<sup>25</sup> que l'expression (II.52), qui satisfont tout aussi bien la relation de conservation

$$\partial_t u + \nabla \cdot \mathbf{S} = 0,$$

<sup>23</sup>... à ne pas confondre avec l'intensité d'un faisceau d'ions, ou d'un courant électrique, qui est une charge électrique par unité de temps.

<sup>24</sup>S'il en est besoin, revoyez pour cela un cours d'électrodynamique classique, en particulier[29, l. II § 27-3].

<sup>25</sup>Difficile de résister au prétexte de citer [32], ne serait-ce que pour son titre.

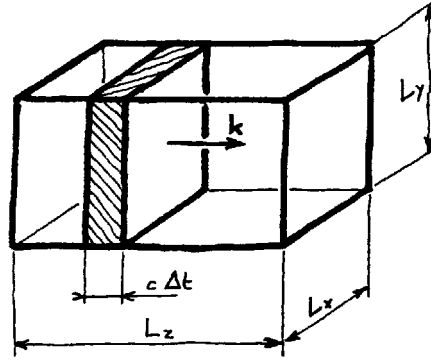


Figure II.2: La boîte à modes et le volume charrié par le faisceau pendant une durée  $\Delta t$ .

et qui ont donc la même valeur intégrale; il suffit pour cela d'ajouter à  $\mathbf{S}$  un champ vectoriel de divergence nulle. Dans le cas d'une onde plane, on peut se permettre d'attribuer à  $\mathbf{S}$  une signification locale car il a alors la même valeur en tous points d'une surface d'onde, et la même valeur moyenne temporelle partout.

## II.6 La règle d'or de Fermi

Sous ce titre évocateur, je vais faire un rappel fébrile de la théorie des perturbations dépendant du temps. Nonobstant l'usage très libéral du mot "théorie" par les physiciens, il s'agit en fait d'une méthode de résolution approchée de l'équation d'évolution

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = (H_0 + H_{\text{int}}(t)) |\psi(t)\rangle \quad (\text{II.53})$$

d'un système conservatif (hamiltonien  $H_0$  indépendant du temps) soumis à une perturbation dépendant du temps  $H_{\text{int}}(t)$ . La question est traitée dans tous les cours de théorie quantique élémentaire<sup>26</sup>. Le recours que nous allons avoir à ses premiers résultats en justifie un bref rappel. Le traitement, pour simple qu'il soit, recèle quelques subtilités, et puis n'est-il pas agréable parfois de lire des choses que l'on connaît déjà?

En adoptant une base d'états stationnaires de  $H_0$  — autrement dit états propres du même,  $H_0 |I\rangle = E_I |I\rangle$  — dont la supposée discrétion allège l'écriture, on peut développer tout état  $|\psi(t)\rangle$  du système perturbé, avec des coefficients qui dépendront

<sup>26</sup>Entre autres, sous une forme particulièrement commode, dans G. Baym, ref.[9].

généralement du temps:

$$|\psi(t)\rangle = \sum_I c_I(t) e^{-i\frac{E_I t}{\hbar}} |I\rangle.$$

On extrait les facteurs exponentiels des coefficients pour une simple raison de confort ultérieur: en l'absence de perturbation, les coefficients (ou *amplitudes*)  $c_I$  seront constants. Le report de ce développement dans l'équation de Schrödinger (II.53) et le produit à gauche par un bra propre  $\langle B|$  de l'hamiltonien non perturbé nous donnent les équations des amplitudes couplées:

$$i\hbar \dot{c}_B(t) = \sum_I e^{\frac{i}{\hbar}(E_B - E_I)t} \langle B| H_{\text{int}}(t) |I\rangle c_I(t). \quad (\text{II.54})$$

Dans le cas, important en pratique parce que facile à traiter formellement, où l'état initial est un état propre de l'hamiltonien non perturbé, soit  $|\psi(0)\rangle = |A\rangle$ , on a  $c_I(0) = \delta_{IA}$ . Dans la mesure où il ne s'est pas encore écoulé beaucoup de temps, les  $c_I(t)$  n'ont guère vieilli et on ne devrait pas trop se tromper en remplaçant les  $c_I(t)$ , sources dans les deuxièmes membres des équations d'évolution (II.54), par leurs valeurs initiales  $c_I(0)$ . La résolution des équations couplées est alors ramenée à une simple quadrature qui nous donne l'approximation, dite du premier ordre,

$$i\hbar c_B(t) \approx i\hbar c_B^{(1)}(t) \stackrel{\text{df}}{=} \int_0^t dt_1 e^{\frac{i}{\hbar}(E_B - E_A)t_1} \langle B| H_{\text{int}}(t_1) |A\rangle. \quad (\text{II.55})$$

Les critères de validité d'une telle approximation sont loins d'être évidents. L'évaluation de l'approximation du deuxième ordre, obtenue en prenant l'estimation du premier ordre (II.55) comme source dans les équations d'évolution (II.54) est trop coûteuse (sauf lorsqu'elle est nécessaire pour cause de nullité de l'élément de matrice  $\langle B| H_{\text{int}}(t_1) |A\rangle$ ), et pas totalement convaincante car la suite perturbative n'est généralement pas convergente. Heureusement, les premiers termes peuvent quand même fournir une estimation car il y a certainement convergence asymptotique, lorsque la perturbation est nulle! Contentons nous d'espérer heuristiquement que l'approximation est satisfaisante tant que l'amplitude de l'état initial  $|A\rangle$  dans l'état  $|\psi(t)\rangle$  ne s'est pas trop vidée au profit des autres états de base.

Pour les mêmes raisons pratiques et théoriques, considérons le cas où la dépendance temporelle de la perturbation  $H_{\text{int}}(t)$  est purement harmonique, avec une pulsation  $\omega$ , soit

$$H_{\text{int}}(t) = W e^{-i\omega t} + W^\dagger e^{i\omega t},$$

l'hermiticité de l'hamiltonien du système,  $H_0 + H_{\text{int}}$ , requérant celle de  $H_{\text{int}}$ . Le calcul de l'approximation du premier ordre (II.55) est alors un jeu d'enfant, avec pour résultat:

$$i\hbar c_B^{(1)}(t) = \langle B| W |A\rangle \frac{e^{\frac{i}{\hbar}(E_B - E_A - \hbar\omega)t} - 1}{\frac{i}{\hbar}(E_B - E_A - \hbar\omega)} + \langle B| W^\dagger |A\rangle \frac{e^{\frac{i}{\hbar}(E_B - E_A + \hbar\omega)t} - 1}{\frac{i}{\hbar}(E_B - E_A + \hbar\omega)}.$$

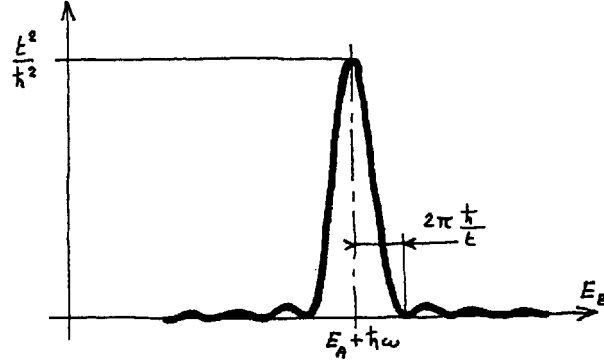


Figure II.3: Le graphe de  $\left( \frac{\sin \left( \frac{E_B - E_A - \hbar\omega}{2\hbar} t \right)}{\frac{E_B - E_A - \hbar\omega}{2}} \right)^2$ .

A pulsation de la perturbation  $\omega$  et état initial  $A$  donnés, le module de cette expression va présenter des maximums notables pour les états  $B$  singularisés par un des deux dénominateurs, c'est-à-dire tels que  $E_B \approx E_A \pm \hbar\omega$ .

Pour ceux de ces états dont l'énergie est proche de  $E_A + \hbar\omega$  par exemple, la lectrice vérifiera facilement que le module du deuxième terme du deuxième membre reste alors inférieur à  $|\langle B|W^+|A\rangle|/\omega$ , ce qui sera loin d'être le cas du premier terme. Lorsqu'on néglige ce deuxième terme (ce que l'on appelle parfois l'*approximation de l'onde tournante*), on obtient pour la probabilité d'avoir l'état  $|B\rangle$  dans l'état  $|\psi(t)\rangle$ , évolué de  $|A\rangle$ , à un instant donné  $t$ :

$$\begin{aligned} |\langle B|\psi(t)\rangle|^2 &= |c_B(t)|^2 \\ &\approx |c_B^{(1)}(t)|^2 \\ &\approx |\langle B|W|A\rangle|^2 \left( \frac{\sin \left( \frac{E_B - E_A - \hbar\omega}{2\hbar} t \right)}{\frac{E_B - E_A - \hbar\omega}{2}} \right)^2. \end{aligned}$$

Cette *probabilité de transition* est le produit du module carré de l'*élément de matrice de la transition* et d'une quantité qui, considérée comme fonction de l'énergie  $E_B$  de l'état de base dont nous cherchons l'amplitude, mérite notre attention. L'allure de son graphe est facile à tracer (fig. II.3) et nous révèle que l'aire totale comprise entre la courbe représentative et l'axe des abscisses est proportionnelle à  $(t^2/\hbar^2)(2\pi\hbar/t)$ , soit  $2\pi t/\hbar$ . Explicitons ce facteur en écrivant plutôt

$$|\langle B|\psi(t)\rangle|^2 \approx \frac{2\pi t}{\hbar} |\langle B|W|A\rangle|^2 f(E_B),$$

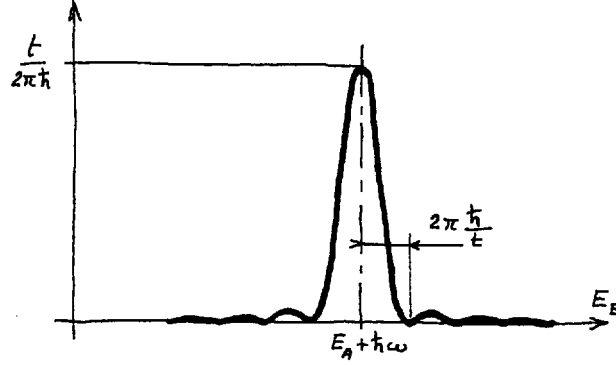


Figure II.4: Le graphe de  $\frac{1}{\pi} \frac{\sin^2\left((E_B - E_A - \hbar\omega)\frac{t}{2\hbar}\right)}{(E_B - E_A - \hbar\omega)^2 \frac{t}{2\hbar}}$ .

avec

$$f(E_B) \stackrel{\text{df}}{=} \frac{1}{\pi} \frac{\sin^2\left((E_B - E_A - \hbar\omega)\frac{t}{2\hbar}\right)}{(E_B - E_A - \hbar\omega)^2 \frac{t}{2\hbar}}. \quad (\text{II.56})$$

Toutes choses égales par ailleurs, ce sont les probabilités des divers états  $B$  que nous analysons, d'où l'intérêt d'étudier le facteur  $f$  en tant que fonction de la variable  $E_B$ , dépendant paramétriquement de  $t$  et de  $E_A + \hbar\omega$ . L'allure du graphe de cette fonction est non moins facile à établir (fig. II.4). La croissance de  $t$  effile le pic et correspond à modifier en proportions inverses les échelles des abscisses et des ordonnées. On en déduit que l'aire totale sous cette courbe reste constante. L'argument dimensionnel est confirmé lorsqu'on écrit l'intégrale correspondante,

$$\mathcal{A} \stackrel{\text{df}}{=} \int_{-\infty}^{\infty} dE_B f(E_B)$$

qui, au moyen du changement de variable  $x \stackrel{\text{df}}{=} (E_B - E_A - \hbar\omega)t/2\hbar$ , se réduit à

$$\mathcal{A} = \frac{1}{\pi} \int_{-\infty}^{\infty} dx \frac{\sin^2 x}{x^2}.$$

La limite d'un pic dont l'acuité croît indéfiniment en gardant une aire  $\mathcal{A}$  constante porte un nom; au facteur  $\mathcal{A}$  près, c'est la "fonction" de Dirac:

$$f(E_B) \underset{t \rightarrow \infty}{\sim} \mathcal{A} \delta(E_B - E_A - \hbar\omega).$$

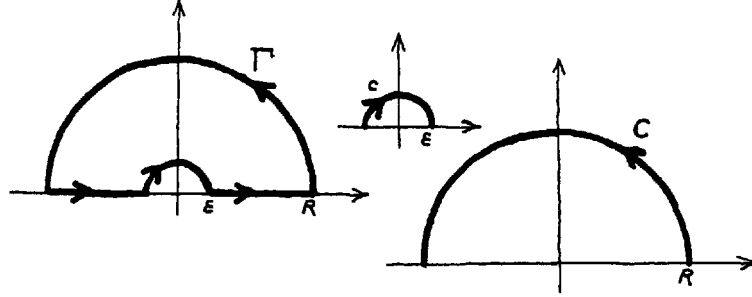


Figure II.5: Les contours d'intégration pour évaluer  $\int_{-\infty}^{\infty} dz \frac{1-e^{iz}}{z^2}$ .

Reste à calculer la constante  $\mathcal{A}$ , c'est à dire l'intégrale

$$\begin{aligned} \int_{-\infty}^{\infty} dx \frac{\sin^2 x}{x^2} &= \int_{-\infty}^{\infty} dx \frac{1 - \cos 2x}{2x^2} = \int_{-\infty}^{\infty} dz \frac{1 - e^{iz}}{z^2} \\ &= \lim_{\substack{\epsilon \rightarrow 0 \\ R \rightarrow \infty}} \left( \int_{\Gamma} \dots - \int_c \dots - \int_C \dots \right), \end{aligned}$$

après passage dans le plan  $z$  complexe, selon les différents contours d'intégration indiqués (fig. II.5).

L'intégrale sur  $\Gamma$  est nulle par absence de pôle intérieur, tandis que le lemme de Jordan nous garantit que l'intégrale sur  $C$  tend vers zéro. Reste l'intégrale sur  $c$  qui se calcule facilement en posant  $z = \epsilon e^{i\theta}$ , soit  $dz = \epsilon e^{i\theta} i d\theta = iz d\theta$ . Il vient alors

$$\int_c dz \frac{1 - e^{iz}}{z^2} = \int_c dz \frac{-iz - (iz)^2/2! - \dots}{z^2} \xrightarrow{\epsilon \rightarrow 0} \int_{\pi}^0 d\theta,$$

d'où  $\mathcal{A} = 1$ . On a ainsi retrouvé une représentation connue de la fonction de Dirac:

$$\frac{1}{\pi} \frac{\sin^2 tx}{tx^2} \underset{t \rightarrow \infty}{\sim} \delta(x). \quad (\text{II.57})$$

En ce qui nous concerne, l'aire sous la courbe représentative de  $f$  reste, en tous temps, égale à un et, lorsque le temps écoulé est assez grand, la probabilité de trouver l'état  $|B\rangle$  est donc donnée par la *règle d'or de Fermi*:

$$|\langle B|\psi(t)\rangle|^2 = \frac{2\pi t}{\hbar} |\langle B|W|A\rangle|^2 \delta(E_B - E_A - \hbar\omega). \quad (\text{II.58})$$

La pudeur me retient de préciser qu'il s'agit d'une limite lorsque  $t$  tend vers l'infini. Les critères de validité de cette expression sont ambigus, car contradictoires entre

les exigences de l'approximation du premier ordre et de l'aiguillage du pic pour en faire une fonction de Dirac.

La raison du succès de cette formule réside dans le fait qu'en pratique on ne mesure évidemment jamais une fonction de Dirac mais plutôt son intégrale. La probabilité de trouver *tous* les états  $B$  est une quantité plus réaliste. Dans la mesure où ceux-ci font partie d'un continuum, ou tout au moins ont des énergies assez voisines par rapport à la largeur  $2\pi\hbar/t$  du pic, on peut remplacer la somme discrète par une intégration dont l'élément différentiel est un nombre d'états,  $dn_B$ . La présence de la fonction  $\delta(E_B - E_A - \hbar\omega)$  suggère alors d'adopter plutôt l'énergie comme variable d'intégration, ce qui fait apparaître la *densité d'états finals*,  $\rho$ , du système:

$$\begin{aligned} dn_B &= \frac{dn_B}{dE_B} dE_B \\ &= \rho(E_B) dE_B. \end{aligned}$$

On obtient ainsi la règle d'or de Fermi sous sa forme intégrée, plus habituelle, qui donne la probabilité de transition totale

$$\text{Pr}_{A \rightarrow B} = \frac{2\pi t}{\hbar} |\langle B|W|A \rangle|^2 \rho(E_B), \quad (\text{II.59})$$

où  $E_B = E_A + \hbar\omega$ .

Je n'épiloguerai pas sur les quelques raffinements techniques, sans difficulté, mais nécessaires, lorsque le niveau  $E_B$  est dégénéré. D'un autre côté, on comprend maintenant que l'approximation par une fonction  $\delta$  (équation (II.58), ou (II.59) sous forme intégrée) est valide dans la mesure où la largeur du pic  $2\pi\hbar/t$  est assez faible pour que les variations du module de l'élément de matrice  $\langle B|W|A \rangle$  et de la densité  $\rho$  en fonction de l'énergie  $E_B$  y soient négligeables.

La simplicité de la dépendance temporelle dans la règle d'or permet de définir un *taux de transition*

$$\Gamma_{A \rightarrow B} \stackrel{\text{df}}{=} \frac{|\langle B|\psi(t)\rangle|^2}{t},$$

soit, pour la forme que nous utiliserons,

$$\Gamma_{A \rightarrow B} = \frac{2\pi}{\hbar} |\langle B|W|A \rangle|^2 \delta(E_B - E_A - \hbar\omega). \quad (\text{II.60})$$

Attention! En dépit de sa dimension (l'inverse d'un temps) et de sa dénomination courante, cette quantité n'a en aucune façon la signification physique d'une "probabilité de transition par unité de temps". Ce n'est pas la transition qui est par unité de temps, comme si après deux unités (par exemple deux heures!) on devait nécessairement avoir une probabilité deux fois plus grande. C'est parce qu'elle est indépendante du temps que cette expression est si commode, mais ce n'est pas vrai



tout le temps: celui-ci ne doit être ni trop court (la présumée fonction de Dirac se répand), ni trop long (tout peut arriver au-delà du premier ordre). Le temps est ici un paramètre. A l'instant  $t$  l'état est  $|\psi(t)\rangle$ , et c'est l'instant choisi pour l'analyse du système qui conditionne la probabilité de transition à l'état  $|B\rangle$  selon la règle orthodoxe (sinon cohérente) de la mesure en théorie quantique. . . mais ceci est une autre histoire abondamment glosée, à défaut d'être élucidée.<sup>27</sup>

## II.7 Taux d'excitation

Revenons enfin à notre “atome” de quantons sans spin en présence d'un noyau infiniment lourd. Soumis à un champ électromagnétique, son hamiltonien est donné par les équations (II.23), (II.24) et (II.25). Le choix de la jauge de radiation permet de négliger précisément la contribution du potentiel coulombien instantané,  $\phi$ , à l'interaction, dans la mesure où l'atome est assez loin des sources du champ électromagnétique. Supposons aussi le champ assez faible pour pouvoir négliger le terme diamagnétique en  $\mathbf{A}^2$ , du second ordre par rapport au terme paramagnétique proportionnel à  $\mathbf{A}$  (donc pas d'effets non linéaires). L'hamiltonien d'interaction, sous la forme factorisée (II.28), (II.29), se réduit alors à

$$H_{\text{int}} \approx -q \int d^3\mathbf{r} \mathbf{j}(\mathbf{r}) \cdot \mathbf{A}(\mathbf{r}, t), \quad (\text{II.61})$$

en représentation de positions des quantons.

Le courant de matière  $\mathbf{j}$  est défini en (II.27); il dépend des grandeurs physiques positions et impulsions des quantons, et du paramètre position  $\mathbf{r}$ . Il est par contre indépendant du champ électromagnétique, tandis que  $\mathbf{A}(\mathbf{r}, t)$ , indépendant des grandeurs physiques des quantons, est un simple facteur réel, pas un opérateur. Ce courant  $\mathbf{j}$  n'a rien à voir avec le courant source du champ électromagnétique, contingence technique dont nous n'aurons plus à nous préoccuper. L'hamiltonien d'interaction  $H_{\text{int}}$  dépend, quant à lui, des positions et impulsions des quantons et du temps, mais pas du paramètre  $\mathbf{r}$  sur lequel on a maintenant intégré.

Réfugions-nous dans notre boîte pour décrire le rayonnement. L'utilisation du développement du potentiel vecteur en modes plans rectilignes (II.46) conduit à l'expression de l'hamiltonien d'interaction

$$H_{\text{int}} = -\frac{q}{\sqrt{\mathcal{V}}} \sum_{\mathbf{k}T} \left( A_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{j}_{-\mathbf{k}} e^{-i\omega t} + A_{\mathbf{k}T}^* \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{j}_{\mathbf{k}} e^{i\omega t} \right),$$

plus simple qu'il n'y paraît, car les seuls opérateurs agissant dans l'espace des états des quantons sont les

$$\mathbf{j}_{\mathbf{k}} \stackrel{\text{df}}{=} \int d^3\mathbf{r} e^{-i\mathbf{k} \cdot \mathbf{r}} \mathbf{j}(\mathbf{r}), \quad (\text{II.62})$$

---

<sup>27</sup>Voir par exemple [26] et [49].

soit, par définition de  $\mathbf{j}$ ,

$$\mathbf{j}_{\mathbf{k}} = \frac{1}{2m} \sum_{i=1}^Z \left( \mathbf{p}_i e^{-i\mathbf{k} \cdot \mathbf{r}_i} + e^{-i\mathbf{k} \cdot \mathbf{r}_i} \mathbf{p}_i \right).$$

Dans le cas d'un faisceau monomode  $\mathbf{k}T$ , il ne reste que

$$H_{\text{int}} = -\frac{q}{\sqrt{\mathcal{V}}} \left( A_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{j}_{-\mathbf{k}} e^{-i\omega t} + A_{\mathbf{k}T}^* \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{j}_{\mathbf{k}} e^{i\omega t} \right). \quad (\text{II.63})$$

Le terme d'interaction a alors une variation temporelle purement harmonique et, suivant le rappel de la section précédente, le taux de transition (on dit d'excitation dans ce cas) d'un atome initialement dans son état fondamental  $A$  est:

$$\Gamma_{A \rightarrow B} = \frac{2\pi q^2}{\hbar \mathcal{V}} |A_{\mathbf{k}T}|^2 |\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{j}_{-\mathbf{k}} | A \rangle|^2 \delta(E_B - E_A - \hbar\omega), \quad (\text{II.64})$$

c'est-à-dire non négligeable seulement vers des états finals  $B$  tels que  $E_B \approx E_A + \hbar\omega$ . Cette formule rend compte à merveille des processus d'absorption au cours desquels un atome passe de son fondamental à un état excité sous l'action d'un rayonnement électromagnétique extérieur pas trop intense. Avec un faisceau de fréquence adéquate, l'état final peut être tellement excité qu'il correspond à l'éjection d'un quanton dans un état non lié. On a donc aussi une description quantique de l'effet photo-électrique.

Tout aussi simplement, en partant de l'état initial excité  $B$ , la contribution du deuxième terme de (II.63) cette fois, permet de calculer le taux de désexcitation  $\Gamma_{B \rightarrow A}$ , d'ailleurs égal à  $\Gamma_{A \rightarrow B}$ .

Tout serait-il donc maintenant pour le mieux dans le meilleur des mondes? Loin s'en faut. Nous allons voir que le modèle semi-classique de l'interaction matière-rayonnement auquel nous sommes parvenus ne fait pas toute la lumière sur les phénomènes observables, et que les tentatives de réparation heuristiques conduisent à des descriptions du rayonnement contradictoires.

## Exercices

1. Feynman [29, l. II § 27-3] recommande un “truc” de calcul fort commode lorsque l'opérateur dérivatif  $\nabla$  s'applique à un produit de fonctions, que celles-ci soient scalaires ou vectorielles. Il suffit de décomposer l'opérateur en somme d'opérateurs analogues agissant respectivement sur chacune des fonctions et indifféremment à gauche ou à droite. Il ne reste plus ensuite qu'à appliquer les règles ordinaires de l'algèbre vectorielle à ces opérateurs spécialisés. Par exemple:

$$\nabla \cdot (\mathbf{A} \wedge \mathbf{B}) = (\nabla_A + \nabla_B) \cdot (\mathbf{A} \wedge \mathbf{B})$$

$$\begin{aligned}
&= \nabla_A \cdot (\mathbf{A} \wedge \mathbf{B}) + \nabla_B \cdot (\mathbf{A} \wedge \mathbf{B}) \\
&= \mathbf{B} \cdot (\nabla_A \wedge \mathbf{A}) + \mathbf{A} \cdot (\mathbf{B} \wedge \nabla_B) \\
&= \mathbf{B} \cdot (\nabla \wedge \mathbf{A}) - \mathbf{A} \cdot (\nabla_B \wedge \mathbf{B}) \\
&= \mathbf{B} \cdot (\nabla \wedge \mathbf{A}) - \mathbf{A} \cdot (\nabla \wedge \mathbf{B}).
\end{aligned}$$

i) Calculez de cette façon  $\nabla(fg)$ , où  $f$  et  $g$  sont des fonctions scalaires. En déduire l'expression développée de  $\nabla(e^{i\omega(\mathbf{r})}f(\mathbf{r}))$ .

ii) Calculez de la même façon  $\nabla \cdot (f\mathbf{A})$ , où  $\mathbf{A}$  est une fonction vectorielle. En déduire les expressions de:

a)  $\nabla \cdot (e^{i\omega(\mathbf{r})}\mathbf{A}(\mathbf{r}))$ ;

b)  $\nabla \cdot (\boldsymbol{\varepsilon} e^{i\mathbf{k} \cdot \mathbf{r}})$ , où  $\boldsymbol{\varepsilon}$  et  $\mathbf{k}$  sont des vecteurs constants;

c)  $\nabla \cdot (f\nabla g)$ , où  $f$  et  $g$  sont des fonctions scalaires.

iii) Calculez  $\nabla \wedge (\mathbf{A}(\mathbf{r}, t)\psi(\mathbf{r}, t))$ . En déduire l'expression de  $\nabla \wedge (\boldsymbol{\varepsilon} e^{i\mathbf{k} \cdot \mathbf{r}})$ .

**2.** La fonction-d'onde  $\psi(\mathbf{r}, t) = e^{-i\frac{q}{\hbar}\omega(\mathbf{r}, t)}\psi'(\mathbf{r}, t)$  évolue selon l'équation de Schrödinger libre

$$i\hbar \frac{\partial}{\partial t} \psi = \frac{\mathbf{P}^2}{2m} \psi.$$

i) Calculez  $\mathbf{P}^2 e^{-i\frac{q}{\hbar}\omega} \psi'$ ,

a) naïvement, en utilisant directement  $\mathbf{P}^2 = -\hbar^2 \Delta$ ;

b) plus astucieusement, en suivant la méthode exposée dans le texte (p. 13) et grâce aux relations trouvées (Exercice 1), c'est à dire en calculant d'abord  $\mathbf{P} e^{-i\frac{q}{\hbar}\omega} \psi'$ , puis l'action de  $\mathbf{P}$  (sous forme de produit scalaire) sur le résultat.

c) Appréciez la simplicité de la deuxième méthode de calcul.

ii) En déduire l'équation d'évolution de  $\psi'$ .

**3.** Vérifiez que les mêmes champs  $\mathbf{E}, \mathbf{B}$  (définis par (II.6–II.7)) dérivent des potentiels  $\phi, \mathbf{A}$  et des potentiels transformés  $\phi', \mathbf{A}'$  (définis par (II.11–II.12)).

**4.** Les fonctions-d'onde  $\psi$  et  $\psi'$  sont reliées par la transformation de jauge locale (équ. (II.3)):

$$\psi'(\mathbf{r}, t) = e^{-i\frac{q}{\hbar}\omega(\mathbf{r}, t)} \psi(\mathbf{r}, t).$$

i) Comparez les éléments de matrice

a)  $\langle \psi(t) | \mathbf{R} | \psi(t) \rangle$  et  $\langle \psi'(t) | \mathbf{R} | \psi'(t) \rangle$ ;

b)  $\langle \psi(t) | \mathbf{P} | \psi(t) \rangle$  et  $\langle \psi'(t) | \mathbf{P} | \psi'(t) \rangle$ .

ii) Vérifiez que  $\langle \psi(t) | (\mathbf{P} - q\mathbf{A}) | \psi(t) \rangle$  est invariant de jauge.

**5.** De l'utilité du tenseur antisymétrique de Levi-Civita (p. 22):

- i) Montrez que  $\varepsilon_{ijk} = \varepsilon_{kij} = \varepsilon_{jki}$ .  
 ii) Vérifiez que la composante  $i$  d'un produit vectoriel peut s'écrire  $(\mathbf{A} \wedge \mathbf{B})_i = \varepsilon_{ijk} A_j B_k$ .  
 iii) Vérifiez que  $\varepsilon_{ijk} \varepsilon_{ilm} = \delta_{jl} \delta_{km} - \delta_{jm} \delta_{kl}$ . En déduire des expressions compactes de  $\mathbf{A} \wedge (\mathbf{B} \wedge \mathbf{C})$ ,  $\nabla \wedge (\nabla \wedge \mathbf{A})$  et  $(\mathbf{A} \wedge \mathbf{B}) \cdot (\mathbf{C} \wedge \mathbf{D})$ .

6. Etant donné une constante réelle  $B \dots$

- i) Déterminez une fonction de jauge  $\omega(\mathbf{r}, t)$  qui permette de passer des potentiels  $\phi = 0$ ,  $\mathbf{A} = (-By/2, Bx/2, 0)$ , aux potentiels  $\phi' = 0$ ,  $\mathbf{A}' = (-By, 0, 0)$ .  
 ii) Vérifiez que les champs  $\mathbf{E}$  et  $\mathbf{B}$  sont encore les mêmes.  
 iii) Mêmes questions pour passer à  $\phi'' = -Vx$ ,  $\mathbf{A}'' = (Vt, Bx, 0)$ , où  $V$  est une constante.

7. A propos des moments angulaire orbital et de spin d'un quanton.

- i) Vérifiez que la propriété  $[R_i, P_j] = i\hbar \delta_{ij}$  et la définition  $\mathbf{L} \stackrel{\text{df}}{=} \mathbf{R} \wedge \mathbf{P}$  impliquent  $[L_i, L_j] = i\hbar \varepsilon_{ijk} L_k$ .  
 ii) Vérifiez que la propriété  $[J_i, J_j] = i\hbar \varepsilon_{ijk} J_k$  et la définition  $\mathbf{S} \stackrel{\text{df}}{=} \mathbf{J} - \mathbf{L}$  impliquent  $[S_i, S_j] = i\hbar \varepsilon_{ijk} S_k$ .  
 iii) Montrez que  $[J_i, R_j] = i\hbar \varepsilon_{ijk} R_k$  et  $[J_i, P_j] = i\hbar \varepsilon_{ijk} P_k$ . (Commutateurs caractéristiques d'opérateurs vectoriels.)  
 iv) Montrez que  $[S_i, R_j] = [S_i, P_j] = 0$ .

8. A propos des diverses formes de l'équation de Schrödinger d'un quanton libre de spin 1/2.

- i) Ecrivez l'équation aux dérivées partielles correspondant à (II.19), dans le cas  $a = 0$ .  
 ii) Déterminez le système différentiel linéaire du premier ordre équivalent à cette équation. (Le "truc" consiste à définir des fonctions à deux composantes auxiliaires  $\chi_x \stackrel{\text{df}}{=} \partial_x \psi$ , etc.) Quel est le rang de ce système?  
 iii) Et dans le cas général ( $a$  différent de 0 et de 1)?

9. Dans le cas d'un quanton de spin 1/2, quelle théorie de jauge pouvez vous fonder sur l'hamiltonien libre ( $a$  est un paramètre réel)

$$H_0 = \frac{1}{2m} (a\mathbf{P}^2 + (1-a)(\boldsymbol{\sigma} \cdot \mathbf{P})^2) ?$$

Quelle valeur du facteur  $g$  cette théorie prédit-elle pour le quanton?

10. A propos des matrices de Pauli.

- i) Calculez les matrices représentatives de  $S_x$ ,  $S_y$  et  $S_z$  dans la base des états propres de  $\mathbf{S}^2$  et  $S_z$ :  $|1/2 \ 1/2\rangle$  et  $|1/2 \ -1/2\rangle$ . Pour cela il faut se souvenir que

pour tout triplet d'opérateurs satisfaisant les relations (II.15) (c'est le cas de  $S_x$ ,  $S_y$  et  $S_z$ ) on a

$$\begin{aligned} J_z |j m\rangle &= \hbar m |j m\rangle \\ J_{\pm} |j m\rangle &= \hbar \sqrt{j(j+1) - m(m \pm 1)} |j m \pm 1\rangle, \end{aligned}$$

avec  $2m$  entier  $\in [-2j, 2j]$ , et  $J_{\pm} \stackrel{\text{df}}{=} J_x \pm iJ_y$ .

- ii) En déduire les expressions des trois matrices de Pauli.
- iii) Calculez les neuf produits  $\sigma_i \sigma_j$  et vérifiez les propriétés (II.17).
- iv) A l'aide des propriétés des  $\varepsilon_{ijk}$  (trouvées dans l'exercice 5), démontrez l'identité (II.18).

### 11. Opérateur vitesse d'un quanton.

i) Montrez que l'évolution de la valeur moyenne d'une grandeur physique d'un système d'hamiltonien  $H$  est donnée par

$$i\hbar \frac{d}{dt} \langle \psi(t) | G | \psi(t) \rangle = \langle \psi(t) | [G, H] | \psi(t) \rangle + \langle \psi(t) | i\hbar \frac{dG}{dt} | \psi(t) \rangle.$$

ii) Calculez le commutateur  $[X_1, H]$  de la composante  $x$  de la position du quanton 1 avec l'hamiltonien (II.22).

iii) En déduire l'expression de l'opérateur vitesse  $\mathbf{V}_i$  du quanton  $i$ . Celle-ci est-elle invariante de jauge?

## Chapitre III

# La théorie quantique du rayonnement

J'espère avoir maintenant suffisamment insisté sur le fait que nous n'allons étudier que des modèles. Un modèle se situe strictement dans le cadre des règles d'une ou plusieurs théories et constitue l'intermédiaire obligé de l'adéquation entre théorie et données empiriques. La lectrice vient d'en voir un exemple avec le modèle semi-classique du rayonnement qui ne remettait nullement en cause — mais au contraire invoquait toutes — les règles de deux théories fermement établies: la théorie quantique et l'électrodynamique classique. Les insuffisances auxquelles j'ai fait allusion ne seront nullement imputables à la théorie quantique. L'électrodynamique classique par contre va rencontrer ses limites.

La nouvelle description du rayonnement viendra se couler dans la théorie quantique pour finalement ne constituer qu'un modèle au sein de celle-ci. J'utiliserai maintenant la dénomination de théorie, habituelle dans le métier (le mot fait tellement mieux!), que ce soit pour les théories semi-classique ou quantique du rayonnement, ou plus tard la théorie quantique des champs. Mais en se souvenant de cette distinction modèle/théorie la lectrice parviendra peut-être à s'économiser quelques blocages conceptuels.<sup>1</sup> Théorie ou modèle, nous allons maintenant élaborer une description unifiée, dans le seul cadre quantique, de processus qui étaient classiquement condamnés au diptyque: mécanique du point et théorie du champ.

---

<sup>1</sup>Le moment est peut-être bien choisi pour tenter de lire, ou relire, le préluce. La lectrice philosophe particulièrement intéressée par cette question tirera profit de l'ouvrage de Bunge[14].

### III.1 Insuffisances et contradictions de la théorie semi-classique

Un modèle fondé à la fois sur un traitement quantique de la charge d'épreuve et sur une description classique du champ électromagnétique ne présente, à première vue, aucune incohérence. Pourquoi la théorie quantique doit-elle alors remettre en jeu les succès acquis par le modèle du quanton, et se risquer sur le territoire mouvant des champs?

#### III.1.1 Le photon entre en scène

La position du pic qui surgit au milieu de la figure II.3, ou la fonction de Dirac à laquelle il se réduit dans le taux d'excitation (II.64), nous proclament la pertinence de la quantité  $\hbar\omega$ , associée au rayonnement monomode de pulsation  $\omega$ . Attachés au principe de conservation de l'énergie<sup>2</sup>, nous attribuons le rôle d'énergie à ces paquets  $\hbar\omega$  qui constituent les éléments d'une décomposition modale du rayonnement. Notre modèle du quanton semblerait donc prédire l'existence des *photons* si l'on oubliait que les quantons de la théorie semi-classique sont inertes, que ce ne sont pas des sources de rayonnement mais de simples charges d'épreuve. Le photon est ici surajouté au nom d'un principe de correspondance aussi vague que le rôle de particule (qui s'avérera en partie usurpé) que se voit ainsi attribuer le paquet d'énergie  $\hbar\omega$ . Sept ans après l'hypothèse granulaire du rayonnement par Einstein, et une douzaine d'années avant l'avènement de la théorie quantique, c'était précisément le "truc" heuristique de Bohr:

- l'atome n'émet de rayonnement que lorsqu'il transite entre des états par ailleurs stationnaires (il ne fallait pas avoir peur des contradictions logiques!);
- la fréquence du rayonnement est donnée par la différence des énergies des états (et se voit donc attribuer un rôle d'énergie).

Notre formule (II.64) pour le taux d'excitation semble suffire pour étudier la photo-ionisation de l'atome et donc l'effet photo-électrique, un des phénomènes à l'origine de l'hypothèse d'Einstein. Elle ne fournit par contre aucune explication pour une expérience de résonance optique (fig. III.1). Non seulement notre modèle semi-classique est incapable de prédire l'émission spontanée, c'est à dire cette lumière, de même fréquence que le faisceau incident, émise dans toutes les directions, mais il n'offre aucune explication pour l'absorption du faisceau. C'est en ce sens que l'appellation d'absorption pour la formule (II.64) était surfaite. Les quantons ne sont encore ni sources ni puits, mais des bouchons agités au gré des ondes, et la théorie semi-classique du rayonnement est un modèle pour le mécanisme de transition des quantons chargés, mais n'a rien à dire sur l'évolution du rayonnement.

---

<sup>2</sup>Relire à ce propos la parabole des cubes et de la maman de Feynman, refs. [29, vol. I, §4-1] ou [28, p. 80].

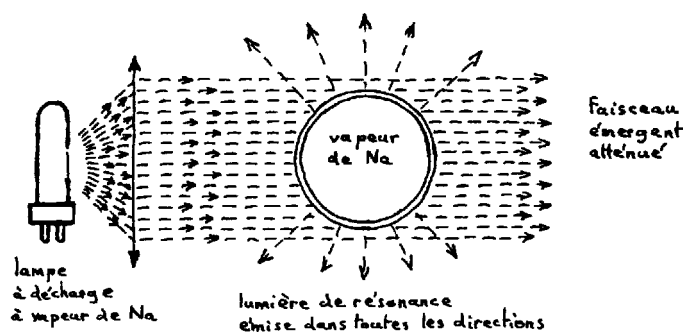


Figure III.1: Expérience de résonance optique.

### III.1.2 Deux sortes de rayonnement?

Pourquoi ne pas se contenter d'une théorie semi-classique et d'un principe de correspondance qui nous assurent qu'un rayonnement monomode ne peut intervenir que par paquets d'énergie  $\hbar\omega$ ? Admettons. On se trouve alors confronté à deux types de rayonnement:

- un rayonnement classique, fabriqué par des sources classiques et les équations de Maxwell;
- un rayonnement plutôt vague, susceptible d'être absorbé par paquets ou mystérieusement créé par des émissions plus ou moins spontanées de la part des quantons (alors même que l'interdiction du rayonnement était la motivation affirmée de la "mécanique" quantique naissante).

J'ai peu parlé des sources, protégées par l'anonymat classique, mais il est temps de réaliser qu'elles ne sont généralement pas plus classiques que les quantons. Le rayonnement incident est le plus souvent lui aussi créé par des émissions spontanées (il suffit de penser à la lampe à décharge de la figure III.1); pourquoi serait-il alors une onde maxwellienne monomode  $\omega$  plutôt qu'un convoi de paquets d'énergie  $\hbar\omega$ ? On ne peut espérer résoudre cette contradiction que par une théorie quantique du rayonnement, reprenant à son compte la structure granulaire de celui-ci, en même temps que sa création et son absorption par des sources quantiques. Restera à retrouver les comportements maxwelliens en allant explorer les limites classiques de cette théorie.

### III.1.3 Les avatars de $\hbar$

Des phénomènes comme l'effet photo-électrique ou la résonance optique sont essentiellement quantiques; leur description exigeait d'une façon ou d'une autre l'in-



tervention de  $\hbar$ , que ce soit — dans l'ordre chronologique — d'abord dans le rayonnement seul (Einstein), puis dans l'atome seul (Bohr, Schrödinger). La théorie quantique du rayonnement va introduire  $\hbar$  dans les deux acteurs de ces processus et, accessoirement, concilier la description du rayonnement et le traitement quantique de l'interaction coulombienne instantanée déjà réalisé dans l'équation de Schrödinger de l'atome non perturbé.

## III.2 L'espace de Fock des photons

Construire un modèle quantique, c'est se donner un espace des états, ou plutôt la structure de celui-ci en spécifiant les relations de commutation entre opérateurs — représentatifs de grandeurs physiques du système — agissant dans cet espace. En repartant de l'idée initiale d'Einstein — à savoir que l'interaction matière-rayonnement suggère la structure granulaire de celui-ci (ou tout au moins de son énergie) par paquets de  $\hbar\omega$  — le rayonnement entre d'ores et déjà dans le domaine quantique. Il ne reste plus qu'à parfaire un modèle compatible avec la théorie semi-classique qui convenait si bien pour rendre compte des seules transitions de la matière. Nous allons voir que ces conditions restreignent extraordinairement l'éventail des possibilités.

### III.2.1 La correspondance

Un faisceau monomode  $\mathbf{k}T$  étudié dans la boîte de volume  $\mathcal{V}$  a pour énergie moyenne

$$U = 2\varepsilon_0\omega^2|A_{\mathbf{k}T}|^2,$$

selon l'expression classique de la densité (II.51), et en revanche

$$U = \hbar\omega n_{\mathbf{k}T}$$

s'il est considéré, quantiquement, comme un ensemble de  $n_{\mathbf{k}T}$  photons. On a donc une relation entre l'amplitude de la description classique et le nouveau concept de nombre de photons:

$$|A_{\mathbf{k}T}|^2 = \frac{\hbar}{2\varepsilon_0\omega} n_{\mathbf{k}T}, \quad (\text{III.1})$$

relation qui bien entendu n'a de sens que dans les situations où la description classique du faisceau est suffisante.<sup>3</sup>

---

<sup>3</sup>Vous verrez que le nombre de photons est alors une variable aléatoire, mais très localisée autour de sa valeur moyenne, comme toujours lorsqu'une grandeur quantique admet une limite classique.

### III.2.2 L'espace de Fock

Pour préciser cette représentation du rayonnement quantique en ensemble de photons, on donne à son espace des états la structure d'un *espace de Fock*, c'est à dire que l'espace des états d'un mode  $\mathbf{k}T$  est sous-tendu par des kets de base correspondant à des nombres donnés de photons. Par exemple:

- $|1\rangle_{\mathbf{k}T}$  est un état à un photon du mode  $\mathbf{k}T$ ;
- $|7\rangle_{\mathbf{k}T}$  est un état à sept photons; *etc.* et même (et surtout)...
- l'état à zéro photon,  $|0\rangle_{\mathbf{k}T}$ , conventionnellement appelé *vide*.

La signification de ces kets n'est pas encore très lumineuse mais, comme toujours dans l'élaboration d'une théorie physique, elle se révélera progressivement par l'utilisation que nous ferons de ceux-ci.

Plus généralement, l'espace des états du rayonnement électromagnétique sera un espace de Fock constitué du produit direct des espaces de Fock de tous les modes du rayonnement dans la boîte. Un ket de base typique de cet espace sera par exemple:

$$|p_{\mathbf{k}_1 T_1}, q_{\mathbf{k}_2 T_2}, 0_{\mathbf{k}_3 T_3}, \dots\rangle \stackrel{\text{df}}{=} |p\rangle_{\mathbf{k}_1 T_1} |q\rangle_{\mathbf{k}_2 T_2} |0\rangle_{\mathbf{k}_3 T_3} \cdots \quad (\text{III.2})$$

Nous décidons évidemment que ces kets de base sont *orthogonaux*: un photon et deux photons, ça n'est pas du tout la même chose — ça s'exclue —, non plus que trois photons dans le mode  $\mathbf{k}T$  et trois photons dans le mode  $\mathbf{k}'T'$ . Pour des raisons plus pratiques que fondamentales, nous les choisissons aussi *normés* à l'unité.

### III.2.3 Excitation de la matière et absorption d'un photon

Reprenons la description du processus d'excitation dans ce nouveau modèle. Nous considérons maintenant un système constitué d'un ensemble de quantons — la matière, par exemple les électrons “sans spin” d'un atome, — et du rayonnement quantique. L'espace des états de ce système est donc le produit de l'espace des états des quantons et de l'espace de Fock du rayonnement dans la boîte.

Commençant avec un état initial du système

$$|A; \dots n_{\mathbf{k}T} \dots\rangle \stackrel{\text{df}}{=} |A\rangle |\dots n_{\mathbf{k}T} \dots\rangle,$$

on se demande quelle peut être, au bout du temps  $t$ , la probabilité de l'état

$$|B; \dots (n-1)_{\mathbf{k}T} \dots\rangle,$$

toutes choses (les points de suspension) égales par ailleurs, afin que ce processus corresponde à l'excitation des quantons de l'état  $A$  à l'état  $B$  et à l'absorption d'un photon du mode  $\mathbf{k}T$ .

Notre plus cher désir est évidemment d'obtenir le même taux de transition que par la théorie semi-classique (équation (II.64)), lorsque celle-ci est en accord avec

les données empiriques. Nos états initial et final sont états propres de l'hamiltonien de la matière prolongé à l'espace produit,

$$\tilde{H}_{\text{mat}} \stackrel{\text{df}}{=} H_{\text{mat}} \otimes 1_{\text{ray}},$$

c'est-à-dire

$$\begin{aligned}\tilde{H}_{\text{mat}}|A; \dots n_{\mathbf{k}_T} \dots\rangle &= E_A|A; \dots n_{\mathbf{k}_T} \dots\rangle \\ \tilde{H}_{\text{mat}}|B; \dots n_{\mathbf{k}_T} \dots\rangle &= E_B|B; \dots n_{\mathbf{k}_T} \dots\rangle.\end{aligned}$$

Une telle évolution ne peut se produire que par l'intercession d'un hamiltonien d'interaction  $\tilde{H}_{\text{int}}(t)$ , à déterminer, agissant dans l'espace produit. Or le résultat semi-classique (II.64) peut se réécrire, moyennant la correspondance (III.1):

$$\Gamma_{A \rightarrow B} = \frac{2\pi}{\hbar} \frac{q^2 \hbar}{2\varepsilon_0 \omega \mathcal{V}} n_{\mathbf{k}_T} |\langle B | \hat{\mathbf{e}}_{\mathbf{k}_T} \cdot \mathbf{j}_{-\mathbf{k}} | A \rangle|^2 \delta(E_B - E_A - \hbar\omega).$$

On ne peut espérer retrouver ce résultat avec la théorie des perturbations appliquée au terme d'interaction  $\tilde{H}_{\text{int}}(t)$  que si celui-ci contient une contribution additive dont la dépendance temporelle soit de la forme  $e^{-i\omega t}$ . Notons cette contribution  $\tilde{W}e^{-i\omega t}$ . La règle d'or (II.60) nous donne alors pour le taux de transition de notre système matière-rayonnement:

$$\Gamma_{A; n_{\mathbf{k}_T} \rightarrow B; (n-1)_{\mathbf{k}_T}} = \frac{2\pi}{\hbar} |\langle B; \dots (n-1)_{\mathbf{k}_T} \dots | \tilde{W} | A; \dots n_{\mathbf{k}_T} \dots \rangle|^2. \quad (\text{III.3})$$

### III.2.4 Opérateurs de création et d'annihilation

L'équivalence des deux expressions du taux d'excitation implique la proportionnalité des éléments de matrice ou, plus précisément, la relation:

$$|\langle B; \dots (n-1)_{\mathbf{k}_T} \dots | \tilde{W} | A; \dots n_{\mathbf{k}_T} \dots \rangle|^2 = \frac{q^2 \hbar}{2\varepsilon_0 \omega \mathcal{V}} n_{\mathbf{k}_T} |\langle B | \hat{\mathbf{e}}_{\mathbf{k}_T} \cdot \mathbf{j}_{-\mathbf{k}} | A \rangle|^2.$$

L'opérateur inconnu  $\tilde{W}$  a donc nécessairement une forme de produit...

- d'un opérateur  $\hat{\mathbf{e}}_{\mathbf{k}_T} \cdot \mathbf{j}_{-\mathbf{k}}$  agissant exclusivement dans l'espace des états de la matière,
- et d'un opérateur, que l'on désignera par  $a_{\mathbf{k}_T}$ , agissant dans l'espace des états du rayonnement,

soit

$$\tilde{W} = -q \sqrt{\frac{\hbar}{2\varepsilon_0 \omega \mathcal{V}}} (\hat{\mathbf{e}}_{\mathbf{k}_T} \cdot \mathbf{j}_{-\mathbf{k}}) \otimes a_{\mathbf{k}_T},$$

avec

$$|\langle \dots (n-1)_{\mathbf{k}_T} \dots | a_{\mathbf{k}_T} | \dots n_{\mathbf{k}_T} \dots \rangle|^2 \stackrel{\text{df}}{=} n_{\mathbf{k}_T}.$$

J'ai défini l'opérateur  $a_{\mathbf{k}T}$  de façon à faire apparaître, pour des raisons de confort ultérieur, le signe "moins" devant l'expression de  $\widetilde{W}$ . Usant encore de notre liberté de choix de phase, nous pouvons prendre

$$\langle \dots (n-1)_{\mathbf{k}T} \dots | a_{\mathbf{k}T} | \dots n_{\mathbf{k}T} \dots \rangle \stackrel{\text{df}}{=} \sqrt{n_{\mathbf{k}T}}.$$

Ayant décidé *a priori* que le processus de transition  $A \rightarrow B$  correspond dans notre modèle quantique à l'absorption d'un, et un seul, photon du mode  $\mathbf{k}T$ , nous complétons la définition de l'opérateur  $a_{\mathbf{k}T}$  par la propriété

$$\langle \dots n'_{\mathbf{k}T} \dots | a_{\mathbf{k}T} | \dots n_{\mathbf{k}T} \dots \rangle \stackrel{\text{df}}{=} 0,$$

- si  $n'_{\mathbf{k}T} \neq n_{\mathbf{k}T} - 1$ ,
- ou si  $n_{\mathbf{k}T} = 0$  (pas de photon à absorber!),
- ou si les états du rayonnement dans les autres modes (représentés par les points de suspension...) ne sont pas les mêmes à gauche et à droite.

En d'autres termes,  $a_{\mathbf{k}T}$  n'agit que dans l'espace de Fock du mode  $\mathbf{k}T$ .

L'opérateur  $a_{\mathbf{k}T}$  est maintenant complètement défini, avec pour conséquence (Exercice 2) que

$$a_{\mathbf{k}T} |n\rangle_{\mathbf{k}T} = \begin{cases} \sqrt{n} |n-1\rangle_{\mathbf{k}T} & \text{si } n \text{ entier } > 0; \\ 0 & \text{si } n = 0. \end{cases} \quad (\text{III.4})$$

(On conçoit ce que peuvent être des états à 3 photons ou à zéro photon, mais des états à  $-1$  ou à  $-4$  photon(s?), ça n'a pas de sens. Moins que rien, c'est vraiment rien!)

Tout est fini... et tout commence. Les états  $|n\rangle_{\mathbf{k}T}$  constituent, par définition, une base de l'espace de Fock du mode; les opérateurs  $a_{\mathbf{k}T}$  sont donc maintenant parfaitement déterminés. L'algèbre de ces opérateurs, et de leurs adjoints  $a_{\mathbf{k}T}^+$ , vous est bien connue puisque c'est en fait celle des *opérateurs d'annihilation* et de *création* d'un oscillateur harmonique. On a par exemple (en omettant les indices  $\mathbf{k}T$  tant que l'on reste dans le même mode):

$$\begin{aligned} \langle p | a a^+ | q \rangle &= \sum_{r=0}^{\infty} \langle p | a | r \rangle \langle r | a^+ | q \rangle \\ &= \langle p | a | p+1 \rangle \langle p+1 | a^+ | q \rangle \\ &= (p+1) \delta_{pq}, \end{aligned} \quad (\text{III.5})$$

et de même  $\langle p | a^+ a | q \rangle = p \delta_{pq}$ , donc  $[a, a^+] = 1$ .

Agissant dans des espaces différents, les opérateurs associés à des modes différents commutent; on a donc plus généralement:

$$[a_{\mathbf{k}T}, a_{\mathbf{k}'T'}^+] = \delta_{\mathbf{k}\mathbf{k}'} \delta_{TT'} \quad (\text{III.6})$$

$$[a_{\mathbf{k}T}, a_{\mathbf{k}'T'}] = [a_{\mathbf{k}T}^+, a_{\mathbf{k}'T'}^+] = 0. \quad (\text{III.7})$$

Notre rayonnement quantique est donc algébriquement équivalent à un ensemble d'oscillateurs; un mode  $\mathbf{k}T$  correspond à un oscillateur à une dimension, et un état du rayonnement à  $n$  photons dans le mode  $\mathbf{k}T$  correspond au  $n$ ème état excité de cet oscillateur.

### III.3 L'opérateur de champ

Nous connaissons maintenant la forme de la perturbation  $\widetilde{W}e^{-i\omega t}$  requise pour décrire un processus d'absorption d'un photon d'un mode quelconque. Comme tout hamiltonien digne de ce nom, l'hamiltonien total de la matière en interaction avec le rayonnement quantique doit être hermitique. Les raisons en remontent à l'exigence de conservation de la norme du vecteur d'état d'un système quantique (c'est cette propriété qui assure la pérennité du quanton par exemple).

La condition d'hermiticité ne peut ici être satisfaite — sans rien changer par ailleurs — qu'en ajoutant l'expression conjuguée  $\widetilde{W}^+e^{i\omega t}$  qui, dans l'approximation de l'onde tournante, n'apporte aucune contribution intempestive au taux d'absorption. D'autre part, il n'y a pas de mode intrinsèquement privilégié; il faut donc qu'une interaction du même type intervienne, de manière additive, quel que soit le mode du faisceau. L'hamiltonien d'interaction matière-rayonnement quantique est alors:

$$\widetilde{H}_{\text{int}}(t) = -q \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \left\{ (\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{j}_{-\mathbf{k}}) \otimes a_{\mathbf{k}T} e^{-i\omega t} + (\hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{j}_{-\mathbf{k}}^+) \otimes a_{\mathbf{k}T}^+ e^{i\omega t} \right\}.$$

L'opérateur  $a_{\mathbf{k}T}$  agissant dans l'espace des états du rayonnement quantique réalise l'absorption d'un photon. La présence du conjugué  $a_{\mathbf{k}T}^+$  nous laisse augurer des processus d'émission de photon sur lesquels nous reviendrons.

Grâce à la définition de  $\mathbf{j}_{\mathbf{k}}$  (eq. (II.62)), et à l'hermiticité de l'opérateur courant de matière, il vient

$$\begin{aligned} \widetilde{H}_{\text{int}}(t) = -q \int d^3\mathbf{r} \mathbf{j}(\mathbf{r}) \cdot \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \\ \times \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} + a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k} \cdot \mathbf{r} - \omega t)} \right\}, \end{aligned} \quad (\text{III.8})$$

expression dans laquelle un produit (scalaire) peut en cacher un autre, direct (d'opérateurs agissant respectivement dans les espaces des états de la matière et du rayonnement). L'analogie formelle avec l'hamiltonien d'interaction (II.61) de la théorie

semi-classique suggère la notation:

$$\mathbf{A}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} + a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}. \quad (\text{III.9})$$

Dans cette expression,  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$  sont des opérateurs agissant dans l'espace des états du rayonnement, chaque  $\hat{\mathbf{e}}_{\mathbf{k}T}$  est un triplet de nombres,  $T$  peut prendre deux valeurs pour chaque valeur d'un vecteur d'onde  $\mathbf{k}$  d'un mode de la boîte de volume  $\mathcal{V} = L_x L_y L_z$ , et  $\omega$  est une fonction de  $|\mathbf{k}|$ , définie par la relation de dispersion

$$\omega(\mathbf{k}) = |\mathbf{k}|c. \quad (\text{III.10})$$

Enfin, l'opérateur  $\mathbf{A}$  n'a qu'une dépendance paramétrique en  $\mathbf{r}$  et  $t$  qui ne sont pas des opérateurs et ne jouent aucun rôle de variables dynamiques, ni des quantons, ni du rayonnement. Le sens physique de  $\mathbf{A}(\mathbf{r}, t)$ , à ne pas confondre avec le potentiel vecteur de l'électrodynamique classique, va se dégager progressivement, et son importance ultérieure lui mérite un nom en propre, à savoir — puisqu'à chaque point de l'espace, à chaque instant, il associe un opérateur — celui d'*opérateur de champ*.

Nous verrons plus tard qu'existent des phénomènes de rayonnement qui ne peuvent être correctement pris en compte par l'interaction (III.8). Effectuant un pas heuristique de plus dans l'élaboration de la théorie quantique, n'hésitons pas à nous inspirer de l'expression complète (II.28) de l'hamiltonien d'interaction de la théorie semi-classique. (Cette forme présente au moins l'avantage esthétique d'être invariante de jauge.) Aussi, on adopte finalement:

$$H_{\text{int}} \stackrel{\text{df}}{=} \int d^3\mathbf{r} \left\{ -q\mathbf{j}(\mathbf{r}) \cdot \mathbf{A}(\mathbf{r}, t) + \frac{q^2}{2m} \rho(\mathbf{r}) \mathbf{A}^2(\mathbf{r}, t) \right\}, \quad (\text{III.11})$$

par extension de l'expression (III.8) que l'on reconnaît dans le premier terme. Maintenant que la lectrice doit avoir acquis l'habitude de l'espace produit de notre modèle, j'abandonne, pour alléger les notations et symboles, tout l'attirail de  $\otimes$  et autres "tildes" qui servaient à identifier les prolongements dans l'espace produit. Au-delà d'une simple analogie formelle avec l'expression classique (II.28), l'hamiltonien de l'interaction de la matière et du rayonnement quantiques (III.11) se justifiera par ses succès et par l'existence de limites classiques qui, inversement, expliquent la forme (II.28).

### III.4 L'opérateur énergie du rayonnement

Disposant de l'opérateur de champ  $\mathbf{A}(\mathbf{r}, t)$ , et en ayant remarqué le lien de cousinage (tout au moins quant à son utilisation dans le calcul des taux de transition)

avec le potentiel vecteur classique, nous pouvons tirer parti de cette parenté et tenter de construire des grandeurs physiques du rayonnement quantique par analogie formelle avec les grandeurs du rayonnement classique. De simples résultats d'un exercice de calcul au départ, ces opérateurs vont progressivement acquérir un sens effectivement physique, à mesure de leur utilisation, que ce soit par des limites classiques, des principes de conservation<sup>4</sup>, ou par leur rôle dans des transformations ou des changements de représentation. L'importance de ces opérateurs nécessite de les baptiser. Un double souci d'économiser les patronymes et l'effort d'apprentissage d'un nouveau vocabulaire, alors que nous avons bien assez à faire avec les nouveaux concepts, conduit à recourir tout simplement aux noms des grandeurs classiques auxquelles ces opérateurs s'apparentent.

### III.4.1 Quelques définitions

On définit ainsi, en s'inspirant de la définition classique (II.36), un *opérateur champ électrique du rayonnement*  $\mathbf{E}(\mathbf{r}, t) \stackrel{\text{df}}{=} -\partial_t \mathbf{A}(\mathbf{r}, t)$ , soit

$$\mathbf{E}(\mathbf{r}, t) = i \sum_{\mathbf{k}T} \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} - a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}, \quad (\text{III.12})$$

ainsi qu'un *opérateur champ magnétique*  $\mathbf{B}(\mathbf{r}, t) \stackrel{\text{df}}{=} \nabla \wedge \mathbf{A}(\mathbf{r}, t)$ . Par calcul direct des composantes, ou en utilisant le “truc” d'analyse vectorielle de Feynman (Exercice I.1), on trouve facilement que

$$\begin{aligned} \nabla \wedge (\boldsymbol{\varepsilon} e^{i\mathbf{k}\cdot\mathbf{r}}) &= i\mathbf{k} \wedge \boldsymbol{\varepsilon} e^{i\mathbf{k}\cdot\mathbf{r}} \\ &= i \frac{\omega}{c} \hat{\mathbf{k}} \wedge \boldsymbol{\varepsilon} e^{i\mathbf{k}\cdot\mathbf{r}}, \end{aligned}$$

et l'on a donc:

$$\mathbf{B}(\mathbf{r}, t) = i \sum_{\mathbf{k}T} \sqrt{\frac{\hbar\omega}{2\varepsilon_0 c^2 \mathcal{V}}} \left\{ a_{\mathbf{k}T} \mathbf{k} \wedge \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} - a_{\mathbf{k}T}^+ \mathbf{k} \wedge \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}. \quad (\text{III.13})$$

A partir de ces opérateurs, on songe naturellement à définir un *opérateur énergie du rayonnement*

$$U \stackrel{\text{df}}{=} \frac{\varepsilon_0}{2} \int_{\mathcal{V}} d^3\mathbf{r} \left( \mathbf{E}^2(\mathbf{r}, t) + c^2 \mathbf{B}^2(\mathbf{r}, t) \right) \quad (\text{III.14})$$

dont le calcul en fonction des opérateurs de création et annihilation de photons nécessite cette fois un peu plus de travail. Je vais en présenter quelques détails pour l'édification de l'arpète.

<sup>4</sup>Si vous ne l'avez déjà fait, voilà une occasion de plus de lire la parabole des cubes de Feynman, références [29, vol. I, §4-1] ou [28, p. 80].

Le développement du champ électrique (III.12) donne

$$\begin{aligned} \mathbf{E}^2(\mathbf{r}, t) = & -\frac{\hbar}{2\varepsilon_0} \sum_{\mathbf{k}T} \sum_{\mathbf{k}'T'} \sqrt{\omega\omega'} \\ & \times \left\{ a_{\mathbf{k}T} a_{\mathbf{k}'T'} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'} \frac{e^{i((\mathbf{k}+\mathbf{k}')\cdot\mathbf{r}-(\omega+\omega')t)}}{\mathcal{V}} \right. \\ & + a_{\mathbf{k}T}^+ a_{\mathbf{k}'T'}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \frac{e^{-i((\mathbf{k}+\mathbf{k}')\cdot\mathbf{r}-(\omega+\omega')t)}}{\mathcal{V}} \\ & - a_{\mathbf{k}T} a_{\mathbf{k}'T'}^+ \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \frac{e^{i((\mathbf{k}-\mathbf{k}')\cdot\mathbf{r}-(\omega-\omega')t)}}{\mathcal{V}} \\ & \left. - a_{\mathbf{k}T}^+ a_{\mathbf{k}'T'} \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'} \frac{e^{-i((\mathbf{k}-\mathbf{k}')\cdot\mathbf{r}-(\omega-\omega')t)}}{\mathcal{V}} \right\}, \end{aligned}$$

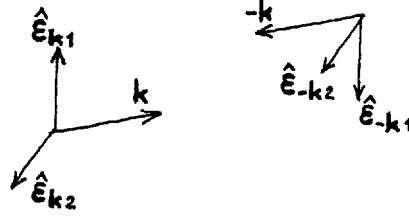
où j'ai posé  $\omega' \stackrel{\text{df}}{=} |\mathbf{k}'|c$ . C'est dans le calcul de l'intégrale de volume de cette quantité qu'apparaît l'utilité pratique de la propriété d'orthonormalité (II.45) de nos exponentielles de base, conséquence de notre choix idoine de conditions aux limites sur les parois de la boîte; ainsi,

$$\begin{aligned} \int_{\mathcal{V}} d^3\mathbf{r} \mathbf{E}^2(\mathbf{r}, t) = & \frac{\hbar}{2\varepsilon_0} \sum_{\mathbf{k}T} \omega \left\{ a_{\mathbf{k}T} a_{\mathbf{k}'T'}^+ \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* + a_{\mathbf{k}T}^+ a_{\mathbf{k}'T'} \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'} \right. \\ & \left. - a_{\mathbf{k}T} a_{-\mathbf{k}T'} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{-\mathbf{k}T'} e^{-2i\omega t} - a_{\mathbf{k}T}^+ a_{-\mathbf{k}T'}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \hat{\mathbf{e}}_{-\mathbf{k}T'}^* e^{2i\omega t} \right\}. \end{aligned}$$

J'avais mentionné (page 38) un choix commode pour les vecteurs polarisation de base. La seule condition qui leur est imposée est d'être orthogonaux au vecteur d'onde  $\mathbf{k}$  du mode; il est par ailleurs tout indiqué de les choisir unitaires et orthogonaux (ça ne change bien entendu rien physiquement, mais les calculs en sont tellement plus simples...), et de prévoir la possibilité de les multiplier (et combiner linéairement) par des nombres complexes (pour décrire la polarisation circulaire). La condition alors adoptée,  $\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* = \delta_{TT'}$ , est par contre muette en ce qui concerne le produit scalaire de vecteurs de polarisation de base associés à des modes de vecteurs d'onde  $\mathbf{k}$  et  $\mathbf{k}'$  différents. Ayant décidé des vecteurs de base du mode  $\mathbf{k}$ , nous pouvons profiter de la latitude qui nous est laissée dans le choix des vecteurs de base du mode  $-\mathbf{k}$ ; leur seule contrainte est d'être orthogonaux à  $-\mathbf{k}$ , donc à  $\mathbf{k}$ . Nous pouvons les choisir parallèles aux précédents, le trièdre restant direct (fig. III.2). La liberté de phase des vecteurs de base nous permet d'en décider les orientations. Le choix

$$\hat{\mathbf{e}}_{-\mathbf{k}T} \stackrel{\text{df}}{=} (-)^T \hat{\mathbf{e}}_{\mathbf{k}T}, \quad (\text{III.15})$$



Figure III.2: Un choix de vecteurs polarisation de base du mode  $-\mathbf{k}$ .

est immédiatement ratifié par la simplification apportée à l'expression de l'intégrale précédente qui devient:

$$\int_{\mathcal{V}} d^3\mathbf{r} \mathbf{E}^2(\mathbf{r}, t) = \frac{\hbar}{2\varepsilon_0} \sum_{\mathbf{k}T} \omega \left\{ a_{\mathbf{k}T} a_{\mathbf{k}T}^+ + a_{\mathbf{k}T}^+ a_{\mathbf{k}T} - (-)^T \left( a_{\mathbf{k}T} a_{-\mathbf{k}T} e^{-2i\omega t} + a_{\mathbf{k}T}^+ a_{-\mathbf{k}T}^+ e^{2i\omega t} \right) \right\},$$

expression indépendante, comme il se devait, des vecteurs polarisation de base, puisque  $\mathbf{E}^2$ , scalaire, ne peut en dépendre.

Le calcul de l'intégrale de volume de  $\mathbf{B}^2$  se conduit de la même façon; la lectrice y parviendra sans difficulté (Exercice 4), après avoir noté que

$$\hat{\mathbf{k}} \wedge \hat{\mathbf{e}}_{\mathbf{k}1} = \hat{\mathbf{e}}_{\mathbf{k}2} \quad (\text{III.16})$$

$$\hat{\mathbf{k}} \wedge \hat{\mathbf{e}}_{\mathbf{k}2} = -\hat{\mathbf{e}}_{\mathbf{k}1} \quad (\text{III.17})$$

C'est alors que le regroupement de ces intégrales pour calculer l'opérateur énergie (III.14) réserve une divine surprise que sa définition ne laissait guère prévoir. Les termes indépendants du temps sont identiques, tandis que les termes dépendants du temps se compensent exactement! Il vient ainsi

$$U = \sum_{\mathbf{k}T} \frac{\hbar\omega}{2} (a_{\mathbf{k}T} a_{\mathbf{k}T}^+ + a_{\mathbf{k}T}^+ a_{\mathbf{k}T}), \quad (\text{III.18})$$

où, rappelons-le,  $\omega$  est une fonction de  $\mathbf{k}$ , à savoir:  $\omega(\mathbf{k}) \stackrel{\text{df}}{=} |\mathbf{k}|c$ .

### III.4.2 Nombre de photons

En vue de son rôle ultérieur, il est utile de définir l'opérateur (hermitique)

$$N_{\mathbf{k}T} \stackrel{\text{df}}{=} a_{\mathbf{k}T}^+ a_{\mathbf{k}T}.$$

Mais quelle est sa signification? Nous avons déjà calculé (équation (III.5)) ses éléments de matrice entre états de base du même mode, et l'on en déduit immédiatement

$$N_{\mathbf{k}T}|n\rangle_{\mathbf{k}T} = n|n\rangle_{\mathbf{k}T}.$$

Les états de base à  $n$  photons sont donc états propres de notre opérateur avec les valeurs propres  $n$ , et nous ne pouvons baptiser celui-ci autrement qu'*opérateur nombre de photons du mode  $\mathbf{k}T$* .

Compte tenu des commutateurs (III.6–III.7), on a

$$a_{\mathbf{k}T}a_{\mathbf{k}T}^+ = N_{\mathbf{k}T} + 1,$$

et l'expression (III.18) de l'opérateur  $U$  devient finalement

$$U = \sum_{\mathbf{k}T} \hbar\omega(N_{\mathbf{k}T} + \tfrac{1}{2}). \quad (\text{III.19})$$

Un état de base du rayonnement (tel que (III.2) par exemple) est état propre simultané de tous les  $N_{\mathbf{k}T}$  (ceux-ci n'agissent effectivement que dans les espaces de leurs modes respectifs, donc ils commutent). Il est donc aussi état propre de  $U$ , avec pour valeur propre la somme des énergies de tous les photons de l'état (soit  $p\hbar\omega_1 + q\hbar\omega_2 + \dots$  dans l'exemple (III.2)).

Notre schéma, d'arbitraire qu'il pouvait paraître au début, commence donc à acquérir une certaine cohérence. Passer d'un état comportant, entre autres,  $n$  photons dans le mode  $\mathbf{k}T$  à un état à  $n+1$  photons dans ce même mode, toutes choses égales par ailleurs, correspond à accroître la valeur propre de  $U$  de la quantité  $\hbar\omega$  (c'est à dire l'énergie d'un photon du mode). L'appellation d'énergie totale du rayonnement pour l'opérateur  $U$  semble donc assez bien choisie. Dans un état propre, la valeur de cette grandeur est (presque, comme on va s'en apercevoir bientôt) la somme des énergies de tous les photons, le nom d'énergie d'un photon ayant lui-même été attribué à la quantité  $\hbar\omega$  à partir de motivations conservatrices appliquées à des processus de transition de la matière tels que l'absorption du rayonnement par un atome (effet photo-électrique). Notre hypothèse de départ de structure granulaire pour l'énergie du rayonnement semble en tous cas parfaitement accommodée par notre construction... qui nous réserve pourtant une surprise de taille!

### III.5 L'énergie du vide

L'existence d'un état à zéro photon dans tous les modes était nécessaire, à la base de notre modèle, pour pouvoir envisager des situations caractérisées par l'absence de tout rayonnement. C'est maintenant aussi une nécessité formelle. Par applications répétées de l'opérateur d'annihilation (III.4) sur un état de base à nombre de

photons donné, on obtient finalement un état propre du nombre de photons pour la valeur propre zéro, sur lequel une application supplémentaire ne conduit qu'au ket nul. Mais la valeur propre de l'opérateur énergie du rayonnement correspondante,

$$\mathcal{E}_0 = \frac{1}{2} \sum_{\mathbf{k}T} \hbar\omega, \quad (\text{III.20})$$

est, non seulement non nulle, mais infinie puisque, *horresco referens*, on a là une série infinie de modes  $\mathbf{k}$ , dont le terme général  $\hbar\omega$  est croissant.

Remarquons tout de même, au passage, que  $N_{\mathbf{k}T}$  étant un opérateur positif (ses états propres sont, ni plus ni moins, les états de base, et son spectre est l'ensemble des entiers naturels), la quantité  $\mathcal{E}_0$ , quoique infinie, est l'énergie minimale du rayonnement, ce qui est heureux pour un état supposément associé à l'absence de tout rayonnement et lui mériterait l'appellation de *fondamental*, plus dépouillée de connotations que "vide".

Formellement, l'irruption de cet infini était inévitable. Vous savez qu'à l'état fondamental d'un oscillateur quantique unidimensionnel de pulsation  $\omega$  est associée une *énergie de point zéro* valant  $\hbar\omega/2$ . Algébriquement, nous avons ici une infinité d'oscillateurs — chacun dans son fondamental — donc une énergie de point zéro de notre rayonnement qui est infinie. De par la structure continue que nous attribuons à l'espace des positions classiques en chaque point duquel nous attachons, dans une théorie de champ, une grandeur physique, nous ne pouvons qu'obtenir une infinité de degrés de liberté, chacun apportant son énergie de point zéro. Par mise en boîte, avec des conditions adéquates sur les parois, nous pouvons réduire cette infinité à être dénombrable, mais aller plus loin (si l'on peut dire) exigerait le pas (littéralement) introduit par une discrétisation de l'espace. C'est d'ailleurs ce que font numériquement les actuelles théories sur réseau: une boîte de points en nombre fini joue les rôles d'un lit de Procuste et d'une guillotine qui sectionnent une tranche de modes entre deux fréquences de coupure, l'une inférieure (équation (II.47)), fonction des dimensions de la boîte, l'autre supérieure,  $\omega_c \stackrel{\text{df}}{=} 2\pi c/\delta$ , selon le pas  $\delta$  entre points du réseau.

Nous allons voir que l'on peut malgré tout, sans remettre en cause la continuité de l'espace, extraire des informations significatives de cette énergie du vide infinie. La présence de tels infinis a été considérée, dès l'avènement de la théorie quantique du rayonnement, comme une tare qui devait condamner celle-ci à brève échéance. Soixante ans après, il n'en est toujours rien. Les physiciens ont appris à cohabiter avec la chose, encore qu'ils soient généralement plutôt honteux de son incongruité. Curieusement, un tel scrupule — après tout honorable, sinon justifié — a presque totalement disparu à l'endroit de l'électrodynamique classique dont l'ingrédient de base, la charge ponctuelle, est aussi responsable d'infinis — d'ailleurs pires<sup>5</sup> — que

<sup>5</sup>Sur cette question vous pouvez (et si l'histoire de la théorie des champs vous intéresse, vous devez) lire les récits de deux personnages principaux, Weinberg [77, p. 171] et Weisskopf [79, p. 69].

chacun a bien refoulés dans son inconscient.

### III.5.1 La soustraction

Le lieu commun qui permet d'éluder le problème consiste à rappeler que l'énergie n'est définie qu'à une constante additive près — seules ses variations importent — et qu'en conséquence un simple décalage (infini quand même) d'origine suffit à ramener l'énergie à des proportions plus rassurantes. Je vais effectivement procéder sur un exemple à cette opération de soustraction de deux quantités infinies pour obtenir une quantité à la fois finie et significative.

Mais l'existence d'un tel algorithme est loin de répondre aux questions que la lectrice einsteinienne ne manque pas de se poser. En relativité, l'énergie a un caractère absolu (!) avec, dans le cas d'une particule, l'existence d'un minimum infranchissable lorsqu'on est dans le repère propre de celle-ci. L'énergie est une composante d'un *quadrivecteur* par rapport aux transformations de Lorentz, et l'on ne peut impunément lui ajouter une constante sans tout gâcher. D'autre part, l'interaction gravitationnelle a pour source l'énergie (communément appelée masse lorsque cette source est de la matière au repos), et un décalage de la valeur de celle-ci ne peut laisser le monde indifférent. Le problème est donc réel, même si la bonne éducation des physiciens leur permet de faire comme s'il n'existait pas; il s'agit là d'un de ces trous noirs de la physique, aux contours plus ou moins délimités, que l'on apprend à contourner, mais qui attendent toujours leur clarification. De celle-ci surgira peut-être une idée "révolutionnaire", comme la relativité restreinte lorsqu'elle a balayé toutes les énigmatiques propriétés que les physiciens se voyaient obligés d'attribuer au vide, alors appelé éther.

Quoi qu'il en soit, le calcul explicite de l'*effet Casimir* va mettre en évidence un effet physique de l'énergie de point zéro qui n'implique en rien la relativité ou la gravitation. Les pulsations  $\omega$  des modes propres de notre boîte dépendent (un peu) de sa forme et (beaucoup) de ses dimensions.<sup>6</sup> Le déplacement d'un objet macroscopique, comme une paroi, modifie les pulsations propres, ce qui entraîne une variation de l'énergie de point zéro, variation à laquelle correspond une force sur la paroi, et ceci en l'absence de toute source de rayonnement. Notre but est de calculer cette force d'interaction mutuelle entre deux grandes plaques conductrices, neutres, parallèles, séparées de la distance  $a$  (fig. III.3).

### III.5.2 Estimation de l'effet

Avant de se plonger dans les détails du calcul, il est de bonne pratique de faire appel à l'analyse dimensionnelle pour tenter une estimation de cette force, si toutefois elle n'est pas nulle... ou infinie. Imaginons les deux plaques tellement grandes qu'elles

---

<sup>6</sup>Un violon ressemble plus à une flûte qu'à une contrebasse, tout au moins pour l'oreille.

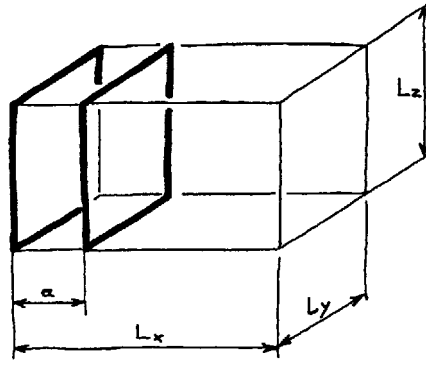


Figure III.3: Les deux plaques distantes de  $a$  et la boîte  $L_x, L_y, L_z$ .

soient infinies. Ce système jouit d'une totale invariance par rapport aux translations parallèles aux plaques. En l'occurrence, la force qui s'exerce sur un élément d'une plaque est la même quelle que soit la position de l'élément. Il suffit alors de juxtaposer deux éléments identiques pour réaliser que la force sur un élément est nécessairement proportionnelle à son aire.

De quelles quantités peut dépendre la force divisée par la surface? Il n'y a ni rayonnement, ni charge électrique, dans le système. L'examen de l'expression (III.20) montre que, pour une forme de boîte donnée, l'énergie de point zéro ne peut dépendre que de:

- la constante  $\hbar$ ;
- la constante  $c$ , puisque celle-ci intervient dans la relation de dispersion qui donne les valeurs de la pulsation en fonction des vecteurs d'onde des modes sur lesquels on doit sommer;
- et des longueurs qui mesurent la boîte et conditionnent les valeurs du vecteur d'onde.

Indépendante de la taille de la boîte, la force par unité de surface qui dérive de cette énergie, ne peut être que

$$\frac{F}{L_y L_z} \propto \frac{\hbar c}{a^4}, \quad (\text{III.21})$$

à un facteur numérique près, voisin de l'unité en vertu du *principe zéro de la physique* fondé, comme tout principe, sur l'expérience:

*Dans toute formule prise au hasard dans un livre de physique pris au hasard, les nombres sans dimension sont voisins de un.*<sup>7</sup>

<sup>7</sup>Ce principe reste à vrai dire de nature statistique. Il souffre des exceptions, nombres de

J'ai choisi un symbole graphique rappelant que ce n'est pas absolument une égalité (donc  $\neq$ ), ni une différence (donc  $\neq$ ), mais plutôt une égalité ( $=$ ) à un facteur ( $\times$ ) près.

### III.5.3 Les modes d'une cavité résonante

Etant donnée la présence de conducteurs, il est tout indiqué d'employer ceux-ci à la construction de notre boîte à modes propres, de façon à n'avoir à considérer qu'un seul type de condition limite sur les parois. Cette boîte a ici la propriété de confiner le rayonnement et n'est pas seulement un artifice de calcul; elle peut être réalisée (tout au moins dans des dimensions finies) et porte alors le nom de *cavité résonante*. La paroi conductrice impose idéalement comme condition que le champ électrique lui soit orthogonal ou, sinon, qu'il soit nul, ce qui va nous suffire pour déterminer les modes propres, sans faire appel au potentiel vecteur.

Le champ électrique dans la cavité est solution de l'équation d'onde homogène, conséquence directe des équations de Maxwell hors des sources,

$$(\Delta - \frac{1}{c^2} \partial_t^2) \mathbf{E}(\mathbf{r}, t) = 0.$$

Etant donnée la forme parallélépipédique des frontières, et les rôles analogues joués par les variables  $x, y, z$  et  $t$  dans l'équation d'onde, cherchons, pour la première composante, des solutions de la forme  $E_x(\mathbf{r}, t) = f_1(x)f_2(y)f_3(z)f_4(t)$ . Après report dans l'équation d'onde, et division par  $E_x$ , on obtient la condition

$$\frac{f_1''}{f_1} + \frac{f_2''}{f_2} + \frac{f_3''}{f_3} + \frac{f_4''}{f_4} = 0. \quad (\text{III.22})$$

Celle-ci ne peut être satisfaite, pour tous  $x, y, z$ , et  $t$ , que si chacun des rapports  $f_i''/f_i$  est constant. Chacune des fonctions  $f_i$  est donc une combinaison de sinus et de cosinus (circulaires ou hyperboliques).

Ce sont les solutions stationnaires qui nous intéressent (par opposition aux transitoires), c'est-à-dire oscillantes avec une pulsation  $\omega$ . D'autre part, la condition d'orthogonalité de  $\mathbf{E}$  sur les parois conductrices impose que  $E_x(\mathbf{r}, t)$  soit nulle lorsque  $y = 0$ ,  $y = L_y$ ,  $z = 0$ , ou  $z = L_z$ . On a donc, moyennant un choix convenable d'origine des temps,

$$E_x(\mathbf{r}, t) = f_1(x) \sin\left(\frac{m\pi}{L_y} y\right) \sin\left(\frac{n\pi}{L_z} z\right) \cos \omega t, \quad m, n \text{ entiers naturels.}$$

Par le même raisonnement, on a, pour les deux autres composantes,

$$E_y(\mathbf{r}, t) = \sin\left(\frac{l\pi}{L_x} x\right) g_2(y) \sin\left(\frac{n'\pi}{L_z} z\right) \cos(\omega' t + \varphi')$$

---

Reynolds ou exposants critiques par exemple, généralement associées à des transitions de phase, autrement dit à un changement de modèle pertinent. Quoi qu'il en soit, il participe du Premier principe moral de Wheeler [71]: Ne jamais commencer un calcul avant d'en connaître le résultat!

$$E_z(\mathbf{r}, t) = \sin\left(\frac{l'\pi}{L_x}x\right) \sin\left(\frac{m'\pi}{L_y}y\right) h_3(z) \cos(\omega''t + \varphi''),$$

où  $l, l', m'$  et  $n'$  sont des entiers naturels.

Mais nous sommes dans le vide, et notre solution de l'équation d'onde doit, pour être solution des équations de Maxwell, satisfaire la condition  $\nabla \cdot \mathbf{E} = 0$  pour tous  $x, y, z, t$  (dans la cavité). Calculez la divergence de la solution proposée et vous verrez facilement (je l'espère) que cette condition n'est réalisable que dans la mesure où les dépendances en  $x, y, z$  et  $t$  dans  $\partial_x E_x$ ,  $\partial_y E_y$  et  $\partial_z E_z$  sont les mêmes. Il s'ensuit que:

$$\omega' = \omega'' = \omega, \quad \varphi' = \varphi'' = 0,$$

$$l' = l, \quad m' = m, \quad n' = n,$$

$$\begin{aligned} f_1'(x) &\propto \sin\left(\frac{l\pi}{L_x}x\right), \\ g_2'(y) &\propto \sin\left(\frac{m\pi}{L_y}y\right), \\ h_3'(z) &\propto \sin\left(\frac{n\pi}{L_z}z\right). \end{aligned}$$

Notre tentative de solution est donc récompensée. On peut écrire ses composantes

$$\mathbf{E}(\mathbf{r}, t) = \begin{cases} E_{0x} \cos(k_x x) \sin(k_y y) \sin(k_z z) \cos \omega t \\ E_{0y} \sin(k_x x) \cos(k_y y) \sin(k_z z) \cos \omega t, \\ E_{0z} \sin(k_x x) \sin(k_y y) \cos(k_z z) \cos \omega t \end{cases}$$

expressions dans lesquelles on a

$$\begin{cases} k_x = (\pi/L_x)n_x \\ k_y = (\pi/L_y)n_y, \\ k_z = (\pi/L_z)n_z \end{cases} \quad \text{avec } n_x, n_y, n_z \text{ entiers naturels,} \quad (\text{III.23})$$

et  $\omega = |\mathbf{k}|c$ , en vertu de la condition (III.22), tandis que la condition de divergence nulle impose finalement  $\mathbf{E}_0 \cdot \mathbf{k} = 0$ .

Ainsi, il existe généralement deux modes (deux vecteurs  $\mathbf{E}_0$  indépendants) pour un vecteur d'onde  $\mathbf{k}$  donné. Mais lorsqu'un indice d'onde ( $n_x$  par exemple) est nul, le vecteur d'onde est parallèle à deux des parois de la cavité (les faces  $x = 0$  et  $x = L_x$ ) et il n'existe plus qu'un mode de polarisation qui soit à la fois orthogonal à ces parois et transverse. Enfin, lorsque deux indices sont nuls, le vecteur d'onde est orthogonal à deux faces en regard et il n'existe aucun mode satisfaisant ces deux conditions.

### III.5.4 L'effet Casimir

Calculons l'énergie de point zéro de la cavité de la figure III.3. La cloison intermédiaire, conductrice, étant étanche au rayonnement, l'énergie de cette cavité est, généralement, la somme des énergies de la cavité d'épaisseur  $a$  et de la cavité de longueur  $L_x - a$ :

$$\mathcal{E}_0 = \mathcal{E}(a) + \mathcal{E}(L_x - a). \quad (\text{III.24})$$

En vertu de notre comptage de modes précédent, l'expression (III.20) donne, pour la première cavité,

$$\mathcal{E}(a) = \sum'_{n_x, n_y, n_z} \hbar c \sqrt{k_x^2 + k_y^2 + k_z^2}$$

où la somme sur trois indices modifiée,  $\sum'$ , comporte par définition un facteur qui vaut...

- 1 si les trois indices sont différents de zéro,
- 1/2 si un (et un seul) indice est nul,
- 0 si deux (ou trois) indices sont nuls.

La composante  $k_z$  varie par pas  $\Delta k_z = \pi/L_z$ . On a donc

$$\begin{aligned} \sum_{n_z} f(k_z) &= \frac{L_z}{\pi} \Delta k_z \sum_{n_z} f(k_z) \\ &\underset{L_z \rightarrow \infty}{\sim} \frac{L_z}{\pi} \int_0^\infty dk_z f(k_z), \end{aligned}$$

et, dans cette hypothèse de quasi continuité de la variable  $k_z$ , on n'a que faire de la distinction du cas  $n_z = 0$ . Ainsi, pour une cavité auxiliaire dont les dimensions  $L_y$  et  $L_z$  sont suffisamment grandes, nous avons:

$$\mathcal{E}(a) = \hbar c \frac{L_y L_z}{\pi^2} \sum'_{n_x=0} \int_0^\infty dk_y \int_0^\infty dk_z \sqrt{\frac{\pi^2 n_x^2}{a^2} + k_y^2 + k_z^2},$$

où la somme simple modifiée affecte un facteur 1/2 à la contribution  $n_x = 0$ . Lorsque la dimension selon la direction  $x$  est grande elle aussi, on a évidemment:

$$\mathcal{E}(L_x - a) = \hbar c \frac{(L_x - a) L_y L_z}{\pi^3} \int_0^\infty dk_x \int_0^\infty dk_y \int_0^\infty dk_z \sqrt{k_x^2 + k_y^2 + k_z^2}.$$

La lectrice circonspecte se sera déjà inquiétée de nos manipulations désinvoltes, d'abord de sommes, puis d'intégrales, manifestement divergentes. Il faut bien voir qu'il ne s'agit là que de divergences de notre modèle et certainement pas de divergences du système physique que nous tentons d'analyser! La cavité résonante perd en réalité de son étanchéité dans les hautes fréquences, où la conductivité



des conducteurs diminue. La cavité tend vers la transparence, et l'idéal stationnaire est de plus en plus éloigné de la réalité pour les modes de fréquence élevée. Tentons d'améliorer le modèle en diminuant, de force, les contributions des hautes fréquences; nous en serons récompensés par une exquise surprise: le résultat final est indépendant des détails — plus culinaires que gastronomiques — du procédé choisi pour assurer la convergence. Il nous suffit de savoir que, d'une façon ou d'une autre, nos sommes et intégrales sont rendues convergentes par l'opération d'un facteur  $g(\omega)$  soumis à la seule condition d'être égal à un à basse fréquence, et nul à haute fréquence.<sup>8</sup>

C'est parce que nous pouvons soustraire de l'énergie  $\mathcal{E}_0$  (équation (III.24)) une constante idoine que nous parvenons finalement à en extraire la dépendance en  $a$  qui nous intéresse — sans préciser autrement l'instrument de convergence. Retranchons de  $\mathcal{E}_0$  l'énergie de point zéro d'une cavité de mêmes dimensions hors tout, exempte de cloison interne (fig. III.4). Cette énergie est bien — à toutes fins utiles — une constante, puisqu'indépendante de  $a$ . Quant à l'énergie décalée, elle vaut:

$$\begin{aligned}\mathcal{E}_{0-} &\stackrel{\text{df}}{=} \mathcal{E}_0 - \mathcal{E}(L_x) \\ &= \hbar c \frac{L_y L_z}{\pi^2} \left\{ \sum'_{n_x=0} \int_0^\infty dk_y \int_0^\infty dk_z \sqrt{\frac{\pi^2 n_x^2}{a^2} + k_y^2 + k_z^2} \right. \\ &\quad \left. - \frac{a}{\pi} \int_0^\infty dk_x \int_0^\infty dk_y \int_0^\infty dk_z \sqrt{k_x^2 + k_y^2 + k_z^2} \right\}.\end{aligned}$$

Pour l'intégration dans le plan  $(k_y, k_z)$ , on passe à des coordonnées polaires réduites  $\rho, \vartheta$ , définies par:

$$\begin{cases} k_y = (\pi/a)\rho \cos \vartheta \\ k_z = (\pi/a)\rho \sin \vartheta. \end{cases}$$

Remplaçons l'intégrale sur  $k_x$  par une intégrale sur la variable réduite  $u$ , telle que  $k_x = \pi u/a$ . Il vient alors:

$$\begin{aligned}\mathcal{E}_{0-} &= \pi \hbar c \frac{L_y L_z}{a^3} \left\{ \sum'_{n_x=0} \int_0^\infty d\rho \rho \int_0^{\pi/2} d\vartheta \sqrt{n_x^2 + \rho^2} \right. \\ &\quad \left. - \int_0^\infty du \int_0^\infty d\rho \rho \int_0^{\pi/2} d\vartheta \sqrt{u^2 + \rho^2} \right\};\end{aligned}$$

---

<sup>8</sup>Vous trouverez un calcul détaillé avec un facteur de convergence explicite du type  $e^{-\lambda\omega}$ , ainsi que des réflexions sur diverses manifestations d'énergies de point zéro et sur le rôle des procédés de soustraction en physique, dans [13]. Pour une revue, pédagogique, historique et à jour sur l'effet Casimir, voir [25].

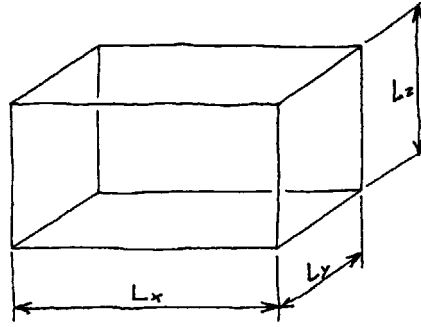


Figure III.4: La cavité sans cloison.

soit, après intégration tous azimuths et un nouveau changement de variable  $v \stackrel{\text{df}}{=} \rho^2$ ,

$$\mathcal{E}_{0-} = \frac{\pi^2}{4} \hbar c \frac{L_y L_z}{a^3} \left\{ \sum'_{n_x=0} f(n_x) - \int_0^\infty du f(u) \right\}, \quad (\text{III.25})$$

avec

$$f(u) \stackrel{\text{df}}{=} \int_0^\infty dv \sqrt{u^2 + v} g(\sqrt{u^2 + v}),$$

en faisant apparaître explicitement le facteur  $g$  de convergence.

Le reste est un jeu d'enfant dans la mesure où celle-ci se souvient de la *formule d'Euler-Maclaurin* (utilisée par exemple pour les corrections à l'estimation d'une intégrale définie par la méthode des trapèzes), que j'écris sous la forme<sup>9</sup>

$$\begin{aligned} \sum_{n=1}^{N-1} f(n) &= \int_0^N du f(u) - \frac{1}{2} [f(0) + f(N)] \\ &\quad - \frac{1}{12} [f^{(1)}(0) - f^{(1)}(N)] + \frac{1}{720} [f^{(3)}(0) - f^{(3)}(N)] - \dots \end{aligned}$$

Grâce à notre recette de convergence, nous sommes assurés que pour  $N$  assez grand,  $f$  et ses dérivées sont nulles. Appliquée à (III.25), cette formule donne

$$\mathcal{E}_{0-} = \frac{\pi^2}{4} \hbar c \frac{L_y L_z}{a^3} \left\{ -\frac{1}{12} f^{(1)}(0) + \frac{1}{720} f^{(3)}(0) - \dots \right\}.$$

<sup>9</sup>Voir le formulaire, commode, d'Abramowitz [1, p. 16]. Pour la démonstration, reportez vous à votre ouvrage d'analyse favori, par exemple [8, vol. I, p. 229]. Ce développement, dont je n'ai donné que les premiers termes, n'est pas convergent. En fait, c'est une série asymptotique. Il y aurait beaucoup à dire sur le rôle fondamental des séries divergentes en physique et en analyse numérique... sans parler de leur goût délicieux de fruit défendu! Mais un éloge des séries divergentes reste à faire, à la suite de Poincaré [61, t. 2, § 118].

L'évaluation des dérivées de  $f$  ne recèle aucune difficulté. On a, après un changement de variable de plus:

$$f(u) = \int_{u^2}^{\infty} dw w^{1/2} g(w^{1/2}),$$

et

$$f(u+du) - f(u) = - \int_{u^2}^{u^2+2u du} dw w^{1/2} g(w^{1/2}) = -2u^2 g(u) du,$$

d'où

$$\begin{aligned} f^{(1)}(u) &= -2u^2 g(u) \\ f^{(2)}(u) &= -4u g(u) - 2u^2 g^{(1)}(u) \\ f^{(3)}(u) &= -4g(u) - 8u g^{(1)}(u) - 2u^2 g^{(2)}(u). \end{aligned}$$

Au passage, vous vérifiez que  $f$  et ses dérivées sont bien nulles pour les grandes valeurs de l'argument<sup>10</sup> et, surtout, que

$$\begin{aligned} f^{(1)}(0) &= 0 \\ f^{(3)}(0) &= -4. \end{aligned}$$

Nous obtenons donc, finalement,

$$\mathcal{E}_{0-} = -\frac{\pi^2}{720} \frac{\hbar c}{a^3} L_y L_z,$$

où se révèle la dépendance en  $a$  qui était infiniment dissimulée dans  $\mathcal{E}_0$ . S'en dérive une force sur la cloison qui tend à diminuer  $a$  pour abaisser cette énergie, à savoir  $F = -\partial \mathcal{E}_0 / \partial a = -\partial \mathcal{E}_{0-} / \partial a$ , soit une *force attractive* entre les plaques en regard qui vaut, par unité de surface,

$$\frac{F}{L_y L_z} = \frac{\pi^2}{240} \frac{\hbar c}{a^4}, \quad (\text{III.26})$$

en accord avec l'expression subodorée (III.21), à un facteur 25 près qui rétrospectivement confère à celle-ci un rapport qualité/prix tout à fait honorable.

L'expression de cette force d'origine purement électromagnétique est, paradoxalement, indépendante de toute charge électrique, élémentaire ou non. On ne pouvait

<sup>10</sup>La dépendance du résultat par rapport au procédé de convergence (la structure de la fonction  $g$  utilisée) n'apparaîtrait en fait que dans les termes ultérieurs du développement d'Euler-Maclaurin, par le jeu des conditions — assez restrictives — de convergence uniforme pour la dérivabilité sous le signe somme.

que s'attendre à cette propriété, puisque j'ai pris soin de considérer l'interaction entre deux plaques neutres; aucune charge électrique n'intervient, comme il se devait à propos d'un phénomène intitulé énergie du vide! Mais il convient de remarquer qu'un tel prodige n'est dimensionnellement possible, alors même que le seul paramètre est une longueur (la distance entre les plaques), que par l'irruption entre celles-ci de la constante fondamentale  $\hbar$ , emblème du territoire quantique.

### III.5.5 Confirmation expérimentale

Dans le Système International,  $\hbar \approx 10^{-34}$  J s,  $c \approx 3 \times 10^8$  m s<sup>-1</sup> et, pour des plaques distantes de  $a = 1 \mu\text{m} = 10^{-6}$  m, l'effet Casimir (III.26) prédit une attraction mutuelle de  $1,2 \times 10^{-3}$  N m<sup>-2</sup>.

L'ordre de grandeur de l'effet est extrêmement faible, comparé par exemple à la pression atmosphérique (mille hectopascals, comme on dit à la météo, soit  $10^5$  N m<sup>-2</sup>), d'autant plus qu'il n'est évidemment pas question d'approcher à un micron deux plaques de un mètre carré! Jusqu'en 1958, les résultats de mesures n'étaient pas incompatibles (sans plus) avec cette prédiction de la théorie quantique du rayonnement. Diminuer la distance  $a$  pour mettre en évidence une force d'attraction plus élevée est impossible dans la mesure où nous avons assimilé les plaques à des conducteurs sans épaisseur : la profondeur  $\delta$  de pénétration du champ électromagnétique (l'effet de peau) atteint  $0,1 \mu\text{m}$  pour les plus basses fréquences mises en jeu qui sont données par la fréquence de coupure correspondant à l'épaisseur  $a$  de la boîte. Pour être valablement négligée, cette épaisseur effective des conducteurs, qui constituerait autrement un paramètre supplémentaire du phénomène, doit satisfaire la condition  $\delta \ll a$ .

On ne peut analyser les choses plus finement entre ces deux plaques qu'en dépassant le simple modèle heuristique — fondé sur un mélange astucieux de conducteurs et modes classiques, et d'énergie quantique — que je viens de décrire. Il nous faut préciser microscopiquement l'interaction entre les plaques, c'est-à-dire l'interaction entre leurs atomes. A courte distance ( $< 150$  Å) l'interaction dipôle-dipôle du type London-van der Waals entre atomes neutres se traduit par une attraction décroissante avec la distance comme  $a^{-7}$ . Au delà, la durée de propagation est supérieure au temps de vie des dipôles; la force entre atomes, dite alors de London-van der Waals retardée, s'en trouve modifiée. (C'est encore Casimir qui l'a calculée.) Il est alors possible, à partir d'un modèle de diélectrique constitué d'atomes neutres polarisables, de calculer l'attraction entre deux plaques diélectriques. Ainsi, Lifschitz a trouvé la même loi (III.26), corrigée par un facteur sans dimension dépendant de la permittivité relative du matériau des plaques.

Avec des plaques diélectriques, il n'y a plus d'effet de peau et il devient possible de tester expérimentalement l'effet Casimir à des courtes distances qui permettent la

manifestation de forces plus importantes.<sup>11</sup> La loi de force a pu être mesurée à courte distance entre deux plaques de mica obtenues par clivage, garantie de régularité moléculaire de leurs surfaces. Cette belle expérience [70] confirme parfaitement la prédiction de London-van der Waals de 50 à 150 Å, et sa version modifiée par le vide de rayonnement entre 150 et 300 Å.

### III.6 L'hamiltonien du rayonnement

Pour construire notre modèle quantique du système matière-rayonnement, nous avons donné à son espace des états la structure d'un produit...

- de l'espace des états de la matière,
- et de l'espace des états du rayonnement avec lequel nous commençons à nous familiariser.

Nous sommes parvenus à décrire des transitions de ce système, c'est à dire certaines évolutions, au moyen d'un hamiltonien constitué par la somme de l'hamiltonien de la matière seule  $H_{\text{mat}}$ , et d'un terme  $H_{\text{int}}(t)$  que j'ai appelé hamiltonien d'interaction.

Curieusement, et nous allons en voir bientôt la raison, nous n'avons eu nul besoin d'un hamiltonien du rayonnement quantique seul, alors qu'en bonne orthodoxie l'hamiltonien d'un système produit (par exemple l'atome d'hydrogène modélisé par le produit d'espaces de deux quantons, un proton et un électron) est en général la somme des hamiltoniens libres des sous-systèmes (les hamiltoniens libres du proton et de l'électron) agissant dans leurs espaces respectifs, et d'un terme d'interaction (l'interaction coulombienne dans le modèle le plus simple). Il n'est quand même pas interdit de se demander comment pourrait évoluer notre rayonnement quantique seul, en absence de toute matière, c'est à dire quel peut être son hamiltonien libre? Nous venons d'ailleurs de voir, avec l'effet Casimir, que ce rayonnement seul, et même dans son état fondamental — le vide de photon — a des propriétés physiques mesurables.

Guidé par une analogie formelle avec l'électrodynamique classique, j'ai déjà défini un opérateur énergie du rayonnement. L'appellation d'énergie n'est, semble-t-il, pas totalement frauduleuse. La dérivée de la valeur propre de cet opérateur associée à l'état fondamental, par rapport à un déplacement virtuel, donne une quantité à laquelle correspond bien, empiriquement, une force. Mais ce genre d'incantation a-t-il un pouvoir suffisant pour en faire l'hamiltonien du rayonnement libre, à savoir, dans un registre plus pédant, le générateur de l'évolution temporelle des états de ce système?

---

<sup>11</sup>De plus, les forces de London-van der Waals, retardées ou non, sont à longue portée par rapport aux dimensions atomiques, et leur étude est aussi appréciable pour la compréhension des phénomènes de floculation dans les colloïdes lyophobes et autres ratages de spécialités béarnaises, hollandaises ou mahonnaises [74].

Tel est bien le cas. Nous sommes partis d'une équation d'évolution pour le système matière-rayonnement (nulle part énoncée, mais implicite dans l'application du traitement perturbatif conduisant à (III.3) par la règle d'or):

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = (H_{\text{mat}} + H_{\text{int}}(t)) |\psi(t)\rangle, \quad (\text{III.27})$$

où  $|\psi(t)\rangle$  est un ket d'état de l'espace produit. Juste pour voir, associons lui un autre ket du même espace:

$$|\psi^{(S)}(t)\rangle \stackrel{\text{df}}{=} e^{-\frac{i}{\hbar}Ut} |\psi(t)\rangle, \quad (\text{III.28})$$

où  $U$  est l'opérateur énergie du rayonnement (III.14), et cherchons l'équation d'évolution de ce ket. Par définition, on a

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = e^{-\frac{i}{\hbar}Ut} \left( i\hbar \frac{d}{dt} + U \right) |\psi(t)\rangle,$$

expression qui peut s'écrire, en vertu de l'équation d'évolution (III.27) et en insérant une astucieuse décomposition de l'unité:

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = e^{-\frac{i}{\hbar}Ut} (H_{\text{mat}} + H_{\text{int}}(t) + U) e^{\frac{i}{\hbar}Ut} e^{-\frac{i}{\hbar}Ut} |\psi(t)\rangle.$$

Comme toute fonction de l'opérateur  $U$ , l'exponentielle  $\exp(iUt/\hbar)$  commute avec  $U$  lui-même. Elle commute également avec  $H_{\text{mat}}$  car ce sont des prolongements d'opérateurs agissant dans des espaces différents. On a donc finalement:

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = (H_{\text{mat}} + U + e^{-\frac{i}{\hbar}Ut} H_{\text{int}}(t) e^{\frac{i}{\hbar}Ut}) |\psi^{(S)}(t)\rangle. \quad (\text{III.29})$$

Cette équation d'évolution est strictement équivalente à l'équation (III.27); on peut utiliser l'une ou l'autre, qu'à passer d'un ket représentatif à l'autre par la relation de définition (III.28) lorsqu'il s'agit de calculer une grandeur physique en valeur moyenne ou propre. Mais sa forme est ici plus éloquente, car en absence de matière — et *a fortiori* d'interaction avec celle-ci — elle se réduit à:

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = U |\psi^{(S)}(t)\rangle.$$

Cette équation d'évolution d'un ket d'état du rayonnement libre a la forme d'une équation de Schrödinger dans laquelle l'opérateur énergie du rayonnement  $U$  joue manifestement le rôle de générateur des translations temporelles habituellement dévolu à l'hamiltonien. Il mérite donc parfaitement le nom d'*hamiltonien du rayonnement* et en conséquence je le noterai dorénavant  $H_{\text{ray}}$ . Pour la suite, il suffit de

se souvenir de son expression la plus commode, en fonction des opérateurs nombres de photons (équation (III.19)):

$$H_{\text{ray}} = \sum_{\mathbf{k}T} \hbar\omega \left( N_{\mathbf{k}T} + \frac{1}{2} \right).$$

Pas plus que  $H_{\text{mat}}$ , cet hamiltonien ne dépend explicitement du temps; on dit qu'il est *conservatif*. Qu'en est-il du troisième terme de l'hamiltonien total qui, dans l'équation (III.29), pilote l'évolution du ket d'état  $|\psi^{(S)}(t)\rangle$  du système matière-rayonnement,

$$e^{-\frac{i}{\hbar} H_{\text{ray}} t} H_{\text{int}}(t) e^{\frac{i}{\hbar} H_{\text{ray}} t} ?$$

Dans  $H_{\text{int}}(t)$  (éq. (III.11)), les seules q-entités qui risquent de ne pas commuter avec  $H_{\text{ray}}$  (et donc avec son exponentielle) sont les opérateurs d'annihilation et de création de photons  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$  qui interviennent dans le développement modal (III.9) de l'opérateur de champ  $\mathbf{A}$ . Etudions donc d'abord l'opérateur

$$a(t) \stackrel{\text{df}}{=} e^{-\frac{i}{\hbar} H_{\text{ray}} t} a_{\mathbf{k}T} e^{\frac{i}{\hbar} H_{\text{ray}} t}, \quad (\text{III.30})$$

le reste s'ensuivra aisément. Dériver cet opérateur par rapport à son paramètre  $t$  ne pose aucun problème (la lectrice incrédule reviendra pour cela à la définition de la dérivée); on a, grâce aux commutateurs (III.6):

$$\begin{aligned} i\hbar \frac{d}{dt} a(t) &= e^{-\frac{i}{\hbar} H_{\text{ray}} t} [H_{\text{ray}}, a_{\mathbf{k}T}] e^{\frac{i}{\hbar} H_{\text{ray}} t} \\ &= \sum_{\mathbf{k}'T'} \hbar\omega' e^{-\frac{i}{\hbar} H_{\text{ray}} t} \left[ a_{\mathbf{k}'T'}^+ a_{\mathbf{k}'T'} + \frac{1}{2}, a_{\mathbf{k}T} \right] e^{\frac{i}{\hbar} H_{\text{ray}} t} \\ &= \sum_{\mathbf{k}'T'} \hbar\omega' e^{-\frac{i}{\hbar} H_{\text{ray}} t} (-a_{\mathbf{k}'T'}) \delta_{\mathbf{k}'\mathbf{k}} \delta_{T'T} e^{\frac{i}{\hbar} H_{\text{ray}} t}. \end{aligned}$$

Notre opérateur  $a(t)$  est donc solution de l'équation différentielle

$$i\hbar \frac{d}{dt} a(t) = -\hbar\omega a(t),$$

assortie de la condition initiale  $a(0) = a_{\mathbf{k}T}$ , par définition. Cette solution est évidemment  $a(t) = a_{\mathbf{k}T} e^{i\omega t}$ . On a donc:

$$e^{-\frac{i}{\hbar} H_{\text{ray}} t} a_{\mathbf{k}T} e^{\frac{i}{\hbar} H_{\text{ray}} t} = a_{\mathbf{k}T} e^{i\omega t},$$

et, par conjugaison,

$$e^{-\frac{i}{\hbar} H_{\text{ray}} t} a_{\mathbf{k}T}^+ e^{\frac{i}{\hbar} H_{\text{ray}} t} = a_{\mathbf{k}T}^+ e^{-i\omega t}.$$

Au moyen de ces deux relations et du développement (III.9) de l'opérateur de champ on obtient maintenant facilement

$$e^{-\frac{i}{\hbar} H_{\text{ray}} t} \mathbf{A}(\mathbf{r}, t) e^{\frac{i}{\hbar} H_{\text{ray}} t} = \mathbf{A}(\mathbf{r}, 0),$$

comme on pouvait, après tout, s'y attendre pour la translation temporelle de la grandeur physique  $\mathbf{A}(\mathbf{r}, 0)$  d'un système dont le générateur est  $H_{\text{ray}}$ .

Par judicieuses introusses de l'unité idoine,

$$1 = e^{\frac{i}{\hbar} H_{\text{ray}} t} e^{-\frac{i}{\hbar} H_{\text{ray}} t},$$

on trouve immédiatement que le même type de relation vaut pour  $\mathbf{A}^2(\mathbf{r}, t)$  et plus généralement pour toute fonction de  $\mathbf{A}(\mathbf{r}, t)$ . En particulier, pour l'hamiltonien d'interaction (III.11), on a

$$\begin{aligned} e^{-\frac{i}{\hbar} H_{\text{ray}} t} H_{\text{int}}(t) e^{\frac{i}{\hbar} H_{\text{ray}} t} &= \int d^3 \mathbf{r} \left\{ -q \mathbf{j}(\mathbf{r}) \cdot \mathbf{A}(\mathbf{r}, 0) + \frac{q^2}{2m} \rho(\mathbf{r}) \mathbf{A}^2(\mathbf{r}, 0) \right\} \\ &= H_{\text{int}}(0). \end{aligned}$$

L'équation d'évolution (III.29) pour le ket d'état  $|\psi^{(S)}(t)\rangle$  peut donc finalement s'écrire

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = (H_{\text{mat}} + H_{\text{ray}} + H_{\text{int}}(0)) |\psi^{(S)}(t)\rangle,$$

c'est-à-dire sous forme d'une équation de Schrödinger dont l'hamiltonien est indépendant du temps. En choisissant de décrire le système matière-rayonnement par le ket d'état  $|\psi^{(S)}(t)\rangle$  (c'est ce que l'on appelle la *représentation de Schrödinger*), nous nous apercevons que ce système est conservatif.

Le ket d'état  $|\psi(t)\rangle$ , avant la transformation (III.28), n'appartient pas tout à fait — comme les bonnes élèves le soupçonnaient — à la *représentation de Heisenberg* puisque, avec l'équation (III.27), ce ket évolue. Mais l'absence dans celle-ci d'un des hamiltoniens libres nous apprend que nous faisons alors, prosaïquement, de la *représentation d'interaction* sans le savoir.

A un changement des kets représentatifs de l'état d'un système il faut bien entendu associer un changement de tout opérateur représentant une grandeur physique, puisque celui-ci est défini par l'ensemble des doublets:

- valeur numérique définie (ou valeur propre);
- ket d'état correspondant (ou ket propre).

Dans les calculs de probabilités (recouvrements d'états), un éventuel changement de représentation ne joue aucun rôle. Pour ce qui est des grandeurs physiques (valeurs propres, valeurs moyennes, dispersions et autres moments), c'est le couple opérateur-ket d'état qui intervient conjointement, et là encore le résultat est indépendant de la représentation.



Si toutes les représentations sont donc équivalentes, il n'en reste pas moins que leurs points de vue sont différents; les perspectives offertes et les détails mis en évidence peuvent varier, et en conséquence l'itinéraire du calcul pour parvenir au résultat cherché peut être plus ou moins commode.

### III.7 L'impulsion du rayonnement

Toujours par le jeu de construction d'analogies formelles avec la théorie classique, nous pouvons nous essayer à définir, à partir de la densité d'impulsion (II.52) — le vecteur de Poynting — du rayonnement classique, un *opérateur impulsion du rayonnement* quantique

$$\mathbf{P}_{\text{ray}} \stackrel{\text{df}}{=} \varepsilon_0 \int d^3\mathbf{r} \mathbf{E}(\mathbf{r}, t) \wedge \mathbf{B}(\mathbf{r}, t), \quad (\text{III.31})$$

où les opérateurs champs électrique et magnétique sont donnés par leurs développements modaux (III.12)-(III.13). Encore une fois, cette grandeur n'a pour l'instant qu'un sens purement rhétorique, et sa signification physique va se dégager progressivement par l'usage que nous parviendrons à en faire. Nous pouvons, comme pour l'opérateur énergie (III.14) qui a fini par accéder à la dignité d'hamiltonien — c'est à dire de générateur de l'évolution du rayonnement — chercher l'expression de cet opérateur impulsion en fonction des opérateurs de création et d'annihilation des photons.

Le calcul se conduit de la même façon, il ne nécessite qu'un peu plus de courage. Grâce à notre choix commode de vecteurs de base de polarisation du mode  $-\mathbf{k}$  en fonction de ceux du mode  $\mathbf{k}$  (équ. (III.15)), par ailleurs réels, nous bénéficions des identités

$$\begin{aligned} \hat{\mathbf{e}}_{\mathbf{k}T} \wedge \hat{\mathbf{e}}_{-\mathbf{k}T} &= \hat{\mathbf{e}}_{\mathbf{k}T} \wedge \hat{\mathbf{e}}_{\mathbf{k}T} = 0 \\ \hat{\mathbf{e}}_{\mathbf{k}1} \wedge \hat{\mathbf{e}}_{-\mathbf{k}2} &= \hat{\mathbf{e}}_{\mathbf{k}1} \wedge \hat{\mathbf{e}}_{\mathbf{k}2} = \hat{\mathbf{k}} \\ \hat{\mathbf{e}}_{\mathbf{k}2} \wedge \hat{\mathbf{e}}_{-\mathbf{k}1} &= -\hat{\mathbf{e}}_{\mathbf{k}2} \wedge \hat{\mathbf{e}}_{\mathbf{k}1} = \hat{\mathbf{k}}, \end{aligned}$$

et la lectrice trouvera aisément que

$$\begin{aligned} \int_V d^3\mathbf{r} \mathbf{E} \wedge \mathbf{B} &= \frac{\hbar}{2\varepsilon_0 c} \sum_{\mathbf{k}} \omega_{\mathbf{k}} \left\{ a_{\mathbf{k}1} a_{\mathbf{k}1}^+ + a_{\mathbf{k}2} a_{\mathbf{k}2}^+ + a_{\mathbf{k}1}^+ a_{\mathbf{k}1} + a_{\mathbf{k}2}^+ a_{\mathbf{k}2} \right. \\ &\quad \left. - (a_{\mathbf{k}1} a_{-\mathbf{k}1} - a_{\mathbf{k}2} a_{-\mathbf{k}2}) e^{-2i\omega t} - (a_{\mathbf{k}1}^+ a_{-\mathbf{k}1}^+ - a_{\mathbf{k}2}^+ a_{-\mathbf{k}2}^+) e^{2i\omega t} \right\}, \end{aligned}$$

où, en vertu de la relation de dispersion entre pulsation et vecteur d'onde, nous avons  $\omega_{\mathbf{k}} = k/c$ . Une somme telle que

$$\sum_{\mathbf{k}} \mathbf{k} a_{\mathbf{k}1} a_{-\mathbf{k}1} e^{-2i\omega t}$$

est nulle, car à toute contribution d'une valeur de  $\mathbf{k}$  dans la somme correspond une contribution opposée de la valeur  $-\mathbf{k}$  du fait de la commutativité de  $a_{\mathbf{k}1}$  et  $a_{-\mathbf{k}1}$  (éq. (III.7)). Reste finalement, grâce aux commutateurs (III.6),

$$\int_{\mathcal{V}} d^3\mathbf{r} \mathbf{E} \wedge \mathbf{B} = \frac{\hbar}{2\varepsilon_0} \sum_{\mathbf{k}T} \mathbf{k} (2a_{\mathbf{k}T}^+ a_{\mathbf{k}T} + 1).$$

Pour la même raison de symétrie, le terme constant n'apporte aucune contribution globale, du type "point zéro", et l'opérateur impulsion peut tout simplement s'écrire, en fonction de l'opérateur nombre de photons:

$$\mathbf{P}_{\text{ray}} = \sum_{\mathbf{k}T} \hbar \mathbf{k} N_{\mathbf{k}T}. \quad (\text{III.32})$$

Nous avons ainsi la surprise — déjà éventée à propos de l'opérateur énergie du rayonnement — d'obtenir un opérateur qui ne dépend pas explicitement du temps. L'analogie ne s'arrête pas là; ces deux opérateurs s'expriment en fonctions additives des opérateurs nombre de photons dans chaque mode, ce qui signifie ici qu'un état de base de l'espace de Fock — à nombres de photons déterminés dans chaque mode — est état propre de  $\mathbf{P}_{\text{ray}}$ , et surtout qu'un état avec un photon de plus dans le mode  $\mathbf{k}T$  est encore état propre avec une valeur propre augmentée de  $\hbar \mathbf{k}$ .

L'expression (III.19) obtenue pour l'énergie du rayonnement était une conséquence de l'origine heuristique du concept de photon (un paquet d'énergie). Inversement, nous voyons ici qu'il nous faut, avec la relation (III.32), attribuer au photon d'un mode  $\mathbf{k}T$  une nouvelle caractéristique individuelle mesurée par la quantité  $\hbar \mathbf{k}$ . Sa dimension, et l'appellation d'impulsion décernée au total  $\mathbf{P}_{\text{ray}}$  (ce n'est encore qu'un présupposé sémantique de notre jeune théorie quantique du rayonnement), nous déconseillent de baptiser cette quantité  $\hbar \mathbf{k}$  autrement qu'*impulsion d'un photon du mode  $\mathbf{k}T$* .

Selon cette terminologie notre photon, originellement *paquet d'énergie*, est aussi *paquet d'impulsion* et présente donc les attributs d'une particule. Tout notre traitement du rayonnement, calqué formellement sur la théorie classique de l'électromagnétisme, jouit, comme celle-ci, de l'invariance par rapport aux transformations de Lorentz même si, dans la formulation maxwellienne d'où nous sommes partis, elle reste dissimulée.<sup>12</sup> Les quantités  $\hbar \omega$  et  $\hbar \mathbf{k}$  ne sauraient donc se montrer dignes de leurs titres présumés d'énergie et d'impulsion d'une particule qu'en étant composantes d'un quadri-vecteur dont le carré, invariant, définit la *masse du photon*:

$$\begin{aligned} (mc^2)^2 &\stackrel{\text{df}}{=} E^2 - \mathbf{p}^2 c^2 \\ &= (\hbar \omega)^2 - (\hbar \mathbf{k})^2 c^2, \end{aligned}$$

<sup>12</sup>La formulation équivalente en termes de tenseur du champ électromagnétique exhibe clairement cette invariance.

soit  $m = 0$ , à cause de la relation de dispersion  $\omega(\mathbf{k}) = |\mathbf{k}|c$ .

Qu'il ait une masse nulle devrait déjà suffire à nous mettre en garde sur la "particularité" du photon. Comme il est doué d'énergie et d'impulsion, imaginons lui naïvement une vitesse classique. Celle-ci vaut, en bonne orthodoxie relativiste:

$$\mathbf{v} = \frac{\mathbf{p}}{E} = \frac{\mathbf{k}}{\omega},$$

d'où  $|\mathbf{v}| = c$ . L'intervalle de temps propre entre deux événements de la vie du photon, séparés dans notre repère par le temps  $dt$  et la distance  $d\mathbf{r} = \mathbf{v}dt$ , vaut donc

$$(d\tau)^2 \stackrel{\text{df}}{=} (dt)^2 - \left(\frac{d\mathbf{r}}{c}\right)^2 = 0.$$

Ainsi, le temps ne s'écoule pas pour notre photon. Mais, loin de lui conférer l'éternité cette propriété le rend, de son point de vue, parfaitement éphémère. Tout lui arrive en même temps: création, vie et annihilation. Ces étranges avatars ne sont que la traduction de l'impossibilité de figer au moyen d'une transformation de Lorentz un "objet" qui se "meut" à la vitesse  $c$ . Voilà une particule bien volage, que l'on ne peut contraindre au repos pour en étudier les propriétés statiques.

Il y a plus grave. Souvenons nous que notre photon désigne en fait un état du rayonnement quantique et comme tel, il n'y a pas de position — au sens de variable dynamique — qui lui soit associée. Les seules positions  $\mathbf{r}$  dont nous parlons à propos de champs sont des repères dans l'espace, ce ne sont pas des positions de quelque chose. Privé de position, le photon se démarque donc d'une particule classique, ce qui n'est pas surprenant puisque  $\hbar$  est au cœur même de sa définition. (Souvenez vous... les paquets d'énergie.) Mais nous verrons que le photon n'est pas plus une particule quantique — un quanton. Il n'existe aucune limite des équations d'évolution du rayonnement quantique, qui ait la forme de l'équation de Schrödinger d'un quanton, où interviendraient des grandeurs physiques  $\mathbf{R}$  et  $\mathbf{P}$ . Telle quelle, l'équation de Schrödinger d'un quanton est d'ailleurs visiblement dénuée de sens lorsque la masse est nulle.

Quelle image nous reste-t-il du photon? Certainement pas celle d'une bille classique: sa masse est nulle et il ne possède pas de variable dynamique position  $\mathbf{r}(t)$  dont on puisse suivre l'évolution au cours du temps. Ce n'est même pas le nuage protéiforme des probabilités de position d'un quanton. Paradoxalement, c'est le souvenir des origines classiques de notre théorie quantique du rayonnement qui va suggérer à notre entendement une représentation harmonieuse du photon.

### III.8 Le photon n'est pas une particule!

Nul ne songerait, en effet, à chercher une interprétation d'un champ classique (comme le champ électromagnétique, la hauteur des vagues ou l'élongation trans-

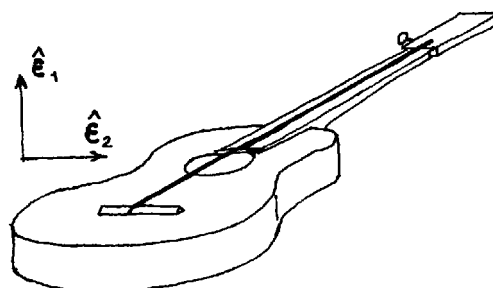
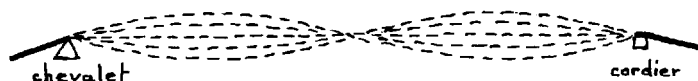


Figure III.5: Une corde de guitare et un choix, parmi d'autres, de vecteurs de base de polarisation.

verse d'une corde de guitare) en termes de particules. La longueur de la corde vibrante, par exemple — comme les dimensions de notre boîte à modes — est un des paramètres qui déterminent les modes, ou les fréquences. La corde est une “boîte” à une dimension. Ses vibrations transverses peuvent être décomposées selon deux vecteurs de base, orthogonaux à la configuration de repos de la corde, exacts analogues des vecteurs de base de polarisation du rayonnement. Pour une corde idéalement simple, les modes sont harmoniques (les fréquences sont multiples d'une fréquence fondamentale), et on peut associer à chaque fréquence un couple de vecteurs de base de polarisation. On n'est pas obligé de prendre les mêmes vecteurs de base pour tous les modes, encore que ce soit une simplification particulièrement indiquée dans ce cas unidimensionnel (fig. III.5).

A titre d'exemple, j'ai représenté sur la figure III.6 une vibration (polarisée dans le plan) verticale dans le premier harmonique, double du fondamental. Octavia — la guitariste — n'a bien entendu aucune raison de dire «à tel instant la vibration est ici, à tel autre instant la vibration est là». La vibration est un état de la corde, comme le photon est un état du rayonnement. Elle n'a pas de position, et une fois précisé son mode (fréquence et polarisation), n'importe plus pour la caractériser que son amplitude. C'est justement ce rôle d'amplitude que joue le nombre de photons dans un mode du rayonnement. Le caractère quantique — ou discret — de ceux-ci se bornant à restreindre l'amplitude — ou plutôt son carré puisqu'il s'agit de l'énergie — aux valeurs discrètes d'une suite arithmétique.

«Mais, objecte la lectrice opiniâtre, le photon est pourtant bien détecté ici, dans ma rétine, après avoir été émis là, dans le filament de l'ampoule». Un état à nombre de photons donné est un état stationnaire du rayonnement libre, sans matière avec laquelle interagir, comme notre vibration harmonieuse est, pour l'instant, le mouvement d'une corde libre. Cette corde ne peut être excitée, ou éventuellement entendue, que par interaction avec le plectre ou avec le chevalet couplé à la table,

Figure III.6: Une vibration selon le mode  $\omega_1, \hat{e}_2$ .

qui jouent respectivement les rôles de créateur et d'annihilateur, analogie qui montre qu'il faut se garder d'attribuer à la vibration, ou au photon, une position qui n'est que celle de la source ou du détecteur.

Bien sûr, l'approche de la réalité nécessite toujours la complication du modèle. Le plectre n'excite pas qu'un seul mode, et la guitariste sait, d'ailleurs, que de sa position sur la corde dépendent la richesse harmonique et la sonorité de la vibration. Il en va de même pour le rayonnement; ce n'est qu'au premier ordre des perturbations que les transitions n'impliquent qu'un photon dans un mode.

Notre image sonore illustre encore d'autres propriétés. Tout comme la probabilité de polarisation du photon émis ou absorbé par un atome dépend de l'état initial de celui-ci, nous remarquons que la polarisation de la vibration de la corde va dépendre de la direction donnée au mouvement du plectre. De même, la géométrie du chevalet, qui n'a absolument pas la symétrie de révolution autour de la corde, différencie sensiblement les taux d'absorption (on dit plutôt d'amortissement) des vibrations selon les modes  $\hat{e}_1$  et  $\hat{e}_2$  (voir la figure III.5).<sup>13</sup>

Après cette excursion dans le paysage audio-visuel, j'espère la lectrice convaincue de l'inutilité de chercher dans le photon une particule au sens classique, ou même quantique. Le photon n'est pas une entité localisée dans l'espace qui puisse être suivie, avec plus ou moins de précision, au cours du temps et, à l'inverse des objets auxquels s'applique le modèle du quanton pendant leur existence, le photon ne peut se manifester que par sa naissance ou sa mort, sa création ou son annihilation.

«D'accord. Mais alors pourquoi se permettre d'attribuer au photon de l'énergie et de l'impulsion (et même une masse, vraisemblablement nulle, et éventuellement un moment angulaire), si ce n'est une particule?» Nous allons voir dans le chapitre suivant que les grandeurs physiques associées au photon satisfont, en cas d'interaction avec la matière, des relations de conservation mettant en jeu énergie et impulsion de celle-ci. Suivant l'enseignement de la parabole des cubes et de la maman de Feynman, il est alors bien naturel de baptiser ces grandeurs énergie et impulsion du photon, comme la perte (localisée) d'énergie et d'impulsion du plectre, lorsqu'il gratte la corde, conduit à attribuer à la vibration de celle-ci une énergie et une impulsion (non localisées) correspondantes. Le rayonnement quantique, quant à

<sup>13</sup>Ce dernier effet conditionne en particulier la sonorité de la corde d'un piano au cours de la tenue après la frappe du marteau; ne subsiste à long terme que le mode de polarisation le moins absorbé par le chevalet. Voir [78].

lui, est non seulement inaudible (ou plutôt invisible), mais discret; les échanges d'énergie et d'impulsion (lorsqu'ils ont des valeurs déterminées) ne procèdent que par multiples d'une seule unité de base, le photon.

Un analogue plus quantique que la corde de guitare est constitué par les vibrations des atomes autour de leurs positions moyennes dans un réseau cristallin. Ces atomes peuvent être modélisés par un ensemble d'oscillateurs harmoniques quantiques dont l'excitation élémentaire s'appelle un *phonon*. Il est évidemment commode de nommer énergie et impulsion du phonon les sous-multiples élémentaires des valeurs des grandeurs physiques du système oscillant qui se trouvent satisfaire des relations de conservation avec l'énergie et l'impulsion du système exciteur (un neutron diffusé, par exemple). Mais, même si son langage peut donner l'impression contraire, la physicienne se représente toujours le phonon comme une vibration. Nul ne songerait à qualifier le phonon de particule (tout au plus ose-t-on parler de "quasiparticule") ou de quanton et ne tenterait encore moins de le représenter comme tels. La possibilité de "voir" la matière atomique vibrer inhibe cette velléité, tandis que dans le cas du rayonnement, l'immatérialité du support de la vibration, l'éther, lache la bride aux visions oniriques les plus extravagantes. Le vide est la page blanche de nos fantasmes.

## Exercices

1. Souvenirs de l'oscillateur harmonique; soit l'opérateur  $a$ , tel que  $[a, a^+] = 1$ .
  - i) Montrer que  $[a, (a^+)^n] = n (a^+)^{n-1}$ . En déduire, pour une fonction  $F$  d'une variable, la relation  $[a, F(a^+)] = F'(a^+)$ .
  - ii) Soit l'opérateur  $N \stackrel{\text{df}}{=} a^+ a$ . Que peut-on dire du spectre de  $N$ ? Que peut-on dire du signe de sa valeur moyenne  $\langle \alpha | a^+ a | \alpha \rangle$  dans un état quelconque? Quelle en est la conséquence pour le spectre de  $N$ ?
  - iii) Soit  $|r\rangle$  un état propre de  $N$  pour la valeur propre  $r$ . Montrer que  $a|r\rangle$  est état propre de  $N$ , pour une valeur propre à déterminer.
  - iv) Calculer  $[N, (a^+)^n]$ .
  - v) Soit  $|0\rangle$  l'état propre de  $N$  pour la valeur propre 0. Calculer le ket  $N (a^+)^n |0\rangle$  et la valeur moyenne  $\langle 0 | a^n (a^+)^n | 0 \rangle$ . En déduire l'expression de l'état propre normalisé  $|n\rangle$  en fonction de  $(a^+)^n |0\rangle$ .
  - vi) Montrer que  $[a, e^{\beta a^+}] = \beta e^{\beta a^+}$ .
  - vii) Montrer que  $a^+$  ne peut avoir de vecteur propre de norme finie.
  - viii) Déterminer le développement, sur les états de base  $|n\rangle$ , du ket propre de l'opérateur  $a$ , normé à l'unité, associé à la valeur propre  $z$ .

2. Partant de la définition de l'opérateur  $a$  par ses éléments de matrice entre les états de base orthonormés  $|n\rangle$ ,  $n$  entier naturel:  $\langle n | a | m \rangle = \sqrt{n} \delta_{n, m-1}$ .

- i) Déterminez les éléments de matrice de l'adjoint  $a^+$ .

- ii) Déterminez les éléments de matrice de  $a^+a$ ,  $aa^+$  et  $[a, a^+]$ .
- iii) Déterminez les kets propres de l'opérateur  $N \stackrel{\text{df}}{=} a^+a$ . Quel est son spectre?
- iv) Calculez les commutateurs  $[N, a]$  et  $[N, a^+]$ . En déduire les résultats de l'action des opérateurs  $Na$  et  $Na^+$  sur le ket  $|n\rangle$ .
- v) Montrez que l'on peut obtenir tous les états de base par actions répétées de l'opérateur  $a^+$  sur le seul état  $|0\rangle$ .

**3.** Juste par curiosité, imaginons des “particules de champ” analogues aux photons, mais avec cette seule différence que leur nombre est sévèrement limité: dans un mode, on ne peut avoir qu’une seule particule, ou pas du tout. L’espace des états du mode est donc sous-tendu par les vecteurs de base  $|0\rangle$  et  $|1\rangle$ . Dans ce cas, la condition définitoire de l’opérateur  $a$  du mode devient  $\langle 0|a|1\rangle = 1$ , les trois autres éléments de matrice étant nuls.

- i) Ecrivez la matrice représentative de  $a$  sur les états de base. En déduire la matrice du conjugué,  $a^+$ .
- ii) Calculez les matrices de  $a^+a$ ,  $aa^+$  et  $[a, a^+]$ .
- iii) Calculez la matrice de l’anticommutateur  $[a, a^+]_+ \stackrel{\text{df}}{=} aa^+ + a^+a$ .
- iv) Quels sont les kets et valeurs propres de  $N \stackrel{\text{df}}{=} a^+a$ ?
- v) Quels sont les résultats de  $a|0\rangle$ ,  $a|1\rangle$ ,  $a^+|0\rangle$  et  $a^+|1\rangle$ ?

**4.** Soit l’opérateur intégral  $\int_V d^3\mathbf{r} \mathbf{B}^2(\mathbf{r}, t)$ , envisagé page 64. Calculer son développement en modes.

## Chapitre IV

# Propriétés des grandeurs physiques du rayonnement quantique

Les états de base de l'espace de Fock du rayonnement quantique pourraient-ils être des états propres des supposées grandeurs physiques de ce rayonnement? En d'autres termes, les états à nombre de photons déterminé, ou états propres des opérateurs de nombre de photon  $N_{\mathbf{k}T}$ , sont-ils aussi états propres des opérateurs de champ  $\mathbf{A}(\mathbf{r}, t)$ ,  $\mathbf{E}(\mathbf{r}, t)$ ,  $\mathbf{B}(\mathbf{r}, t)$  et des opérateurs qui s'en déduisent,  $H_{\text{ray}}$  et  $\mathbf{P}_{\text{ray}}$ ? Pour ces derniers, nous savons déjà que la réponse est affirmative, et qu'elle nous a apporté des précisions sur la signification, tout à la fois, de ces grandeurs et du concept de photon. L'ignorance dans laquelle nous sommes vis-à-vis des opérateurs de champ motive maintenant la recherche de toute information supplémentaire sur leur compte.

Toutes les grandeurs physiques que nous avons vues sont finalement des combinaisons des opérateurs d'annihilation  $a_{\mathbf{k}T}$  et de création  $a_{\mathbf{k}T}^+$ . Ces opérateurs ne commutent généralement pas avec les nombres de photons; par exemple:

$$[N_{\mathbf{k}T}, a_{\mathbf{k}'T'}] = [a_{\mathbf{k}T}^+ a_{\mathbf{k}T}, a_{\mathbf{k}'T'}] = -\delta_{\mathbf{k}\mathbf{k}'} \delta_{TT'} a_{\mathbf{k}T}.$$

Le nombre total de photons  $\sum N_{\mathbf{k}T}$  ne saurait donc commuter avec un opérateur tel que la composante  $E_x(\mathbf{r}, t)$  du champ électrique. Nos états de base ne seront pas états propres des combinaisons linéaires d'opérateurs  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$  que sont les opérateurs de champ. Pour dégager la signification de ces derniers, nous allons en être réduits, faute de valeurs déterminées en même temps que le nombre de photons, à calculer les moments les plus significatifs de leur distribution, à savoir



leurs valeurs moyennes et dispersions. La compatibilité de ces grandeurs est aussi une question physique importante, et nous aurons donc à évaluer les commutateurs des opérateurs correspondants.

Des fonctions bilinéaires des opérateurs de création et d'annihilation peuvent commuter avec le nombre de photons; c'est justement le cas des opérateurs énergie et impulsion — linéaires par rapport aux nombres de photons  $a_{\mathbf{k}T}^+ a_{\mathbf{k}T}$  — et la raison pour laquelle ils admettent des états propres communs, ce qui permet d'attribuer énergie et impulsion au photon. Reste à justifier ces appellations, ce que nous allons faire en invoquant le principe de conservation de l'impulsion et en montrant que l'opérateur impulsion a bien les propriétés attendues d'un générateur des translations. En ce qui concerne l'énergie, j'ai déjà montré que cet opérateur était effectivement un générateur de l'évolution temporelle du système, ce qui lui avait valu, de plein droit, le titre d'hamiltonien.

## IV.1 Valeurs moyennes et dispersions

Rappelons l'expression modale de l'opérateur dit champ électrique:

$$\mathbf{E}(\mathbf{r}, t) = i \sum_{\mathbf{k}T} \sqrt{\frac{\hbar \omega}{2\varepsilon_0 V}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} - a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k} \cdot \mathbf{r} - \omega t)} \right\}, \quad (\text{IV.1})$$

Nous allons étudier cet opérateur, plutôt que  $\mathbf{A}(\mathbf{r}, t)$  dont la pertinence physique est, après tout, sujette à caution si l'on s'en tient au principe d'invariance de jauge. En particulier, dans le vide (ou plus exactement dans l'état fondamental, ou à zéro photon dans tous les modes) nous avons  $\langle 0 | a_{\mathbf{k}T} | 0 \rangle = \langle 0 | a_{\mathbf{k}T}^+ | 0 \rangle = 0$ , et donc:

$$\langle 0 | \mathbf{E}(\mathbf{r}, t) | 0 \rangle = 0.$$

Dans ce cas, en tous lieux  $\mathbf{r}$ , tous instants  $t$  (qui, rappelons le, ne sont que des paramètres, pas des grandeurs physiques), la valeur moyenne du champ électrique est nulle, ce qui est rien moins que satisfaisant pour la cohérence d'une description dans laquelle  $\mathbf{E}(\mathbf{r}, t)$  est censé représenter un champ électrique, et  $|0\rangle$  l'absence de tout rayonnement. . . encore que l'existence de l'effet Casimir nous permette déjà de nourrir quelque doute sur ce dernier point. Les catastrophes commencent d'ailleurs à se produire dès que nous nous interrogeons sur la même valeur moyenne dans un état excité. On a en effet:

$$\mathbf{k}T \langle n | a_{\mathbf{k}T} | n \rangle_{\mathbf{k}T} = 0,$$

et la valeur moyenne  $\langle \mathbf{E}(\mathbf{r}, t) \rangle$  dans un état excité — c'est à dire à nombres de photons déterminés non tous nuls — est aussi nulle!

Essayons maintenant un opérateur bilinéaire par rapport aux opérateurs de création/annihilation, à savoir la valeur moyenne  $\langle \mathbf{E}^2(\mathbf{r}, t) \rangle_0$  du carré de l'opérateur

champ électrique dans l'état vide, dans l'espoir qu'elle ne soit pas nulle. Je laisse le soin à la lectrice de faire ce calcul, en me bornant à remarquer que parmi les quatre types de produits d'opérateurs de création/annihilation qui apparaissent, un seul admet des valeurs moyennes non nulles dans le vide,  $\langle 0|a_{\mathbf{k}T}a_{\mathbf{k}'T'}^+|0\rangle = \delta_{\mathbf{k}\mathbf{k}'}\delta_{TT'}$ , ce qui nous vaut finalement

$$\langle \mathbf{E}^2(\mathbf{r}, t) \rangle_0 = \frac{1}{2\varepsilon_0\mathcal{V}} \sum_{\mathbf{k}T} \hbar\omega = \infty,$$

et récompense — trop largement — notre attente! Cela signifie, pour ce qui est de la dispersion du champ électrique en tous lieux et instants:

$$(\Delta \mathbf{E})_0^2 \stackrel{\text{df}}{=} \langle (\mathbf{E} - \langle \mathbf{E} \rangle_0)^2 \rangle_0 = \langle \mathbf{E}^2 \rangle_0 - \langle \mathbf{E} \rangle_0^2 = \infty;$$

la *fluctuation du vide* est infinie! Encore un de ces étranges infinis qui avaient déjà commencé à se manifester dans notre embryonnaire théorie quantique d'un champ. Que  $(\Delta \mathbf{E})_0$  soit différent de zéro n'est pas pour surprendre, puisque l'état  $|0\rangle$  est un état propre du nombre total de photons, opérateur qui ne commute pas avec  $\mathbf{E}(\mathbf{r}, t)$ . L'important, physiquement, est que cette dispersion ne soit pas nulle, et il ne faut (littéralement) pas se formaliser d'une valeur infinie. Cet excès est lié à la définition adoptée pour caractériser la largeur du spectre de valeurs de  $\mathbf{E}(\mathbf{r}, t)$ . J'ai choisi la dispersion — ou écart-type, ou deuxième moment centré — mais il existe bien d'autres façons d'évaluer la largeur d'une courbe "piquée": largeur à mi-hauteur, largeur aux points d'inflexion, *etc.*<sup>1</sup>

Dans ce cas particulier, la solution a été apportée par Bohr et Rosenfeld<sup>2</sup> en s'interrogeant sur le processus de mesure du champ électrique, qui ne peut, en dernière analyse, s'effectuer que par l'intervention de *charges d'épreuves* d'extension finie ne serait-ce qu'à cause de la dispersion quantique qui affecte leur position. Le champ électrique  $\mathbf{E}(\mathbf{r}, t)$  n'intervient donc que par sa moyenne spatiale  $\overline{\mathbf{E}}$  sur une zone caractérisée par la largeur  $b$ . Cette opération de moyenne ne touche pas à la dépendance opératorielle en  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$ , et l'on a donc encore  $\langle \overline{\mathbf{E}} \rangle_0 = 0$ .

Par contre, mais pour la même raison, la valeur moyenne quantique de  $\overline{\mathbf{E}}^2$  ne sera pas nulle et nous pouvons même prédire son expression... si elle n'est pas infinie. L'examen du développement modal de  $\mathbf{E}$  (eq. (IV.1)) nous révèle que  $\overline{\mathbf{E}}^2$  est proportionnel à  $\hbar/\varepsilon_0$ . Le produit  $\varepsilon_0\overline{\mathbf{E}}^2$  a la dimension d'une densité d'énergie, autrement dit d'une énergie divisée par le cube d'une longueur. Or nous disposons dans notre patrimoine de la combinaison  $\hbar c$  qui a la dimension d'une énergie multipliée par une longueur. Si tout va bien, la moyenne  $\langle \overline{\mathbf{E}}^2 \rangle_0$  devrait être indépendante

<sup>1</sup>Par exemple, l'écart-type d'une lorentzienne — un pic d'allure benoîte — s'avère infini! C'est un artefact, et la largeur à mi-hauteur convient tout aussi bien, et même mieux dans ce cas, pour caractériser une lorentzienne.

<sup>2</sup>Voir la discussion dans [35, p. 81].

du volume  $\mathcal{V}$  de la boîte à modes, surtout lorsque celui-ci croît indéfiniment. Ne nous reste comme paramètre de longueur que la largeur  $b$  de la fenêtre dans laquelle nous calculons notre moyenne. En vertu de notre principe zéro de la physique, la moyenne quantique quadratique du champ moyen (!) est donc

$$\langle \overline{\mathbf{E}}^2 \rangle_0 \propto \frac{\hbar c}{\varepsilon_0 b^4} . \quad (\text{IV.2})$$

On comprend bien, avec ce résultat, qu'il était impossible d'obtenir une dispersion finie sans faire intervenir un nouveau paramètre de longueur  $b$  caractéristique du procédé de mesure. La lectrice qui manque d'exercice pourra vérifier (Exercice 2) que, dans le cas d'un mécanisme de moyenne "à la Gauss" de portée  $b$ ,

$$\overline{\mathbf{E}}(t) \stackrel{\text{df}}{=} \frac{\int d^3\mathbf{r} e^{-\frac{r^2}{2b^2}} \mathbf{E}(\mathbf{r}, t)}{\int d^3\mathbf{r} e^{-\frac{r^2}{2b^2}}} ,$$

on obtient, dans une grande boîte,

$$\langle \overline{\mathbf{E}}(t)^2 \rangle_0 = \frac{1}{4\pi^2} \frac{\hbar c}{\varepsilon_0 b^4} ,$$

résultat finalement voisin de notre estimation (IV.2)... à un facteur de l'ordre de 40 près.<sup>3</sup>

## IV.2 Commutateurs des opérateurs de champ

Passons à la seconde partie de notre programme et demandons nous dans quelle mesure nos opérateurs de champ sont éventuellement compatibles, c'est-à-dire si ils admettent une base de vecteurs propres communs — ils ont alors respectivement des valeurs parfaitement définies, dites propres — et donc commutent. Les expressions de quelques commutateurs, choisis plus ou moins au hasard, seront pleines d'enseignements.

Commençons par un commutateur de deux composantes de l'opérateur de champ  $\mathbf{A}(\mathbf{r}, t)$  en des lieux distincts, mais au même instant. Un tel commutateur, qui réapparaîtra souvent dans la théorie, est appelé *commutateur aux temps égaux* ou *commutateur isochrone*. À l'aide du développement modal (III.9), de vecteurs de base de polarisation réels, et des commutateurs (III.6–III.7), il vient

$$[A_x(\mathbf{r}, t), A_y(\mathbf{r}', t)]$$

---

<sup>3</sup>On aurait d'ailleurs pu aider la chance en adjoignant au dénominateur de (IV.2) un facteur  $4\pi$  sans lequel la permittivité du vide  $\varepsilon_0$  est rarement rencontrée. L'estimation dimensionnelle ne différerait plus alors du résultat gaussien que par un facteur 3.

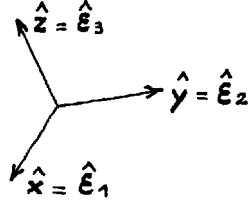


Figure IV.1: Un trièdre de projection  $\hat{x}, \hat{y}, \hat{z}$  particulier pour des vecteurs orthonormés  $\hat{\epsilon}_1, \hat{\epsilon}_2$  et  $\hat{\epsilon}_3$ !

$$\begin{aligned}
 &= \frac{\hbar}{2\varepsilon_0\mathcal{V}} \sum_{\mathbf{k}T} \sum_{\mathbf{k}'T'} \frac{1}{\sqrt{\omega\omega'}} (\hat{\epsilon}_{\mathbf{k}T})_x (\hat{\epsilon}_{\mathbf{k}'T'})_y \\
 &\quad \times [a_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} + a_{\mathbf{k}T}^+ e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)}, a_{\mathbf{k}'T'} e^{i(\mathbf{k}'\cdot\mathbf{r}'-\omega' t)} + a_{\mathbf{k}'T'}^+ e^{-i(\mathbf{k}'\cdot\mathbf{r}'-\omega' t)}] \\
 &= \frac{\hbar}{2\varepsilon_0\mathcal{V}} \sum_{\mathbf{k}T} \frac{1}{\omega} (\hat{\epsilon}_{\mathbf{k}T})_x (\hat{\epsilon}_{\mathbf{k}T})_y \left( e^{i\mathbf{k}\cdot(\mathbf{r}-\mathbf{r}')} - e^{-i\mathbf{k}\cdot(\mathbf{r}-\mathbf{r}')} \right),
 \end{aligned}$$

soit,

$$[A_x(\mathbf{r}, t), A_y(\mathbf{r}', t)] = i \frac{\hbar}{\varepsilon_0\mathcal{V}} \sum_{\mathbf{k}} \frac{1}{\omega} \sin(\mathbf{k} \cdot (\mathbf{r} - \mathbf{r}')) \sum_T (\hat{\epsilon}_{\mathbf{k}T})_x (\hat{\epsilon}_{\mathbf{k}T})_y. \quad (\text{IV.3})$$

Avec le choix (III.15) adopté pour les vecteurs de base de polarisation du mode  $-\mathbf{k}$  relativement à ceux du mode  $\mathbf{k}$ , on voit immédiatement que le résultat de la somme sur l'indice  $T$  est une fonction impaire de  $\mathbf{k}$ . Cette propriété peut se montrer plus formellement au moyen d'une astuce de la profession qu'il n'est pas inutile, une fois n'est pas coutume, de détailler. Soit trois vecteurs unitaires  $\hat{\epsilon}_T$ , avec  $T = 1, 2, 3$ , constituant un trièdre orthonormé (fig. IV.1) et demandons nous quelle peut être la valeur de  $\sum_{T=1}^3 (\hat{\epsilon}_T)_x (\hat{\epsilon}_T)_y$ .

Si nous faisons choix d'un trièdre de projection dont les axes  $\hat{x}, \hat{y}$  et  $\hat{z}$  sont suivant  $\hat{\epsilon}_1, \hat{\epsilon}_2$  et  $\hat{\epsilon}_3$  respectivement, nous avons évidemment

$$\begin{aligned}
 \sum_{T=1}^3 (\hat{\epsilon}_T)_x (\hat{\epsilon}_T)_y &= 0, \\
 \sum_{T=1}^3 (\hat{\epsilon}_T)_x (\hat{\epsilon}_T)_x &= 1,
 \end{aligned}$$

et des expressions analogues pour les autres combinaisons de composantes, tant et

si bien que

$$\sum_{T=1}^3 (\hat{\mathbf{e}}_T)_i (\hat{\mathbf{e}}_T)_j = \delta_{ij}.$$

Mais une quantité telle que  $(\hat{\mathbf{e}}_1)_i (\hat{\mathbf{e}}_1)_j$  est composante d'un tenseur à deux indices, car  $\hat{\mathbf{e}}_1$  est un vecteur par rapport aux rotations. En tant que somme de composantes de tenseur de mêmes indices, le premier membre de cette équation est donc composante d'un tenseur, ainsi que le second membre comme composante du tenseur de Kronecker. Les indices sont les mêmes à gauche et à droite; nous disposons donc d'une équation tensorielle équilibrée, dont la forme est invariante quel que soit le trièdre de projection. Dans le cas de nos vecteurs de base de polarisation d'un mode  $\mathbf{k}$ , nous avons  $\hat{\mathbf{e}}_3 = \mathbf{k}/k$ , et finalement

$$\sum_{T=1}^2 (\hat{\mathbf{e}}_T)_i (\hat{\mathbf{e}}_T)_j = \delta_{ij} - \frac{k_i k_j}{k^2},$$

qui est bien une fonction paire de  $\mathbf{k}$ .

Revenons à l'expression (IV.3), dans laquelle  $\omega$  est une fonction paire de  $\mathbf{k}$ , tandis que le sinus est une fonction impaire. Le terme général de la somme sur  $\mathbf{k}$  est donc une fonction impaire et nous obtenons

$$[A_x(\mathbf{r}, t), A_y(\mathbf{r}', t)] = 0.$$

Ce résultat est tout à fait satisfaisant d'un point de vue einsteinien: deux supposées grandeurs physiques, en des lieux différents au même instant, donc en des événements séparés par un intervalle du genre espace, ne devraient pas interférer. Que ce point de vue prévale ici n'est pas surprenant pour une théorie du rayonnement fondée au départ sur des équations de Maxwell qui ont — quoiqu'elles le cachent bien — une forme invariante par rapport aux transformations de Lorentz. Notre rayonnement — mais pas encore notre matière — est donc, par essence, relativiste einsteinien. En l'occurrence, les deux grandeurs  $A_x(\mathbf{r}, t)$  et  $A_y(\mathbf{r}', t)$  sont compatibles. On peut leur trouver des états propres communs et elles n'ont à subir aucune corrélation du type inégalité de Heisenberg, inévitable pour deux grandeurs qui ne commutent pas.<sup>4</sup>

Les grandeurs précédentes, même attachées au même point  $\mathbf{r} = \mathbf{r}'$ , commutent. Pour donner un exemple de grandeurs incompatibles lorsqu'associées au même événement, mais qui commutent dès lors qu'elles sont considérées en des lieux différents, évaluons le commutateur isochrone d'une composante de l'opérateur champ électrique et d'une composante de l'opérateur champ magnétique. Les

---

<sup>4</sup>Les corrélations (ou plutôt l'absence de corrélation) dont il s'agit ici concernent les dispersions des grandeurs physiques, à ne pas confondre avec les éventuelles corrélations entre valeurs moyennes des grandeurs physiques, mises en question par les inégalités de Bell.

développements en modes (III.12) et (III.13) — encore sur des vecteurs de base de polarisation rectiligne, c'est-à-dire réels, pour alléger l'écriture — donnent :

$$[E_i(\mathbf{r}, t), B_j(\mathbf{r}', t)] = \sum_{\mathbf{k}T} \frac{\hbar\omega}{2\varepsilon_0 c\mathcal{V}} (\hat{\boldsymbol{\epsilon}}_{\mathbf{k}T})_i (\mathbf{k} \wedge \hat{\boldsymbol{\epsilon}}_{\mathbf{k}T})_j \left( e^{i\mathbf{k}\cdot(\mathbf{r}-\mathbf{r}')} - e^{-i\mathbf{k}\cdot(\mathbf{r}-\mathbf{r}')} \right).$$

Grâce aux produits vectoriels (III.16–III.17), nous avons, dans un mode  $\mathbf{k}$  donné :

$$\sum_{T=1}^2 (\hat{\boldsymbol{\epsilon}}_T)_i (\mathbf{k} \wedge \hat{\boldsymbol{\epsilon}}_T)_j = (\hat{\boldsymbol{\epsilon}}_1)_i (\hat{\boldsymbol{\epsilon}}_2)_j - (\hat{\boldsymbol{\epsilon}}_1)_j (\hat{\boldsymbol{\epsilon}}_2)_i,$$

soit encore une composante d'un tenseur, cette fois-ci antisymétrique. On en déduit immédiatement, par exemple,

$$[E_x(\mathbf{r}, t), B_x(\mathbf{r}', t)] = 0$$

quelles que soient les positions  $\mathbf{r}$  et  $\mathbf{r}'$ , et surtout

$$\begin{aligned} [E_x(\mathbf{r}, t), B_y(\mathbf{r}', t)] &= i \frac{\hbar}{\varepsilon_0 c\mathcal{V}} \sum_{\mathbf{k}} \omega \sin(\mathbf{k} \cdot (\mathbf{r} - \mathbf{r}')) \frac{k_z}{k} \\ &= -i \frac{\hbar}{\varepsilon_0} \partial_z \frac{1}{\mathcal{V}} \sum_{\mathbf{k}} \cos(\mathbf{k} \cdot (\mathbf{r} - \mathbf{r}')) \\ &\underset{\nu \rightarrow \infty}{\sim} -i \frac{\hbar}{\varepsilon_0} \partial_z \frac{1}{(2\pi)^3} \int d^3\mathbf{k} e^{i\mathbf{k}\cdot(\mathbf{r}-\mathbf{r}')}, \end{aligned}$$

soit

$$[E_x(\mathbf{r}, t), B_y(\mathbf{r}', t)] = -i \frac{\hbar}{\varepsilon_0} \partial_z \delta^3(\mathbf{r} - \mathbf{r}'). \quad (\text{IV.4})$$

Ce commutateur, quoique très singulier pour  $\mathbf{r} = \mathbf{r}'$ , est donc à nouveau nul pour  $\mathbf{r} \neq \mathbf{r}'$ , au même instant  $t$ , en accord avec les prescriptions relativistes. Mais, en un événement  $\mathbf{r}, t$  donné, nous voyons que si  $E_x$  est compatible avec  $B_x$ , par contre il ne l'est pas avec  $B_y$ .

Plus tard — dans une présentation formalisée de la théorie quantique des champs dont l'invariance de Lorentz sera explicite — ce seront des commutateurs isochrones que l'on prendra comme point de départ pour établir, en bonne orthodoxie quantique, la structure de l'espace des états du système étudié. De la nullité du commutateur isochrone on déduira l'algèbre des opérateurs de création et annihilation et, par transformation de Lorentz, la compatibilité de grandeurs physiques en tous événements séparés par un intervalle du genre espace, que ces événements soient simultanés ou non.

Nous pouvons d'ailleurs déjà trouver ce résultat directement — mais lourdement, car l'invariance de Lorentz de notre théorie n'est pas explicite — en calculant par

exemple le commutateur de  $E_x$  et  $B_y$  en des événements quelconques. La lectrice obtiendra, par la même méthode que précédemment, et si elle accepte de travailler dans une grosse boîte:

$$[E_x(\mathbf{r}, t), B_y(\mathbf{r}', t')] = -i \frac{\hbar}{2\varepsilon_0} \partial_z \frac{1}{(2\pi)^3} \int d^3\mathbf{k} \left( e^{i(\mathbf{k} \cdot (\mathbf{r} - \mathbf{r}') - \omega(t - t'))} + \text{c.c.} \right),$$

où l'abréviation "c.c." désigne l'expression conjuguée complexe de la précédente. L'intégrale se traite sans difficultés (tout au moins formellement, car elle est divergente!) en calculant

$$\begin{aligned} \int d^3\mathbf{k} e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} &= \int_0^\infty dk k^2 \int_{-1}^1 d(-\cos \theta) e^{i(kr \cos \theta - \omega t)} 2\pi \\ &= 2\pi \int_0^\infty dk k^2 e^{-i\omega t} \frac{e^{-ikr} - e^{ikr}}{-ikr} \\ &= \frac{i\pi}{r} \int_{-\infty}^\infty dk k \left( e^{-ik(r+ct)} - e^{ik(r-ct)} \right). \end{aligned}$$

On peut encore transformer cette expression, pour s'apercevoir que l'on avait affaire à une représentation intégrale de la fonction de Dirac:

$$\begin{aligned} \int d^3\mathbf{k} e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} &= \frac{i\pi}{c} \frac{1}{-i\dot{c}} \partial_t \int_{-\infty}^\infty dk \left( e^{-ik(r+ct)} - e^{ik(r-ct)} \right) \\ &= -\frac{2\pi^2}{rc} \partial_t (\delta(r+ct) + \delta(r-ct)), \end{aligned}$$

et l'on obtient finalement

$$\begin{aligned} [E_x(\mathbf{r}, t), B_y(\mathbf{r}', t')] &= i \frac{\hbar}{4\pi\varepsilon_0 c} \\ &\times \partial_z \partial_t \frac{1}{|\mathbf{r} - \mathbf{r}'|} \left\{ \delta(|\mathbf{r} - \mathbf{r}'| + c(t - t')) - \delta(|\mathbf{r} - \mathbf{r}'| - c(t - t')) \right\}. \quad (\text{IV.5}) \end{aligned}$$

Ce commutateur est nul pour tout couple d'événements qui ne remplissent pas l'une des conditions

$$|\mathbf{r} - \mathbf{r}'| = \pm c|t - t'|,$$

éminemment satisfaisantes pour une théorie de la lumière: des grandeurs physiques en des événements qui n'appartiennent pas à leurs cônes de lumière respectifs — c'est-à-dire qui ne peuvent échanger du champ électromagnétique — sont compatibles et leurs dispersions ne sont donc soumises *a priori* à aucune corrélation.

J'ai maintenant tiré l'essentiel de ces calculs de commutateurs, tout en profitant de l'occasion pour indiquer quelques trucs techniques.<sup>5</sup> Remarquons, pour finir, que

---

<sup>5</sup>La lectrice qui le désire pourra calculer bien d'autres commutateurs que ceux qui ont été présentés ici... et vérifier ses résultats dans [35, app. 2].

dans les commutateurs non nuls tels que (IV.4) et (IV.5), et plus généralement dans toute la théorie quantique du rayonnement que nous sommes en train d'élaborer, il n'apparaît aucune masse ou charge élémentaire. Les seules constantes fondamentales qui interviennent sont  $\hbar$  et  $c$ ,  $\varepsilon_0$  n'étant là que pour des raisons de choix d'unités. Aucune référence n'est faite à des quantons élémentaires et cette théorie du rayonnement est construite en toute indépendance de ses sources.

### IV.3 L'impulsion et le générateur des translations

Reprenons la prétendue impulsion du rayonnement (III.31). Nous avons vu que cet opérateur a finalement pour expression

$$\mathbf{P}_{\text{ray}} = \sum_{\mathbf{k}T} \hbar \mathbf{k} N_{\mathbf{k}T} \quad (\text{IV.6})$$

en fonction des nombres de photons. Impulsion et nombre de photons sont donc compatibles,

$$[\mathbf{P}_{\text{ray}}, N_{\mathbf{k}T}] = 0,$$

ce qui nous a d'ailleurs permis de parler d'impulsion des photons en même temps que de leur énergie, puisque l'hamiltonien du rayonnement,

$$H_{\text{ray}} = \sum_{\mathbf{k}T} \hbar \omega (N_{\mathbf{k}T} + \frac{1}{2}),$$

jouit de la même propriété et commute avec  $\mathbf{P}_{\text{ray}}$ .

Mais pourquoi avoir baptisé cette grandeur impulsion du rayonnement? Nous allons voir que ce n'est pas seulement à cause d'une simple analogie formelle avec l'impulsion du rayonnement classique mais que la raison en est, de fait, exactement la même qu'en électrodynamique classique.

#### IV.3.1 Conservation de l'impulsion

Revenons à un système de quantons chargés (la matière) en interaction avec le rayonnement quantique, dont l'évolution est générée par l'hamiltonien (§ III.6):

$$H^{(S)} = H_{\text{mat}} + H_{\text{ray}} + H_{\text{int}}(0),$$

en représentation de Schrödinger. Considérons en particulier un système de quantons libres; alors l'impulsion totale des quantons,  $\mathbf{P}_{\text{mat}} \stackrel{\text{df}}{=} \sum_{i=1}^Z \mathbf{P}_i$ , et leur hamiltonien,  $H_{\text{mat}} \stackrel{\text{df}}{=} \sum_{i=1}^Z \mathbf{P}_i^2 / 2m_i$ , commutent:

$$[\mathbf{P}_{\text{mat}}, H_{\text{mat}}] = 0.$$



En absence de couplage avec le rayonnement, l'impulsion totale des quantons est donc bien une constante du mouvement. Mais qu'en est-il de l'impulsion totale  $\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}$  du système couplé? Nous avons:

$$[\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}, H^{(S)}] = [\mathbf{P}_{\text{mat}}, H_{\text{int}}(0)] + [\mathbf{P}_{\text{ray}}, H_{\text{int}}(0)],$$

commutateurs qui restent à calculer car le terme d'interaction,  $H_{\text{int}}(0)$ , éq. (III.11), dépend à la fois des grandeurs physiques positions et impulsions des quantons par l'intermédiaire des densités de matière  $\rho$  et  $\mathbf{j}$ , et des grandeurs physiques du rayonnement par l'intermédiaire de l'opérateur de champ  $\mathbf{A}(\mathbf{r}, 0)$ .

D'une part, on a, en représentation de positions des quantons,  $(\mathbf{P}_i)_x = \frac{\hbar}{i} \partial_{x_i}$ , et donc:

$$[(\mathbf{P}_{\text{mat}})_x, H_{\text{int}}(0)] = \frac{\hbar}{i} \sum_{i=1}^Z \partial_{x_i} H_{\text{int}}(0).$$

Pour ce qui est du rayonnement, d'autre part, les seuls opérateurs qui interviennent dans  $\mathbf{A}(\mathbf{r}, 0)$  sont les  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$ . De l'expression (IV.6), la lectrice déduira immédiatement que  $[\mathbf{P}_{\text{ray}}, a_{\mathbf{k}T}] = -\hbar \mathbf{k} a_{\mathbf{k}T}$ , et  $[\mathbf{P}_{\text{ray}}, a_{\mathbf{k}T}^+] = \hbar \mathbf{k} a_{\mathbf{k}T}^+$ . On peut faire encore mieux en remarquant que dans l'opérateur de champ

$$\mathbf{A}(\mathbf{r}, t) = \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0 \omega \mathcal{V}}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} + a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k} \cdot \mathbf{r} - \omega t)} \right\},$$

les opérateurs d'annihilation  $a_{\mathbf{k}T}$  (ou de création  $a_{\mathbf{k}T}^+$ ) n'interviennent que flanqués d'un facteur  $e^{i\mathbf{k} \cdot \mathbf{r}}$  (respectivement  $e^{-i\mathbf{k} \cdot \mathbf{r}}$ ) où, rappelons nous,  $\mathbf{r}$  est un paramètre de repérage dans l'espace, destiné à intégration, et absolument pas la valeur propre d'une grandeur physique position du quanton. Il est alors particulièrement intéressant de noter que ce sont toujours des commutateurs du même type qui interviennent, que l'on peut mettre sous la forme:

$$\begin{aligned} [(\mathbf{P}_{\text{ray}})_x, a_{\mathbf{k}T} e^{i\mathbf{k} \cdot \mathbf{r}}] &= \sum_{\mathbf{k}'T'} \hbar k'_x [N_{\mathbf{k}'T'}, a_{\mathbf{k}T}] \\ &= -\hbar k_x a_{\mathbf{k}T} e^{i\mathbf{k} \cdot \mathbf{r}} \\ &= i\hbar \partial_x a_{\mathbf{k}T} e^{i\mathbf{k} \cdot \mathbf{r}}, \end{aligned}$$

ou une forme analogue pour l'opérateur de création. Ainsi, toute fonction  $F$  qui ne dépend de la position  $\mathbf{r}$  que par l'intermédiaire de l'opérateur de champ, satisfait:

$$[(\mathbf{P}_{\text{ray}})_x, F(\mathbf{A}(\mathbf{r}, t))] = i\hbar \partial_x F(\mathbf{A}(\mathbf{r}, t)). \quad (\text{IV.7})$$

Or l'hamiltonien  $H_{\text{int}}(0)$  est justement obtenu par produit de telles fonctions ( $\mathbf{A}$  et  $\mathbf{A}^2$ ) avec les densités  $\rho$  et  $\mathbf{j}$ , équations (II.26–II.27), qui jouent le rôle de peignes

et, après intégration sur le paramètre  $\mathbf{r}$ , assignent à ce dernier les diverses valeurs propres de positions des quantons,  $\mathbf{r}_i$ . On a donc:

$$[(\mathbf{P}_{\text{ray}})_x, H_{\text{int}}(0)] = i\hbar \sum_{i=1}^Z \partial_{x_i} H_{\text{int}}(0),$$

et enfin, résultat inespéré,

$$[\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}, H^{(S)}] = 0.$$

En présence de rayonnement, l'impulsion totale des quantons n'est donc certes plus conservée,

$$[\mathbf{P}_{\text{mat}}, H^{(S)}] = [\mathbf{P}_{\text{mat}}, H_{\text{int}}(0)] \neq 0,$$

mais par contre nous venons de découvrir que la quantité  $\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}$  est conservée. L'impulsion  $\mathbf{P}_{\text{mat}}$  est une grandeur physique de la matière tandis que  $\mathbf{P}_{\text{ray}}$  ne concerne que le rayonnement. Considérant qu'ici le défaut de conservation de l'impulsion de la matière est parfaitement compensé par un gain de  $\mathbf{P}_{\text{ray}}$ , si nous tenons à un principe de conservation de l'impulsion il nous faut baptiser  $\mathbf{P}_{\text{ray}}$  *impulsion du rayonnement*, tout comme Feynman et sa maman à propos des différents termes de leur formule de conservation du nombre de cubes.<sup>6</sup> La lectrice, qui n'a pas oublié son électrodynamique classique, se souvient que ce n'est pas autrement que l'on peut attribuer au vecteur de Poynting la signification d'une densité d'impulsion du champ électromagnétique.

Rappelons ce que l'on entend en théorie quantique par les expressions *constante du mouvement* ou *grandeur conservée*. L'équation d'évolution du vecteur d'état est, en représentation de Schrödinger,

$$i\hbar \frac{d}{dt} |\psi^{(S)}(t)\rangle = H^{(S)} |\psi^{(S)}(t)\rangle,$$

d'où l'on déduit, pour la valeur moyenne d'une grandeur physique ne dépendant pas explicitement du temps, comme c'est le cas de notre impulsion totale,

$$\begin{aligned} i\hbar \frac{d}{dt} \langle \psi^{(S)}(t) | (\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}) | \psi^{(S)}(t) \rangle &= \langle \psi^{(S)}(t) | [\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}, H^{(S)}] | \psi^{(S)}(t) \rangle \\ &= 0. \end{aligned}$$

La valeur moyenne de l'impulsion totale du système matière-rayonnement reste constante au cours du temps. Il en est de même pour toute fonction de l'impulsion. Les divers moments de la distribution des valeurs de l'impulsion, et par conséquent cette distribution elle-même, sont donc des constantes. Il s'ensuit en particulier qu'un état propre de l'impulsion le reste.

---

<sup>6</sup>Toujours [29, vol. I, §4-1] ou [28, p. 80].

C'est en nous plaçant en représentation de Schrödinger que nous avons pu attribuer à  $\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}}$  la signification d'une impulsion totale du système. Quel peut être l'opérateur représentant la même grandeur physique lorsqu'on utilise la représentation d'interaction (§ III.6), souvent plus commode? Les changements de représentation pour lesquels nous avons toute latitude doivent néanmoins respecter les spectres des grandeurs physiques et en particulier leurs valeurs moyennes. On repasse à la représentation d'interaction par la substitution (III.28),

$$|\psi^{(S)}(t)\rangle = e^{-\frac{i}{\hbar}H_{\text{ray}}t}|\psi(t)\rangle,$$

qui permet d'écrire, pour la valeur moyenne de l'impulsion:

$$\langle\psi^{(S)}(t)|(\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}})|\psi^{(S)}(t)\rangle = \langle\psi(t)|e^{\frac{i}{\hbar}H_{\text{ray}}t}(\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}})e^{-\frac{i}{\hbar}H_{\text{ray}}t}|\psi(t)\rangle.$$

D'une part,  $H_{\text{ray}}$  et  $\mathbf{P}_{\text{mat}}$  commutent car ils agissent respectivement dans l'espace des états du rayonnement et dans l'espace des états des quantons. Mais il en va de même pour  $H_{\text{ray}}$  et  $\mathbf{P}_{\text{ray}}$  qui s'expriment tous deux en fonctions des opérateurs nombre de photons  $N_{\mathbf{k}T}$ . Nous avons donc:

$$\langle\psi^{(S)}(t)|(\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}})|\psi^{(S)}(t)\rangle = \langle\psi(t)|(\mathbf{P}_{\text{mat}} + \mathbf{P}_{\text{ray}})|\psi(t)\rangle.$$

C'est le même opérateur qui joue le rôle d'impulsion du système aussi bien dans la représentation de Schrödinger que dans la représentation d'interaction.

Notre retour sur la justification de la dénomination de l'opérateur  $\mathbf{P}_{\text{ray}}$  par sa relation avec  $\mathbf{P}_{\text{mat}}$  a peut être déclenché une interrogation retrospective dans l'esprit de la lectrice. Comment en était-on venu à attribuer à un opérateur relatif à un quanton le nom d'impulsion, représentée classiquement par une fonction à valeurs bien déterminées associée à une particule? Une raison (parmi d'autres) en est dans les équations d'évolution des valeurs moyennes de  $\mathbf{R}$  et  $\mathbf{P}$  pour un quanton soumis à l'hamiltonien

$$H = \frac{\mathbf{P}^2}{2m} + V(\mathbf{R}).$$

Ces équations (le théorème d'Ehrenfest),

$$\begin{aligned}\frac{d}{dt}\langle\mathbf{R}\rangle &= \left\langle\frac{\mathbf{P}}{m}\right\rangle \\ \frac{d}{dt}\langle\mathbf{P}\rangle &= \langle-\nabla V(\mathbf{R})\rangle,\end{aligned}$$

ressemblent aux équations d'évolution de la position et de l'impulsion (les équations de Hamilton) d'une particule classique, d'où l'adoption des mêmes patronymes pour les opérateurs  $\mathbf{R}$  et  $\mathbf{P}$ . Cette simple ressemblance, car en général  $\langle\nabla V(\mathbf{R})\rangle \neq \nabla V(\langle\mathbf{R}\rangle)$ , devient une équivalence lorsque le quanton a une fonction d'onde bien localisée, ce qui est somme toute rassurant: dans ce cas limite, la valeur moyenne de l'impulsion quantique évolue comme une impulsion classique.

### IV.3.2 Invariance et générateur

La raison fondamentale des similitudes formelles ou réelles entre théories classique et quantique réside finalement dans les structures identiques que nous y attribuons à l'espace. Dans ces deux cadres théoriques, nous croyons en effet au même principe d'invariance par translation qui permet à une physicienne et un physicien n'usant pas de la même origine. . .

- d'avoir des lois physiques de même forme,
- et d'établir un dictionnaire simple de correspondance entre valeurs qu'ils attribuent respectivement aux grandeurs physiques.

En physique classique, l'impulsion d'un système de particules invariant par translation est une constante du mouvement et, en ce qui nous concerne, tout est dit. En physique quantique la question est plus subtile, et potentiellement plus riche. L'opérateur qui permet d'établir la correspondance entre vecteurs d'état pour les deux protagonistes a une importance qui lui mérite un nom, celui de générateur des translations. Lorsque le système de quantons — ou l'hamiltonien qui en dicte l'avenir — jouit de l'invariance par translation, le générateur est une constante du mouvement et, à cette grandeur constante associée à la même invariance qu'en théorie classique, on donne encore le nom d'impulsion. Mais surtout, ce schéma logique éclaire le véritable rôle fonctionnel de l'impulsion quantique: c'est l'opérateur qui engendre le comportement des vecteurs d'état et grandeurs physiques du système au cours d'une translation spatiale. Nous allons voir que notre impulsion du rayonnement joue également ce rôle.<sup>7</sup>

Reprenons une esquisse plus formelle du rôle joué par une transformation et son générateur vis-à-vis d'un système quantique. Par exemple, Juliette et Roméo étudient le même système, mais leurs origines spatiales diffèrent. A un état du système que Juliette représente par le ket  $|\psi_J\rangle$ , Roméo assigne le ket  $|\psi_R\rangle$ . Le système fait l'objet d'une expérience de mesure de probabilité c'est-à-dire, pour Juliette, de recouvrement de deux états  $|\psi_J\rangle$  et  $|\phi_J\rangle$  par exemple. Pour Roméo, il s'agit donc d'une mesure du recouvrement des deux états correspondants  $|\psi_R\rangle$  et  $|\phi_R\rangle$ . Nos personnages croient à l'invariance par translation, qui se traduit ici par l'égalité de leurs probabilités et donc des modules des recouvrements:

$$|\langle\phi_J|\psi_J\rangle| = |\langle\phi_R|\psi_R\rangle|.$$

Cette condition concernant la correspondance bijective entre les kets attribués par Juliette et par Roméo aux mêmes états physiques du système, jointe au fait — apparemment anodin — que ces kets sont des éléments d'un espace de Hilbert (c'est-à-dire de normes finies), suffit pour que l'application satisfasse une propriété

---

<sup>7</sup>Par goût, et donc par plaisir, je préfère, comme vous n'avez pu manquer de vous en apercevoir, un exposé de la construction des modèles quantiques fondé sur des principes d'invariance. Ce point de vue est parfaitement relaté dans [51].

très spécifique: il existe un choix de phase pour tous les kets d'état tel que

$$\begin{aligned} |\psi_R\rangle &= U|\psi_J\rangle, \\ |\phi_R\rangle &= U|\phi_J\rangle, \\ &\text{etc.} \end{aligned}$$

où l'opérateur  $U$  est, de deux choses l'une, ou bien unitaire, ou bien anti-unitaire. C'est le *théorème de Wigner*.<sup>8</sup>

Le produit de deux translations est encore une translation, représentée par le produit des opérateurs  $U_1$  et  $U_2$  correspondants. Le produit de trois translations est indépendant d'éventuelles associations deux-à-deux de ces translations. L'outil mathématique qui satisfait ces *a priori* physiques est la structure de groupe. Toute translation est continûment connectée à l'identité par une suite de translations infinitésimales, identité représentée par l'opérateur unité et unitaire. Les translations (comme toutes les transformations continues) sont donc représentées par des opérateurs unitaires. Permettant d'associer à toute translation un opérateur linéaire dont la loi de multiplication est isomorphe de la loi de composition des translations, on dit que notre système quantique, ou ses kets d'état, engendrent une *représentation linéaire* du groupe des translations.

La translation de Juliette à Roméo est spécifiée par le vecteur  $\mathbf{a}$  dont dépend donc, paramétriquement, son opérateur représentatif  $U$ . Dans le cas d'une translation infinitésimale, disons parallèle à l'axe des  $x$  et de longueur  $da$ , l'opérateur représentatif doit être voisin de 1, et son écart à l'unité un opérateur (puisque'il en faut bien un) proportionnel au paramètre numérique infinitésimal  $da$ . On peut donc écrire, en introduisant des facteurs  $i$  et  $\hbar$  pour notre confort ultérieur,

$$U(\hat{\mathbf{x}} da) = 1 + \frac{i}{\hbar} P_x da. \quad (\text{IV.8})$$

L'opérateur  $P_x$  ainsi défini — toute ressemblance avec un opérateur existant ou ayant existé ne serait, pour l'instant, qu'une coïncidence! — agit dans l'espace des états du système quantique considéré et caractérise le comportement dudit système lors des translations, tout au moins infinitésimales; on l'appelle *générateur* des translations du système suivant l'axe  $\hat{\mathbf{x}}$ . L'unitarité de la représentation (IV.8) — le contenu du théorème de Wigner — impose évidemment l'hermiticité du générateur. La condition de linéarité par rapport au paramètre infinitésimal est essentielle si l'on veut que le produit de deux translations infinitésimales puisse encore s'écrire sous la même forme.

En cette fin de XX<sup>e</sup> siècle, notre vision physicienne de la nature est encore analytique. Concrètement, nous croyons que toute translation finie peut être impunément analysée en suite de translations infinitésimales, autrement dit que le

---

<sup>8</sup>La démonstration de ce théorème, purement mathématique mais non triviale, ne recèle en elle-même aucune signification physique. Rarement exposée, on la trouve dans [53, p. 540] ou [51].

comportement du système dans une translation finie est entièrement dicté par ses propriétés lors des translations infinitésimales. Si l'on veut “faire bien”, on prononce le mot “groupe” et l'on demande, dans ces conditions, que le groupe des translations — ici selon la direction  $\hat{\mathbf{x}}$  — soit intégrable, c'est-à-dire que sa structure soit entièrement dictée par ses propriétés au voisinage de l'identité. Encore autrement dit, il faut que les paramètres de la transformation résultant du produit de deux transformations successives soient des fonctions analytiques des paramètres de ces deux transformations.

L'outil mathématique adapté à cette vision de la nature s'appelle *groupe de Lie*. Dans le cas des translations dans l'espace vital à trois dimensions, ce groupe sera *commutatif*, ou *abelien*, puisque les translations dans cet espace sont, tant qu'on se le figure euclidien, des opérations qui commutent. Il n'en va pas de même pour les rotations dans ce même espace car à l'évidence — comme toute passionnée du “Rubik's cube” l'a expérimenté — le groupe des rotations est *non abelien*.

L'étude mathématique des groupes de Lie permet d'en démontrer de multiples propriétés. La plus utile pour nous est très simple: les groupes de Lie admettent des représentations exponentielles. Cette propriété est évidente pour ce qui est des translations à un paramètre: l'intégration de (IV.8) donne

$$U(a \hat{\mathbf{x}}) = e^{\frac{i}{\hbar} P_x a}.$$

Il en va bien sûr de même à trois dimensions:

$$U(\mathbf{a}) = e^{\frac{i}{\hbar} \mathbf{P} \cdot \mathbf{a}},$$

puisque les générateurs  $P_x$ ,  $P_y$  et  $P_z$ , comme les translations qu'ils génèrent, doivent commuter. Dans le cas des groupes de Lie non abeliens, la démonstration de cette propriété d'exponentiation n'est absolument pas triviale. Rares sont les physicien(ne)s qui prennent la peine de l'examiner. Pour les autres c'est une évidence! Vous pouvez donc être rassurée: il n'est nullement nécessaire de s'infliger l'étude de la “théorie des groupes” pour profiter de leurs propriétés.

### IV.3.3 Translations et quanton

Après ces considérations sur la correspondance entre kets utilisés par Juliette et par Roméo pour représenter le même état du même système, qu'en est-il des grandeurs physiques? Imaginons d'abord la plus simple des situations:

- Juliette, Roméo, et le système quantique objet de leur attention, vivent (effectivement) dans un espace à une dimension;
- le système n'est pas indifférent aux translations, il admet des grandeurs physiques qui, dans un état donné, n'ont pas les mêmes valeurs (moyennes par exemple) pour Juliette et Roméo; alors le système admet un générateur des translations,  $P$ , non trivial;

- parmi les grandeurs physiques, il en est une, disons  $X$ , qui mérite le nom de position (quantique) pour la simple raison que sa valeur moyenne se transforme comme une position classique:  $\langle \psi_R | X | \psi_R \rangle = \langle \psi_J | X | \psi_J \rangle - a$ .

On a alors, en cas de Juliette et Roméo très rapprochés:

$$\langle \psi_J | (1 - \frac{i}{\hbar} P da) X (1 + \frac{i}{\hbar} P da) | \psi_J \rangle = \langle \psi_J | X | \psi_J \rangle - da,$$

et ceci quel que soit l'état du système. On en déduit nécessairement que:

$$[X, P] = i\hbar.$$

Appréciez: nous pouvons déjà compter sur la cohérence d'un système quantique minimal, n'ayant que deux grandeurs physiques indépendantes,  $X$  et  $P$ . Pour la lectrice qui s'est infligée le prologue (chap. I): avec l'ensemble irréductible d'opérateurs  $X, P$  et leur commutateur, nous disposons maintenant du modèle du quanton (à une dimension).

Mais ces considérations n'interdisent nullement l'existence d'autres grandeurs physiques indépendantes. On peut maintenant imaginer, par produits d'espaces, toutes sortes de modèles ayant d'autres propriétés, en plus de celle du modèle minimal précédent. Rien de plus facile, par exemple, que de construire un modèle candidat pour décrire plusieurs quantons à une dimension (Exercice ??). D'un intérêt plus vital, nous pouvons par le même procédé fabriquer un nouveau modèle, dans le monde à trois dimensions, doué d'un générateur des translations  $\mathbf{P}$ , d'une grandeur position  $\mathbf{R}$  telle que (figure IV.2):

$$\langle \psi_R | \mathbf{R} | \psi_R \rangle = \langle \psi_J | \mathbf{R} | \psi_J \rangle - \mathbf{a}. \quad (\text{IV.9})$$

Je vous laisse le soin de montrer que les composantes de ces grandeurs physiques ont nécessairement la propriété

$$[R_i, P_j] = i\hbar \delta_{ij}. \quad (\text{IV.10})$$

Vous vous retrouvez donc en possession de votre bon vieux modèle du quanton à trois dimensions dont le générateur des translations n'est autre que ce que vous appeliez impulsion.<sup>9</sup>

Il convient de noter qu'en écrivant la propriété de transformation (IV.9), nous avons implicitement décidé que Juliette et Roméo utilisent le même opérateur position  $\mathbf{R}$ . En théorie quantique ce ne sont ni les opérateurs seuls, ni les kets d'état seuls qui déterminent les valeurs des grandeurs physiques (valeurs propres, valeurs

---

<sup>9</sup>Pour quelles raisons au fait? De vagues incantations de "correspondance"? Plutôt une croyance affirmée en un principe de conservation de l'impulsion (Feynman, sa maman et ses cubes). Mieux: l'évolution de la valeur moyenne de  $\mathbf{P}$  en vertu du théorème d'Ehrenfest (page 98). Encore mieux: le comportement classique de cette valeur moyenne au cours d'une transformation de Galilée, réf. [51].

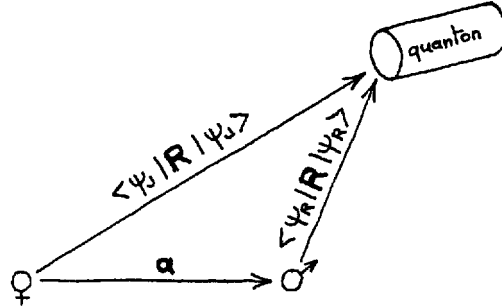


Figure IV.2: Juliette, Roméo et le quanton.

moyennes, ou plus généralement moments des distributions de ces valeurs) mais, indissolublement, les couples “opérateur-ket”. Nous avons donc l’embarras du choix pour répartir la transformation d’une valeur moyenne par exemple. Toute la responsabilité peut en être rejetée sur le vecteur d’état, l’opérateur restant alors impavide, c’est ce que j’ai fait ici, mais vous auriez pu, à l’inverse, fixer les kets d’état et modifier les opérateurs, ou encore répartir la transformation en proportions quelconques entre ceux-ci et ceux-là. La description de l’évolution temporelle d’un système quantique nous offre exactement les mêmes choix entre la représentation de Schrödinger, la représentation de Heisenberg, et les multiples représentations d’interaction.

Dernière remarque, dire que l’environnement d’un système est invariant par translation, c’est exprimer que la translation n’altère en rien la forme de l’équation d’évolution temporelle du système, et donc que le générateur correspondant, ou hamiltonien  $H$ , est représenté par le même opérateur pour Juliette et pour Roméo. Ceci implique en particulier  $\langle \psi_R | H | \psi_R \rangle = \langle \psi_J | H | \psi_J \rangle$ . En considérant le cas de la translation infinitésimale selon  $\hat{x}$ , on en déduit immédiatement  $[H, P_x] = 0$ . Comme promis, à une invariance de la nature correspond un générateur du système, et à la même invariance de l’environnement (triviale si le système est isolé) correspond une constante du mouvement, le générateur lui-même.

#### IV.3.4 Translations et champ

Venons en enfin au système isolé constitué par le rayonnement libre. Les choses sont un peu plus subtiles — il s’agit d’un champ et non plus d’un quanton — et il nous faut d’abord bien nous entendre sur les implications de la translation de Juliette à Roméo lorsqu’ils étudient un champ classique vectoriel par exemple,  $\mathbf{A}(\mathbf{r})$ , que ce soit un potentiel vecteur, un champ électrique ou l’elongation d’une corde de guitare.

Au point  $Q$  de l’espace est attaché une valeur du champ, disons  $\mathbf{v}$  (voir fig. IV.3).



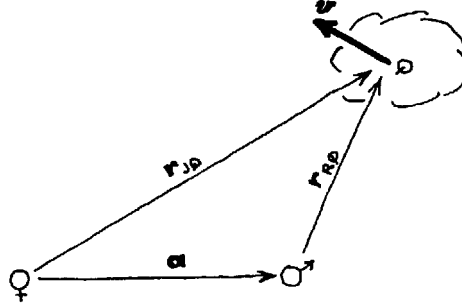


Figure IV.3: Juliette, Roméo et le champ.

Croire à l'invariance par translation de ce champ, c'est dire que Juliette et Roméo attribuent la même valeur  $\mathbf{v}$  au champ au point  $Q$ . Mais, pour des raisons pratiques évidentes, chacun de nos deux personnages préfère repérer le point  $Q$  par la distance de celui-ci à sa propre origine (sinon il n'y aurait pas translation des points de vue!). Il en résulte — bien que la grandeur  $\mathbf{v}$  soit la même — deux dépendances fonctionnelles distinctes que nous devons — une fois n'est pas coutume chez les physicien(ne)s<sup>10</sup> — noter différemment, tant et si bien que l'expression de cette invariance est  $\mathbf{v} = \mathcal{A}_R(\mathbf{r}_{RQ}) = \mathcal{A}_J(\mathbf{r}_{JQ})$ , soit  $\mathcal{A}_R(\mathbf{r}_{JQ} - \mathbf{a}) = \mathcal{A}_J(\mathbf{r}_{JQ})$ , ou encore:

$$\mathcal{A}_R(\mathbf{r}) = \mathcal{A}_J(\mathbf{r} + \mathbf{a}), \quad \text{pour tout } \mathbf{r}. \quad (\text{IV.11})$$

Il ne nous reste plus qu'à imposer la même loi de transformation aux valeurs moyennes de l'opérateur de champ, si nous ne voulons pas abandonner tout espoir que la théorie quantique du rayonnement puisse éventuellement admettre comme limite la théorie classique. Encore une fois, nous disposons d'une certaine latitude, entre les kets et les opérateurs, pour représenter la transformation. Mais les héros shakespeariens eux-mêmes doivent finir par prendre une décision, la moins coûteuse étant, bien sûr, identique à la précédente: les kets d'état de Juliette et de Roméo sont différents tandis que les opérateurs représentant les grandeurs physiques ne changent pas. La transposition quantique de (IV.11) est donc, dans ces conditions,

$$\langle \psi_R | \mathbf{A}(\mathbf{r}, t) | \psi_R \rangle = \langle \psi_J | \mathbf{A}(\mathbf{r} + \mathbf{a}, t) | \psi_J \rangle, \quad (\text{IV.12})$$

soit:

$$U^+(\mathbf{a}) \mathbf{A}(\mathbf{r}, t) U(\mathbf{a}) = \mathbf{A}(\mathbf{r} + \mathbf{a}, t).$$

<sup>10</sup>Pour une discussion des habitudes respectives et justifiées des mathématicien(ne)s et physicien(ne)s dans ce domaine, lire les remarques de Weyl [81, p. 85].

En se rapprochant, Juliette et Roméo sont à même d'utiliser l'expression (IV.8) et nous en arrivons tous et toutes à la conclusion:

$$[P_x, \mathbf{A}(\mathbf{r}, t)] = i\hbar \partial_x \mathbf{A}(\mathbf{r}, t). \quad (\text{IV.13})$$

Cette relation, caractéristique du générateur des translations, est identique à la relation (IV.7), caractéristique de l'impulsion, et les deux concepts ne font qu'un, comme nous l'avons bien anticipé par l'adoption d'une notation unique.

#### IV.3.5 Encore quelques considérations sur les translations

La présentation que j'ai donnée (plus exactement, vendue) des relations structurelles (IV.10) et (IV.13) caractérisant les espaces des états de deux modèles quantiques, le quanton et le rayonnement, semblerait remiser le traditionnel *principe de correspondance* «écrivons les équations classiques, remplaçons  $\mathbf{p}$  par  $(\hbar/i)\nabla$ ... et toutes ces sortes de choses», au rang de vestige archéologique. Il est certain que l'on parvient à une plus grande cohérence de l'exposé, et que, surtout, l'on épargne à notre théorie quantique un grave défaut en évitant de laisser sous-entendre qu'elle puisse être déduite de la mécanique classique. Mais il ne serait pas question de croire qu'une théorie s'élabore sans présupposés; les relations (IV.9) ou (IV.12), en attribuant implicitement la même structure à l'espace dans les théories classiques et quantiques, constituent un *super-principe de correspondance*, point de départ rationnel pour obtenir les relations (IV.10) et (IV.13).

La lectrice a peut-être déjà complété d'elle-même l'expression (IV.8) pour envisager le cas d'une translation infinitésimale  $d\mathbf{a}$  de direction quelconque. La représentation  $U(d\mathbf{a})$  dépend du paramètre vectoriel  $d\mathbf{a}$ . Dire que  $d\mathbf{a}$  est un vecteur, c'est sous-entendre un certain comportement canonique de ses composantes au cours des rotations de point de vue. Tout autant qu'à l'invariance par translation de nos descriptions, nous croyons à leur invariance par rotation. En l'occurrence  $U(d\mathbf{a})$  doit donc avoir un comportement de scalaire par rapport aux rotations, ce qui ne peut être réalisé qu'au moyen d'un générateur vectoriel  $\mathbf{P}$ , c'est à dire un triplet d'opérateurs dont les lois de transformations au cours d'une rotation sont celles des composantes d'un vecteur. Dans ces conditions,

$$U(d\mathbf{a}) = 1 + \frac{i}{\hbar} \mathbf{P} \cdot d\mathbf{a}$$

et

$$[P_i, A_j(\mathbf{r}, t)] = i\hbar \partial_i A_j(\mathbf{r}, t).$$

Une autre question plus subtile se pose peut-être à la lectrice cultivée, ou perspicace: «Alors que le signe dans la relation définitoire du générateur  $P_x$ , éq. (IV.8), est complètement arbitraire, comment se fait-il que seul le signe effectivement choisi conduise à une relation (IV.13) équivalente à (IV.7)? D'ailleurs, pourquoi prendre

des conventions différentes pour les translations spatiales et les translations temporelles? Je sais bien que l'équation de Schrödinger

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = H(t) |\psi(t)\rangle \quad (\text{IV.14})$$

implique un opérateur d'évolution infinitésimale

$$U(t + dt) = 1 - \frac{i}{\hbar} H(t) dt \quad (\text{IV.15})$$

où  $H(t)$  joue son rôle de générateur. Mais si je décide d'une expression analogue pour les translations:

$$U(\mathbf{\hat{x}} da) = 1 - \frac{i}{\hbar} P_x da,$$

je trouve:  $[P_x, \mathbf{A}] = -i\hbar \partial_x \mathbf{A}$ , irréconciliable avec (IV.7)! »

La raison en est dans l'invariance (bien cachée) de notre théorie quantique du rayonnement par rapport aux transformations de Lorentz, invariance qui remonte à l'origine formelle de la théorie dans les équations de Maxwell. Au cours d'une transformation de Lorentz, les intervalles de temps et d'espace obéissent à des lois de transformation canoniques, dites d'un *quadrivecteur*. L'énergie et l'impulsion totales du champ libre jouissent elles aussi, de par leurs origines, de ces propriétés de transformation. Exactement comme dans le cas des rotations, pour préserver l'invariance de Lorentz de notre représentation des translations spatio-temporelles nous ne devons faire appel qu'au scalaire (vis-à-vis de ces transformations)

$$\begin{aligned} \underline{P} \cdot d\underline{a} &= \left( \frac{H}{c}, \mathbf{P} \right) \cdot (c dt, d\mathbf{a}) \\ &= H dt - \mathbf{P} \cdot d\mathbf{a}. \end{aligned}$$

J'ai implicitement convenu d'un signe pour le générateur des translations temporelles, ou hamiltonien, lorsque j'ai écrit l'équation de Schrödinger sous la forme habituelle — éq. (IV.14) par exemple — qui conduit à l'expression (IV.15) pour la représentation d'une translation temporelle infinitésimale. L'invariance de Lorentz — encore une fois nullement évidente, mais bien réelle — de notre théorie impose donc d'écrire pour la représentation d'une translation d'espace-temps infinitésimale

$$\begin{aligned} U(d\underline{a}) &= 1 - \frac{i}{\hbar} \underline{P} \cdot d\underline{a} \\ &= 1 - \frac{i}{\hbar} H dt + \frac{i}{\hbar} \mathbf{P} \cdot d\mathbf{a}. \end{aligned}$$

Ainsi se trouve expliquée la mystérieuse exigence de cohérence qui était apparue pour imposer le choix de signe dans la définition (IV.8) du générateur des translations spatiales.

La lectrice opiniâtre qui se demandait pourquoi de telles questions ne se posent pas à propos du modèle du quanton aura maintenant répondu d'elle-même: «Elémentaire! Ce modèle n'est pas soumis aux rigueurs de l'invariance de Lorentz.

Translations de temps et translations d'espace n'ont à satisfaire aucune exigence commune.» Seul le conformisme nous conduit à choisir, généralement, le même signe devant les générateurs...

- dans l'opérateur d'évolution (IV.15),
- dans les représentations des translations

$$U(d\mathbf{a}) = 1 - \frac{i}{\hbar} \mathbf{P} \cdot d\mathbf{a},$$

- et même dans les représentations des rotations

$$U(d\boldsymbol{\omega}) = 1 - \frac{i}{\hbar} \mathbf{J} \cdot d\boldsymbol{\omega}.$$

Cette lectrice a tout à fait raison, encore que l'invariance galiléenne implique une forme générale (beaucoup mieux connue que l'implication elle-même) de  $H$  en fonction de  $\mathbf{R}$  et  $\mathbf{P}$ :

$$H(t) = \frac{1}{2m} \left( \mathbf{P} - \mathbf{A}(\mathbf{R}, t) \right)^2 + V(\mathbf{R}, t),$$

mais dans laquelle justement, le signe devant  $\mathbf{P}$  n'intervient qu'en relation avec le signe du potentiel vecteur  $\mathbf{A}$ , parfaitement arbitraire... mais ceci est une autre histoire.<sup>11</sup>

Profitons encore de l'exemple du quanton pour remarquer que la notion de position y intervient à deux titres différents:

- comme paramètre des translations qui séparent Juliette et Roméo,
- et comme grandeur physique du système, ce qui lui vaut d'être conjuguée — par la relation (IV.10) — du générateur des translations.

Le temps, en revanche, n'a que le seul rôle de paramètre; il n'est jamais une grandeur physique. Une transformation de Lorentz venant combiner intervalles de temps et d'espace, il est clair que ceux-ci doivent jouer dans la théorie le même rôle fonctionnel, ce qui interdit à la position d'être une grandeur physique du système. Il ne peut y avoir de grandeur physique position compatible avec l'invariance de Lorentz; il faut abandonner tout espoir d'élaborer un modèle du quanton relativiste. Il n'y a pas de particule quantique relativiste au sens où l'on aimerait l'entendre, c'est-à-dire où l'on pourrait suivre les évolutions au cours du temps d'une concentration spatiale d'énergie et d'impulsion. Dans un modèle quantique relativiste, temps et espace doivent se borner au même rôle de paramètres de repérage, il ne peut donc y avoir de tels modèles que de champs. C'est le cas de notre théorie quantique du rayonnement, prototype de théorie quantique d'un champ, qui semble bien s'obstiner à satisfaire l'invariance de Lorentz.

---

<sup>11</sup> Je ne connais guère d'autres références sur ce sujet que [38] et [51].

## IV.4 Le moment angulaire et les rotations

L'accumulation de renseignements maintenant obtenus concernant une impulsion du rayonnement quantique née — analogie purement formelle au départ — du rayonnement classique, laisse présager des considérations de même teneur en ce qui concerne le *moment angulaire* total du rayonnement, défini à partir de la densité d'impulsion correspondant aux opérateurs champs électrique et magnétique:

$$\mathbf{J}_{\text{ray}} \stackrel{\text{df}}{=} \varepsilon_0 \int d^3\mathbf{r} \, \mathbf{r} \wedge (\mathbf{E}(\mathbf{r}, t) \wedge \mathbf{B}(\mathbf{r}, t)).$$

Le caractère vectoriel de cette définition augure quelques complications de calculs que j'omettrai, me contentant de mentionner les nouveautés conceptuelles liées à cette grandeur physique.

On obtient une décomposition instructive de cette grandeur en exprimant le champ magnétique en fonction du potentiel vecteur  $\mathbf{B} = \nabla \wedge \mathbf{A}$ . De la propriété de transversalité du champ électrique de rayonnement  $\nabla \cdot \mathbf{E} = 0$ , on déduit alors la décomposition  $\mathbf{J}_{\text{ray}} = \mathbf{L}_{\text{ray}} + \mathbf{S}_{\text{ray}}$ , avec

$$\mathbf{L}_{\text{ray}} \stackrel{\text{df}}{=} \varepsilon_0 \int_{\mathcal{V}} d^3\mathbf{r} \sum_{j=1}^3 E_j \mathbf{r} \wedge \frac{1}{i} \nabla A_j, \quad (\text{IV.16})$$

$$\mathbf{S}_{\text{ray}} \stackrel{\text{df}}{=} \varepsilon_0 \int_{\mathcal{V}} d^3\mathbf{r} \, \mathbf{E} \wedge \mathbf{A}, \quad (\text{IV.17})$$

et, rappelons le,  $\mathbf{E} = -\partial_t \mathbf{A}$ .<sup>12</sup>

On conçoit que la définition de  $\mathbf{L}_{\text{ray}}$  puisse être transposée au cas d'un champ scalaire. Effectivement, appelons  $\varphi$  ce champ d'une théorie hypothétique, qui serait lui aussi doué d'énergie et d'impulsion; il n'est pas impossible, et c'est d'ailleurs le cas<sup>13</sup>, qu'il s'avère posséder une grandeur analogue

$$\mathbf{L}_{\text{ray}} \propto \varepsilon_0 \int_{\mathcal{V}} d^3\mathbf{r} \, (-\partial_t \varphi)(\mathbf{r} \wedge \nabla \varphi).$$

Par contre, le produit vectoriel  $(\partial_t \mathbf{A}) \wedge \mathbf{A}$  qui figure dans la définition de  $\mathbf{S}_{\text{ray}}$  est inconvertible au cas d'un champ scalaire. La contribution  $\mathbf{S}_{\text{ray}}$  au moment angulaire total du champ est intrinsèque — liée au caractère vectoriel de ce champ —, moyennant quoi nous appellerons naturellement  $\mathbf{L}_{\text{ray}}$  le *moment orbital du rayonnement* et  $\mathbf{S}_{\text{ray}}$  le *spin du rayonnement*.

<sup>12</sup>Les conditions aux limites interviennent évidemment dans les manipulations qui aboutissent à ces résultats et que vous trouverez détaillées dans [53, § XXI-23] ou [20, p. 33]. Le moment angulaire ne peut évidemment être une constante du mouvement dans un monde non invariant par rotation tel qu'une boîte parallélépipédique. Il faut recourir à une boîte sphérique, ou à un espace infini et un développement en modes continu. Ce problème ne se posait pas dans le cas de l'impulsion car une boîte parallélépipédique dont les dimensions croissent tend vers un monde invariant par rapport à toute translation.

<sup>13</sup>Voir par exemple [53, § XXI-7].

### IV.4.1 Représentations des rotations

Remarquons que, jusqu'à présent, rien dans notre démarche ne repose sur le fait que  $\mathbf{A}$ ,  $\mathbf{E}$  ou  $\mathbf{B}$  soient des opérateurs. La décomposition  $\mathbf{L}_{\text{ray}} + \mathbf{S}_{\text{ray}}$  s'applique aussi bien au rayonnement classique. La lectrice ingénue risque d'éprouver quelque surprise de se trouver confrontée à l'idée que le spin puisse être une grandeur classique! L'énigme se résout sans déception en pénétrant le monde merveilleux des relations entre le moment angulaire et les transformations de rotation. L'invariance par rotation de nos points de vue est encore une invariance de la nature à laquelle tout un chacun croit implicitement (tout au moins dans la communauté physicienne).<sup>14</sup>

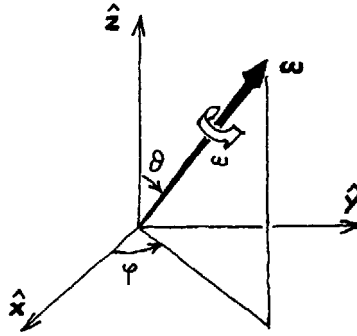
Nous savons que les composantes des vecteurs se transforment, au cours des rotations de points de vue, à l'aide de matrices orthogonales (donc unitaires, mais réelles) dont les coefficients dépendent de la rotation spécifique envisagée. Nous avons ainsi un moyen d'associer à toute rotation une matrice unitaire de rang trois. Ces matrices sont des opérateurs linéaires dans l'espace des vecteurs colonnes à trois lignes, et leur loi de multiplication est isomorphe de la loi de composition des rotations; elles en constituent une *représentation*. Les vecteurs engendrent une représentation unitaire de rang trois des rotations (qui bien entendu constituent un groupe).<sup>15</sup>

«Et alors, qu'y a-t-il de nouveau par rapport aux translations?» C'est que, pour celles-ci, les représentations unitaires (p. 100) engendrées par les kets d'état avaient un rang infini, car l'espace de Hilbert des états des systèmes que nous avons envisagés (quants et rayonnement quantique) est de dimension infinie. Ici, nous avons une représentation de rang fini — et même modéré — qui autorise toutes sortes de manipulations interdites aux représentations de rang infini, et qui jouit en conséquence de propriétés originales. Les rotations forment, elles aussi, un groupe de Lie: de deux rotations successives nous désirons obtenir une rotation dont les paramètres soient des fonctions analytiques des paramètres de celles-là.

«Est-ce à dire que l'on pourrait aussi trouver des représentations unitaires de rang fini pour le groupe des translations? Les vecteurs sont tellement commodes... »

<sup>14</sup>En fait, la notion d'invariance est paradoxalement beaucoup plus concrète, plus facile à expliquer au profane — si même elle ne lui paraît pas d'une trivialité indigne de son attention — que les constantes du mouvement qui peuvent s'ensuivre: les principes de conservation — pensez à l'énergie — sont terriblement abstraits et difficiles à faire admettre. Il en est ainsi parce que les propriétés de l'espace-temps des physiciennes ne sont que la formalisation de ce que tout un chacun ressent dans la vie... justement peu courante, car dès que le mouvement intervient, l'invariance galiléenne, elle, se heurte au “bon sens”; nous sommes tous des aristotéliens. Par contre, l'invariance par translation des points de vue, qui “tombe sous le sens”, n'est pas satisfaite dans certaines théories: si Juliette et Roméo n'ont pas le même horizon cosmologique, des parties de l'univers de celle-là peuvent être inobservables pour celui-ci.

<sup>15</sup>Si vous êtes curieuse de la chose, vous trouverez une élégante introduction à l'utilisation de la théorie des groupes en théorie quantique dans [53, app. D]. Les groupes qui nous intéressent le plus ici, à savoir les groupes continus, sont très bien exposés dans [33, chap. 8], avec (pour la plupart) démonstrations des théorèmes pertinents.

Figure IV.4: Les paramètres d'une rotation  $\omega$ .

Et bien non! D'abord cela se saurait, depuis le temps que l'on fait de la physique. Il n'y a pas d'objets jouant vis-à-vis des translations un rôle analogue à celui des vecteurs vis-à-vis des rotations. Il y a une distinction entre les deux groupes de transformations en question: alors que les paramètres des translations varient dans un domaine illimité — pour une translation  $\hat{x}a$ , le paramètre  $a$  peut prendre toute valeur dans l'intervalle ouvert  $]-\infty, \infty[$  — les paramètres des rotations eux sont restreints à des domaines finis et, surtout, fermés. Par exemple, pour une rotation paramétrée par un vecteur  $\omega$ , l'axe de rotation  $\hat{\omega}$  est repéré par la longitude  $\varphi$  qui varie de 0 à  $2\pi$  ( $180^\circ\text{E}$  à  $180^\circ\text{W}$  pour les marins), par la colatitude  $\vartheta$  entre 0 et  $\pi$  (ou la latitude de  $90^\circ\text{S}$  à  $90^\circ\text{N}$ ), tandis que l'angle de rotation propre  $\omega$  varie de 0 à  $\pi$  (figure IV.4). Il en va de même si la rotation est paramétrée par les angles d'Euler.

Entendons nous bien: quels qu'ils soient, les paramètres peuvent fort bien prendre des valeurs hors des intervalles précités, mais ils ne désignent pas alors de nouvelles rotations. Pour cette raison, les rotations sont qualifiées de *groupe compact* et admettent des représentations unitaires de rang fini. Comme pour tout théorème, la démonstration de cette assertion est, par définition, purement mathématique et ne présente aucun intérêt physique, argument péremptoire qui donne bonne conscience au physicien moyen pour ignorer cette démonstration d'existence... d'autant plus qu'elle est compliquée!

Les matrices associées aux rotations des vecteurs ont une propriété supplémentaire, elles constituent une *représentation irréductible*. Dans l'ensemble des rotations, les composantes de vecteurs se transforment linéairement de façon absolument conviviale, sans laisser à l'écart aucun membre de la famille: quelle que soit la composante, il existe évidemment une rotation qui la mélange à n'importe quelle autre. Deux sous-familles dont la cohésion se maintiendrait, envers et contre toute rotation, sans intrusion d'élément allogène, permettraient d'engendrer deux représentations

de rangs plus petits. Outre les vecteurs qui engendrent une représentation de rang trois pour les rotations, la lectrice connaît bien les scalaires qui, banalement, engendrent une représentation de rang un, ou les tenseurs cartésiens à deux indices qui engendrent une représentation unitaire de rang neuf, mais réductible car les neuf composantes peuvent être regroupées en sous-familles qui préservent leur intimité au cour des rotations:

- la trace (un scalaire),
- la partie antisymétrique (trois composantes vectorielles),
- et la partie symétrique de trace nulle (six composantes moins une condition) qui, engendrant une représentation unitaire et irréductible de rang cinq, est appelée pour cette raison tenseur irréductible de rang cinq.

Dans cette discussion, la lectrice aura peut-être ressenti, plus ou moins consciemment, la confusion terminologique propre à l'usage libéral du mot vecteur dans la langue vernaculaire des physiciens. Les vecteurs sont ici des triplets jouissant d'une certaine loi de transformation par rotation, et là des éléments d'un espace de Hilbert des états, de dimension infinie! Pour dissiper cette ambiguïté, il est bon d'avoir présent à l'esprit — à défaut d'employer — le volapük mathématique, et de réserver l'appellation de tenseur, ou de suite tensorielle, aux multiplets qui engendrent des représentations d'un groupe, ici les rotations, tandis que "vecteur" reste le terme générique désignant les éléments d'un espace vectoriel, sans préjuger de leur comportement au cours des transformations. Il faut également remarquer que cette notion de tenseur est relative à un certain groupe de transformations compact. Ainsi les vecteurs — en fait tenseurs de rang trois — par rapport au groupe des rotations, n'en sont pas par rapport au groupe des transformations de Lorentz, domaine où la notion correspondante est celle de tenseur de rang quatre, ou tenseur à un indice, ou (quadri)vecteur, puisque ce groupe contient le sous-groupe des rotations mais pratique l'échangisme familial entre les trois composantes spatiales et la composante temporelle.

Pour clore ce panorama superficiel, l'honnêteté m'oblige à ajouter que se contenter d'une simple allusion à la compacité d'un groupe de transformations — si elle a peut-être permis de faire sentir physiquement l'enjeu — laisse subsister une ambiguïté. D'un paramètre variant dans tout le domaine des nombres réels, on peut très bien, ne serait-ce qu'en en prenant la tangente hyperbolique, passer à une paramétrisation dans un domaine fini... sans que le groupe devienne pour cela compact. Il est essentiel que le paramètre prenne ses valeurs dans un intervalle fermé.

Après cette démonstration(!) de l'existence de représentations des rotations, unitaires, irréductibles de rang un, trois et cinq, il n'est pas difficile d'imaginer, au prix de complications croissantes, des procédés de construction pour des rangs impairs plus élevés, à partir des multiplets engendrant la représentation d'ordre trois. L'intérêt de la notion se manifeste dans l'abondante utilisation que font les physicien(ne)s, depuis la fin du XIX<sup>e</sup> siècle, des objets qui engendrent ces représentations,



pour écrire leurs équations sous une forme qui exhibe clairement leur invariance par rotation.

Mais la lectrice férue du modèle du quanton sait bien qu'existent aussi des représentations de rang pair. La fonction d'onde d'un quanton de spin  $1/2$ , pour ne citer que ce cas, a deux composantes et se transforme linéairement au cours d'une rotation. Elle engendre une représentation (irréductible) de rang deux en associant à toute rotation spatiale  $\boldsymbol{\omega}$ , la matrice de transformation des vecteurs de l'espace des états du spin:

$$\begin{aligned} e^{-\frac{i}{\hbar} \mathbf{S} \cdot \boldsymbol{\omega}} &= e^{-\frac{i}{2} \boldsymbol{\sigma} \cdot \boldsymbol{\omega}} \\ &= \cos \frac{\omega}{2} - i \boldsymbol{\sigma} \cdot \hat{\boldsymbol{\omega}} \sin \frac{\omega}{2}, \end{aligned}$$

où  $\sigma_x$ ,  $\sigma_y$  et  $\sigma_z$  sont les matrices de Pauli (page 24).<sup>16</sup> Ainsi, la valeur du spin  $s$  d'un objet (ici la fonction-d'onde) est une mesure du rang  $2s+1$  de la représentation irréductible des rotations que cet objet engendre. Dans ces conditions le spin n'a plus rien de spécifiquement quantique, et nous pouvons dire que notre champ vectoriel, qu'il soit classique ou quantique, a un spin 1.<sup>17</sup> Disons quand même, par précaution, et j'aurai l'occasion d'y revenir, que le caractère transverse du champ de rayonnement n'en fait pas un tenseur à part entière dans l'espace à trois dimensions (dans un mode  $\mathbf{k}$ , il est confiné au plan perpendiculaire à ce vecteur) et l'empêche de jouer pleinement son rôle de géniteur de représentations.

La lectrice concevra en revanche qu'à propos du rayonnement quantique on puisse répéter *mutatis mutandis* les arguments qui ont jalonné la discussion de la section précédente sur l'impulsion du rayonnement: le moment angulaire total du rayonnement commute avec l'hamiltonien du rayonnement, c'est une constante du mouvement du rayonnement libre, tandis qu'en interaction avec une matière quantique dont le moment angulaire est séparément conservé, c'est la somme des moments angulaires du rayonnement et de la matière qui est une constante du mouvement.

#### IV.4.2 Polarisation circulaire

Plus intéressantes vont être les considérations liées à la polarisation du rayonnement. Nous avons jusqu'à présent utilisé effectivement comme vecteurs de base de polarisation d'un mode  $\mathbf{k}$ , les vecteurs de polarisation rectiligne  $\hat{\mathbf{e}}_1$  et  $\hat{\mathbf{e}}_2$  qui, pour cause

<sup>16</sup>Cette représentation est unitaire, mais maintenant complexe, car elle agit sur des objets, les deux composantes du vecteur d'état du spin  $1/2$ , qui sont complexes.

<sup>17</sup>Sous ce rapport, la seule originalité de la physique quantique — cela pourrait presque en constituer une définition — est de ne pouvoir se passer des représentations de rang pair pour décrire la Nature, alors que les représentations impaires suffisent à la physique classique (même si les représentations de rang pair peuvent y être parfois utiles). On peut lire l'histoire de la découverte, par Hamilton, des représentations de rang deux des rotations dans [57, § 41.1].

de transversalité, devaient être orthogonaux à  $\mathbf{k}$  et, par raison de commodité, constituaient avec ce dernier un trièdre orthogonal direct. Mais nous pouvons essayer de profiter de la latitude, que nous nous étions réservée, d'utiliser des vecteurs de base complexes, et recourir plutôt aux vecteurs de base

$$\hat{\mathbf{e}}_{\pm} \stackrel{\text{df}}{=} \mp \frac{1}{\sqrt{2}} (\hat{\mathbf{e}}_1 \pm i\hat{\mathbf{e}}_2).$$

La lectrice clairvoyante aura reconnu les expressions des vecteurs de polarisation circulaire de l'optique ondulatoire classique; mais que signifient-ils pour nous? Toujours dans le mode  $\mathbf{k}$  (sous-entendu), ce sont des vecteurs de base tout aussi légitimes, auxquels seraient associés, en rayonnement classique de nouvelles amplitudes et, en rayonnement quantique, de nouveaux opérateurs d'annihilation, le tout décrivant le même champ et satisfaisant donc la condition définitoire:

$$a_+ \hat{\mathbf{e}}_+ + a_- \hat{\mathbf{e}}_- \stackrel{\text{df}}{=} a_1 \hat{\mathbf{e}}_1 + a_2 \hat{\mathbf{e}}_2.$$

On en déduit (Exercice 3) l'expression de ces nouveaux opérateurs en fonction des opérateurs de photons polarisés rectilignes,

$$a_{\pm} = \mp \frac{1}{\sqrt{2}} (a_1 \mp ia_2),$$

et les expressions conjuguées pour les opérateurs de création:

$$a_{\pm}^+ = \mp \frac{1}{\sqrt{2}} (a_1^+ \pm ia_2^+).$$

L'action de ces derniers sur le vide de photons crée des états

$$|1\rangle_{\mathbf{k}\pm} \stackrel{\text{df}}{=} a_{\mathbf{k}\pm}^+ |0\rangle, \quad (\text{IV.18})$$

dits respectivement à un photon *droit* (+), ou à un photon *gauche* (−). Il ne s'agit pour l'instant que d'une dénomination motivée par (sinon fondée sur) l'analogie formelle classique et dont la signification physique finira par se révéler (§ V.8). De plus, les conventions concernant la gauche et la droite diffèrent entre physicien(ne)s des “particules” et opticien(ne)s... ce qui facilite remarquablement les erreurs! Quoi qu'il en soit, un calcul dont la complication me dispense de le détailler ici, permet de trouver des résultats particulièrement simples pour l'action sur ces états des composantes du moment orbital (IV.16) et du spin (IV.17) du rayonnement selon l'impulsion du photon:

$$\begin{aligned} \mathbf{k} \cdot \mathbf{L}_{\text{ray}} |1\rangle_{\mathbf{k}\pm} &= 0 \\ \mathbf{k} \cdot \mathbf{S}_{\text{ray}} |1\rangle_{\mathbf{k}\pm} &= \pm \hbar k |1\rangle_{\mathbf{k}\pm}. \end{aligned}$$

*exercice,  
voir dans  
Lurié*

Encore une fois, le photon conspire pour se parer des plumes du quanton. Etat du rayonnement, d'impulsion  $\hbar \mathbf{k}$ , il s'arrange pour que, de façon bien ordinaire, la

composante du moment orbital sur cette impulsion soit nulle tandis que la composante du spin prend des valeurs caractéristiques d'un quanton de spin 1. On en déduit d'ailleurs, pour la composante  $J_{\mathbf{k}} \stackrel{\text{df}}{=} \hat{\mathbf{k}} \cdot \mathbf{J}_{\text{ray}}$  du moment angulaire le long de l'impulsion:

$$J_{\mathbf{k}}|1\rangle_{\mathbf{k}\pm} = \pm\hbar|1\rangle_{\mathbf{k}\pm}. \quad (\text{IV.19})$$

Chose curieuse, cette composante ne semble avoir que deux valeurs propres. J'entends déjà la lectrice: «Ne s'agit-il pas tout simplement d'un spin 1/2 déguisé? Nous avons peut-être oublié un facteur deux quelque part dans la définition du spin du champ ce qui rétablirait des valeurs propres plus orthodoxes pour  $J_{\mathbf{k}}$ . Il est d'ailleurs déjà arrivé que des facteurs deux surgissent mystérieusement à propos du spin 1/2, c'est le facteur  $g$  de l'électron.» Et bien non, il n'y a pas d'erreur ou de définition mal choisie. Aucun facteur ne manque à l'appel, et le photon a bien un spin 1, comme il sied à l'unité d'excitation d'un rayonnement vectoriel. Il semble que la nature s'interdise bien le troisième état correspondant à une valeur propre nulle pour  $J_{\mathbf{k}}$ , et ceci à cause de la condition de transversalité qui ne nous laisse que deux modes de polarisation indépendants.

Examinons l'effet d'une rotation d'axe  $\hat{\mathbf{k}}$  et d'angle  $\omega$  sur les états propres de  $J_{\mathbf{k}}$ . En vertu de l'équation aux valeurs propres (IV.19), on a

$$e^{-\frac{i}{\hbar}J_{\mathbf{k}}\omega}|1\rangle_{\mathbf{k}\pm} = e^{\mp i\omega}|1\rangle_{\mathbf{k}\pm}.$$

Ces états sont donc modifiés par des facteurs de phase qui, pour être simples, ne sont pas triviaux, car ils sont différents. Une combinaison linéaire de ces états subit donc, elle, une véritable transformation. Par contre, pour l'hypothétique état manquant au tableau du spin 1, un état à un photon, état propre de  $J_{\mathbf{k}}$  pour la valeur propre zéro, on aurait

$$e^{-\frac{i}{\hbar}J_{\mathbf{k}}\omega}|1\rangle_{\mathbf{k}0} = |1\rangle_{\mathbf{k}0},$$

et un tel état serait donc réellement invariant par rotation. Les seuls vecteurs de l'espace vital qui soient invariants par rotation d'axe  $\hat{\mathbf{k}}$  étant parallèles à cet axe, l'état  $|1\rangle_{\mathbf{k}0}$  ne pourrait correspondre qu'à l'excitation d'un mode de rayonnement parallèle à  $\mathbf{k}$ , ou longitudinal, précisément interdit par la condition de transversalité. L'absence de photon longitudinal est donc liée à la transversalité du rayonnement. Nous verrons plus tard que cette transversalité est à associer à la nullité de la masse du photon, au sens de la relation de dispersion (III.10). Un "rayonnement" vectoriel massif (comme celui dont les bosons intermédiaires  $W_{\pm}$  et  $Z_0$  sont les "photons") admet parfaitement des états longitudinaux. Inversement, il ne faudra pas être surpris qu'un champ longitudinal comme le champ électrique coulombien puisse être associé aux modes longitudinaux. Enfin, ultime retour en arrière, la masse nulle du photon sera attribuée à l'invariance de jauge, origine de toute cette épopée.

Les états à un photon  $|1\rangle_{\mathbf{k}\pm}$  définis par (IV.18) sont états propres de la composante  $J_{\mathbf{k}}$  du moment angulaire. Il importe de remarquer qu'ils sont états propres de l'impulsion, mais pas de l'opérateur  $\mathbf{J}_{\text{ray}}^2$ , ni même d'une composante de  $\mathbf{J}_{\text{ray}}$  sur un axe indépendant d'un mode particulier. Il est en fait possible de construire des états à un photon, états propres de  $\mathbf{P}_{\text{ray}}^2$  (ou de l'énergie  $H_{\text{ray}}$ ),  $\mathbf{J}_{\text{ray}}^2$ ,  $(\mathbf{J}_{\text{ray}})_z$  et de la réflexion, ensemble de grandeurs physiques compatibles, et en nombre suffisant pour étiqueter de façon unique un état par leurs valeurs propres. De tels états ne peuvent être créés par des opérateurs associés à de nouveaux états de base de polarisation. Il est clair qu'il faut aussi de nouvelles fonctions modales de base, douées de dépendances spatiales en  $\mathbf{r}$  différentes de  $e^{i\mathbf{k}\cdot\mathbf{r}}$ . Il faut pour cela recourir à une forme de boîte différente, bien évidemment sphérique, mais qui interdit l'usage d'une condition de périodicité sur la paroi qui nous limiterait par trop à des fonctions de base isotropes. La seule condition utilisable est celle de nullité, satisfaisante pour la description de phénomènes stationnaires mais qui — pour décrire des phénomènes de propagation pendant un temps raisonnable avant de se heurter aux parois — exigera le passage final à la limite d'un rayon infini. On peut ainsi obtenir — en lieu et place du développement (III.9) de l'opérateur de champ sur les modes transverses d'une boîte parallélépipédique — un développement multipolaire sur les modes transverses d'une boîte sphérique<sup>18</sup>,

$$\mathbf{A}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sum_{kjm\pi} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \left\{ a_{kjm\pi} e^{-i\omega t} \mathbf{Y}_{kjm\pi}(\mathbf{r}) + a_{kjm\pi}^+ e^{i\omega t} \mathbf{Y}_{kjm\pi}^*(\mathbf{r}) \right\}. \quad (\text{IV.20})$$

Cette représentation n'implique aucune nouveauté conceptuelle en ce qui concerne notre théorie quantique du rayonnement, sinon une complication évidente. Elle est en revanche un outil essentiel dans l'étude du système matière-rayonnement, lorsque cette matière est constituée par les électrons d'un atome ou par les nucléons d'un noyau, dont l'hamiltonien non perturbé — sans couplage avec le rayonnement — est invariant par rotation. Les états stationnaires du système non perturbé sont alors états propres du moment angulaire (ce ne serait pas le cas avec les électrons d'un cristal par exemple), et le calcul des taux de transition entre ces états stationnaires par la théorie des perturbations est considérablement simplifié par les règles de sélection, évidentes, des éléments de matrices des opérateurs multipolaires du développement (IV.20). Les exemples les plus courants que je traiterai ne nécessiteront pas cet arsenal.

## Exercices

### 1. Rafraîchissement quantique: l'inégalité de Heisenberg.

---

<sup>18</sup>Voir [53, § XXI-29] ou [20, p. 47]

i) Soit  $A$  un opérateur hermitique. Qu'en est-il de  $A^2$ ? Que peut-on dire du signe de la valeur moyenne  $\langle A^2 \rangle_\psi \stackrel{\text{df}}{=} \langle \psi | A^2 | \psi \rangle$  de l'opérateur  $A^2$  dans l'état  $|\psi\rangle$ ?

ii) Soit deux opérateurs hermitiques  $A$  et  $B$ , un état  $|\psi\rangle$ , et un nombre réel  $\lambda$ . Calculer la norme de  $|\varphi\rangle \stackrel{\text{df}}{=} (i\lambda A + B)|\psi\rangle$ , et en déduire la relation entre valeurs moyennes dans l'état  $|\psi\rangle$ :

$$\langle A^2 \rangle_\psi \langle B^2 \rangle_\psi \geq \frac{1}{4} \left\langle \frac{1}{i} [A, B] \right\rangle_\psi^2.$$

iii) A la grandeur représentée par l'opérateur  $A$ , et à l'état  $|\psi\rangle$ , on associe la dispersion quantique définie par son carré:  $(\Delta A)_\psi^2 \stackrel{\text{df}}{=} \langle (A - \langle A \rangle_\psi)^2 \rangle_\psi$ . Soit l'opérateur  $\hat{A}_\psi \stackrel{\text{df}}{=} A - \langle A \rangle_\psi$ , et  $\hat{B}_\psi$  son analogue. Calculer le commutateur  $[\hat{A}, \hat{B}]$ . Montrer que

$$(\Delta A)_\psi (\Delta B)_\psi \geq \frac{1}{2} \left| \left\langle \frac{1}{i} [A, B] \right\rangle_\psi \right|.$$

## 2. Champ moyen.

i) Calculer l'expression de l'opérateur champ électrique moyen  $\bar{\mathbf{E}}(t)$  au sens de Gauss (page 90).

ii) Montrer que le carré de celui-ci a pour valeur moyenne dans le vide

$$\langle 0 | \bar{\mathbf{E}}(t) \cdot \bar{\mathbf{E}}(t) | 0 \rangle = \frac{\hbar c}{\varepsilon_0} \frac{1}{V} \sum_{\mathbf{k}} k e^{-k^2 b^2}.$$

iii) Calculer cette valeur dans la limite d'une grande boîte.

3. En fonction des vecteurs de base de polarisation rectiligne dans un mode  $\mathbf{k}$ , on définit les vecteurs (complexes!)  $\hat{\mathbf{e}}_\sigma$ ,  $\sigma \in \{+, -\}$ :

$$\hat{\mathbf{e}}_\pm \stackrel{\text{df}}{=} \mp \frac{1}{\sqrt{2}} (\hat{\mathbf{e}}_1 \pm i \hat{\mathbf{e}}_2).$$

i) Vérifier que ces vecteurs sont orthonormés (au sens d'Hermite), c'est-à-dire que  $\hat{\mathbf{e}}_\sigma^* \cdot \hat{\mathbf{e}}_{\sigma'} = \delta_{\sigma\sigma'}$ .

ii) On peut choisir de développer l'opérateur de champ  $\mathbf{A}(\mathbf{r}, t)$  sur cette nouvelle base. Associés à ces nouveaux vecteurs de base, on a donc de nouveaux opérateurs d'annihilation  $a_\sigma$  tels que, par définition,

$$\sum_{\sigma=+}^{-} a_\sigma \hat{\mathbf{e}}_\sigma \stackrel{\text{df}}{=} \sum_{T=1}^2 a_T \hat{\mathbf{e}}_T.$$

En déduire les expressions de  $a_+$  et  $a_-$  en fonction de  $a_1$  et  $a_2$ .

iii) Montrer, à l'aide des commutateurs des  $a_T$  et  $a_T^+$ , que  $[a_\sigma, a_{\sigma'}^+] = \delta_{\sigma\sigma'}$ , et  $[a_\sigma, a_{\sigma'}] = 0$ .

## Chapitre V

# Quelques états du rayonnement quantique intéressants

Après avoir étudié des grandeurs physiques du rayonnement quantique telles que les opérateurs de champs  $\mathbf{A}(\mathbf{r}, t)$ ,  $\mathbf{E}(\mathbf{r}, t)$  et les opérateurs hamiltonien  $H_{\text{ray}}$ , ou impulsion  $\mathbf{P}_{\text{ray}}$ , nous allons examiner quelques kets particuliers de l'espace des états du rayonnement quantique libre.

Dans ce chapitre, nous nous placerons généralement, et sauf avis contraire, dans l'espace de Fock des états d'un seul mode polarisé rectiligne, dont les paramètres caractéristiques,  $\mathbf{k}$  et  $\hat{\mathbf{e}}_{\mathbf{k}T}$ , seront généralement sous-entendus. Cet espace sera donc sous-tendu par les kets propres du nombre de photons dans le mode,  $|n\rangle$  mis pour  $|n\rangle_{\mathbf{k}T}$ . Il suffit pour cela d'imaginer que nous nous restreignons au sous-espace des états du rayonnement correspondant au vide de photons dans tous les autres modes que  $\mathbf{k}T$ . Les opérateurs effectifs dans ce sous-espace seront obtenus en tirant, sans sommation (sur les modes), la contribution pertinente du développement modal de l'opérateur complet; ainsi, par exemple:

$$\begin{aligned} N &\stackrel{\text{df}}{=} N_{\mathbf{k}T}, \\ E(\mathbf{r}, t) &\stackrel{\text{df}}{=} i\sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} \left\{ a e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} - a^\dagger e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}, \\ &\text{etc.} \end{aligned} \tag{V.1}$$

Ce faisant, nous allégeons évidemment l'écriture, mais cette opération simple dissimule quand même l'omission des contributions infinies que les opérateurs complets donnent en général dans l'ensemble des modes vides!

L'étude de la polarisation circulaire sera une exception à cette simplification. Nous serons alors conduits à envisager simultanément deux polarisations rectilignes de base associées au mode  $\mathbf{k}$ , orthogonales, notées  $T = 1$  et  $T = 2$ . D'autre part, la lectrice concevra sans peine que l'on puisse ensuite compliquer les choses à loisir en étendant les résultats trouvés à tout l'espace des états avec des nombres de photons non nuls dans les différents modes, et même à des ensembles statistiques impurs en adjoignant à notre arsenal formel un opérateur densité du rayonnement quantique.

Ces considérations sur les états du rayonnement vont nous apporter des éclaircissements supplémentaires, et d'ailleurs nécessaires, sur la signification de nos prétendues grandeurs physiques. En effet, si nous avons maintenant quelque idée des rôles assurés par l'hamiltonien et par l'impulsion du rayonnement (générateurs et constantes du mouvement), notre compréhension des opérateurs de champ, elle, est encore très indirecte. Ces derniers interviennent, bien sûr, dans l'interaction entre matière et rayonnement (III.11), dans la définition de l'énergie du rayonnement (III.14) qui conduit à l'hamiltonien, *etc.* Mais méritent-ils leurs appellations de potentiel vecteur, champ électrique, *etc.* autrement que par une analogie formelle de définition avec leurs homonymes classiques?

Que le problème se pose est évident si l'on songe que les seuls états considérés pour l'instant — les états  $|n\rangle$  qui précisément définissent en le sous-tendant l'espace des états — correspondent à une valeur déterminée de l'énergie du rayonnement,

$$H_{\text{ray}}|n\rangle = \hbar\omega\left(n + \frac{1}{2}\right)|n\rangle,$$

alors qu'ils se révèlent impuissants à donner des valeurs, ne serait-ce que moyennes, à notre prétendu champ électrique:  $\langle n|E(\mathbf{r}, t)|n\rangle = 0$ . Cela ne signifie pas pour autant que le champ électrique soit un opérateur nul. D'ailleurs — nous l'avons vu (page 89), et le reverrons (page 123) — sa dispersion dans ces états n'est pas nulle. La difficulté est circonscrite lorsqu'elle est nommée. L'énergie  $H_{\text{ray}}$  (ou le nombre de photons  $N$ ) et le champ électrique  $E(\mathbf{r}, t)$  sont deux grandeurs incompatibles (on disait complémentaires), car leurs opérateurs ne commutent pas:

$$[N, E(\mathbf{r}, t)] = -i\sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} \left\{ ae^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} + a^\dagger e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}.$$

Pour s'assurer de la signification de  $E(\mathbf{r}, t)$  il va donc nous falloir trouver ses états propres — qui ne peuvent être des états à nombre de photons déterminé — ou, à tout le moins, des états dans lesquels on puisse espérer que la valeur moyenne de  $E$  présente, ne serait-ce qu'à la limite, un comportement qui rappelle celui de son éponyme classique. A mode donné — spécifié par un vecteur d'onde  $\mathbf{k}$  et un vecteur de polarisation  $\hat{\mathbf{e}}_{\mathbf{k}T}$  —, ce comportement ne peut être qu'une onde plane (monochromatique, polarisée rectiligne) du type  $\hat{\mathbf{e}}_{\mathbf{k}T} E_0 \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \varphi)$ . Ainsi, en adoptant une écriture plus symétrique, notre programme est la recherche d'un

état, noté  $|?\rangle$  pour l'instant, tel que l'opérateur champ électrique effectif

$$E(\mathbf{r}, t) = i \left\{ \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} a e^{-i(\omega t - \mathbf{k}\cdot\mathbf{r})} - \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} a^+ e^{i(\omega t - \mathbf{k}\cdot\mathbf{r})} \right\}, \quad (\text{V.2})$$

ait une valeur moyenne de la forme

$$\langle ? | E(\mathbf{r}, t) | ? \rangle = i \left\{ \frac{E_0}{2} e^{i\varphi} e^{-i(\omega t - \mathbf{k}\cdot\mathbf{r})} - \frac{E_0}{2} e^{-i\varphi} e^{i(\omega t - \mathbf{k}\cdot\mathbf{r})} \right\}, \quad (\text{V.3})$$

et une dispersion  $(\triangle E(\mathbf{r}, t))_?$  nulle, ne serait-ce qu'à la(?) limite.

Dans l'expression de la valeur moyenne apparaissent les grandeurs physiques classiques amplitude  $E_0$  et phase  $\varphi$ . On va donc commencer par chercher, dans l'expression de l'opérateur champ effectif, des opérateurs qui joueraient, ne serait-ce qu'en valeurs moyennes, ces rôles d'amplitude et de phase.

## V.1 Opérateurs amplitude et phase

Comparant les expressions de l'opérateur champ électrique et de sa valeur moyenne désirée, on peut tenter de poser

$$a = A e^{i\Phi}, \quad (\text{V.4})$$

où les opérateurs  $A$  et  $\Phi$ , hermitiques comme il se doit pour des grandeurs physiques qui classiquement ont des valeurs réelles, seraient les opérateurs amplitude (à un facteur numérique près) et phase cherchés. Mais cette équation définit-elle effectivement les opérateurs  $A$  et  $\Phi$ ? Elle se conjugue en tout cas sans difficulté,  $a^+ = e^{-i\Phi} A$ , pour donner:

$$\begin{aligned} aa^+ &= A e^{i\Phi} e^{-i\Phi} A, \\ a^+ a + 1 &= A^2, \end{aligned}$$

soit enfin  $A^2 = N + 1$ . L'opérateur  $N + 1$  est défini positif (son spectre est l'ensemble des entiers positifs), sa racine carrée existe:

$$\sqrt{N + 1} \stackrel{\text{df}}{=} \sum_{n=0}^{\infty} |n\rangle \sqrt{n + 1} \langle n|. \quad (\text{V.5})$$

L'opérateur  $A$  est donc parfaitement défini:

$$A \stackrel{\text{df}}{=} \sqrt{N + 1}, \quad (\text{V.6})$$



et nous l'appellerons dorénavant, à toute fin pratique, l'*opérateur amplitude* (effectif, dans le mode  $\mathbf{k}T$ ). Cet opérateur n'ayant pas de valeur propre nulle, il a un inverse:

$$A^{-1} = (N+1)^{-1/2} = \sum_{n=0}^{\infty} |n\rangle (n+1)^{-1/2} \langle n|.$$

L'opérateur  $e^{i\Phi} = A^{-1}a = (N+1)^{-1/2}a$ , est donc, lui aussi, défini. Remarquez que si l'on avait, par malchance, posé au départ  $a = e^{i\Phi}A$ , les choses ne se seraient pas si bien arrangées; on aurait certes  $A = \sqrt{N}$ , mais pas d'inverse.

Avons-nous maintenant bien défini un opérateur  $\Phi = -i \ln((N+1)^{-1/2}a)$ ? Certainement pas: il n'est pas hermitique (en contradiction avec les manipulations effectuées pour en arriver là), et on peut impunément lui ajouter un multiple quelconque de  $2\pi$ . Force nous est donc, par prudence, d'en rester à l'opérateur  $e^{i\Phi}$  ou, pour user d'une notation moins compromettante, à l'opérateur  $F$ , comme (facteur de) phase,

$$F \stackrel{\text{df}}{=} (N+1)^{-1/2}a, \quad (\text{V.7})$$

autrement dit:  $a = \sqrt{N+1}F$ .

Quelles sont les propriétés de cet *opérateur phase*  $F$ ? Son adjoint est facile à calculer car,  $N$  étant hermitique,

$$F^+ = a^+(N+1)^{-1/2}.$$

En vertu du commutateur de  $a$  et  $a^+$ , on a donc

$$\begin{aligned} FF^+ &= (N+1)^{-1/2}aa^+(N+1)^{-1/2} \\ &= (N+1)^{-1/2}(N+1)(N+1)^{-1/2} \\ &= 1, \end{aligned} \quad (\text{V.8})$$

en accord avec la substitution  $e^{i\Phi}e^{-i\Phi} = 1$ , que nous nous étions permise plus haut pour arriver à l'expression de l'opérateur  $A$ . Mais serait-ce à dire que  $F$  est unitaire? On a en fait  $F^+F = a^+(N+1)^{-1}a$ , et en particulier:

$$\langle 0|F^+F|0\rangle = 0. \quad (\text{V.9})$$

Le produit  $F^+F$  ne peut donc être l'opérateur unité et nous sommes en présence de cette espèce rare, dans un domaine quantique aussi élémentaire que l'oscillateur harmonique, d'un opérateur  $F^+$  isométrique — il transforme un système orthonormé en un autre système orthonormé — mais non unitaire.<sup>1</sup> L'action de  $F$  sur les états de base est facile à calculer,

$$F|n\rangle = \begin{cases} |n-1\rangle & \text{si } n \neq 0 \\ 0 & \text{si } n = 0 \end{cases}, \quad (\text{V.10})$$

---

<sup>1</sup>On rencontre des opérateurs isométriques beaucoup plus sophistiqués — y compris au sens français de l'adjectif — en théorie quantique de la diffusion.

ainsi que celle de son adjoint,

$$F^+|n\rangle = |n+1\rangle. \quad (\text{V.11})$$

L'opérateur  $F$  n'est donc pas hermitique. De par sa définition (V.7), il a peu de chances de commuter avec l'opérateur nombre de photons. La lectrice trouvera d'ailleurs, d'elle-même, que

$$\begin{aligned} [N, F] &= -F, \\ [N, F^+] &= F^+. \end{aligned} \quad (\text{V.12})$$

L'opérateur  $F$  ne commute pas avec  $N$ , et n'est ni hermitique, ni unitaire. Que lui reste-t-il donc pour nous plaire?

Puisque  $F$  est construit sur une réminiscence de  $e^{i\varphi}$ , étudions — par analogie avec  $\cos \varphi$  et  $\sin \varphi$  — les opérateurs

$$\begin{aligned} C &\stackrel{\text{df}}{=} \frac{1}{2}(F + F^+), \\ S &\stackrel{\text{df}}{=} \frac{1}{2i}(F - F^+). \end{aligned} \quad (\text{V.13})$$

Ces opérateurs sont hermitiques et ont donc plus de chances d'être pertinents, encore que cela ne soit pas nécessaire pour ceci.<sup>2</sup> Ils ne commutent pas. Vous pouvez vous assurer qu'il en va ainsi parce que  $F$  n'est pas un *opérateur normal*, autrement dit  $[F, F^+] \neq 0$  (Exercice 1). Mais nos “cosinus” et “sinus” de la “phase” satisfont néanmoins d'élégantes relations de commutation avec l'opérateur nombre de photons:

$$\begin{aligned} [N, C] &= -iS, \\ [N, S] &= iC. \end{aligned} \quad (\text{V.14})$$

Celles-ci nous valent des inégalités de Heisenberg corrélatant les dispersions et valeurs moyennes de ces grandeurs avec la dispersion de  $N$  (lui-même apparenté au carré de l'amplitude). Dans un état  $|\psi\rangle$ , on a:

$$\begin{aligned} (\Delta N)_\psi (\Delta C)_\psi &\geq \frac{1}{2} |\langle S \rangle_\psi|, \\ (\Delta N)_\psi (\Delta S)_\psi &\geq \frac{1}{2} |\langle C \rangle_\psi|. \end{aligned} \quad (\text{V.15})$$

## V.2 Les états de phase

Les opérateurs  $C$  et  $S$  n'admettent pas de système complets de vecteurs propres communs, car ils ne commutent pas. Et pourtant nous allons trouver des états

---

<sup>2</sup>La lectrice inquiétée par la prétendue nécessité de l'hermiticité, et intriguée par les angles et les phases en théorie quantique, trouvera rassurante, voire divertissante, la lecture de [46].

qui s'en approchent étrangement. Etant donnés les paramètres  $\varphi$ , réel, et  $s$ , entier positif, définissons le ket d'état

$$|\varphi, s\rangle \stackrel{\text{df}}{=} \frac{1}{\sqrt{s+1}} \sum_{n=0}^s e^{in\varphi} |n\rangle. \quad (\text{V.16})$$

La lectrice vérifiera d'abord sans peine que ce ket est normé à l'unité (Exercice 2). Calculant ensuite la valeur moyenne de  $F$ , on trouve que

$$\langle \varphi, s | F | \varphi, s \rangle = \frac{s}{s+1} e^{i\varphi},$$

et donc

$$\langle F \rangle_{\varphi, s} \xrightarrow{s \rightarrow \infty} e^{i\varphi}, \quad (\text{V.17})$$

surprise certes agréable, peut-être simple coïncidence.

Mais encore plus heureux se révèle le calcul de la valeur moyenne,

$$\langle FF^+ \rangle_{\varphi, s} = \frac{s}{s+1}$$

qui, poussée à la limite

$$\langle FF^+ \rangle_{\varphi, s} \xrightarrow{s \rightarrow \infty} 1, \quad (\text{V.18})$$

a de quoi surprendre car, rappelons-le,  $FF^+$  n'est pas l'opérateur unité. Ce n'est qu'en prenant un nombre infini de composantes dans la définition de  $|\varphi, s\rangle$  que la contribution (V.9) du fondamental — insignifiante mais intempestive lorsqu'on voudrait faire de  $F$  un opérateur unitaire — parvient à passer inaperçue. L'opérateur  $F$  n'en est bien entendu pas plus unitaire pour cela; les états  $|\varphi, s\rangle$  ne jouissent pas de l'indépendance que leur conférerait l'orthogonalité et, surtout, ne constituent pas une suite de Cauchy (ils ne tendent vers aucune limite lorsque  $s$  croît indéfiniment).

La limite (V.17) a pour conséquence, maintenant triviale,

$$\begin{aligned} \langle C \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \cos \varphi, \\ \langle S \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \sin \varphi, \end{aligned} \quad (\text{V.19})$$

tandis que de (V.18) et ses analogues (Exercice 3), on déduit

$$\begin{aligned} \langle C^2 \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \cos^2 \varphi, \\ \langle S^2 \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \sin^2 \varphi, \end{aligned} \quad (\text{V.20})$$

d'où cette propriété, inattendue, des dispersions de nos “cosinus” et “sinus” dans les états  $|\varphi, s\rangle$ :

$$\lim_{s \rightarrow \infty} (\Delta C)_{\varphi, s} = \lim_{s \rightarrow \infty} (\Delta S)_{\varphi, s} = 0. \quad (\text{V.21})$$

La contradiction apparente de ce résultat avec les inégalités de Heisenberg (V.15) n'est rien de plus qu'un paradoxe car, chaque état  $|n\rangle$  pesant du même poids dans la définition de  $|\varphi, s\rangle$ , la limite de la dispersion  $(\Delta N)_{\varphi, s}$  ne peut être, elle, qu'infinie (Exercice 4).

En vertu de (V.19) et (V.21), les kets  $|\varphi, s\rangle$  pour  $s$  assez grand se trouvent constituer quasiment des *états de phase*. Cette phase est d'autant mieux déterminée qu'en fait tous les moments (au delà du premier) des distributions des valeurs de  $C$  et  $S$  associées à l'état  $|\varphi, s\rangle$  ont une limite nulle. Mais les états  $|\varphi, s\rangle$ , pas plus que leur inexistante limite, ne sont des états propres de  $C$  et de  $S$  (Exercice 5).

### V.3 Propriétés des états de nombre de photons

Les états  $|n\rangle$ , états propres du nombre de photons  $N$ , sont bien sûr états propres de l'amplitude  $\sqrt{N+1}$  dont les dispersions, dans ces conditions, sont nulles. Mais qu'en est-il des opérateurs de phase?

Faisant jouer nos relations (V.10) et (V.11), nous obtenons facilement (Exercice 6) les valeurs moyennes dans un état  $|n\rangle$ ,

$$\begin{aligned}\langle C \rangle_n &= \langle S \rangle_n = 0, \\ \langle C^2 \rangle_n &= \langle S^2 \rangle_n = \frac{1}{2} - \frac{1}{4}\delta_{n0},\end{aligned}\tag{V.22}$$

d'où l'on déduit immédiatement les dispersions:

$$(\Delta C)_n = (\Delta S)_n = \begin{cases} \frac{1}{\sqrt{2}} & \text{si } n \neq 0, \\ \frac{1}{2} & \text{si } n = 0. \end{cases}\tag{V.23}$$

Ces valeurs anodines, qui n'ont rien d'énorme au premier coup d'œil, sont en fait considérables pour des grandeurs dont on pouvait espérer le spectre limité à  $[-1, 1]$ , à l'instar de cosinus ou sinus dignes de ce nom. Pour comparaison, l'écart-type d'une distribution équiprobable sur  $[-1, 1]$  vaut  $1/\sqrt{3}$  (Exercice 7). La distribution de probabilité des valeurs de la phase  $\varphi$  correspondant à un état  $|n\rangle$  est donc certainement très étalée et quasi équiprobable (Exercice 8).

En ce qui concerne l'opérateur champ électrique effectif dans le mode, défini en (V.2), nous avons évidemment

$$\langle E(\mathbf{r}, t) \rangle_n = 0,$$

décourageant tout espoir qu'un état  $|n\rangle$  ait quelque chose à voir avec un état quasi classique dans lequel le champ électrique aurait — exigence minimale — une valeur moyenne en forme d'onde. Mais l'histoire n'est pas close car la dispersion dans l'état  $|n\rangle$  vaut (Exercice 9):

$$(\Delta E(\mathbf{r}, t))_n = \sqrt{\frac{\hbar\omega}{\varepsilon_0\mathcal{V}}} \sqrt{n + \frac{1}{2}}.\tag{V.24}$$

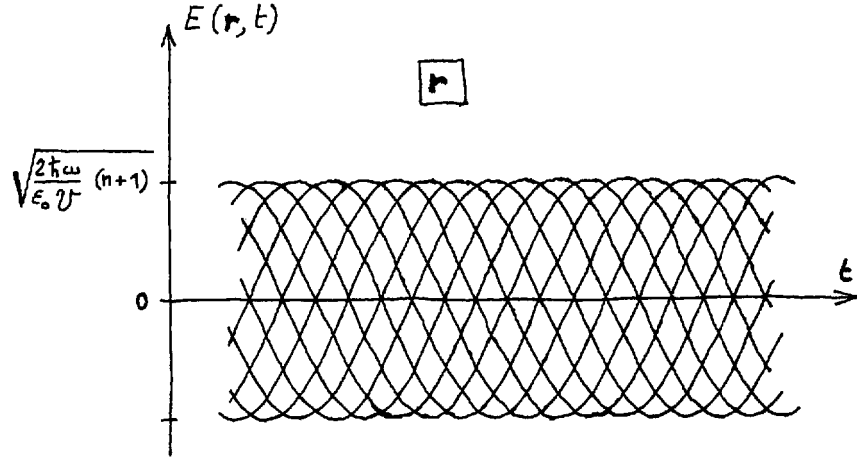


Figure V.1: Dans un état de nombre de photons, quelques unes des équiprobables valeurs du champ électrique en un point; une amplitude, des phases...

Nous pouvons interpréter qualitativement ces valeurs moyennes et dispersion du champ électrique en nous souvenant que, dans un état  $|n\rangle$ , l'amplitude a une valeur parfaitement définie,  $\sqrt{n+1}$ , tandis que la phase n'a qu'une distribution de valeurs étalée. Or, le champ électrique effectif (V.2), réécrit sous la forme

$$E(\mathbf{r}, t) = \sqrt{\frac{2\hbar\omega}{\varepsilon_0 V}} \frac{a^+ e^{i(\omega t - \mathbf{k} \cdot \mathbf{r})} - a e^{-i(\omega t - \mathbf{k} \cdot \mathbf{r})}}{2i},$$

suggère l'expression, très symbolique:

$$E(\mathbf{r}, t) \text{ " = " } \sqrt{\frac{2\hbar\omega}{\varepsilon_0 V}} \times \text{"ampl."} \times \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \text{"phase"}). \quad (\text{V.25})$$

Puisque dans un état  $|n\rangle$  l'amplitude a la valeur précise  $\sqrt{n+1}$ , tandis que la phase n'a qu'un histogramme désespérément plat, toutes les valeurs de phase sont démocratiquement représentées et il faut s'attendre à obtenir également un histogramme plat pour  $E(\mathbf{r}, t)$ , dans lequel toutes ces valeurs contribuent également. C'est ce que l'on peut tenter de représenter, en un lieu  $\mathbf{r}$  donné, par le schéma de la figure V.1<sup>3</sup>. Au vu d'un tel fouillis on conçoit que, *hic et nunc*, en  $\mathbf{r}$  et  $t$ , la valeur moyenne du champ électrique puisse être nulle et la largeur de sa distribution de l'ordre de  $\sqrt{n\hbar\omega/\varepsilon_0 V}$ .

<sup>3</sup>Calquée sur [52, p. 146].

## V.4 Propriétés des états de phase

Nous avons vu que les états  $|\varphi, s\rangle$  parviennent, tout au moins à la limite, à doter les opérateurs de phase de valeurs quasi déterminées:

$$\begin{aligned}\langle C \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \cos \varphi, \\ \langle S \rangle_{\varphi, s} &\xrightarrow{s \rightarrow \infty} \sin \varphi, \\ (\Delta C)_{\varphi, s} \text{ et } (\Delta S)_{\varphi, s} &\xrightarrow{s \rightarrow \infty} 0.\end{aligned}$$

En revanche, et à cause des inégalités de Heisenberg (V.15), il faut s'attendre à ce que la dispersion de  $N$  dans ces états ait une limite infinie. Qu'en est-il au juste? A partir de la définition (V.16), on trouve notamment (Exercice 4):

$$\begin{aligned}\langle N \rangle_{\varphi, s} &= \frac{s}{2}, \\ (\Delta N)_{\varphi, s}^2 &= \frac{s(s+2)}{12},\end{aligned}\tag{V.26}$$

expressions qui croissent indéfiniment avec  $s$ . Aux états de phase  $|\varphi, s\rangle$  qui, en dotant celle-ci d'une valeur précise, méritent bien leur nom, est en revanche associée une distribution de valeurs d'amplitude extrêmement étalée puisque de largeur infinie.

Alors que dans un état  $|n\rangle$  la valeur moyenne du champ est nulle, nous trouvons dans les états de phase une moyenne de l'"amplitude" qui croît indéfiniment. Partant de l'expression (V.25) souhaitée pour le champ électrique, ce comportement peut être illustré par la figure V.2. Au vu de celle-ci, on réalise très bien que la moyenne quantique du champ électrique vaille  $\pm\infty$ , selon le lieu et l'instant considérés. Les états de phase sont prétextes à effets graphiques de moiré, mais restent par ailleurs irréalisables: à une moyenne du champ électrique, ou du nombre de photons d'énergie  $\hbar\omega$ , correspond une énergie moyenne infinie qui rend ces états par trop dispendieux pour nos pauvres moyens, en ces temps de crise de l'énergie<sup>4</sup>. Ces états sont de toute façon bien loin de donner au champ électrique des valeurs approchant l'onde souhaitée, et nous en sommes toujours...

## V.5 A la recherche d'états quasi classiques

Ni les états  $|n\rangle$ , ni les états  $|\varphi, s\rangle$  ne réussissent à donner une valeur moyenne  $\langle E(\mathbf{r}, t) \rangle$  qui ressemble, ne serait-ce que de loin, au champ électrique classique d'un mode polarisé rectiligne,  $E_0 \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \varphi)$ . Et pourtant, ce cas idéal est fort souvent en rapport avec des situations réelles. Notre théorie quantique du rayonnement gagnerait donc en vraisemblance si nous parvenions à trouver des états  $|\varphi\rangle$

<sup>4</sup>...ou plutôt de crise locale de l'entropie. L'énergie, elle, est toujours conservée!

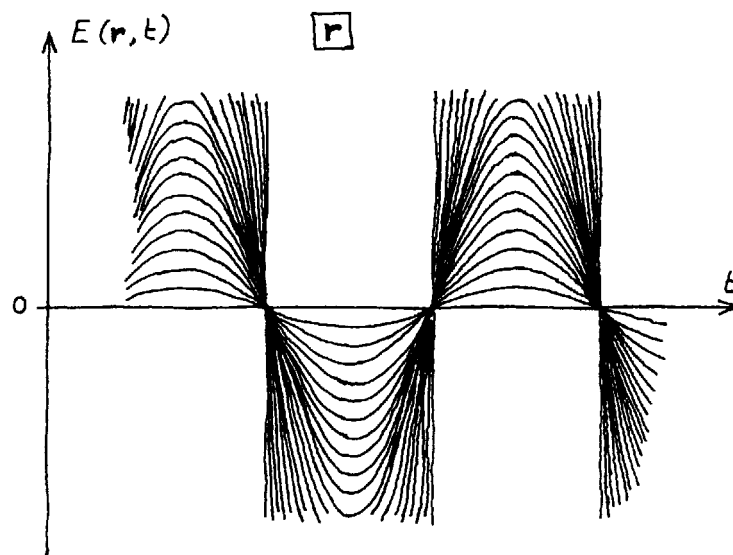


Figure V.2: Dans un état de phase, quelques unes des équiprobables valeurs du champ électrique en un point; une phase, des amplitudes...

dans lesquels amplitude et phase seraient déterminées au mieux, autrement dit, des états tels que les dispersions  $(\Delta N)_?$ ,  $(\Delta C)_?$  et  $(\Delta S)_?$  soient minimales.

### V.5.1 Les états cohérents

Réexaminons encore les expressions (V.2) et (V.3) de l'opérateur champ électrique et de sa valeur désirée avec la meilleure précision possible. Tout irait bien si les opérateurs  $a$  et  $a^+$  prenaient simultanément des valeurs précises, complexes conjuguées, or c'est chose impossible à des opérateurs qui ne commutent pas. Mais exploitons néanmoins cette idée et, faute de mieux, intéressons nous aux états propres de  $a$ .

Les valeurs propres de  $a$  sont *a priori* complexes car  $a$  n'est pas hermitique. Soit  $z$  l'une de ces valeurs propres et  $|z\rangle$  le ket propre correspondant. Autrement dit,  $a|z\rangle = z|z\rangle$ . Déterminer le ket  $|z\rangle$ , c'est trouver son développement sur une base de l'espace des états, par exemple l'ensemble des  $|n\rangle$ , kets propres du nombre de photons  $N = a^+a$ . Introduisant ce développement,  $|z\rangle = \sum_{n=0}^{\infty} c_n |n\rangle$ , dans l'équation aux valeurs propres, on a nécessairement:

$$\sum_{n=1}^{\infty} c_n \sqrt{n} |n-1\rangle = z \sum_{n=0}^{\infty} c_n |n\rangle.$$

Nous en déduisons la récurrence  $c_n = (z/\sqrt{n})c_{n-1}$ , pour  $n > 0$ , qui se résout explicitement en  $c_n = (z^n/\sqrt{n!})c_0$ . Les kets propres de  $a$  sont donc de la forme:

$$|z\rangle = c_0 \sum_{n=0}^{\infty} \frac{z^n}{\sqrt{n!}} |n\rangle.$$

La norme d'un ket d'état n'ayant aucune signification physique, on préfère, pour des raisons purement pratiques, utiliser des kets normés à l'unité. On doit donc avoir:

$$1 = \langle z|z\rangle = |c_0|^2 \sum_{n=0}^{\infty} \frac{|z|^{2n}}{n!} = |c_0|^2 e^{|z|^2}.$$

La phase d'un ket d'état n'ayant pas plus de signification, l'argument du coefficient  $c_0$  est quelconque. Convenons, comme tout le monde, de choisir  $c_0$  réel positif. Nous obtenons ainsi l'expression du ket propre de  $a$  associé à la valeur propre  $z$ :

$$|z\rangle = e^{-\frac{1}{2}|z|^2} \sum_{n=0}^{\infty} \frac{z^n}{\sqrt{n!}} |n\rangle. \quad (\text{V.27})$$

Ces états propres de l'opérateur d'annihilation sont appelés *états cohérents*. Imaginés par Schrödinger [65], dans les tout débuts de la théorie quantique de l'oscillateur harmonique, c'est leur redécouverte à propos du rayonnement qui, par référence à l'optique, leur a valu leur nom.



### V.5.2 Propriétés algébriques

La littérature consacrée aux états cohérents est considérable, et la lectrice qui ne se satisferait pas des propriétés essentielles que je vais passer en revue pourra consulter les références [19, Compl. G<sub>V</sub>] ou, à un niveau plus élevé, [20] et [40]. Par ailleurs, les revues publient continuellement des articles consacrés aux états cohérents — ou quasi-classiques, ou d’“incertitude” minimale — de systèmes autres que l’oscillateur harmonique dont les opérateurs  $a$  et  $a^+$  ont l’algèbre.

A tout nombre complexe  $z$ , on peut donc associer l’état cohérent  $|z\rangle$  donné par son développement (V.27) sur les kets de base de l’espace de Fock. Première propriété — la propriété définitoire —, le ket  $|z\rangle$  est ket propre de l’opérateur d’annihilation  $a$ , pour la valeur propre  $z$ :

$$a|z\rangle = z|z\rangle.$$

Cette relation s’avérera commode tant pour certaines démonstrations que pour les calculs de valeurs moyennes de grandeurs physiques.

Autre propriété — un tantinet inhabituelle encore qu’il ne faille pas s’en offusquer de la part de kets propres d’un opérateur non hermitique —, les états cohérents ne sont pas orthogonaux. A l’aide du développement (V.27), on trouve en effet:

$$\begin{aligned}\langle z'|z\rangle &= e^{-\frac{1}{2}(|z'|^2+|z|^2)} e^{z'^*z} = e^{-\frac{1}{2}(|z'-z|^2-z'^*z+z'z^*)} \\ &= e^{i\Im(z'^*z)} e^{-\frac{1}{2}|z'-z|^2},\end{aligned}$$

et donc, en particulier,

$$|\langle z'|z\rangle|^2 = e^{-|z'-z|^2}.$$

Le même développement (V.27) toujours, nous donne directement la probabilité de trouver l’état à  $n$  photons dans l’état cohérent:

$$|\langle n|z\rangle|^2 = \frac{|z|^{2n}}{n!} e^{-|z|^2}. \quad (\text{V.28})$$

Nul doute, les ichtyophiles hument là une distribution de leur met favori.

Enfin, l’état cohérent  $|z\rangle$  s’obtient aussi par l’action d’un opérateur idoine sur le vide  $|0\rangle$ . Pour cela, souvenons nous que les états de base  $|n\rangle$  peuvent être édifiés à partir du vide par actions successives de l’opérateur de création  $a^+$ :

$$\begin{aligned}|1\rangle &= a^+ |0\rangle, \\ |2\rangle &= \frac{a^+}{\sqrt{2}} |1\rangle = \frac{(a^+)^2}{\sqrt{2!}} |0\rangle, \\ &\dots \\ |n\rangle &= \frac{a^+}{\sqrt{n}} |n-1\rangle = \frac{(a^+)^n}{\sqrt{n!}} |0\rangle.\end{aligned}$$

Le développement (V.27) peut donc aussi bien s'écrire:

$$\begin{aligned} |z\rangle &= e^{-\frac{1}{2}|z|^2} \sum_{n=0}^{\infty} \frac{(za^+)^n}{n!} |0\rangle, \\ &= e^{-\frac{1}{2}|z|^2} e^{za^+} |0\rangle. \end{aligned}$$

Mais ceci peut se mettre sous une forme plus élégante. Grâce à la nullité résultant de l'action de  $a$  sur l'état vide, on a:

$$\begin{aligned} |z\rangle &= e^{-\frac{1}{2}|z|^2} e^{za^+} \left\{ 1 + \sum_{n=1}^{\infty} \frac{(-z^*a)^n}{n!} \right\} |0\rangle, \\ &= e^{-\frac{1}{2}|z|^2} e^{za^+} e^{-z^*a} |0\rangle. \end{aligned}$$

Comme les opérateurs  $a$  et  $a^+$  commutent avec leur commutateur (qui est un nombre), la formule de Campbell, Baker, Hausdorf, Glauber, *etc.* (Exercice 10) permet d'écrire

$$e^{za^+} e^{-z^*a} = e^{za^+ - z^*a + \frac{1}{2}[za^+, -z^*a]} = e^{za^+ - z^*a + \frac{1}{2}|z|^2},$$

et l'on a donc enfin:

$$|z\rangle = D(z) |0\rangle,$$

avec

$$D(z) \stackrel{\text{df}}{=} e^{za^+ - z^*a}. \quad (\text{V.29})$$

L'opérateur déplacement  $D(z)$  a quelques propriétés remarquables, faciles à établir (Exercice 11),

$$\begin{aligned} D(0) &= 1, \\ D^+(z) &= D(-z), \\ D(z')D(z) &= e^{i\Im(z'z^*)} D(z' + z), \\ D(z)D(-z) &= 1 = D(-z)D(z), \end{aligned}$$

qui permettent en particulier de montrer que

$$D(z)D^+(z) = D^+(z)D(z) = 1;$$

autrement dit, l'opérateur déplacement est unitaire.

### V.5.3 Propriétés physiques

Un état cohérent  $|z\rangle$  étant état propre de  $a$ , on a, par conjugaison,  $\langle z|a^+ = z^*\langle z|$ . Le calcul de la valeur moyenne de l'opérateur champ électrique effectif (V.2) dans cet état est alors immédiat, soit

$$\langle z|E(\mathbf{r}, t)|z\rangle = i \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} (ze^{-i(\omega t - \mathbf{k}\cdot\mathbf{r})} - z^*e^{i(\omega t - \mathbf{k}\cdot\mathbf{r})}),$$

ou encore

$$\langle E(\mathbf{r}, t) \rangle_z = \sqrt{\frac{2\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z| \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z), \quad (\text{V.30})$$

donc l'expression en onde plane tant espérée. Mais il n'y a pas encore lieu de triompher car, après tout, n'importe quel état du rayonnement donne à l'opérateur (V.2) une valeur moyenne de cette forme (peut-être nulle, cas des états  $|n\rangle$ , ou infinie, cas limite des états de phase  $|\varphi, s\rangle$ ), puisque les valeurs moyennes de  $a$  et de  $a^+$  sont toujours conjuguées!

Il nous reste donc à aller voir si le champ électrique est bien déterminé autour de sa valeur moyenne et, pour cela, à calculer les dispersions de l'amplitude et de la phase. On a, tout d'abord,

$$\langle N \rangle_z = \langle z | a^+ a | z \rangle = |z|^2,$$

et

$$\langle N^2 \rangle_z = \langle z | a^+ a a^+ a | z \rangle = \langle z | a^+ (1 + a^+ a) a | z \rangle = |z|^4 + |z|^2,$$

d'où la variance du nombre de photons,

$$\begin{aligned} (\Delta N)_z^2 &\stackrel{\text{df}}{=} \langle z | (N - \langle N | z \rangle)^2 | z \rangle, \\ &= \langle N^2 \rangle_z - \langle N \rangle_z^2, \\ &= |z|^2, \end{aligned}$$

et la dispersion,  $(\Delta N)_z = |z|$ . Remarquez au passage que l'on a bien  $(\Delta N)_z^2 = \langle N \rangle_z$ , comme il se devait pour la distribution pisciforme (V.28). La dispersion du nombre de photons relative à sa valeur moyenne,

$$\frac{(\Delta N)_z}{\langle N \rangle_z} = \frac{1}{|z|} = \frac{1}{\sqrt{\langle N \rangle_z}},$$

tend vers zéro lorsque le paramètre  $z$  de l'état cohérent croît. Ceci est encourageant pour l'amplitude qui est fonction de  $N$ . L'état cohérent n'est pas état propre de  $N$  (la dispersion de  $N$  n'est pas nulle) ni de l'amplitude, mais celle-ci se trouvera quand même de mieux en mieux déterminée lorsque l'état représentera un nombre moyen de photons, ou une amplitude moyenne, élevés.

Et qu'en est-il de la phase? Plutôt que l'opérateur  $F$ , utilisons les opérateurs  $S$  et  $C$  qui, étant hermitiques, présentent l'avantage pratique de conduire à des inégalités de Heisenberg. A défaut d'expressions explicites, il est néanmoins possible d'obtenir les développements asymptotiques, pour  $|z|$  grand:

$$\begin{aligned} \langle S \rangle_z &= \sin(\arg z) \left\{ 1 - \frac{1}{8|z|^2} + \dots \right\}, \\ \langle S^2 \rangle_z &= \sin^2(\arg z) - \frac{\sin^2(\arg z) - \frac{1}{2}}{2|z|^2} + \dots, \end{aligned}$$

et des développements analogues pour l'opérateur  $C$  en remplaçant les sinus par des cosinus. Je renonce à vous en présenter les démonstrations, mais les fanatiques d'analyse pourront nourrir leur passion dans l'exercice **13**. Quoi qu'il en soit, nous avons donc

$$\begin{aligned}\langle S \rangle_z &\underset{|z| \rightarrow \infty}{\sim} \sin(\arg z), \\ \langle \Delta S \rangle_z &\underset{|z| \rightarrow \infty}{\sim} \frac{|\cos(\arg z)|}{2|z|},\end{aligned}$$

et analogues pour l'opérateur  $C$ . Ainsi, les états cohérents à nombre moyen de photons élevé déterminent de plus en plus précisément la phase autour de la valeur moyenne  $\arg z$ . Avec ce résultat, les états cohérents commencent à devenir réellement intéressants: la dispersion de la phase (ou plus exactement de ses "sinus" et "cosinus"), comme la dispersion relative de l'amplitude, tendent vers zéro lorsque l'état correspond à un nombre moyen de photons  $|z|^2$  croissant.<sup>5</sup>

Mais encore mieux, les produits des dispersions trouvées pour  $N$ ,  $S$  et  $C$  satisfont les relations:

$$\begin{aligned}(\Delta N)_z(\Delta S)_z &\underset{|z| \rightarrow \infty}{\sim} \frac{1}{2} |\langle C \rangle_z|, \\ (\Delta N)_z(\Delta C)_z &\underset{|z| \rightarrow \infty}{\sim} \frac{1}{2} |\langle S \rangle_z|.\end{aligned}$$

Les états cohérents à  $|z| \gg 1$  saturent donc les inégalités de Heisenberg (V.15). Non seulement les dispersions de l'amplitude et de la phase dans ces états sont contrôlées, mais ce sont des dispersions minimales compatibles avec les inégalités de Heisenberg. Remarquons qu'un contournement de la première inégalité, par exemple, en choisissant la valeur  $\arg z = \pi/2$  serait illusoire; on se retrouverait alors avec une contrainte maximale pour le produit des dispersions  $(\Delta N)_z(\Delta S)_z$  dans la deuxième inégalité. On ne pouvait espérer mieux, même si l'on peut trouver autre avec les *états comprimés* (voir § V.9), en jouant sur la répartition de la borne de Heisenberg entre les dispersions du nombre de photons et de la phase.

Nonobstant ces considérations plutôt formelles, il est temps d'en revenir aux conséquences plus concrètes (!) de ces idéales propriétés des états cohérents en ce qui concerne le champ électrique lui-même. Nous en avons déjà calculé la valeur moyenne (V.30). Qu'en est-il de sa dispersion? Nous avons, d'après (V.2),

$$\begin{aligned}E^2(\mathbf{r}, t) &= -\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}} \left\{ (a^2 e^{-2i(\omega t - \mathbf{k} \cdot \mathbf{r})} + a^{+2} e^{2i(\omega t - \mathbf{k} \cdot \mathbf{r})}) - aa^+ - a^+a \right\}, \\ &= \frac{\hbar\omega}{\varepsilon_0\mathcal{V}} \left\{ N + \frac{1}{2} - \frac{1}{2} (a^2 e^{-2i(\omega t - \mathbf{k} \cdot \mathbf{r})} + a^{+2} e^{2i(\omega t - \mathbf{k} \cdot \mathbf{r})}) \right\},\end{aligned}$$

---

<sup>5</sup>Les états cohérents doivent leur qualificatif à cette propriété, analogue de la cohérence en optique classique.

d'où, toujours grâce à l'équation aux valeurs propres,  $a|z\rangle = z|z\rangle$ , et à sa conjuguée,

$$\begin{aligned}\langle E^2(\mathbf{r}, t) \rangle_z &= \frac{\hbar\omega}{\varepsilon_0\mathcal{V}} (|z|^2 + \frac{1}{2} - |z|^2 \cos 2(\omega t - \mathbf{k} \cdot \mathbf{r})), \\ &= \frac{\hbar\omega}{2\varepsilon_0\mathcal{V}} + \frac{2\hbar\omega}{\varepsilon_0\mathcal{V}} |z|^2 \sin^2(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z), \\ &= \frac{\hbar\omega}{2\varepsilon_0\mathcal{V}} + \langle E(\mathbf{r}, t) \rangle_z^2,\end{aligned}$$

et enfin la dispersion,

$$(\Delta E(\mathbf{r}, t))_z = \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}},$$

indépendante du lieu,  $\mathbf{r}$ , de l'instant,  $t$ , et même du paramètre  $z$  de l'état cohérent!

En fin de compte, l'état cohérent  $|z\rangle$  assure au champ électrique une valeur moyenne en forme d'onde plane  $E_0 \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \varphi)$ , d'amplitude

$$E_0 = \sqrt{\frac{2\hbar\omega}{\varepsilon_0\mathcal{V}}} |z| = \sqrt{\frac{2\hbar\omega}{\varepsilon_0\mathcal{V}}} \langle N \rangle_z,$$

et de phase  $\varphi = \arg z$ , valeurs affectées de dispersions certes, mais la dispersion du champ électrique lui-même vaut

$$\Delta E = \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}},$$

et en valeur relative:

$$\frac{\Delta E}{E_0} = \frac{1}{2|z|} = \frac{1}{2\sqrt{\langle N \rangle_z}}.$$

Nous sommes enfin parvenus au but visé. Avec les états cohérents  $|z\rangle$ , lorsque le module de  $z$  croît, l'amplitude du champ moyen croît également, mais la dispersion du champ reste constante. La grandeur  $E(\mathbf{r}, t)$  est donc, relativement, de mieux en mieux définie, la distribution de ses valeurs est de plus en plus "piquée" autour d'une valeur moyenne qui évolue comme un champ électrique de rayonnement classique monochromatique, plan, polarisé rectiligne.

La figure V.3 représente la moyenne de l'opérateur champ électrique en un point donné, dans divers états cohérents  $|z\rangle$  correspondant à la même phase,  $\arg z$ , mais à des nombres moyens,  $|z|^2$ , de 4, 36 et 100 photons respectivement. Un écart de une dispersion à la moyenne se traduit par un décalage indépendant du nombre moyen de photons, et les bandes de largeur  $\pm \Delta E$  autour de la valeur moyenne illustrent bien la précision relative croissante du champ lorsque son amplitude moyenne augmente. Remarquons enfin que, contrairement à ce que l'on peut souvent lire, il ne suffit pas d'être dans un état à grand nombre de photons pour assister à un comportement quasi-classique. Ce dernier ne peut s'obtenir qu'au prix d'une judicieuse

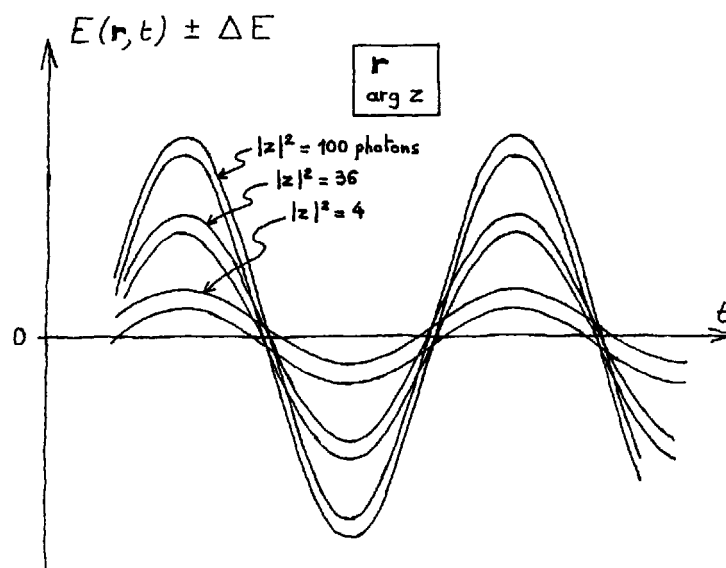


Figure V.3: Le champ électrique moyen en un lieu  $r$  donné, dans des états cohérents de même  $\varphi = \arg z$ , correspondant à différents nombres moyens de photons  $|z|^2$ . La demi-épaisseur en ordonnée de chaque courbe est égale à une dispersion.

superposition d'états correspondants à différents nombres de photons qui, il est vrai, déterminent d'autant mieux ce nombre que sa valeur moyenne est élevée.

Récapitulons enfin les résultats de cette section en faisant réapparaître explicitement les indices des modes,  $\mathbf{k}T$ . Lorsque le rayonnement se trouve dans un état cohérent du mode  $\mathbf{k}T$ , ou état propre de l'opérateur d'annihilation  $a_{\mathbf{k}T}$ , vide dans tous les autres modes, soit

$$|z\rangle_{\mathbf{k}T} = D_{\mathbf{k}T}(z) |0\rangle,$$

avec

$$D_{\mathbf{k}T}(z) = \exp(z a_{\mathbf{k}T}^+ - z^* a_{\mathbf{k}T}),$$

l'opérateur champ électrique complet (III.12) prend, lorsque  $|z| \gg 1$ , des valeurs bien concentrées autour de la valeur moyenne:

$$\mathbf{k}T \langle z | \mathbf{E}(\mathbf{r}, t) | z \rangle_{\mathbf{k}T} = \hat{\mathbf{e}}_{\mathbf{k}T} \sqrt{\frac{2\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z| \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z). \quad (\text{V.31})$$

Il en va évidemment de même pour l'opérateur de champ (III.9),

$$\mathbf{k}T \langle z | \mathbf{A}(\mathbf{r}, t) | z \rangle_{\mathbf{k}T} = \hat{\mathbf{e}}_{\mathbf{k}T} \sqrt{\frac{2\hbar}{\varepsilon_0 \omega \mathcal{V}}} |z| \cos(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z),$$

et pour l'opérateur champ magnétique (III.13):

$$\begin{aligned} \mathbf{k}T \langle z | \mathbf{B}(\mathbf{r}, t) | z \rangle_{\mathbf{k}T} &= \mathbf{k} \wedge \hat{\mathbf{e}}_{\mathbf{k}T} \sqrt{\frac{2\hbar}{\varepsilon_0 \omega \mathcal{V}}} |z| \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z), \\ &= \frac{\hat{\mathbf{k}}}{c} \wedge \mathbf{k}T \langle z | \mathbf{E}(\mathbf{r}, t) | z \rangle_{\mathbf{k}T}. \end{aligned}$$

Ces expressions étant d'inoubliables solutions des équations de Maxwell de l'électrodynamique classique, l'état  $|z\rangle_{\mathbf{k}T}$  présente, à  $|z|$  élevé, toutes les propriétés d'une onde électromagnétique classique de caractéristiques  $\mathbf{k}$  et  $\hat{\mathbf{e}}_{\mathbf{k}T}$ . Ainsi, les appellations de champs électrique et magnétique du rayonnement quantique, proposées (page 62) pour les opérateurs  $\mathbf{E}(\mathbf{r}, t)$  et  $\mathbf{B}(\mathbf{r}, t)$ , n'étaient pas si mal choisies.

#### V.5.4 La fabrication des états cohérents

Rendue à ce stade, la lectrice qui m'a accompagné ne se pose probablement plus qu'une (?) question: «Bien. La théorie accomode des états du rayonnement quantique indistincts d'ondes électromagnétiques classiques. Mais notre théorie contient-elle pour autant les équations de Maxwell de l'électrodynamique classique? En d'autres termes, comment peut-on fabriquer ces états cohérents? Mieux même pour ma tranquillité d'esprit: ces états cohérents sont-ils bien créés par des sources

classiques, puisque, après tout, je contemple jour après jour, ou écoute à longueur de nuit, des ondes électromagnétiques créées par des sources classiques.»

Notre “théorie” quantique du rayonnement décrit l'évolution d'un ket d'état d'un espace produit {états du quanton}  $\otimes$  {états du rayonnement quantique}, évolution régie par une équation du type (III.28), en représentation intermédiaire par exemple. Férue de modèle du quanton, vous savez qu'il admet des états dans lesquels les valeurs de ses diverses grandeurs physiques, d'aléatoires qu'elles étaient, deviennent quasi certaines, avec des distributions concentrées autour de leurs valeurs moyennes. Les dites valeurs moyennes évoluent selon le théorème d'Ehrenfest (page 98) qui prend, dans ce cas, la forme d'équations du mouvement classiques. Dans un tel état, l'hamiltonien des quantons jouit d'une valeur numérique certaine. Ce n'est plus un opérateur, il devient une simple constante ne jouant plus aucun rôle actif dans l'évolution du système. Autrement dit, l'évolution des quantons — ou plutôt des particules maintenant — est découplée du rayonnement, et n'est soumise qu'à la limite classique d'un éventuel potentiel extérieur qui figurait dans  $H_{\text{mat}}$ . Le système peut alors être décrit par un ket d'un espace des états effectif, réduit à l'espace des états du rayonnement, dont l'évolution est régie par l'équation

$$i\hbar \frac{d}{dt} |\psi_{\text{ray}}(t)\rangle = V(t) |\psi_{\text{ray}}(t)\rangle, \quad (\text{V.32})$$

où  $V(t)$  est la version “semi quantique” (par opposition à semi classique) de l'hamiltonien d'interaction matière-rayonnement (III.11), à savoir: courant et densité de matière classiques, opérateur de champ quantique. L'opérateur courant  $\mathbf{j}(\mathbf{r}, \mathbf{R}_1, \dots)$  par exemple est remplacé par sa valeur moyenne (dans un état qui dépend du temps),  $\mathbf{j}(\mathbf{r}, t, \mathbf{r}_1, \dots)$ .

Dans le cadre semi quantique maintenant défini, nous pouvons maintenant envisager le plus simple des modèles, ça allège l'écriture sans restreindre la généralité. Le rayonnement est assez faible pour pouvoir négliger la contribution diamagnétique, non linéaire, dans l'interaction (III.11). Les particules évoluent sous influence de forces extérieures dont l'effet se réduit à un courant effectif dans un seul mode  $\mathbf{k}T$ . Ce sont par exemple des charges dans un milieu conducteur soumises à une différence de potentiel alternative, mais il va sans dire que par superposition de courants (II.62) on peut reproduire n'importe quelle distribution.<sup>6</sup> Usant du développement en modes de l'opérateur de champ, on a enfin l'opérateur de l'interaction du rayonnement quantique avec les particules classiques:

$$V(t) = -q \sqrt{\frac{\hbar}{2\varepsilon_0 \omega \mathcal{V}}} \left\{ \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{j}_{\mathbf{k}T}^*(t) a_{\mathbf{k}T} e^{-i\omega t} + \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{j}_{\mathbf{k}T}(t) a_{\mathbf{k}T}^+ e^{i\omega t} \right\}.$$

<sup>6</sup>Remarquez: maintenant le rayonnement (quantique) est encore couplé à la matière (classique), mais la matière n'est plus couplée au rayonnement! Elle n'est plus que source de celui-ci, classiquement animée par des forces extérieures.



Dans cette expression,  $a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T}^+$  sont des opérateurs, le reste n'est que nombres. Il s'ensuit — remarque importante pour la suite — que généralement les opérateurs  $V(t_1)$  et  $V(t_2)$  à deux instants différents ne commutent pas, mais que leur commutateur n'est qu'un simple nombre, imaginaire pur ( $V$  est hermitique), dépendant de  $t_1$  et  $t_2$ .

Reste à prédire l'évolution du ket d'état du rayonnement, évolution régie par l'équation (V.32). Introduisons pour cela l'*opérateur d'évolution* du rayonnement dans cet environnement, défini par l'équation différentielle

$$i\hbar \frac{\partial}{\partial t} U(t, t_0) = V(t) U(t, t_0),$$

flanquée de la condition initiale  $U(t, t_0) = 1$ . Vous n'aurez aucune peine à vérifier que le ket

$$|\psi_{\text{ray}}(t)\rangle \stackrel{\text{df}}{=} U(t, t_0) |\psi_{\text{ray}}(t_0)\rangle$$

est bien solution de l'équation du mouvement (V.32). Et maintenant — miracle des opérateurs  $V(t)$  dont les commutateurs sont toujours des nombres — nous allons pouvoir trouver une expression explicite de l'opérateur d'évolution. Pour un laps de temps  $\Delta t$  petit, on a par définition

$$\begin{aligned} U(t_0 + \Delta t, t_0) &= U(t_0, t_0) + \left. \frac{\partial U(t, t_0)}{\partial t} \right|_{t_0} \Delta t + \dots \\ &= 1 + \frac{1}{i\hbar} V(t_0) U(t_0, t_0) \Delta t + \dots \\ &= 1 - \frac{i}{\hbar} V(t_0) \Delta t + \dots \\ &\underset{\Delta t \rightarrow 0}{\sim} e^{-\frac{i}{\hbar} \Delta t V(t_0)}, \end{aligned}$$

et donc, pour un laps deux fois plus grand,

$$\begin{aligned} U(t_0 + 2\Delta t, t_0) &= U(t_0 + 2\Delta t, t_0 + \Delta t) U(t_0 + \Delta t, t_0) \\ &\sim e^{-\frac{i}{\hbar} \Delta t V(t_0 + \Delta t)} e^{-\frac{i}{\hbar} \Delta t V(t_0)}, \end{aligned}$$

tant et si bien que:

$$U(t_0 + n\Delta t, t_0) \sim e^{-\frac{i}{\hbar} \Delta t V(t_0 + (n-1)\Delta t)} e^{-\frac{i}{\hbar} \Delta t V(t_0 + (n-2)\Delta t)} \dots e^{-\frac{i}{\hbar} \Delta t V(t_0)}.$$

Grâce aux propriétés de commutation des opérateurs  $V$ , la formule de Campbell *et al.* (Exercice 10), on peut maintenant écrire:

$$U(t_0 + n\Delta t, t_0) \sim e^{i\varphi} e^{-\frac{i}{\hbar} \Delta t \sum_{p=0}^{n-1} V(t_0 + p\Delta t)},$$

où la phase  $\varphi$ , un nombre sans conséquence physique, est tout ce qui reste des commutateurs des opérateurs  $V$ . On a donc enfin, à la limite  $\Delta t \rightarrow 0$ ,

$$U(t, 0) = e^{i\varphi(t)} e^{-\frac{i}{\hbar} \int_0^t dt_1 V(t_1)}.$$

Etant donnée l'expression de  $V(t)$ , son intégrale est de la forme

$$-\frac{i}{\hbar} \int_0^t dt_1 V(t_1) = \alpha_{\mathbf{k}T}(t) a_{\mathbf{k}T}^+ - \alpha_{\mathbf{k}T}^*(t) a_{\mathbf{k}T},$$

où  $\alpha_{\mathbf{k}T}(t)$  est un pur nombre (complexe). L'opérateur d'évolution peut donc s'écrire, en fonction de l'opérateur déplacement (V.29),

$$U(t, 0) = e^{i\varphi(t)} D(\alpha_{\mathbf{k}T}(t)).$$

Ainsi, partant d'un état initial vide, le rayonnement quantique se retrouve, par l'effet de particules chargées classiques animées du mouvement  $\mathbf{j}_{\mathbf{k}}$ , dans l'état:

$$|\psi_{\text{ray}}(t)\rangle = e^{i\varphi(t)} D(\alpha_{\mathbf{k}T}(t)) |0\rangle,$$

c'est-à-dire un état cohérent du mode  $\mathbf{k}T$ , qui, à la limite des grandes excitations, a un comportement d'onde plane électromagnétique. Notre théorie quantique du rayonnement rend compte de toutes les propriétés classiques, aussi bien du rayonnement que de la matière et de leur couplage.

## V.6 Pour les amateurs de classique

Le réémetteur de France-Musique situé à Chamrousse (Isère) a une puissance  $\mathcal{P} = 2 \text{ kW}$ , à la fréquence  $\nu = 91,8 \text{ MHz}$ . La mélomane de l'Institut des Sciences Nucléaires de Grenoble, à  $d = 18 \text{ km}$  de là, a-t-elle besoin de la théorie quantique du rayonnement pour capter son émission favorite?

Dans une description classique de la situation, la puissance rayonnée nous permet d'évaluer l'amplitude du champ électrique à l'emplacement du récepteur. Ce champ électrique nous dira ensuite le nombre moyen de photons correspondant à l'état du rayonnement dans la boîte du récepteur, si l'on se place dans le cadre d'une description quantique par un état cohérent. Grâce à la discussion de la section V.5.3, nous saurons alors si les dispersions quantiques de l'énergie, ou du champ électrique, dans le récepteur, sont négligeables ou non.

La puissance totale rayonnée est égale au flux du vecteur de Poynting. Nous avons donc  $\mathcal{P} \propto Sd^2$ , en fonction du module  $S$  de l'amplitude de du vecteur de Poynting à l'emplacement du récepteur.<sup>7</sup> Dans la région du récepteur, le champ

<sup>7</sup>Ce calcul en ordre de grandeur nous permet de supposer l'émission isotrope, chose à proprement parler impossible (il ne peut y avoir d'émission monopolaire), d'autant que l'antenne de l'émetteur est directionnelle (l'indice d'écoute étant par trop bas chez les marmottes du massif).

électromagnétique a pratiquement la forme d'une onde plane, caractérisée par la relation  $|\mathbf{B}| = |\mathbf{E}|/c$ . Ainsi, le vecteur de Poynting,  $\mathbf{S} = \varepsilon_0 c^2 \mathbf{E} \wedge \mathbf{B}$ , a pour module  $S = \varepsilon_0 c E^2$ . On peut donc exprimer l'amplitude du champ électrique au récepteur en fonction de la puissance de l'émetteur:

$$E^2 \approx \frac{\mathcal{P}}{\varepsilon_0 c d^2}.$$

Quantiquement, la description du rayonnement dans la boîte du récepteur nécessite un état cohérent  $|z\rangle$ , correspondant à un nombre moyen de photons  $\langle N \rangle = |z|^2$ , et à une amplitude moyenne du champ électrique

$$E = \sqrt{\frac{2\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z|.$$

On en déduit le nombre moyen de photons du rayonnement dans la boîte

$$\langle N \rangle \approx \frac{\mathcal{P}\mathcal{V}}{\hbar\omega c d^2}.$$

Nous ne pouvons maintenant éviter de nous demander de quelle boîte il s'agit au juste. Pour étudier un rayonnement de longueur d'onde  $\lambda \approx c/\omega = 3\text{ m}$  ici — notre boîte, avec ses conditions de périodicité, doit avoir une longueur  $L = \lambda$ , ou  $2\lambda$ , ou  $3\lambda$ , *etc.* Autrement dit, la pulsation du rayonnement doit être supérieure à sa pulsation de coupure  $\omega_0$  (équation II.47), et l'on doit prendre nécessairement

$$L \geq \lambda \approx \frac{c}{\omega}.$$

Remarquons que cette valeur est également la dimension de l'antenne du récepteur, et donc la longueur de la cavité à laquelle nous pouvons assimiler cette antenne si nous préférons décrire le rayonnement dans une cavité physique, plutôt que dans une boîte à périodes formelle. Quoi qu'il en soit, nous obtenons

$$\langle N \rangle \geq \frac{\mathcal{P}c^2}{\hbar\omega^4 d^2} = 10^{14} \text{ photons.} \quad (\text{V.33})$$

C'est beaucoup, et les dispersions quantiques des grandeurs physiques dans cet état sont absolument négligeables. Rappelons que les dispersions relatives du champ électrique, du nombre de photons, ou de l'énergie, sont, d'après la section précédente, de l'ordre de  $\langle N \rangle^{-1/2}$  — soit  $10^{-7}$  ici — indécélables sur un champ électrique, par exemple, qui n'est déjà lui-même que de  $10^{-2} \text{ V m}^{-1}$ . La description quantique n'est donc pas nécessaire dans ce cas et, bien entendu(e), la classique suffit au bonheur des auditrices de France-Musique.

Il en va évidemment de même pour l'émetteur, ce que nous pouvions trouver directement car ses seules caractéristiques qui interviennent ici sont sa puissance  $\mathcal{P}$

(qui a la dimension d'une énergie divisée par un temps), et sa pulsation  $\omega$  (l'inverse d'un temps). La valeur de l'action (une énergie multipliée par un temps) caractéristique du fonctionnement de cet émetteur est, en vertu du principe zéro,

$$\frac{\mathcal{P}}{\omega^2} = 10^{-13} \text{ J s} \gg \hbar = 10^{-34} \text{ J s}.$$

L'action de l'émetteur est beaucoup plus grande que l'action caractéristique des phénomènes quantiques, et les techniciens de Radio-France n'ont nul besoin de théorie quantique... tout au moins pour traiter de l'émission de rayonnement proprement dite, car il n'en va plus ainsi s'ils veulent tenter de comprendre la raison des différences de comportement entre isolants et conducteurs!

A quelle distance  $d'$  de l'émetteur devrait se trouver le récepteur pour que la description quantique du rayonnement dans sa cavité n'exige plus qu'un nombre moyen de photons minimum  $\langle N \rangle' = 100$  photons, pour lequel les fluctuations quantiques commencent à être visibles, sur la figure V.3 par exemple? Nous savons maintenant (voir l'équation (V.33)) que le nombre moyen de photons de l'état du rayonnement dans une cavité donnée est inversement proportionnel au carré de la distance à l'émetteur, soit

$$\frac{\langle N \rangle}{\langle N \rangle'} = \left( \frac{d'}{d} \right)^2.$$

Ainsi, dans le cas envisagé de  $\langle N \rangle = 10^{14}$  photons à la distance  $d = 18 \text{ km}$ , nous obtenons  $d' = 10^7 \text{ km}$ , distance bien supérieure à la distance Terre-Lune ( $3,8 \times 10^5 \text{ km}$ ) et plutôt comparable à la distance Soleil-Terre ( $1,5 \times 10^8 \text{ km}$ ). La distance proprement astronomique à laquelle les effets quantiques peuvent commencer à être sensibles dans la réception de France-Musique nous laisse subodorer plutôt leur intervention dans la réception des émissions radio des sources stellaires très lointaines (leur puissance est supérieure à 2 kW!).

## V.7 Ondes et particules, le crépuscule des deux

Nous sommes maintenant armés pour reprendre, en l'élargissant, la discussion sur la nature du photon amorcée dans la section III.8. Comment se fait-il qu'historiquement, dans la période classique — ou plutôt préquantique — qui commence avec le XIX<sup>e</sup> siècle, il ait pu y avoir adéquation, sans recouvrement, des champs et des particules pour décrire le rayonnement et la matière du monde physique? Jusqu'en 1905, un modèle ondulatoire (ou de champ) suffisait à faire la lumière, tandis que les modèles corpusculaires, directement issus de la mécanique du point convenaient, jusqu'au milieu des années vingt, pour expliquer les comportements des molécules, atomes, noyaux, électrons. La théorie quantique vient ensuite mettre fin à cette séparation nette avec la fameuse dualité onde-corpuscule, confusion involontairement entretenue par les pères fondateurs qui avaient, eux, certainement les idées

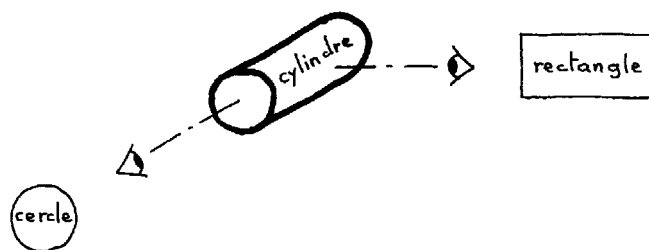


Figure V.4: Le cylindre et la pseudo-dualité cercle-rectangle

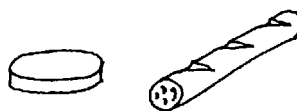


Figure V.5: Le camembert et la baguette. Toujours cylindre, souvent (jamais) cercle, jamais (souvent) rectangle.

beaucoup plus claires sur la question. Disons plutôt que dans cette période, la référence de la théorie quantique est un nouvel objet, que nous avons appelé quanton, dont l'aspect — onde ou particule — peut différer selon le point de vue, tout comme un cylindre qui, pour parfois présenter une apparence de cercle ou de rectangle (fig. V.4), n'en reste pas moins un cylindre plutôt qu'une mystérieuse dualité cercle-rectangle.

Mais alors pourquoi, s'il n'y a là que différence de points de vue, la lumière n'est elle jamais apparue — avant 1905 — comme particules, et l'électron — avant Davisson et Germer — jamais comme onde? Notre analogie conceptuelle avec le cylindre peut fournir un embryon de réponse: il existe des cylindres, type boîte de camembert, qui ne présentent pratiquement jamais l'aspect — sauf point de vue très particulier et fallacieux — d'un rectangle, tandis que d'autres, type baguette de pain, ne peuvent avoir l'apparence d'un cercle (fig. V.5). Ainsi, nous sommes en présence de deux êtres quantiques d'essences différentes pour l'instant, l'un — le photon — dont les manifestations classiques seront des ondes électromagnétiques, tandis que l'autre — auquel nous réserverons, comme nous l'avons fait implicitement jusqu'à présent, le nom de quanton — revêt lorsqu'il le faut l'habit classique de la particule.

«Qu'est-ce à dire? L'électron ne se présenterait jamais comme une onde, et pourtant son comportement ondulatoire est à l'origine de la mécanique du même nom!» Ne nous laissons pas abuser par les pauvres mots des physiciens. La fonction-d'onde de l'électron, porteuse de ses caractéristiques stochastiques, a en commun

avec une onde d'être un champ au sens mathématique — une valeur attachée en chaque point de l'espace à chaque instant —, et d'avoir une évolution régie par une équation linéaire, dont les solutions sont superposables. Pour le reste, les différences sont évidentes; les valeurs des grandeurs physiques de deux ondes classiques sont additives — deux champs électriques s'ajoutent pour donner un champ électrique — contrairement aux valeurs des grandeurs physiques correspondant à des fonctions-d'onde différentes. L'utilisation, pour calculer des grandeurs physiques, de la fonction-d'onde résultant de la superposition exige d'ailleurs l'opération préalable, hautement non linéaire, de normalisation. La distinction entre fonction-d'onde et onde trouve son expression la plus paradoxale dans l'invariance galiléenne de la longueur d'(une) onde (la distance entre les crêtes de vagues est la même pour le gardien de phare et le pilote du Concorde) contrairement à la longueur (d'une fonction-d'onde) dite de de Broglie, puisque l'impulsion dépend du référentiel (voir [47]). L'usage est maintenant trop bien ancré pour tenter de donner un autre nom à la fonction-d'onde, mais il importe de se souvenir que ce n'est pas une onde au sens ordinaire.

La section précédente nous a montré que le photon-baguette de pain, ou plutôt le système quantique constitué par le rayonnement, admet des états quasi-classiques du type onde électromagnétique. «Mais pourquoi l'électron, par exemple, ne dispose-t-il pas de tels états? Pourquoi reste-t-il camembert? Pourquoi n'y a-t-il pas d'ondes électroniques comme il y a des ondes électromagnétiques?» La simple exigence de cohérence de la théorie quantique avec la description de la nature qu'elle nous offre déjà, suffit à expliquer cette absence. Nous savons que pour obtenir un état du rayonnement du type onde électromagnétique, ou état cohérent, il faut une nombre moyen de photons élevé. Cette condition est pratiquement irréalisable pour des électrons que l'on ne saurait accumuler sans voir la charge électrique croître dans des proportions alarmantes.

«Mais pourquoi pas des ondes de neutron?» C'est que pour le neutron, comme pour l'électron d'ailleurs, des contingences encore plus fondamentale en interdisent l'accumulation. Tout d'abord, l'invariance par rotation exclut l'association de quantons de spin demi-entier en nombre quelconque pour constituer un état de moment angulaire donné. Selon l'intégrité (ou la demi-intégrité) de celui-ci, on ne doit superposer que des nombres pairs (ou que des nombres impairs) de quantons. Et surtout, le principe (ou, plus justement, le théorème) de Pauli condamne les fermions — les quantons de spin demi-entier — à la solitude dans leurs modes respectifs. Bien sûr, un principe n'apporte aucune explication, ce n'est après tout qu'une loi empirique, formulée pour les besoins de l'effet, mais bonne, sinon belle, fille qui peut donner plus que ce qu'elle semble avoir et vient, en particulier ici, nous rassurer en démystifiant logiquement l'absence de théorie classique des champs de fermions.

Dans le cas de l'électron, ce sont des paires électron-positron que l'on peut associer en conservant la charge électrique (et l'intégrité du moment angulaire), mais

le principe de Pauli interdit toujours la création de deux fermions identiques dans le même mode, et donc la possibilité d'ondes classiques. Par parenthèse, cette faculté de multiplication d'un électron, par l'intermédiaire de son rayonnement, lorsqu'il est accéléré (autrement dit, qu'il interagit) et qu'il dispose d'une énergie suffisante, marque des limites que le modèle du quanton, individu éternel, ne peut franchir; il faudra pour cela un modèle plus élaboré, un champ quantique des électrons et positrons. En revanche, le modèle du quanton comporte des états quasi-classiques du type particule, c'est-à-dire des états dans lesquels les distributions des valeurs de position et d'impulsion sont suffisamment concentrées<sup>8</sup>. Le théorème d'Ehrenfest dote alors les valeurs moyennes de ces grandeurs d'équations d'évolution classiques, de forme hamiltonienne. D'un autre côté, à trop vouloir préciser la position de l'électron, on augmente la dispersion de son impulsion et donc de son énergie. La création de paire, qui devient alors possible, traduit l'impossibilité d'étendre le concept de fonction-d'onde au domaine relativiste einsteinien. Les mêmes difficultés se présentaient à Dirac lorsqu'il cherchait l'interprétation des solutions de l'équation relativiste qu'il venait de découvrir.

«Tout ceci est bel et bien, mais il existe d'autres bosons que le photon, que personne, lorsqu'ils sont neutres, ni Pauli, ni Coulomb, n'empêche de créer dans le même mode pour en faire des états quasi-classiques du type onde de boson. Pourquoi n'en parle-t-on jamais?» En étendant ces considérations à des bosons matériels nous sommes en train de sortir du cadre des deux seuls modèles que nous connaissons pour l'instant — le quanton et le rayonnement quantique — et nous anticipons sur nos futurs champs quantiques de bosons. Mais, de toute façon, deux raisons empêchent la manifestation de cet aspect des bosons massifs. De tous les bosons massifs connus à l'heure actuelle, les moins éphémères, les  $K_L^0$ , n'ont qu'une durée de vie moyenne de  $5 \times 10^{-8}$  s ce qui laisserait bien peu de temps pour la fabrication et l'observation de ce genre d'onde classique. Et surtout, si, comme pour le photon (page 81), la pulsation du mode est associée à l'énergie du boson — et nous verrons que c'est effectivement le cas —, cette pulsation se trouve bornée inférieurement par l'énergie de repos, ou la masse. Alors que pour le photon cette borne est nulle (voir la relation de dispersion (II.39)), pour le plus léger des bosons massifs neutres, le  $\pi^0$  dont la durée de vie n'est que de  $0,8 \times 10^{-16}$  s, la contrainte est  $\hbar\omega \geq m_\pi c^2$ , soit

$$\omega \geq \frac{m_\pi c^2}{\hbar}.$$

La masse du pion,  $m_\pi c^2 \approx 140$  MeV, et les valeurs des constantes fondamentales,  $\hbar c \approx 200$  MeV fm et  $c \approx 3 \times 10^8$  m s<sup>-1</sup> =  $3 \times 10^{23}$  fm s<sup>-1</sup>, imposent dans ce cas  $\omega \geq 10^{23}$  rd s<sup>-1</sup>. Même dans le mode le plus bas, la fréquence de  $10^{22}$  Hz est déjà beaucoup trop élevée pour que l'onde que l'on pourrait en faire puisse avoir

<sup>8</sup>Dans le cas d'un électron relativiste dans un potentiel harmonique, ces états quasi-classiques sont justement les états cohérents (voir [65]). Ainsi, les états cohérents peuvent aussi bien servir à représenter quantiquement une particule classique qu'une onde classique.

une quelconque signification opérationnelle. Une telle fréquence est inobservable. Il faudrait disposer d'une grandeur physique oscillant avec une fréquence du même ordre pour mettre en évidence des battements; et il n'en va pas mieux pour la phase de cette onde.

Je me suis déjà étendu à loisir, dans la section III.8, sur le fait que le photon n'est pas un quanton, et encore moins une particule. Au delà de la simple constatation, existe-t-il à cela une raison plus profonde, du même niveau que celle qui interdit les ondes de fermions? J'anticipe la formalisation invariante de la théorie quantique des champs qui justifiera quelques détails techniques manquants, mais c'est maintenant que la question se pose des diverses limites classiques des théories quantiques, et nous pouvons y répondre qualitativement<sup>9</sup>.

Que le photon puisse être un quanton, et éventuellement une particule classique, impliquerait l'existence d'une grandeur physique "position" de ce photon et, dans un état donné, d'une densité de probabilité  $\rho(\mathbf{r}, t)$  positive pour les valeurs de cette position. La permanence du quanton implique que cette densité obéisse à une équation de continuité  $\partial_t \rho + \nabla \cdot \mathbf{j} = 0$ . Mais, bien que cela ne soit pas apparent dans la jauge que nous avons adoptée, la théorie quantique du rayonnement a une forme invariante par rapport aux transformations de Lorentz, ne serait-ce que parce qu'elle est calquée sur le modèle des équations de Maxwell. L'hypothétique équation de continuité qui s'en déduirait devrait donc avoir aussi cette invariance — le photon, de par sa masse nulle, ne saurait avoir de limite galiléenne —, ce qui nécessite que la densité  $\rho$  soit la composante temporelle d'un quadrivecteur. Pour que cette composante soit, à tous coups, positive, il faut que le quadrivecteur soit un carré construit avec les tenseurs qui décrivent le rayonnement, à savoir le quadripotential  $A_\mu$  ou le tenseur du champ électromagnétique  $F_{\mu\nu}$ . Ces derniers sont des tenseurs à un ou deux indices, dont les carrés, en dépit de toutes les contractions imaginables, ne pourront jamais former un quadrivecteur (un tenseur à un indice).

Il en serait évidemment de même avec un champ scalaire  $\phi$ ; la composante temporelle de  $(\partial_t \phi, \nabla \phi)$ , seul quadrivecteur que l'on puisse en tirer, n'a aucune chance de rester positive. Les champs scalaire et vectoriel ont en commun la canonicité de leurs lois de transformations au cours d'un changement de repère inertiel, mais leur caractère distinctif le plus évident est le rang de la représentation irréductible du groupe des rotations spatiales qu'ils engendrent... autrement dit leur nombre de composantes, ou leur spin (voir page 112). Nous comprenons ainsi qu'avec des champs de spin entier — c'est-à-dire qui engendrent des représentations de rang impair — on ne puisse songer obtenir de comportement limite du type quanton, et encore moins l'aspect classique d'une particule<sup>10</sup>.

<sup>9</sup>C'est dès le (et surtout au) stade initiatique que se présentent les interrogations les plus fondamentales, et en remettre la discussion à plus tard n'est trop souvent qu'un moyen commode de les oublier, au profit d'une confortable immersion dans les délices du formalisme.

<sup>10</sup>Cette impossibilité fût l'une des motivations qui conduisirent Dirac à envisager des représentations paires (spin demi-entier) pour obtenir une densité positive.



«Et pourtant, des bosons — les mésons  $K$  chargés par exemple — sont bien traités comme des particules par les ingénieurs chargés d'en préparer des faisceaux, autour des grands accélérateurs. Et d'ailleurs, l'équation de Schrödinger convient parfaitement pour décrire le pion, durant son éphémère existence alentour d'un noyau avec lequel il constitue un atome pionique.» Ces bosons, contrairement au photon, sont massifs. Le photon — avec une masse empiriquement inférieure à  $10^{-21}$  MeV — est le plus léger des bosons connus. Le suivant immédiat, le  $\pi^0$ , accuse gaillardement 135 MeV! La création de telles masses est trop onéreuse pour permettre des fluctuations du nombre de bosons lorsque l'énergie éventuellement disponible par interaction est limitée. Dans le cas de l'atome pionique, où l'énergie est peu supérieure à la masse, partis avec un boson massif on reste avec ce boson (tout au moins tant qu'une interaction de la nature ne l'a pas poussé à se désintégrer en produits plus légers, mais de même énergie totale). La densité d'énergie  $\mathcal{H}(\mathbf{r}, t)$  du champ dans cet état, celle dont l'intégrale de volume donne l'hamiltonien du champ de boson libre, permet de définir — comme pour un fermion massif de basse énergie — une densité de position

$$\rho(\mathbf{r}, t) \stackrel{\text{df}}{=} \frac{\langle \mathcal{H}(\mathbf{r}, t) \rangle}{m_\pi c^2}$$

qui reste positive, normée à l'unité (ou presque), et qui évolue selon le modèle du quanton.

Ces considérations qui, d'une seule pierre, portent le coup de grâce à la dualité onde-corpuscule sont empruntées à Pauli et à Peierls [60, §1.3]. La distinction entre photon et particule, en particulier, est fondamentale pour les physiciennes qui exercent leur activité dans le domaine de l'optique quantique, et l'oublier les priverait d'explications cohérentes pour les expériences actuelles d'interférences (voir [41, p. 75], [64, p. 228] et [67]). Curieusement, ces questions, qui se présentent à l'esprit de toute arête en théorie quantique, sont beaucoup moins explicitées chez les professionnelles de la physique dite des "particules". Il faut sans doute en chercher la raison dans l'extrême division des tâches, face au gigantisme des moyens mis en œuvre. Il s'ensuit une compartimentation intellectuelle, chacun des protagonistes — technicien, physicienne expérimentatrice, ou physicien qui ne fait pas d'expériences — se contentant de la représentation de ces particules nécessaire à sa propre activité.<sup>11</sup>

Récapitulons, sous forme de maximes, les conclusions de cette section à caractère quelque peu polémique:

- La fonction-d'onde n'est pas une onde!
- Il ne peut y avoir de théorie classique des champs de fermions.

---

<sup>11</sup>Les soviétiques (voir [2, §2.2] et [44, §1 et 4]) ont fait montre de plus de réalisme pour ce qui est du photon et de la particule, mais on assiste à présent au témoignage épistolier [3, 37, 69] d'une reprise de conscience dans la classe des physiciens occidentaux.

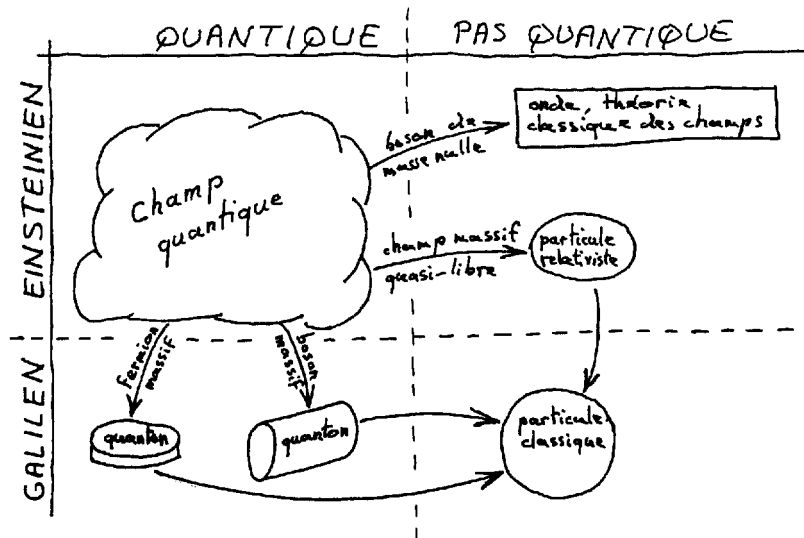


Figure V.6: Les comportements marginaux de la théorie quantique des champs.

- Seuls des champs quasi libres peuvent admettre une limite du type dynamique corpusculaire relativiste.
- A basse énergie, proche de l'énergie de repos, le modèle du quanton et, éventuellement, la dynamique corpusculaire sont applicables.
- Le photon a une masse nulle donc d'une part il est toujours relativiste, d'autre part ses fluctuations de nombre sont peu coûteuses, ce qui permet d'en faire des états cohérents, éventuellement des ondes *stricto sensu* et une théorie classique du champ électromagnétique.
- Une théorie classique des champs de bosons massifs est sans intérêt pratique.
- Toute interaction avec un quanton faisant évoluer celui-ci vers un état dont la dispersion de position est inférieure à la longueur d'onde Compton  $\hbar/mc$  implique une dispersion d'impulsion supérieure à  $mc$  qui signale l'entrée dans le domaine relativiste. S'ensuit une perte d'individualité du quanton, pour cause de création de bosons et de paires de fermions.

L'allégorie de la figure V.6 illustre les divers aspects attendus de notre future théorie quantique des champs.

## V.8 Photons droits et onde tournante

Nous nous sommes, dans ce chapitre, cantonnés jusqu'à présent au sous-espace des états d'un seul mode du rayonnement, plan, polarisé rectiligne. Enhardis par les succès obtenus grâce aux états cohérents, et à la lumière ainsi dispensée sur la théorie quantique du rayonnement, la signification de ses concepts, et sa limite classique, il est temps de nous risquer hors de ce sous-espace. Toujours pour un vecteur d'onde  $\mathbf{k}$  donné, passons à l'espace produit des espaces des états des deux modes de polarisation rectiligne de base,  $\mathcal{E}_{\mathbf{k}} \stackrel{\text{df}}{=} \mathcal{E}_{\mathbf{k}_1} \otimes \mathcal{E}_{\mathbf{k}_2}$ , sous-tendu par les kets  $|p, q\rangle_{\mathbf{k}} \stackrel{\text{df}}{=} |p\rangle_{\mathbf{k}_1} |q\rangle_{\mathbf{k}_2}$ , produits des kets propres des opérateurs nombre de photons  $N_{\mathbf{k}_1}$  et  $N_{\mathbf{k}_2}$ . J'omettrai en général la mention explicite de l'indice  $\mathbf{k}$  pour alléger l'écriture. Ces kets produits sont donc kets propres de la somme des (prolongements des) opérateurs  $N_1$  et  $N_2$ , autrement dit  $N|p, q\rangle = (p + q)|p, q\rangle$ , avec  $N \stackrel{\text{df}}{=} N_1 + N_2$ .

Nous avons eu l'occasion (page 113) d'envisager d'autres vecteurs de base de polarisation du même mode  $\mathbf{k}$ , les  $\hat{\epsilon}_{\pm} \stackrel{\text{df}}{=} \mp \frac{1}{\sqrt{2}} (\hat{\epsilon}_1 \pm i\hat{\epsilon}_2)$ , ainsi que les opérateurs associés

$$a_{\pm} = \mp \frac{1}{\sqrt{2}} (a_1 \mp ia_2). \quad (\text{V.34})$$

Ces opérateurs agissent dans l'espace produit  $\mathcal{E}_{\mathbf{k}}$ . Ce sont en fait des sommes des prolongements de  $a_1$ ,  $a_2$ ,  $a_1^+$  et  $a_2^+$  à l'espace produit et la lectrice consciencieuse a déjà vérifié (Exercice III.2) qu'ils ont les mêmes commutateurs que ceux-ci, à savoir

$$\begin{aligned} [a_{\sigma}, a_{\sigma'}^+] &= \delta_{\sigma\sigma'}, \\ [a_{\sigma}, a_{\sigma'}] &= 0, \end{aligned}$$

avec  $\sigma, \sigma' \in \{+, -\}$ , ce qui mérite à la transformation (V.34) le qualificatif de *canonique*. Jouissant de la même algèbre d'oscillateur harmonique, les opérateurs  $a_{\pm}$  ont toutes les propriétés des opérateurs  $a_1$ ,  $a_2$  déjà étudiées. En particulier, l'action de  $a_{\pm}^+$  sur le vide crée un état à un photon droit (ou gauche),

$$|1\rangle_{\pm} \stackrel{\text{df}}{=} a_{\pm}^+ |0\rangle = \mp \frac{1}{\sqrt{2}} (|1, 0\rangle \pm i|0, 1\rangle),$$

dont j'ai déjà mentionné le comportement vis-à-vis du moment angulaire du rayonnement. Toutes les constructions d'objets associés à un mode polarisé rectiligne  $\hat{\epsilon}_{\mathbf{k}T}$  que nous avons réalisées jusqu'à maintenant peuvent être répétées telles quelles avec le mode droit (ou le mode gauche). On a ainsi des états à  $n$  photons droits (ou gauches),

$$|n\rangle_{\pm} = \frac{1}{\sqrt{n!}} (a_{\pm}^+)^n |0\rangle,$$

états propres normés à un de l'opérateur nombre de photons droits (ou gauches)  $N_{\pm} \stackrel{\text{df}}{=} a_{\pm}^+ a_{\pm}$  et, poursuivant la similitude algébrique, nous pouvons construire des

états cohérents de ces photons, droits par exemple:

$$|z\rangle_+ \stackrel{\text{df}}{=} D_+(z) |0\rangle = e^{za_+^\dagger - z^* a_+} |0\rangle,$$

évidemment états propres de  $a_+$  avec la valeur propre  $z$ .

Ces états bien définis ont-ils pour autant une signification physique? Autrement dit confèrent-ils une valeur moyenne remarquable à l'opérateur champ électrique? Le calcul de cette valeur moyenne exige d'abord l'écriture de l'état  $|z\rangle_+$  en termes d'états de photons polarisés rectilignes puisque c'est sur ces modes que nous connaissons l'opérateur  $\mathbf{E}(\mathbf{r}, t)$  par son développement. C'est facile:

$$|z\rangle_+ = e^{-z \frac{1}{\sqrt{2}}(a_1^\dagger + ia_2^\dagger) + z^* \frac{1}{\sqrt{2}}(a_1 - ia_2)} |0\rangle,$$

et comme les opérateurs  $a_1^\dagger$  et  $a_1$  commutent avec les opérateurs  $a_2^\dagger$  et  $a_2$ , on a

$$\begin{aligned} |z\rangle_+ &= e^{-\frac{z}{\sqrt{2}}a_1^\dagger + \frac{z^*}{\sqrt{2}}a_1} e^{-i\frac{z}{\sqrt{2}}a_2^\dagger - i\frac{z^*}{\sqrt{2}}a_2} |0\rangle, \\ &= D_1\left(-\frac{z}{\sqrt{2}}\right) D_2\left(-i\frac{z}{\sqrt{2}}\right) |0\rangle, \\ &= \left|-\frac{z}{\sqrt{2}}\right\rangle_1 \left|-i\frac{z}{\sqrt{2}}\right\rangle_2. \end{aligned}$$

Merveille: un état cohérent de photons droits est un produit d'états cohérents de photons polarisés rectilignes selon deux directions orthogonales et déphasés de  $\pi/2$ . Comme un état cohérent de photons du mode  $\hat{\mathbf{e}}_1$  parvenait à représenter un champ électrique oscillant, vous vous doutez bien de la fin de cette histoire. Effectivement, utilisant le résultat (V.31), nous obtenons ici pour valeur moyenne du champ électrique

$$\begin{aligned} {}_+\langle z | \mathbf{E}(\mathbf{r}, t) | z \rangle_+ &= {}_2\langle -i\frac{z}{\sqrt{2}} | {}_1\langle -\frac{z}{\sqrt{2}} | \mathbf{E}(\mathbf{r}, t) | -\frac{z}{\sqrt{2}} \rangle_1 \left| -i\frac{z}{\sqrt{2}} \right\rangle_2, \\ &= \hat{\mathbf{e}}_1 \sqrt{\frac{\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z| \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg(-z)) \\ &\quad + \hat{\mathbf{e}}_2 \sqrt{\frac{\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z| \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg(-iz)), \end{aligned}$$

soit finalement (et après adaptation du même calcul au cas de photons gauches):

$$\begin{aligned} {}_\pm \langle z | \mathbf{E}(\mathbf{r}, t) | z \rangle_\pm &= \sqrt{\frac{\hbar\omega}{\varepsilon_0 \mathcal{V}}} |z| \left( \pm \hat{\mathbf{e}}_1 \cos(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z + \frac{\pi}{2}) \right. \\ &\quad \left. + \hat{\mathbf{e}}_2 \sin(\omega t - \mathbf{k} \cdot \mathbf{r} - \arg z + \frac{\pi}{2}) \right). \end{aligned}$$

Victoire, ça tourne! Nous voyons, sur la figure V.7, que dans un état cohérent de photons droits (ou gauches), la valeur moyenne du champ électrique en un point

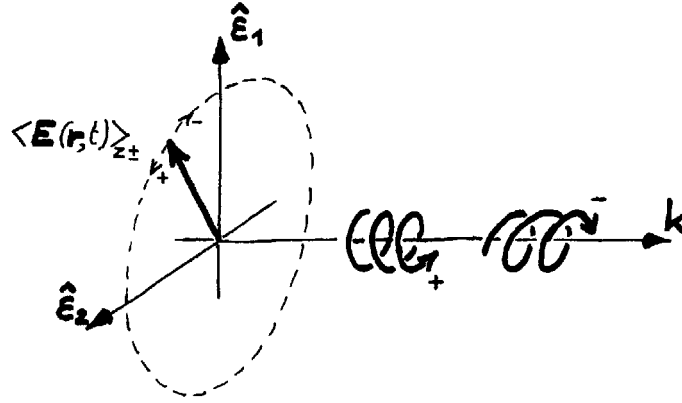


Figure V.7: La valeur moyenne du champ électrique dans un état cohérent de photons droits (ou gauches).

se visse à droite (ou à gauche) selon la convention universellement adoptée, tout au moins chez les serreurs de boulons et les sommeliers... droitiers. En revanche, comme je l'avais signalé (page 113), ce comportement est qualifié d'onde polarisée circulaire gauche (ou droite) par les opticiennes et opticiens.

## V.9 Les états comprimés

Avec les états cohérents, nous avons donc réussi à trouver que les ondes planes du rayonnement classique (elles aussi dites cohérentes, dans la mesure où on ne leur trouve qu'une seule valeur de phase et une seule valeur d'amplitude, et ceci qu'elles soient polarisées rectilignes ou circulaires) figurent bien au palmarès de la théorie quantique du rayonnement. En parlant sur les ondes — électromagnétiques — nous faisons, sans le savoir, de la théorie quantique avec des états cohérents.

Mais il existe encore d'autres états remarquables dans l'espace des états du rayonnement quantique. Pour faire simple, réfugions nous à nouveau dans le sous espace des états du rayonnement associés au mode  $\mathbf{k}, \hat{\mathbf{e}}_{\mathbf{k}T}$ , et reprenons, une fois de plus, l'expression de l'opérateur champ électrique effectif (V.2), spécialisé, toujours pour alléger l'écriture, au point  $\mathbf{r} = 0$ , soit

$$E(t) = i \sqrt{\frac{\hbar\omega}{2\varepsilon_0 V}} \{a e^{-i\omega t} - a^\dagger e^{i\omega t}\}.$$

Au lieu d'analyser, comme nous l'avons fait, le comportement de cette grandeur physique en termes d'opérateurs amplitude et phase, nous pouvons circonvier la

complication liée à la non hermiticité de  $a$  et  $a^+$  en essayant une description en fonction des opérateurs

$$\begin{aligned} P &\stackrel{\text{df}}{=} \frac{1}{2}(a + a^+), \\ Q &\stackrel{\text{df}}{=} \frac{1}{2i}(a - a^+), \end{aligned}$$

qui, eux, ont le mérite d'être hermitiques. Leur commutateur vaut

$$[P, Q] = \frac{i}{2},$$

et ils constituent un ensemble irréductible d'opérateurs tout aussi valide que  $a$  et  $a^+$  pour définir la structure de l'espace des états — et donc le modèle — du rayonnement quantique (dans le mode  $\mathbf{k}T$ ). Remarquez au passage l'identité algébrique du dit modèle avec celui de l'oscillateur harmonique à une dimension d'impulsion  $P$  et de position  $Q$  (à des facteurs numériques près, voir l'exercice 14).

Nous avons maintenant  $a = P + iQ$ , et le champ électrique peut s'écrire:

$$E(t) = \sqrt{\frac{2\hbar\omega}{\varepsilon_0\mathcal{V}}} (P \sin \omega t - Q \cos \omega t).$$

Cette expression illustre la difficulté déjà constatée qu'il y a à trouver un état propre de  $E(t)$  à tout instant  $t$  car, de par leur commutateur, les grandeurs  $P$  et  $Q$  sont affectées, quel que soit l'état du système, de dispersions soumises à l'inégalité de Heisenberg:

$$\Delta P \Delta Q \geq \frac{1}{4}.$$

Récapitulons d'abord les propriétés des états cohérents dans ce cadre renouvelé. Dans un état propre de  $a$  avec la valeur propre  $z$ ,  $a|z\rangle = z|z\rangle$ , nous avons évidemment  $\langle P \rangle_z = (z + z^*)/2$ , et  $\langle Q \rangle_z = (z - z^*)/2i$ . Pour calculer les dispersions, nous avons

$$P^2 = \frac{1}{4}(a^2 + aa^+ + a^+a + a^{+2}) = \frac{1}{4}(a^2 + a^{+2} + 2N + 1),$$

d'où

$$\langle P^2 \rangle_z = \frac{1}{4}(z^2 + z^{*2} + 2|z|^2 + 1),$$

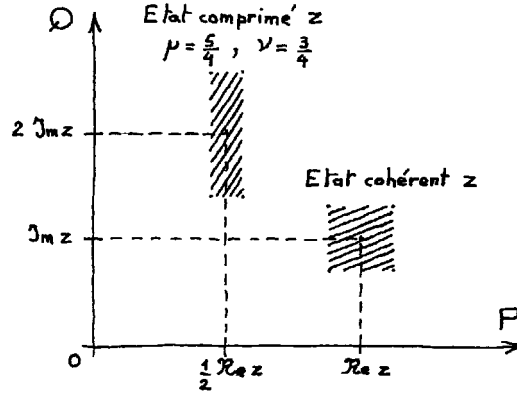
tandis que

$$\langle P \rangle_z^2 = \frac{1}{4}(z^2 + z^{*2} + 2|z|^2).$$

La différence de ces deux expressions et un calcul analogue pour  $Q$  nous donnent finalement

$$(\Delta P)_z = (\Delta Q)_z = \frac{1}{2}.$$

Nous retrouvons la propriété de saturation de l'inégalité de Heisenberg par l'état cohérent, avec un égalitarisme déjà suggéré (page 131) mais maintenant manifeste

Figure V.8: Deux distributions des valeurs de  $P$  et  $Q$ .

dans la description en termes de  $P$  et  $Q$ . Les dispersions de  $P$  et  $Q$  constituent, sous un autre nom, le *bruit quantique* qui limite la précision du champ électrique et donc certaines de ses applications pratiques. Dans un état cohérent, ce bruit affecte dans la même mesure la composante du champ dite *en phase* (proportionnelle à  $\sin \omega t$ ) et la composante *en quadrature* (proportionnelle à  $\cos \omega t$ ). Schématiquement, on peut représenter (figure V.8) les distributions de  $P$  et  $Q$  par un carré (dispersions égales) centré sur leurs valeurs moyennes.

Mais le champ électrique peut être pratiquement analysé, par un filtre convenable, en ses composantes en phase et en quadrature. On peut alors songer à ruser avec la limite quantique imposée par l'inégalité de Heisenberg en reportant la majeure part dans l'une des composantes pour ne plus utiliser que l'autre, presque débarrassée de ses fluctuations quantiques. En d'autres termes, existe-t-il des états de précision maximale — qui saturent l'inégalité de Heisenberg — mais d'inégales dispersions? Effectivement ces états existent, et leur construction s'apparente à celle des états cohérents.

Etant donnés deux nombres complexes  $\mu$  et  $\nu$ , définissons l'opérateur:

$$b \stackrel{\text{df}}{=} \mu a + \nu a^+.$$

Cette transformation est canonique,  $[b, b^+] = 1$ , dans la mesure où les paramètres  $\mu$  et  $\nu$  satisfont la condition

$$|\mu|^2 - |\nu|^2 = 1.$$

L'idée est maintenant, par analogie avec les états cohérents, de s'intéresser aux kets propres de  $b$ :

$$b|\tilde{z}\rangle = z|\tilde{z}\rangle.$$

Dans le cas  $\mu = 1$ ,  $\nu = 0$ , le ket  $|\tilde{z}\rangle$  est tout simplement le ket d'état cohérent  $|z\rangle$ . Dans le cas général, il est facile de s'assurer de l'existence d'un ket  $|\tilde{z}\rangle$ , pour tout  $z$  donné, en établissant la relation de récurrence des coefficients de son développement sur les états de base  $|n\rangle$  (Exercice **16**). Alternativement, on peut profiter astucieusement de la complète analogie algébrique avec les états cohérents pour déterminer le développement de  $|\tilde{z}\rangle$  sur les kets propres de l'opérateur  $\tilde{N} \stackrel{\text{df}}{=} b^+b$  (Exercice **17**). Mais nous n'avons même pas besoin de ces développements pour établir les propriétés des états  $|\tilde{z}\rangle$  qui nous intéressent!

Qu'en est-il donc des valeurs moyennes et dispersions de  $P$  et  $Q$  dans l'état  $|\tilde{z}\rangle$ ? Considérant la propriété définitoire de cet état, tout calcul de valeur moyenne d'une grandeur sera facilité si l'on exprime d'abord la dite grandeur en fonction des opérateurs  $b$  et  $b^+$ . Compte tenu de la condition de canonicité, on a

$$a = \mu^*b - \nu b^+,$$

d'où

$$\begin{aligned} P &= \frac{1}{2}((\mu - \nu)^*b + (\mu - \nu)b^+), \\ Q &= \frac{1}{2i}((\mu + \nu)^*b - (\mu + \nu)b^+), \\ P^2 &= \frac{1}{4}((\mu - \nu)^{*2}b^2 + (\mu - \nu)^2b^{+2} + |\mu - \nu|^2(2\tilde{N} + 1)), \\ Q^2 &= \frac{1}{4}(-(\mu + \nu)^{*2}b^2 - (\mu + \nu)^2b^{+2} + |\mu + \nu|^2(2\tilde{N} + 1)), \end{aligned}$$

expressions qui nous permettent d'obtenir les valeurs moyennes

$$\begin{aligned} \langle P \rangle_{\tilde{z}} &= \frac{1}{2}((\mu - \nu)^*z + (\mu - \nu)z^*), \\ \langle Q \rangle_{\tilde{z}} &= \frac{1}{2i}((\mu + \nu)^*z - (\mu + \nu)z^*), \end{aligned}$$

et les dispersions

$$\begin{aligned} (\Delta P)_{\tilde{z}} &= \frac{1}{2}|\mu - \nu|, \\ (\Delta Q)_{\tilde{z}} &= \frac{1}{2}|\mu + \nu|. \end{aligned}$$

Succès: les dispersions de  $P$  et  $Q$  sont modulables, au gré des paramètres (pas indépendants)  $\mu$  et  $\nu$ , leur produit restant égal à  $1/4$ , la borne de Heisenberg. Les états  $|\tilde{z}\rangle$  sont précisément ce que nous cherchions. On les appelle les *états comprimés*. A titre d'exemple, si le système est dans l'état comprimé  $|\tilde{z}\rangle$  avec  $\mu = 5/4$  et  $\nu = 3/4$ , on trouve les valeurs moyennes et dispersions:

$$\begin{aligned} \langle P \rangle_{\tilde{z}} &= \frac{1}{2} \frac{z + z^*}{2}, & (\Delta P)_{\tilde{z}} &= \frac{1}{4}, \\ \langle Q \rangle_{\tilde{z}} &= 2 \frac{z - z^*}{2i}, & (\Delta Q)_{\tilde{z}} &= 1, \end{aligned}$$



représentées sur la figure V.8 par un rectangle de même aire que le carré de l'état cohérent, et dont les côtés sont dans le rapport 1/4 des dispersions.

Dans un état cohérent, le rayonnement quantique acquiert, lorsque  $|z|$  est grand, toutes les apparences d'un rayonnement classique. Cet état peut alors être qualifié de quasi classique. Dans un état comprimé, par contre, le champ électrique reste affecté de fluctuations différentes sur ses composantes en phase et en quadrature, propriété spécifiquement quantique. L'état comprimé n'est pas un état quasi classique. On peut espérer des applications pratiques des états comprimés en spectroscopie de haute précision (diminution de la largeur naturelle des niveaux atomiques) et en télécommunications (la réduction du bruit dans un canal permet d'en accroître le débit), voir pour cela la référence [76]. Mais le problème de la préparation de ces états n'est pas encore résolu; celle-ci nécessite un couplage non linéaire entre rayonnement et système matériel. Vous trouverez des modèles de fabrication d'états comprimés décrits dans [20, p. 250], [7].

## Exercices

1. Soit l'opérateur  $F \stackrel{\text{df}}{=} (N + \frac{1}{2})^{-1/2}a$ , où  $N \stackrel{\text{df}}{=} a^+a$ , et  $[a, a^+] = 1$ .

i) Calculer les éléments de matrice de l'opérateur  $F^+F$  sur la base des états propres de  $N$ , et en déduire que  $F^+F = 1 - |0\rangle\langle 0|$ .

ii) Montrer que  $[N, F] = -F$ , et que  $[N, F^+] = F^+$ . En déduire que la valeur moyenne du commutateur  $[F, F^+]$  dans l'état de base  $|n\rangle$  vaut  $\delta_{n0}$ . L'opérateur  $F$  est-il normal?

iii) Calculer le commutateur de  $C \stackrel{\text{df}}{=} (F + F^+)/2$  et  $S \stackrel{\text{df}}{=} (F - F^+)/2i$ .

iv) Montrer que  $[C, S] = (a^+(N+1)^{-1}a - 1)/2i$ . En déduire que les éléments de matrice de ce commutateur dans la base des états  $|n\rangle$  sont tous nuls à l'exception de  $\langle 0|[C, S]|0\rangle = i/2$ .

v) Montrer que  $[N, C] = -iS$  et  $[N, S] = iC$ .

2. Soit l'état  $|\varphi, s\rangle$  défini par l'équation (V.16), et l'opérateur  $F$  étudié dans l'exercice précédent.

i) Montrer que  $\langle \varphi, s | \varphi, s \rangle = 1$ .

ii) Montrer que  $\langle \varphi, s | F | \varphi, s \rangle = e^{i\varphi} s / (s + 1)$ .

iii) Montrer que  $\langle \varphi, s | F F^+ | \varphi, s \rangle = s / (s + 1)$ .

3. Toujours avec les états  $|\varphi, s\rangle$ , et l'opérateur  $F$ ...

i) Montrer que  $\langle \varphi, s | F^2 | \varphi, s \rangle = e^{2i\varphi} (s - 2) / (s + 1)$ .

ii) En déduire les valeurs moyennes  $\langle \varphi, s | C^2 | \varphi, s \rangle$  et  $\langle \varphi, s | S^2 | \varphi, s \rangle$ , où  $C$  et  $S$  sont les opérateurs déjà étudiés dans l'exercice 1. Quelles sont les limites de ces valeurs moyennes lorsque  $s \rightarrow \infty$ ?

iii) En déduire les limites des dispersions de  $C$  et de  $S$  dans l'état  $|\varphi, s\rangle$  lorsque le paramètre  $s \rightarrow \infty$ .

4. La dispersion du nombre de photons dans un état de phase.

i) Montrez que  $\langle \varphi, s | N | \varphi, s \rangle = s/2$ , et que  $\langle \varphi, s | N^2 | \varphi, s \rangle = s(2s+1)/6$ , en vous remémorant au besoin que  $\sum_{n=0}^s n = s(s+1)/2$ , et  $\sum_{n=0}^s n^2 = s(s+1)(2s+1)/6$ .

ii) En déduire la valeur de la dispersion de  $N$  dans un état  $|\varphi, s\rangle$ , puis sa limite lorsque  $s \rightarrow \infty$ ?

5. Montrer que

$$C|\varphi, s\rangle = \cos \varphi |\varphi, s\rangle - \frac{1}{\sqrt{s+1}} (e^{-i\varphi} |0\rangle + e^{i(s+1)\varphi} |s\rangle - e^{is\varphi} |s+1\rangle).$$

Que se passe-t-il lorsque  $s \rightarrow \infty$ ?

6. Propriétés des états de nombre de photons.

i) Montrer que  $\langle n | F | n \rangle = 0$ . En déduire la valeur moyenne  $\langle n | C | n \rangle$ .

ii) Montrer que  $\langle n | F^2 | n \rangle = 0$ , que  $\langle n | F F^+ | n \rangle = 1$ , et que  $\langle n | F^+ F | n \rangle = 1 - \delta_{n0}$ .

En déduire la valeur moyenne  $\langle n | C^2 | n \rangle$

iii) En déduire la valeur de la dispersion de la grandeur  $C$  dans l'état  $|n\rangle$ .

7. Calculer l'écart-type d'une variable aléatoire équiprobable sur  $[-1, 1]$ .

8. La variable aléatoire  $\varphi$  étant équiprobable sur  $[0, 2\pi]$ , calculer les valeurs moyennes de  $\cos \varphi$  et de  $\cos^2 \varphi$ . En déduire que l'écart-type de  $\cos \varphi$  vaut  $\frac{1}{\sqrt{2}}$ .

9. Montrer que la valeur moyenne du carré de l'opérateur champ électrique effectif dans le mode  $\mathbf{k}T$  (définition (V.2)), dans un état propre de l'opérateur nombre de photons effectif (V.1), vaut

$$\langle n | E^2(\mathbf{r}, t) | n \rangle = \frac{\hbar \omega}{\varepsilon_0 \mathcal{V}} \left( n + \frac{1}{2} \right).$$

En déduire la dispersion de  $E(\mathbf{r}, t)$  dans l'état  $|n\rangle$ .

10. Etant donné la fonction exponentielle, définie par  $e^x = \sum_{n=0}^{\infty} x^n/n!$ , et deux opérateurs  $X_1$  et  $X_2$ , on se propose d'évaluer le produit  $e^{X_1} e^{X_2}$ .

i) Première tâche, "montrer" que

$$e^{X_1} e^{X_2} = e^{\sum_i d_i X_i + \sum_{i,j} d_{ij} X_i X_j + \sum_{i,j,k} d_{ijk} X_i X_j X_k + \dots}.$$

Utiliser pour cela les développements des exponentielles jusqu'aux monômes du troisième degré compris, et calculer, par identification, les valeurs de chacun des coefficients  $d_i$ , puis  $d_{ij}$ , et enfin  $d_{ijk}$ .

ii) En déduire la relation

$$e^{X_1} e^{X_2} = e^{X_1 + X_2 + \frac{1}{2}[X_1, X_2] + \frac{1}{12}\{[X_1, [X_1, X_2]] + [X_2, [X_2, X_1]]\} + \dots}$$

(début de la relation de Campbell, Baker, Hausdorff, Glauber *etc.*).

**11.** Etant donné l'opérateur déplacement,  $D(z) \stackrel{\text{df}}{=} e^{za^+ - z^*a}$  pour  $z$  complexe et  $[a, a^+] = 1$ , montrer à l'aide de la formule de Campbell, *etc.* que l'on a  $D(z)D(y) = e^{i\Im(z y^*)} D(z+y)$ .

**12.** Une inégalité de Heisenberg pour le "sinus" et le "cosinus".

i) Calculer le commutateur  $[S, C]$  en fonction de  $F$  et  $F^+$ , puis en fonction du projecteur  $|0\rangle\langle 0|$ .

ii) Ecrire l'inégalité de Heisenberg correspondante.

iii) Que prédit cette inégalité dans le cas d'un état cohérent  $|z\rangle$ ? Que se passe-t-il lorsque  $|z| \rightarrow \infty$ ?

**13.** Les "sinus" et "cosinus" de la phase dans un état cohérent (d'après [17]).

i) A l'aide du développement de l'état cohérent  $|z\rangle$  sur les états  $|n\rangle$ , calculer la valeur moyenne  $\langle F \rangle_z$ . En déduire

$$\langle S \rangle_z = e^{-|z|^2} |z| \sin(\arg z) \sum_{n=0}^{\infty} \frac{|z|^{2n}}{n! \sqrt{n+1}},$$

et l'expression analogue pour  $\langle C \rangle_z$ .

ii) Montrer que

$$\frac{1}{\sqrt{n+1}} = \frac{1}{\Gamma(\frac{1}{2})} \int_0^\infty dt \frac{e^{-(n+1)t}}{\sqrt{t}}.$$

En déduire

$$\sum_{n=0}^{\infty} \frac{x^n}{n! \sqrt{n+1}} = \frac{1}{\Gamma(\frac{1}{2})} \frac{1}{\sqrt{x}} \int_0^\infty du \frac{e^{-\frac{u}{x} + x e^{-\frac{u}{x}}}}{\sqrt{u}} = \frac{e^x}{\sqrt{x}} \left(1 - \frac{1}{8x} + O(x^{-2})\right),$$

ainsi que les comportements asymptotiques de  $\langle S \rangle_z$  et  $\langle C \rangle_z$  pour  $|z| \rightarrow \infty$ .

iii) Calculer  $\langle F^2 \rangle_z$ . En déduire

$$\langle S^2 \rangle_z = \frac{1}{2} - \frac{1}{4} e^{-|z|^2} - \frac{z^2 + z^{*2}}{4} e^{-|z|^2} \sum_{n=0}^{\infty} \frac{|z|^{2n}}{n! \sqrt{(n+1)(n+2)}},$$

et l'expression analogue pour  $\langle C^2 \rangle_z$ .

iv) Montrer que

$$\begin{aligned}\sum_{n=0}^{\infty} \frac{x^n}{(n+1)!} &\underset{x \rightarrow \infty}{\sim} \frac{e^x}{x}, \\ \sum_{n=0}^{\infty} \frac{x^n}{(n+2)!} &\underset{x \rightarrow \infty}{\sim} \frac{e^x}{x^2}, \\ \sum_{n=0}^{\infty} \frac{x^n}{(n+2)!(n+2)} &\underset{x \rightarrow \infty}{\sim} \frac{e^x}{x^3}.\end{aligned}$$

En déduire

$$\begin{aligned}\sum_{n=0}^{\infty} \frac{x^n}{n! \sqrt{(n+1)(n+2)}} &= \sum_{n=0}^{\infty} \frac{x^n}{(n+1)!} \left(1 - \frac{1}{2} \frac{1}{n+2} - \frac{1}{8} \frac{1}{(n+2)^2} + \dots\right), \\ &= \frac{e^x}{x} \left(1 - \frac{1}{2x} - \frac{1}{8x^2} + \dots\right),\end{aligned}$$

ainsi que les comportements asymptotiques de  $\langle S^2 \rangle_z$ ,  $\langle C^2 \rangle_z$  et  $\langle S^2 \rangle_z + \langle C^2 \rangle_z$ .

v) Calculer les dispersions  $(\Delta S)_z$  et  $(\Delta C)_z$  pour  $|z|$  grand.

**14.** Retour sur l'oscillateur harmonique à une dimension, d'hamiltonien  $H = P^2/2m + m\omega^2 X^2/2$ .

i) On introduit les opérateurs  $\pi \stackrel{\text{df}}{=} P/\sqrt{m\hbar\omega}$  et  $\chi \stackrel{\text{df}}{=} \sqrt{m\omega/\hbar}X$ . Etablir l'expression de  $H$  en fonction de  $\pi$  et  $\chi$ . Calculer le commutateur  $[\chi, \pi]$ , et énoncer l'inégalité de Heisenberg correspondante.

ii) On introduit l'opérateur  $a \stackrel{\text{df}}{=} (\chi + i\pi)/\sqrt{2}$ . Calculer le commutateur  $[a, a^+]$  et établir l'expression de  $H$  en fonction de  $a$  et  $a^+$ .

iii) Soit  $|z\rangle$  un ket propre de  $a$ , dit état cohérent, correspondant à la valeur propre  $z$  complexe. Est-ce un état stationnaire?

iv) On souhaite étudier l'évolution des valeurs moyennes de grandeurs physiques dans l'état  $|\psi(t)\rangle$ , initialement  $|\psi(0)\rangle = |z\rangle$ . L'évolution temporelle de  $|\psi(t)\rangle$  lui-même ayant quelque chance d'être compliquée (pour ne pas dire complexe), la représentation de Heisenberg s'impose. Désirant calculer  $\langle \psi(t) | \chi | \psi(t) \rangle$  par exemple, pourquoi et comment définit-on la représentation de Heisenberg  $\chi(t)$  de l'opérateur  $\chi$ ? Calculer  $\chi$  en fonction de  $a$  et  $a^+$ , et montrer enfin que l'on a  $\chi(t) = (a e^{-i\omega t} + a^+ e^{i\omega t})/\sqrt{2}$ .

v) Montrer que la valeur moyenne de  $\chi$  et sa dispersion, dans l'état  $|\psi(t)\rangle$  évolué de l'état cohérent à  $t = 0$ , valent respectivement  $\sqrt{2}|z| \cos(\omega t - \arg z)$  et  $1/\sqrt{2}$ .

vi) Transposer ces résultats au cas de l'opérateur  $\pi$  et conclure.

**15.** Le modèle de l'oscillateur harmonique à deux dimensions, isotrope, peut être construit par produit direct de deux oscillateurs unidimensionnels, de même constante, associés respectivement aux directions perpendiculaires  $\hat{\mathbf{x}}$  et  $\hat{\mathbf{y}}$ , avec pour hamiltonien:  $H = H_x + H_y$ . Vous vous rappelez (cf. exercice précédent) qu'un état quasi-classique pour la coordonnée  $X$ , qui dote cette dernière de la valeur moyenne

$\langle X \rangle = \sqrt{2}|z|\cos(\omega t - \arg z)$ , est obtenu par l'action sur le vide de l'opérateur déplacement  $D_x(z) = e^{za_x^+ - z^*a_x}$ .

i) Quelle doit être alors la valeur moyenne de  $Y$  dans un état quasi-classique correspondant à une orbite circulaire dans le sens direct? Quel est l'opérateur  $D_y$  dont l'action sur le vide, concurremment à  $D_x$ , crée cet état?

ii) Montrez que le produit (direct) de ces deux opérateurs déplacement peut s'écrire, moyennant la définition convenable d'un opérateur  $a_+(a_x, a_y)$ , comme un opérateur déplacement  $D_+$  d'argument à préciser.

iii) Soit le ket  $|1\rangle_+ \stackrel{\text{df}}{=} a_+^+|0\rangle$ . Déterminez le développement de ce ket sur la base des états  $|n\rangle_x |m\rangle_y$ .

iv) Exprimez l'opérateur  $L \stackrel{\text{df}}{=} XP_y - YP_x$  en fonction des opérateurs  $a_x, a_x^+, a_y$  et  $a_y^+$ . En déduire le résultat de l'action de  $L$  sur l'état  $|1\rangle_+$ .

v) Quelle doit être la valeur moyenne de  $Y$  dans un état quasi-classique correspondant à une orbite circulaire rétrograde? Quel est l'opérateur  $D_y$  qui, toujours associé au même  $D_x$ , crée cet état?

vi) Montrez que ce produit peut lui aussi se ramener à un opérateur  $D_-$ , moyennant la définition d'un opérateur  $a_-(a_x, a_y)$  qui parachève la canonisation de la transformation  $a_x, a_y \rightarrow a_+, a_-$ .

vii) Exprimez les opérateurs  $H$  et  $L$  en fonctions de  $a_+, a_+^+, a_-$  et  $a_-^+$ , ou, encore mieux, de  $N_+$  et  $N_-$ . Calculez le commutateur  $[H, L]$ . Quel sont les résultats des actions respectives de  $H$  et de  $L$  sur un ket  $(1/\sqrt{n!})(a_\pm^+)^n|0\rangle$ ?

**16.** Deux nombres complexes  $\mu$  et  $\nu$  étant donnés, soit l'opérateur  $b \stackrel{\text{df}}{=} \mu a + \nu a^+$ .

i) Quelle condition  $\mu$  et  $\nu$  doivent-ils vérifier pour que l'on ait  $[b, b^+] = 1$ ?

ii) Soit un ket  $|\tilde{z}\rangle$ , ket propre de  $b$  pour la valeur propre  $z$ . Montrer que les coefficients du développement  $|\tilde{z}\rangle = \sum_{n=0}^{\infty} c_n |n\rangle$  sur la base des kets propres de  $N = a^+a$  s'obtiennent, en fonction du coefficient  $c_0$  sur le vide, par la relation de récurrence:

$$c_n = \frac{z}{\mu\sqrt{n}} c_{n-1} - \frac{\nu}{\mu} \sqrt{\frac{n-1}{n}} c_{n-2}, \quad n \geq 1.$$

**17.** Soit la transformation canonique  $b \stackrel{\text{df}}{=} \mu a + \nu a^+$ ,  $\mu$  et  $\nu$  complexes.

i) Montrer que le ket  $|\tilde{0}\rangle$ , tel que  $b|\tilde{0}\rangle = 0$ , admet un développement sur les états propres de  $N = a^+a$  de la forme:

$$|\tilde{0}\rangle = c_0 \left\{ |0\rangle - \sqrt{\frac{1}{2}} \frac{\nu}{\mu} |2\rangle + \sqrt{\frac{1 \times 3}{2 \times 4}} \left(\frac{\nu}{\mu}\right)^2 |4\rangle - \dots + \sqrt{\frac{(2n-1)!!}{2^n n!}} \left(-\frac{\nu}{\mu}\right)^n |2n\rangle + \dots \right\},$$

le coefficient  $c_0$  étant éventuellement déterminé (plus facile à dire qu'à faire!) pour que le ket  $|\tilde{0}\rangle$  soit normé à un.

ii) Quelle est la nature du spectre de  $\tilde{N} \stackrel{\text{df}}{=} b^+b$ ? Indiquez un procédé de construction de ses kets propres fondé sur le ket  $|\tilde{0}\rangle$ .

iii) Soit un ket  $|\widetilde{z}\rangle$ , ket propre de  $b$  pour la valeur propre  $z$ . Dédurre directement l'expression de son développement sur les kets propres de  $\widetilde{N}$ :

$$|\widetilde{z}\rangle = e^{-\frac{1}{2}|z|^2} \sum_{n=0}^{\infty} \frac{z^n}{\sqrt{n!}} |\widetilde{n}\rangle.$$

iv) Quelle est l'expression de l'opérateur déplacement  $\widetilde{D}(z)$  dont l'action sur le ket  $|\widetilde{0}\rangle$  produit le ket  $|\widetilde{z}\rangle$ ?



## Chapitre VI

# Emission et absorption des photons

Les processus d'émission et d'absorption que nous allons étudier dans le cadre de la théorie quantique du rayonnement, concernent des systèmes aussi divers que les atomes (transitions dipolaires électriques, résonance de fluorescence, lasers, *etc.*), les noyaux (émissions multipolaires), ou les particules plus ou moins élémentaires (désintégration  $\Sigma^0 \rightarrow \Lambda + \gamma$  par exemple). S'intéresser à ces modestes mécanismes avant de se lancer dans une théorie quantique des champs — bénéficiant plus du prestige de la mode —, permettra un retour instructif sur quelques vieux problèmes comme le rayonnement et la stabilité d'un atome, ou le spectre d'un corps noir, et des spectacles aussi délavés que le bleu du ciel.

Comme exemple de matière quantique concrétisant un système émetteur-absorbeur auquel je me référerai souvent, on peut songer aux électrons d'un atome (dont le noyau serait infiniment lourd, pour simplifier un peu l'étude du mouvement), en interaction avec le rayonnement quantique. Souvenons nous de l'hamiltonien de ce système matière-rayonnement en représentation d'interaction,

$$H(t) = H_{\text{él.}} + H_{\text{int.}}(t),$$

expression dans laquelle on trouve...

— l'hamiltonien des électrons inertes,

$$H_{\text{él.}} = \sum_{i=1}^Z \frac{\mathbf{P}_i^2}{2m} + V(\mathbf{R}_1, \mathbf{R}_2, \dots, \mathbf{R}_Z), \quad (\text{VI.1})$$

qui comprend l'interaction coulombienne instantanée électron-noyau et électron-électron,



- et, en jauge de Coulomb, où  $\nabla \cdot \mathbf{A} = 0$  implique  $\mathbf{P} \cdot \mathbf{A} = \mathbf{A} \cdot \mathbf{P}$ , l'interaction avec le rayonnement

$$H_{\text{int.}}(t) = - \sum_i \frac{q}{m} \mathbf{A}(\mathbf{R}_i, t) \cdot \mathbf{P}_i + \sum_i \frac{q^2}{2m} \mathbf{A}^2(\mathbf{R}_i, t) - \sum_i g \frac{q}{2m} \mathbf{S}_i \cdot (\nabla_i \wedge \mathbf{A}(\mathbf{R}_i, t)). \quad (\text{VI.2})$$

où le dernier terme représente l'interaction du spin du quanton avec le champ magnétique (hamiltonien de Pauli, Section II.2.2), justifiée empiriquement, et théoriquement comprise — ou tout au moins nécessitée — par l'invariance de jauge, additionnée du principe de linéarisabilité de l'équation de Schrödinger dictant la valeur  $g = 2$  qui convient (presque) parfaitement à l'électron.

## VI.1 Absorption d'un photon

Nous n'avons à vrai dire pas grand chose de nouveau à attendre de l'absorption d'un photon du mode  $\mathbf{k}T$  par les électrons d'un atome, puisque notre théorie quantique du rayonnement a été bâtie (Chapitre II) précisément pour que ses résultats coïncident, dans ce cas, avec ceux du modèle semi-classique, tout au moins au premier ordre de perturbation.

Amorçant l'évolution du système atome-rayonnement en partant un état, dit "initial",

$$|A\rangle |\dots, n_{\mathbf{k}T}, \dots\rangle,$$

très idéalisé<sup>1</sup>, on se demande quelle est la probabilité d'avoir, dans l'état du système au temps  $t$ , l'état, dit "final",

$$|B\rangle |\dots, n_{\mathbf{k}T} - 1, \dots\rangle,$$

où tous les points de suspension sont mis pour des états identiques: l'état final du rayonnement qui nous intéresse ici ne diffère de l'état initial que par un nombre de photons dans le mode  $\mathbf{k}T$  diminué de une unité.

Au temps  $t$  après l'instant initial, la probabilité de l'état final, évaluée au premier ordre de perturbation suivant l'équation (II.58), a pour expression<sup>2</sup>

$$|c_{B, n_{\mathbf{k}T}-1}^{(1)}(t)|^2 = \frac{2\pi t}{\hbar} |\langle B; n_{\mathbf{k}T} - 1 | W | A; n_{\mathbf{k}T} \rangle|^2 \delta(E_B - E_A - \hbar\omega),$$

<sup>1</sup>Un état initial plus réaliste pour des atomes et, surtout, pour le rayonnement, devrait plutôt être décrit par un opérateur densité.

<sup>2</sup>Il s'agit déjà d'une expression effective. J'ai négligé les contributions des autres modes que  $\mathbf{k}T$ , ainsi que la contribution du terme en  $\mathbf{A}^2$ , fort heureusement car toutes ces contributions sont infinies! C'est un procédé de soustraction qui justifie leur élimination au premier ordre des perturbations.

où  $E_A$  et  $E_B$  sont les énergies des états atomiques initial  $|A\rangle$ , et final  $|B\rangle$ , et  $W$  la partie de l'hamiltonien d'interaction dont la dépendance temporelle est en  $e^{-i\omega t}$ . De l'expression (VI.2) de cet hamiltonien, compte tenu du développement modal (III.9), nous ne retenons que le seul terme de  $\mathbf{A} \cdot \mathbf{P}$  qui contribue. Quant à  $\mathbf{A}^2$ , il n'apporte aucune contribution car son terme en  $e^{-i\omega t}$  implique, beaucoup moins probablement, deux photons. Enfin nous négligeons la contribution du terme de spin; nous verrons (page 203) que cette approximation est loisible dans la mesure où l'élément de matrice de la contribution principale n'est pas nul pour cause de sélection. Reste:

$$\begin{aligned} \langle B; n_{\mathbf{k}T} - 1 | W | A; n_{\mathbf{k}T} \rangle \approx \\ - \langle B; n_{\mathbf{k}T} - 1 | \sum_{i=1}^Z \frac{q}{m} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P}_i e^{i\mathbf{k} \cdot \mathbf{R}_i} | A; n_{\mathbf{k}T} \rangle. \end{aligned}$$

L'élément de matrice de l'opérateur d'annihilation vaut  $\sqrt{n_{\mathbf{k}T}}$ , et nous en déduisons le taux d'absorption (II.60):

$$\begin{aligned} \Gamma_{A; n_{\mathbf{k}T} \rightarrow B; n_{\mathbf{k}T} - 1}^{\text{abs}} = \\ \frac{2\pi q^2}{\hbar} \frac{\hbar}{\mathcal{V} 2\varepsilon_0\omega} n_{\mathbf{k}T} \left| \langle B | \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \frac{\mathbf{P}_i}{m} e^{i\mathbf{k} \cdot \mathbf{R}_i} | A \rangle \right|^2 \delta(E_B - E_A - \hbar\omega). \quad (\text{VI.3}) \end{aligned}$$

La désinvolture avec laquelle j'ai commuté les opérateurs  $\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P}_i$  et  $e^{i\mathbf{k} \cdot \mathbf{R}_i}$  n'est qu'apparente: la transversalité du champ, c'est-à-dire ici l'orthogonalité de  $\hat{\mathbf{e}}_{\mathbf{k}T}$  et de  $\mathbf{k}$ , nous en donne le droit. Nous retrouvons là le résultat de la page 58, identique par construction quantique au résultat du modèle semi-classique (II.64) dans la mesure où l'on prend pour amplitude du rayonnement classique

$$A_{\mathbf{k}T} = \sqrt{\frac{\hbar n_{\mathbf{k}T}}{2\varepsilon_0\omega}}, \quad (\text{VI.4})$$

en fonction de  $n_{\mathbf{k}T}$  le nombre initial de photons dans le mode.

## VI.2 Emission d'un photon

Sous son apparence anodine — par sa parenté avec le mécanisme d'absorption de la section précédente —, l'étude de la transition atomique avec émission d'un photon nous réserve une surprise. En nous limitant toujours au sous-espace des états du rayonnement du mode  $\mathbf{k}T$ , partons de l'état initial  $|B; n_{\mathbf{k}T}\rangle$  de l'atome et du rayonnement, et demandons nous quelle est, après le temps  $t$ , la probabilité d'avoir, dans l'état du système, l'état "final"  $|A; n_{\mathbf{k}T} + 1\rangle$ .

En vertu de la règle d'or de Fermi, la probabilité de transition est proportionnelle à  $\delta(E_A - E_B + \hbar\omega)$  et seul le terme de l'hamiltonien d'interaction dont la dépendance temporelle est en  $e^{i\omega t}$  va contribuer à l'élément de matrice de transition. Dans les mêmes hypothèses que pour le processus d'absorption, nous avons cette fois,

$$\begin{aligned} \langle A; n_{\mathbf{k}_T} + 1 | W^+ | B; n_{\mathbf{k}_T} \rangle \approx \\ - \langle A; n_{\mathbf{k}_T} + 1 | \sum_{i=1}^Z \frac{q}{m} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} a_{\mathbf{k}_T}^+ \hat{\mathbf{e}}_{\mathbf{k}_T}^* \cdot \mathbf{P}_i e^{-i\mathbf{k} \cdot \mathbf{R}_i} | B; n_{\mathbf{k}_T} \rangle, \end{aligned}$$

avec

$$\langle n_{\mathbf{k}_T} + 1 | a_{\mathbf{k}_T}^+ | n_{\mathbf{k}_T} \rangle = \sqrt{n_{\mathbf{k}_T} + 1},$$

d'où le taux d'émission,

$$\begin{aligned} \Gamma_{B; n_{\mathbf{k}_T} \rightarrow A; n_{\mathbf{k}_T} + 1}^{\text{ém}} &= \frac{2\pi}{\hbar} \frac{q^2}{\mathcal{V}} \frac{\hbar}{2\varepsilon_0\omega} \\ &\times (n_{\mathbf{k}_T} + 1) \left| \langle A | \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}_T}^* \cdot \frac{\mathbf{P}_i}{m} e^{-i\mathbf{k} \cdot \mathbf{R}_i} | B \rangle \right|^2 \delta(E_A - E_B + \hbar\omega). \quad (\text{VI.5}) \end{aligned}$$

Surprise (annoncée): cette expression ne donne pas tout à fait le même résultat qu'un calcul de la transition  $B \rightarrow A$  induite par le rayonnement classique d'amplitude (VI.4) correspondant à l'état initial à  $n_{\mathbf{k}_T}$  photons. On ne peut trouver des résultats identiques qu'en prenant une amplitude classique proportionnelle à  $\sqrt{n_{\mathbf{k}_T} + 1}$ , plutôt que  $\sqrt{n_{\mathbf{k}_T}}$ , qui ne correspond donc pas à l'état initial du rayonnement.

Cette différence s'estompe en cas de rayonnement intense: lorsque le nombre de photons est élevé,  $\sqrt{n_{\mathbf{k}_T} + 1}$  est équivalent à  $\sqrt{n_{\mathbf{k}_T}}$ . Les traitements quantique et (semi) classique du processus d'émission sont donc équivalents dans la limite des grands nombres de photons (ou des grandes intensités). Notons toutefois qu'à lui seul un état de nombre de photons élevé ne constitue pas en général une limite classique du rayonnement; il faut pour cela un état cohérent (V.27), c'est-à-dire une superposition d'états de nombres de photons différents.

La distinction est en revanche fondamentale lorsque le nombre de photons est faible. En particulier, en absence de rayonnement le taux de transition prédit par le modèle semi-classique est nul, contrairement au résultat de la théorie quantique (VI.5) lorsque  $n_{\mathbf{k}_T} = 0$ . Ainsi, la théorie quantique rend compte, sans aucun effort supplémentaire, du phénomène d'*émission spontanée*. D'un état initial sans photon on peut parvenir à un état final à un photon, nouvelle manifestation de la présence du rayonnement quantique même dans son état fondamental à zéro photon, que nous avons déjà mise en évidence par les fluctuations du point-zéro (§ IV.1). Notre description quantique ne fait finalement aucune distinction entre les phénomènes d'*émission stimulée* (état initial  $n_{\mathbf{k}_T} \neq 0$ ) et d'émission spontanée

( $n_{\mathbf{k}T} = 0$ ), celle-ci n'étant plus qu'un cas particulier de celle-là; il n'y a qu'un seul mécanisme d'émission. Dans la description semi-classique, le rayonnement  $\mathbf{A}(\mathbf{r}, t)$  agit sur les quantons chargés, sans réciprocité: les quantons ne sont pas source du rayonnement. La théorie quantique s'est d'ailleurs historiquement développée sur cette condition pour empêcher l'effondrement de l'atome d'hydrogène classique. Tout au plus, le modèle de Bohr est venu prendre en compte phénoménologiquement la possibilité, empiriquement avérée, d'émission spontanée à partir de niveaux excités. Remarquons de plus que le photon créé par émission stimulée dans le cadre de la théorie quantique du rayonnement est dans le même mode  $\mathbf{k}T$  que les photons — lorsqu'il y en a — de l'état initial; il a mêmes fréquence, direction et polarisation. Encore une fois, avec un rayonnement classique on ne pouvait rendre compte de cette caractéristique que phénoménologiquement en déclarant que le rayonnement émis était cohérent avec le rayonnement incident.

En rayonnement intense ( $n_{\mathbf{k}T}$ , ou  $A_{\mathbf{k}T}$ , grand), les deux descriptions sont équivalentes, une contribution d'un seul photon, qu'il soit émis ou absorbé, ne faisant guère de différence pour le rayonnement. Celui-ci est alors source intarissable et puits insatiable de photons. Quoi qu'il en soit, la méthode de calcul semi-classique donne le résultat correct à condition de l'appliquer,

- dans le cas de l'absorption  $E_A \rightarrow E_B = E_A + \hbar\omega$ , à un potentiel vecteur classique

$$\mathbf{A}_{\text{cl.}}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sqrt{n_{\mathbf{k}T}} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \hat{\mathbf{e}}_{\mathbf{k}T},$$

- et, dans le cas de l'émission  $E_B \rightarrow E_A = E_B - \hbar\omega$ , au potentiel

$$\mathbf{A}_{\text{cl.}}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sqrt{(n_{\mathbf{k}T} + 1)} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \hat{\mathbf{e}}_{\mathbf{k}T}^*,$$

où  $n_{\mathbf{k}T}$  est toujours le nombre de photons de l'état initial du rayonnement.

### VI.3 L'émission spontanée

Le cas particulier de l'émission spontanée mérite une étude plus détaillée, pour plusieurs raisons. Tout d'abord, il s'agit d'un mécanisme visible à longueur de jour: nous verrons que dans la lumière émise par le Soleil, la contribution du  $n_{\mathbf{k}T}$  dans la formule (VI.5) est négligeable par rapport à celle du +1. «L'émission spontanée est si fréquente que de nombreux noms sont associés à ce qui est fondamentalement la même chose. Si les atomes (ou les molécules) sont excités autrement que par chauffage, l'émission spontanée est appelée luminescence. Les lucioles sont luminescentes. Différents noms sont associés à la luminescence selon le mode de production spécifique des atomes excités (l'électroluminescence, la chimioluminescence, *etc.*). Si l'excitation est effectuée par absorption de rayonnement, l'émission spontanée est

appelée phosphorescence. Parfois les molécules ont un état métastable et continuent à luire longtemps après l'extinction du rayonnement excitateur. Cela s'appelle la phosphorescence. Les figurines qui luisent magiquement dans l'obscurité sont phosphorescentes. Les lasers, bien sûr, produisent leur lumière par émission stimulée. Mais lorsqu'un laser est mis en route, les photons qui entament la stimulation sont eux-mêmes produits par émission spontanée. » [55] Ajoutons à ce catalogue l'émission — moins visible mais tout aussi spontanée — des états excités des noyaux, toujours dominante au regard d'une émission difficile à stimuler par les faibles intensités de rayonnement dont on dispose pratiquement à ces énergies/fréquences.

D'un autre côté, par manque de théorie quantique du rayonnement, le calcul de l'émission spontanée n'est généralement pas décrit dans les cours de théorie quantique élémentaires. Seule l'approche statistique est alors présentée, au prix d'un curieux retournement historique car c'est en adoptant ce point de vue qu'Einstein avait prédit l'existence du mécanisme d'émission stimulée qui devait, un demi-siècle plus tard, aboutir au laser. Pour péremptoire qu'il soit, l'argument statistique ne clôt pas la question. L'étudiante en est plus ou moins consciemment satisfaite et attend généralement une description d'un mécanisme détaillé de l'interaction entre matière et rayonnement qui incorpore, de manière cohérente, l'émission spontanée à partir des niveaux excités et l'existence d'un niveau fondamental de la matière qui, lui, n'émet pas. Enfin, l'émission spontanée va nous fournir une illustration des techniques de calcul utilisées pour l'interaction matière-rayonnement.

Pour calculer le processus

$$\left\{ \begin{array}{l} \text{atome } |B\rangle \\ 0 \text{ photon} \end{array} \right\} \longrightarrow \left\{ \begin{array}{l} \text{atome } |A\rangle \\ 1 \text{ photon } \mathbf{k}T \end{array} \right\},$$

reprenons l'expression (VI.5) du taux d'émission qui, dans ce cas  $n_{\mathbf{k}T} = 0$ , donne:

$$\Gamma_{B \rightarrow A + \gamma_{\mathbf{k}T}}^{\text{ém}} = \frac{2\pi}{\hbar} |\langle A | W_{\text{él.}} | B \rangle|^2 \delta(E_A - E_B + \hbar\omega), \quad (\text{VI.6})$$

où  $W_{\text{él.}}$  est la contribution effective de la partie en  $e^{i\omega t}$  de l'hamiltonien d'interaction dans l'espace des états des électrons, à savoir

$$\begin{aligned} W_{\text{él.}} &\stackrel{\text{df}}{=} \langle 1_{\mathbf{k}T} | W^+ | 0 \rangle, \\ &= -\frac{q}{m} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{p}_i e^{-i\mathbf{k} \cdot \mathbf{R}_i}. \end{aligned} \quad (\text{VI.7})$$

Souvenons nous que le symbole  $\delta$  dans l'équation (VI.6) est mis pour une fonction de  $\hbar\omega$  dont la forme était représentée sur la figure II.4. Or, par réalisme, nous devons envisager un détecteur de photons (nous verrons cela dans la section VII.1) dont la résolution en énergie,  $\Delta(\hbar\omega)$ , est finie. Le détecteur est également sensible à

toutes les transitions dans cet intervalle, et compte donc tous les photons dont le mode appartient au domaine  $[\omega, \omega + \Delta\omega]$ , affectés de leurs probabilités respectives. Soit  $\rho(\hbar\omega) d(\hbar\omega)$  le nombre de modes dans l'intervalle infinitésimal  $[\omega, \omega + d\omega]$ . Lors de l'intégration sur le domaine  $\Delta(\hbar\omega)$  réalisée par le détecteur, le pic de la figure II.4 joue le rôle d'une "fonction" de Dirac, dans la mesure où :

- $2\pi\hbar/t \ll \Delta(\hbar\omega)$ , c'est-à-dire si le détecteur a une *bande large* par rapport à la dispersion des fréquences émissibles,
- l'élément de matrice carré,  $|\langle A|W_{\text{él.}}|B\rangle|^2$ , est une fonction de  $\hbar\omega$  dont la variation est douce par rapport à celle du pic.

Par exemple, pour un temps caractéristique des vies moyennes des niveaux atomiques,  $t = 10^{-8}$  s, la largeur du pic est de l'ordre de

$$\frac{\hbar}{t} = \frac{\hbar c}{ct} = \frac{200 \times 10^2 \times 10^6}{3 \times 10^8 \times 10^{15} \times 10^{-8}} = 10^{-7} \text{ eV},$$

et ces conditions sont satisfaites sans difficulté. Après intégration sur la largeur de bande du détecteur, nous obtenons pour le taux total d'émission spontanée :

$$\Gamma_{B \rightarrow A+\gamma}^{\text{ém}} = \frac{2\pi}{\hbar} |\langle A|W_{\text{él.}}|B\rangle|^2 \rho(\hbar\omega), \quad (\text{VI.8})$$

où  $\rho$  est la *densité de modes* du rayonnement — souvent, et improprement, appelée densité d'états — évaluée en  $\hbar\omega = E_B - E_A$ .

- Nous obtenons la densité de modes en comptant le nombre de modes...
- dont la direction est dans un angle solide  $d^2\hat{\mathbf{k}}$  autour de la direction  $\hat{\mathbf{k}}$ ,
  - dont la pulsation est dans l'intervalle  $[\omega, \omega + d\omega]$ ,
  - et dont la polarisation est  $\hat{\mathbf{e}}_{\mathbf{k}T}$ .

Une composante de  $\mathbf{k}$ , par exemple  $k_x$ , n'est pas une variable continue mais prend des valeurs discrètes distantes de  $2\pi/L_x$  où  $L_x$  est la longueur de la boîte à modes (équation (II.49)). Cette discontinuité des valeurs de  $\mathbf{k}$  est compatible avec son assimilation à une variable continue (par exemple lorsque l'on use d'un intervalle infinitésimal  $d\omega$ ) dans la mesure où la boîte est assez grande pour que les fonctions de  $\mathbf{k}$  qui interviennent dans le problème varient peu dans l'intervalle compris entre deux modes successifs. Le nombre de modes dans le pavé  $d^3\mathbf{k}$  est, dans une notation qui anticipe le fait que la densité est indépendante de la direction et de la polarisation :

$$\begin{aligned} \rho(\hbar\omega) d(\hbar\omega) &= \frac{dk_x}{2\pi/L_x} \frac{dk_y}{2\pi/L_y} \frac{dk_z}{2\pi/L_z}, \\ &= \frac{\mathcal{V}}{(2\pi)^3} k^2 dk d^2\hat{\mathbf{k}}, \\ &= \frac{\mathcal{V}}{(2\pi)^3} \left(\frac{\omega}{c}\right)^2 d\left(\frac{\omega}{c}\right) d^2\hat{\mathbf{k}}. \end{aligned}$$

D'où la densité de modes cherchée:

$$\rho(\hbar\omega) = \frac{1}{(2\pi)^3} \frac{\mathcal{V}}{\hbar c^3} \omega^2 d^2\hat{\mathbf{k}}. \quad (\text{VI.9})$$

Le taux d'émission spontanée d'un photon autour de la direction  $\hat{\mathbf{k}}$ , dans l'angle solide  $d^2\hat{\mathbf{k}}$  et avec la polarisation  $\hat{\mathbf{e}}_{\mathbf{k}T}$  s'écrit ainsi:

$$\Gamma_{B \rightarrow A+\gamma}^{\text{ém}} = |\langle A | W_{\text{él.}} | B \rangle|^2 \frac{\mathcal{V}\omega^2}{(2\pi)^2 \hbar^2 c^3} d^2\hat{\mathbf{k}}. \quad (\text{VI.10})$$

Ce taux est évidemment proportionnel à l'angle solide  $d^2\hat{\mathbf{k}}$ , par un facteur, communément appelé *distribution angulaire* des photons, qui peut finalement s'écrire

$$w_{\mathbf{k}T} \stackrel{\text{df}}{=} \frac{1}{2\pi} \frac{q^2/4\pi\epsilon_0}{\hbar c} \omega \left| \langle A | \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \frac{\mathbf{P}_i}{mc} e^{-i\mathbf{k}\cdot\mathbf{R}_i} | B \rangle \right|^2, \quad (\text{VI.11})$$

avec  $\omega = (E_B - E_A)/\hbar$  et  $|\mathbf{k}| = \omega/c$ .

Que la lectrice se rassure: elle n'a pas à se culpabiliser si la relation entre taux de transition — la probabilité d'avoir l'état “final”, après le temps  $t$ , divisée par ce temps — et taux de comptage dans le détecteur ne lui semble pas parfaitement évidente. Bien que l'on n'ait jamais pu trouver de circonstance expérimentale mettant en défaut cette relation, son obtention reste en tout état de cause peu claire: à la question primordiale de la mesure en théorie quantique vient s'ajouter le problème plus technique de la description de la désintégration sur une durée illimitée — c'est-à-dire sans recourir au traitement perturbatif —, problème qui sera effleuré dans la section VI.5.

## VI.4 L'émission dipolaire électrique

Nous pouvons préciser les caractéristiques de l'émission spontanée en recourant à une approximation dont le domaine de validité est quand même considérable. Par définition, l'exponentielle d'opérateur qui figure dans la distribution angulaire (VI.11) est une série:

$$e^{-i\mathbf{k}\cdot\mathbf{R}} \stackrel{\text{df}}{=} 1 - i\mathbf{k}\cdot\mathbf{R} - \frac{1}{2}(\mathbf{k}\cdot\mathbf{R})^2 + \dots \quad (\text{VI.12})$$

Une longueur d'onde typique du rayonnement émis lors d'une transition atomique correspond à la lumière visible, soit  $\lambda \approx k^{-1} \approx 5 \times 10^3 \text{ \AA}$ . D'autre part, la probabilité est faible de trouver des valeurs de distances d'électrons au noyau supérieures à  $1 \text{ \AA}$ . Appliqué au ket initial  $|B\rangle$ , l'opérateur  $\mathbf{k}\cdot\mathbf{R}$  aura une contribution  $5 \times 10^3$  fois plus petite (en norme) que celle de l'opérateur identité 1 de la même série, en sorte que:

$$e^{-i\mathbf{k}\cdot\mathbf{R}} |B\rangle \approx |B\rangle. \quad (\text{VI.13})$$

Cette approximation, dite *de grande longueur d'onde*, est équivalente à prendre la valeur de l'opérateur de champ  $\mathbf{A}(\mathbf{r}, t)$  à l'emplacement du noyau,  $\mathbf{r} = 0$ , où se trouve centrée la distribution de matière, et à négliger ses variations spatiales (pour les modes qui nous concernent) dans le domaine où évoluent les électrons:

$$\left(\mathbf{A}(\mathbf{r}, t)\right)_{\text{effectif}} \approx \mathbf{A}(0, t).$$

Dans le cas de l'émission spontanée d'un noyau, la longueur d'onde typique est donnée par l'énergie de la transition, de l'ordre du MeV,

$$\lambda \approx \frac{c}{\omega} = \frac{\hbar c}{E_B - E_A} = \frac{200}{1} = 200 \text{ fm},$$

tandis que les nucléons ne s'éloignent guère plus que de un fermi. La même approximation est, peut-être, encore justifiée, mais la contribution du premier terme négligé n'est que 200 fois plus petite. Les diverses contributions de la série décroissent plus faiblement; elles se manifesteront plus intensément.

A toutes fins pratiques, la distribution angulaire (VI.11) est donc dominée par le premier terme du développement de l'exponentielle, tout au moins lorsque l'élément de matrice de l'opérateur de transition correspondant, entre les états initial et final, n'est pas nul. Dans cette approximation, dite *dipolaire électrique*, la distribution angulaire a pour expression

$$w_{\mathbf{k}T} = \frac{1}{2\pi} \frac{q^2/4\pi\epsilon_0}{\hbar c} \omega \left| \langle A | \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \frac{\mathbf{P}_i}{mc} | B \rangle \right|^2. \quad (\text{VI.14})$$

Partis d'un état excité  $|B\rangle$  de l'atome, et en absence de tout rayonnement, la probabilité pour avoir, au bout du temps  $t$ , l'atome dans l'état  $|A\rangle$  et un photon d'un mode dirigé dans l'angle solide  $(\mathbf{k}, d^2\hat{\mathbf{k}})$  et polarisé selon  $\hat{\mathbf{e}}_{\mathbf{k}T}$ , vaut:  $t w_{\mathbf{k}T} d^2\hat{\mathbf{k}}$ . Mais n'oublions pas que cette expression de la probabilité n'a de validité que dans le cadre du premier ordre de perturbation, c'est-à-dire durant un laps de temps assez court pour que l'amplitude de l'état initial dans l'état du système ne se soit pas notablement vidée. Les regroupements de facteurs opérés dans l'écriture de la distribution angulaire en font clairement apparaître la dimension: l'inverse d'un temps. En particulier, dans le cas le plus fréquent où les quantons chargés sont des électrons (charge  $q = -|e|$ ), ou des protons à part entière (charge  $q = |e|$ ) plutôt que dotés de charges effectives (comme les protons d'un noyau décrit par un modèle collectif), on reconnaît la constante de structure fine, sans dimension, et la distribution angulaire s'écrit

$$w_{\mathbf{k}T} = \frac{\alpha}{2\pi} \omega \left| \langle A | \sum_{i=1}^Z \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \frac{\mathbf{P}_i}{mc} | B \rangle \right|^2. \quad (\text{VI.15})$$



### VI.4.1 Pourquoi dipolaire électrique?

La lectrice peut légitimement s'interroger sur la raison de ce qualificatif attribué au traitement approximatif (VI.13) d'un mécanisme de transition qui se traduit par la formule (VI.15) dans laquelle on ne distingue ni champ, ni dipôle électrique.

Avant d'éclaircir les origines de cette dénomination, bornons nous au cas d'un seul quanton afin d'alléger l'écriture des formules illustrant toutes les discussions à venir. La distribution angulaire s'écrit alors

$$w_{\mathbf{k}T} = \frac{\alpha}{2\pi} \omega \left| \langle A | \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \frac{\mathbf{P}}{mc} | B \rangle \right|^2, \quad (\text{VI.16})$$

et convient telle quelle (toujours dans le cadre de l'approximation dipolaire électrique) à l'atome d'hydrogène et, à une très bonne approximation liée à la validité du modèle en couches, à un atome dans lequel un seul électron participe effectivement à la transition (cas d'un cœur fermé plus un électron de valence). Nous nous affranchissons ainsi de toute la quincaillerie du "problème à  $N$ -corps" (deuxième quantification, ou coefficients de parenté fractionnelle), rameau de la théorie quantique des champs qui ne nous intéresse pas directement ici.

Calculons le commutateur d'une composante de l'opérateur position de l'électron avec son hamiltonien (équation (VI.1) réduite au cas  $Z = 1$ ):

$$\begin{aligned} [H_{\text{él.}}, X] &= \frac{1}{2m} [P_x^2, X] \\ &= \frac{1}{2m} (P_x [P_x, X] + [P_x, X] P_x) \\ &= -i \frac{\hbar}{m} P_x. \end{aligned}$$

On en déduit l'élément de matrice de la composante correspondante de l'impulsion entre deux états propres du même hamiltonien:

$$\begin{aligned} \langle A | P_x | B \rangle &= i \frac{m}{\hbar} \langle A | (H_{\text{él.}} X - X H_{\text{él.}}) | B \rangle \\ &= i \frac{m}{\hbar} (E_A - E_B) \langle A | X | B \rangle \\ &= -im\omega \langle A | X | B \rangle. \end{aligned} \quad (\text{VI.17})$$

Attention, ceci n'est qu'une égalité entre éléments de matrice particuliers dans laquelle intervient un facteur,  $\omega = (E_B - E_A)/\hbar$ , qui dépend des états et qui interdit absolument de conclure à une équivalence entre opérateurs, du (mauvais) genre " $\mathbf{P} = -im\omega \mathbf{R}$ ". Revenons au développement modal de l'opérateur de champ:

$$\mathbf{A}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r} - \omega t)} + a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k}\cdot\mathbf{r} - \omega t)} \right\},$$

dans lequel seul le terme créant un photon dont l'énergie  $\hbar\omega$  vaut  $E_B - E_A$  contribue à la transition  $|B; n_{\mathbf{k}T}\rangle \rightarrow |A; n_{\mathbf{k}T} + 1\rangle$ . On peut ainsi écrire l'élément de matrice afférent au calcul perturbatif de la transition:

$$\begin{aligned}
 \langle A; n_{\mathbf{k}T} + 1 | -\frac{q}{m} \mathbf{A}(0, t) \cdot \mathbf{P} | B; n_{\mathbf{k}T} \rangle \\
 &= \langle A; n_{\mathbf{k}T} + 1 | iq\omega \mathbf{A}(0, t) \cdot \mathbf{R} | B; n_{\mathbf{k}T} \rangle \\
 &= \langle A; n_{\mathbf{k}T} + 1 | q \frac{\partial \mathbf{A}(0, t)}{\partial t} \cdot \mathbf{R} | B; n_{\mathbf{k}T} \rangle \\
 &= \langle A; n_{\mathbf{k}T} + 1 | -q\mathbf{R} \cdot \mathbf{E}(0, t) | B; n_{\mathbf{k}T} \rangle, \quad (\text{VI.18})
 \end{aligned}$$

en fonction du champ électrique de rayonnement. Ainsi, l'hamiltonien d'interaction a la forme effective  $-\mathbf{P} \cdot \mathbf{E}(0, t)$  qui rappelle la classique énergie potentielle d'interaction d'une distribution de charge dont le comportement lointain est caractérisé par son *moment dipolaire électrique* — l'opérateur  $\mathbf{P} \stackrel{\text{df}}{=} \sum_{i=1}^Z q_i \mathbf{R}_i$ , dans le cas général — avec le champ électrique  $\mathbf{E}(0, t)$  au centre de la distribution, où se trouvent les charges positives qui en annulent le premier moment c'est-à-dire la charge totale.

Dans l'approximation utilisée, il est indifférent de calculer la probabilité de transition avec les perturbations  $-(q/m)\mathbf{A}(0, t) \cdot \mathbf{P}$  ou  $-\mathbf{P} \cdot \mathbf{E}(0, t)$ . Mais, je persiste, ces opérateurs ne sont absolument pas équivalents; seuls leurs éléments de matrice sont égaux...

- entre des états particuliers, à savoir des solutions stationnaires de l'hamiltonien non perturbé satisfaisant la condition de résonance  $E_B = E_A + \hbar\omega$ ,
- dans une jauge particulière pour exprimer le potentiel vecteur, à savoir une jauge de Coulomb.

A propos de ce dernier point, remarquons que, depuis la page 31, nous nous sommes placés dans la jauge de radiation (une jauge de Coulomb particulière). Au cours d'un changement de jauge, le potentiel vecteur est modifié, ainsi que les fonctions-d'onde initiale et finale de l'électron par un facteur de phase local (voir (II.3)) que l'action de  $\mathbf{P}$ , opérateur dérivatif, ne laisse pas indifférent. Il est donc fort rassurant qu'après bien des calculs intermédiaires effectués dans une jauge spécifique, nous soyons finalement parvenus, avec (VI.18), à une expression de l'élément de matrice de transition qui soit clairement indépendante de la jauge — en accord avec notre principe affiché dès le chapitre I —, puisqu'elle n'implique que l'opérateur champ électrique et aucun opérateur dérivatif. La généralisation de la formule fondamentale (VI.18) aux dérivées de tous ordres du potentiel, qui interviennent dans le développement de Taylor de celui-ci au voisinage du centre de la distribution, constitue le *théorème de Siegert*.

La question de l'apparente dépendance de jauge des probabilités de transition calculées perturbativement avec l'hamiltonien d'interaction (VI.2) a déjà fait couler pas mal d'encre. Le choix de jauge adapté au type d'approximation utilisée pour le

calcul d'un processus donné est un problème complexe, sujet à polémiques, trame de deux ouvrages récents [20],[21] qui en explorent tenants et aboutissants.

Cette même question se complique encore dans son avatar relativiste pour la matière: la théorie quantique des champs en général, et l'électrodynamique quantique pour l'interaction qui nous occupe ici. L'invariance de jauge est liée à l'invariance de Lorentz, comme d'ailleurs à l'invariance galiléenne de l'équation de Schrödinger (voir la note page 16). D'autre part, des potentiels  $\mathbf{A}$  et  $\phi$  qui, pour une observatrice, satisfont la condition de jauge de Coulomb deviennent, par transformation de Lorentz, des potentiels  $\mathbf{A}'$  et  $\phi'$ , pour un observateur tout aussi inertiel, mais qui ne satisfont pas la condition de Coulomb. Seule une condition dont la forme est invariante — la condition de jauge de Lorenz<sup>3</sup> — peut être consensuelle. Des calculs qui soient explicitement invariants de Lorentz à tous les stades exigent donc plutôt une formulation dans la jauge de Lorenz. On ne peut plus alors se contenter d'un opérateur de champ correspondant au potentiel vecteur seul. Il faut les opérateurs de champ  $\mathbf{A}$  et  $\phi$ , tous deux indissolublement liés par la condition de jauge et mélangés par les changements de repères inertiels. Le potentiel vecteur — qui n'est plus transverse — et le potentiel scalaire créent des photons de nouveaux types (longitudinal et temporel) que l'on associe au champ coulombien, longitudinal et instantané, même si leur proportion varie au cours d'une transformation de Lorentz ce qui nécessite d'éliminer certains états non physiques au moyen d'une condition supplémentaire.<sup>4</sup> Heureusement, nous verrons que les calculs effectués dans la jauge de Coulomb — dont la signification est plus évidente — conduisent, même s'ils ne sont pas explicitement invariants de Lorentz aux stades intermédiaires, à des résultats physiques finals qui, eux, le sont bien, tout comme nous venons de les trouver indépendants du choix de jauge.

### VI.4.2 Règles de sélection

Revenons à notre taux d'émission dipolaire électrique spontanée que nous pouvons maintenant écrire, d'après (VI.16) et (VI.17),

$$w_{\mathbf{k}T} = \frac{\alpha}{2\pi} \frac{\omega^3}{c^2} |\langle A | \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{R} | B \rangle|^2, \quad (\text{VI.19})$$

où l'opérateur  $\hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{R}$  désigne la somme des trois opérateurs  $X$ ,  $Y$  et  $Z$  multipliés respectivement par des coefficients purement numériques (les trois cosinus directeurs

---

<sup>3</sup>Non, ce n'est pas une coquille. Tels les Dupondt, ont sévis un Lorenz, danois, père des potentiels retardés qui satisfont la condition de jauge du même, et un Lorentz, néerlandais qui, lui, s'est illustré dans une théorie de l'électron et une transformation des coordonnées d'espace et de temps qui laissait invariante la forme des équations de Maxwell. Mieux: la condition de jauge de Lorenz est invariante de Lorentz!

<sup>4</sup>Tout cela parce que la condition de jauge de Lorenz est moins contraignante que celle de Coulomb.

du vecteur de la polarisation  $\hat{\mathbf{e}}_{\mathbf{k}T}$  à laquelle est sensible le détecteur, lorsque la polarisation est rectiligne, sinon ces coefficients sont complexes, voir page 113). On peut trouver des règles de sélection, c'est-à-dire des cas de nullité de l'élément de matrice, en remarquant que l'hamiltonien atomique de l'électron commute avec toutes les composantes du moment angulaire total de l'électron, traduction du fait que, pour l'interaction responsable de la liaison d'un atome isolé, toutes les directions se valent. Il existe donc une base d'états propres  $|\beta jm\rangle$  communs à  $H_{\text{él.}}$ ,  $\mathbf{J}^2$  et  $J_z$ , parmi lesquels se trouvent les états initial et final.

Il est préférable, dans ces conditions, d'exprimer l'opérateur de transition

$$\hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{R} = (\hat{\mathbf{e}}_{\mathbf{k}T}^*)_x X + (\hat{\mathbf{e}}_{\mathbf{k}T}^*)_y Y + (\hat{\mathbf{e}}_{\mathbf{k}T}^*)_z Z$$

en fonction des *composantes sphériques*

$$\begin{aligned} R_{\pm 1}^{(1)} &\stackrel{\text{df}}{=} \mp \frac{1}{\sqrt{2}}(X \pm iY) &= R \sqrt{\frac{4\pi}{3}} Y_{\pm 1}^{(1)}(\hat{\mathbf{R}}) \\ R_0^{(1)} &\stackrel{\text{df}}{=} Z &= R \sqrt{\frac{4\pi}{3}} Y_0^{(1)}(\hat{\mathbf{R}}) \end{aligned} \quad (\text{VI.20})$$

où les  $Y_m^{(l)}$  sont les fonctions harmoniques sphériques. Le théorème de Wigner-Eckart [53, § XIII-32] permet de décomposer l'élément de matrice d'une composante sphérique en produit...

- d'un facteur géométrique (un coefficient 3-j, ou un coefficient de Clebsch-Gordan, et un *élément de matrice réduit* d'harmonique sphérique),
- et d'un facteur spécifique (une intégrale radiale) qui seul dépend de l'interaction de liaison de l'électron dans l'atome (ou plus exactement, qui dépend de ses caractéristiques autres que sa banale invariance par rotation),

à savoir:

$$\begin{aligned} \langle \beta_A j_A m_A | R_q^{(1)} | \beta_B j_B m_B \rangle &= \\ &(-)^{j_A - m_A} \begin{pmatrix} j_A & 1 & j_B \\ -m_A & q & m_B \end{pmatrix} \langle \beta_A j_A \| \mathbf{Y}^{(1)}(\hat{\mathbf{R}}) \| \beta_B j_B \rangle \\ &\times \sqrt{\frac{4\pi}{3}} \int_0^\infty dr r^2 R_{\beta_A}^*(r) r R_{\beta_B}(r). \end{aligned} \quad (\text{VI.21})$$

L'apparition du coefficient 3-j nous suffit pour affirmer, en vertu des règles de couplage des moments angulaires, que les éléments de matrice des composantes de  $\mathbf{R}$  seront nuls à moins que  $j_A$ , 1 et  $j_B$  satisfassent l'inégalité triangulaire:

$$j_A = j_B, |j_B \pm 1|. \quad (\text{VI.22})$$

Ici, dans le cas d'un électron d'un atome (ou d'un proton d'un noyau), la valeur  $j_A = j_B = 0$  est exclue puisque le moment angulaire total est demi-entier. Il va également presque sans dire que la condition

$$m_A = m_B + q \quad (\text{VI.23})$$

doit être vérifiée, ce qui sera automatiquement le cas pour une des trois composantes  $R_q^{(1)}$  si l'inégalité triangulaire est satisfaite.

Il est empiriquement confirmé que l'interaction électromagnétique (responsable de la liaison de l'atome), comme l'interaction forte (qui lie le noyau), sont invariantes par réflexion. A une très bonne approximation qui consiste à négliger la contribution de l'interaction dite (à bon droit ici) faible, dont le caractère distinctif est de ne pas avoir cette invariance, l'hamiltonien de l'électron commute avec l'opérateur de réflexion qui représente cette transformation spatiale dans l'espace des états. L'hamiltonien et l'opérateur de réflexion admettent donc une base d'états propres, communs aussi à  $\mathbf{J}^2$  et  $J_z$ , caractérisés par les valeurs propres de l'opérateur de réflexion que l'exemple de discrétion de ce groupe de transformations — qui n'a que deux éléments — limite à  $\pm 1$ , selon que la fonction-d'onde est paire ou impaire. Revenant — pour les besoins de cette discussion — aux composantes cartésiennes, il est clair que des éléments de matrice du type

$$\langle A|X|B\rangle = \int_{-\infty}^{\infty} dy \int_{-\infty}^{\infty} dz \int_{-\infty}^{\infty} dx \psi_A^*(\mathbf{r}) x \psi_B(\mathbf{r})$$

ne peuvent être non nuls que dans la mesure où les états  $A$  et  $B$  ont des parités opposées. Qu'elle soit purement coulombienne, ou qu'elle dépende du spin par un terme du type  $\mathbf{L} \cdot \mathbf{S}$ , l'interaction subie par l'électron commute — pour raison d'invariance par rotation — avec  $\mathbf{L}^2$ , si ce n'est avec une quelconque de ses composantes. Les états propres de  $H_{\text{él.}}$  peuvent donc être choisis également comme états propres de  $\mathbf{L}^2$  (en même temps que de  $\mathbf{J}^2$ ,  $J_z$ ,  $\mathbf{S}^2$  et de l'opérateur de réflexion) et avoir pour fonctions-d'onde des combinaisons d'harmoniques sphériques  $Y_m^{(l)}(\hat{\mathbf{r}})$  pour une valeur de  $l$  donnée. La parité d'une harmonique sphérique est  $(-)^l$ , d'où l'on conclut à la nullité de la probabilité de transition à moins que  $(-)^{l_A} = -(-)^{l_B}$ . De cette condition de parité, du couplage  $\mathbf{J} = \mathbf{L} + \mathbf{S}$  qui se traduit par  $l = j \pm \frac{1}{2}$  et de la condition triangulaire (VI.22), on déduit la règle de sélection supplémentaire:

$$l_A = |l_B \pm 1|. \quad (\text{VI.24})$$

La lectrice au fait des développements techniques du théorème de Wigner-Eckart [53, app. C] pouvait tout aussi bien arriver à cette conclusion en évaluant directement l'élément de matrice réduit de l'opérateur  $\mathbf{Y}^{(1)}(\hat{\mathbf{R}})$  qui figure dans l'équation (VI.21), compte tenu du fait que cet opérateur n'agit pas dans l'espace des états de spin.

### VI.4.3 Taux de transition

Indépendamment des règles de sélection que nous venons d'établir et qui permettent d'éliminer d'emblée, sans calcul, une multitude de cas pour lesquels la probabilité de transition dipolaire électrique est nulle, nous pouvons sans difficulté pousser un

peu plus loin nos prédictions quantitatives concernant la distribution angulaire des photons émis et la durée de vie du niveau excité initial.

Ecrivons l'élément de matrice de transition intervenant dans la distribution angulaire (VI.19) sous la forme:

$$\langle A | \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{R} | B \rangle = \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \langle A | \mathbf{R} | B \rangle,$$

où  $\langle A | \mathbf{R} | B \rangle$  est mis pour les trois éléments de matrice des composantes de  $\mathbf{R}$ , qu'un choix convenable des phases des vecteurs d'états stationnaires de l'atome permet toujours de rendre réels (grâce à l'invariance de l'hamiltonien atomique par inversion du sens du temps [53, problème XV-P12]). L'ensemble est scalaire (heureusement: le taux de transition ne doit quand même pas dépendre du choix des axes!) et  $\hat{\mathbf{e}}_{\mathbf{k}T}$  est un vecteur; le triplet d'éléments de matrice  $\langle A | X | B \rangle$ ,  $\langle A | Y | B \rangle$  et  $\langle A | Z | B \rangle$  constitue donc lui-même un vecteur de l'espace vital. Pour les transitions entre deux états atomiques  $|B\rangle$  et  $|A\rangle$  donnés, ces trois nombres sont calculables, le vecteur  $\langle A | \mathbf{R} | B \rangle$  est déterminé et, en désignant par  $\theta_T$  l'angle qu'il forme avec  $\hat{\mathbf{e}}_{\mathbf{k}T}$  — la polarisation rectiligne à laquelle le détecteur est supposément réceptif —, on a:

$$\langle A | \hat{\mathbf{e}}_{\mathbf{k}T}^* \cdot \mathbf{R} | B \rangle = |\langle A | \mathbf{R} | B \rangle| \cos \theta_T.$$

Cette relation traduit l'évidence que, pour des états atomiques donnés, la distribution angulaire ne dépend pas de la direction d'émission du photon, ou plutôt n'en dépend que d'une façon indirecte à travers sa direction de polarisation rectiligne seulement.

Mais ce genre de considération est plutôt académique car ce n'est que rarement que l'on décide d'observer l'émission spontanée avec un détecteur sensible à la polarisation, et tout aussi rarement que l'ensemble des atomes sur lesquels on fait l'expérience est préparé dans un seul et unique état excité initial  $|B\rangle$  (cas d'une source polarisée par relaxation dans un champ magnétique externe ou, de façon plus pratique, par le mécanisme de production de l'état initial), et que la détection ne soit sensible qu'à un état final  $|A\rangle$  particulier (cas où l'état  $|A\rangle$  émet lui-même spontanément, en *cascade*, un deuxième photon dont on mesure aussi la polarisation). Si le détecteur réagit à un photon nonobstant la polarisation dudit photon, il suffit d'additionner les probabilités correspondant à deux états de polarisation indépendants. On obtient ainsi (voir la figure VI.1) la distribution angulaire des photons non polarisés d'émission dipolaire électrique spontanée entre les états  $|B\rangle$  et  $|A\rangle$ :

$$\begin{aligned} w_{\mathbf{k}} &= w_{\mathbf{k}_1} + w_{\mathbf{k}_2} \\ &= \frac{\alpha}{2\pi} \frac{\omega^3}{c^2} |\langle A | \mathbf{R} | B \rangle|^2 (\cos^2 \theta_1 + \cos^2 \theta_2) \\ &= \frac{\alpha}{2\pi} \frac{\omega^3}{c^2} |\langle A | \mathbf{R} | B \rangle|^2 \sin^2 \theta. \end{aligned}$$

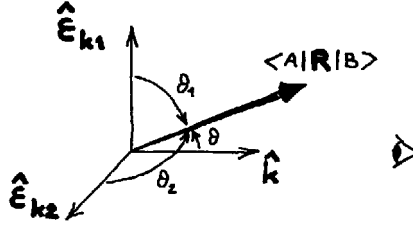


Figure VI.1: La direction  $\hat{\mathbf{k}}$  d'émission des photons, deux polarisations indépendantes, et le vecteur  $\langle A|\mathbf{R}|B\rangle$ .

La probabilité pour que au temps  $t$ , ni trop long ni trop court, l'atome transitant de l'état  $B$  à l'état  $A$ , un photon soit émis dans l'angle solide  $d^2\hat{\mathbf{k}}$  dans la direction  $\hat{\mathbf{k}}$  est  $t w_{\mathbf{k}} d^2\hat{\mathbf{k}}$ . On obtient le taux de transition  $B \rightarrow A$  de l'atome, sans détection du photon, en intégrant sur toutes les directions d'émission:

$$\Gamma_{B \rightarrow A} = \int d^2\hat{\mathbf{k}} w_{\mathbf{k}}.$$

L'élément de matrice  $\langle A|\mathbf{R}|B\rangle$  est caractérisé par les états atomiques mais ne dépend aucunement de la direction d'émission  $\hat{\mathbf{k}}$ . Prenant la direction de  $\langle A|\mathbf{R}|B\rangle$  comme axe  $\hat{\mathbf{z}}$  d'un nouveau trièdre de projection, moyennant

$$\begin{aligned} \int d^2\hat{\mathbf{k}} \sin^2 \theta &= \int d^2\hat{\mathbf{k}} \frac{k_x^2 + k_y^2}{|\mathbf{k}|^2} = \frac{2}{3} \int d^2\hat{\mathbf{k}} \frac{k_x^2 + k_y^2 + k_z^2}{|\mathbf{k}|^2} = \frac{2}{3} \int d^2\hat{\mathbf{k}} \\ &= \frac{8\pi}{3}, \end{aligned}$$

on a immédiatement le taux de transition spontanée:

$$\Gamma_{B \rightarrow A} = \frac{4}{3} \alpha \frac{\omega^3}{c^2} |\langle A|\mathbf{R}|B\rangle|^2. \quad (\text{VI.25})$$

Rappelons qu'il s'agit là d'une expression doublement approchée puisqu'elle résulte du premier ordre de perturbation et de la troncation du développement de l'exponentielle (VI.12) au premier terme!

Pour calculer le taux de transition correspondant à la situation expérimentale la plus fréquente, il est commode de se souvenir qu'il est toujours possible de choisir des états stationnaires qui sont aussi états propres du moment angulaire  $\mathbf{J}^2$  de l'électron et d'une de ses composantes, par exemple  $J_z$ . Partant d'un tel état  $|\beta_B j_B m_B\rangle$ , on ne mesure généralement pas la valeur de  $J_z$  de l'état final. La probabilité correspondante est obtenue en additionnant les probabilités de transition (VI.25) tous  $m_A$

confondus, et de même pour le taux:

$$\Gamma_{\beta_B j_B m_B \rightarrow \beta_A j_A} \stackrel{\text{df}}{=} \sum_{m_A} \Gamma_{\beta_B j_B m_B \rightarrow \beta_A j_A m_A}.$$

Enfin, l'expérience est réalisée non pas avec un seul atome, mais un ensemble d'atomes. S'il n'y a pas de champ magnétique extérieur statique (pas d'effet Zeeman), les différentes valeurs de  $m_B$  sont dégénérées en énergies et statistiquement équiprobables dans notre population d'atomes. Le taux de transition moyen, ramené à un atome de la population, est alors:

$$\Gamma_{\beta_B j_B \rightarrow \beta_A j_A} \stackrel{\text{df}}{=} \frac{1}{2j_B + 1} \sum_{m_B, m_A} \Gamma_{\beta_B j_B m_B \rightarrow \beta_A j_A m_A}, \quad (\text{VI.26})$$

avec

$$\Gamma_{\beta_B j_B m_B \rightarrow \beta_A j_A m_A} = \frac{4}{3} \alpha \frac{\omega^3}{c^2} \left| \langle \beta_A j_A m_A | \mathbf{R} | \beta_B j_B m_B \rangle \right|^2. \quad (\text{VI.27})$$

Il ne reste plus qu'à calculer explicitement l'élément de matrice figurant dans cette expression, pour chacune des composantes de  $\mathbf{R}$  et des valeurs de  $m_A$  et  $m_B$ . Mais la lectrice qui connaît son théorème de Wigner-Eckart est particulièrement aidée en cela car la dépendance en  $m_A$ ,  $m_B$  et les composantes de  $\mathbf{R}$  se trouve factorisée en facteurs géométriques dont la sommation, dans (VI.26), est élémentaire. Je ne détaille pas ce calcul qui nous éloignerait de la théorie quantique du rayonnement, mais qui permet d'obtenir le facteur sans dimension dans l'expression finale du taux moyen en fonction de l'élément de matrice réduit  $\langle \beta_A j_A \| \mathbf{R} \| \beta_B j_B \rangle$ . Le calcul s'achève en prenant les parties radiales explicites des fonctions-d'onde (par exemple  $R_{n_B l_B}(r)$  et  $R_{n_A l_A}(r)$  dans le cas d'un atome d'hydrogène) pour évaluer l'intégrale radiale.

#### VI.4.4 Ordres de grandeur

Les valeurs des éléments de matrice des composantes de  $\mathbf{R}$  qui interviennent dans le taux de transition dipolaire électrique sont de l'ordre de la dimension spatiale  $a$  du système qui subit la transition, et l'on a donc

$$\Gamma \propto \alpha \frac{\omega^3}{c^2} a^2. \quad (\text{VI.28})$$

Les énergies de transition  $\hbar\omega$  sont (en les surestimant) de l'ordre de l'énergie de liaison qui, pour un système comme l'atome principalement lié par une interaction coulombienne centralisatrice, est elle-même de l'ordre de l'énergie potentielle électrostatique (en vertu par exemple du théorème du viriel pour un potentiel en  $1/r$ ):  $\hbar\omega \propto q_e^2/4\pi\epsilon_0 a$ . On a donc  $a \propto \alpha c/\omega$ , et ainsi:

$$\Gamma \propto \alpha^3 \omega = \alpha^3 \frac{c \hbar \omega}{\hbar c}.$$



On en tire une estimation du taux, en (seconde)<sup>-1</sup>, en fonction de l'énergie de la transition  $\hbar\omega$  en Mev, grâce aux valeurs...

- de la constante de structure fine,  $\alpha \approx 1/137$ ,
- du produit de constantes fondamentales  $\hbar c \approx 2 \times 10^2 \text{ MeV fm}$ ,
- et de la constante (dite “vitesse de la lumière”)  $c \approx 3 \times 10^8 \times 10^{15} \text{ fm s}^{-1}$ ,

soit:

$$(\Gamma)_{\text{s}^{-1}} \approx 10^{15} (\hbar\omega)_{\text{MeV}}. \quad (\text{VI.29})$$

Pour des photons émis dont l'énergie va de 1 eV à 1 keV, c'est-à-dire de l'infrarouge au rayonnement X en passant par le visible, on obtient des taux de  $10^9 \text{ s}^{-1}$  à  $10^{12} \text{ s}^{-1}$ , typiques des transitions atomiques.

Pour des énergies de 1 à 10 MeV (autrement dit le rayonnement gamma), caractéristiques des transitions nucléaires, l'expression (VI.25), indépendante de la masse du quanton chargé, qu'il soit électron ou proton, conduit donc à des taux de l'ordre de  $10^{15}$  à  $10^{16} \text{ s}^{-1}$ . Mais l'utilisation de cette expression est alors abusive, car ce n'est pas du tout l'interaction coulombienne qui lie les nucléons constituant le noyau. Il n'y a pas de relation simple entre énergies de transition et rayon du noyau. Nous pouvons néanmoins utiliser la taille connue des noyaux pour en estimer directement les taux de transition dipolaire électrique à partir de la formule (VI.28), soit

$$\Gamma \propto (\hbar\omega)^3 \frac{c}{(\hbar c)^3} a^2,$$

et donc

$$(\Gamma)_{\text{s}^{-1}} \approx 3 \times 10^{14} ((\hbar\omega)_{\text{MeV}})^3 ((a)_{\text{fm}})^2.$$

La taille des noyaux étant de l'ordre du fermi, on retrouve en fait les ordres de grandeur annoncés sur la base de l'interaction coulombienne.

Remarquons que nous n'avons analysé que le cas où un seul quanton chargé participe à la transition, et que l'expression générale du taux de transition dipolaire électrique, déduite de (VI.25), doit comporter une somme des contributions de tous les quantons chargés du système:

$$\Gamma_{B \rightarrow A} = \frac{4}{3} \alpha \frac{\omega^3}{c^2} \left| \langle A | \sum_{i=1}^Z \mathbf{R}_i | B \rangle \right|^2.$$

## VI.5 Largeur naturelle

Ce n'est pas par hasard que je suis parvenu, jusqu'à présent, à éviter de faire allusion à la *durée de vie*, ou *vie moyenne*, d'un niveau atomique. Cette notion découle en effet de la loi exponentielle de désintégration en fonction du temps, observée certes, mais dont nous sommes encore loin de rendre compte théoriquement: partis d'un état initial du système atome-rayonnement,  $|\psi(t=0)\rangle = |B;0\rangle$ , le calcul

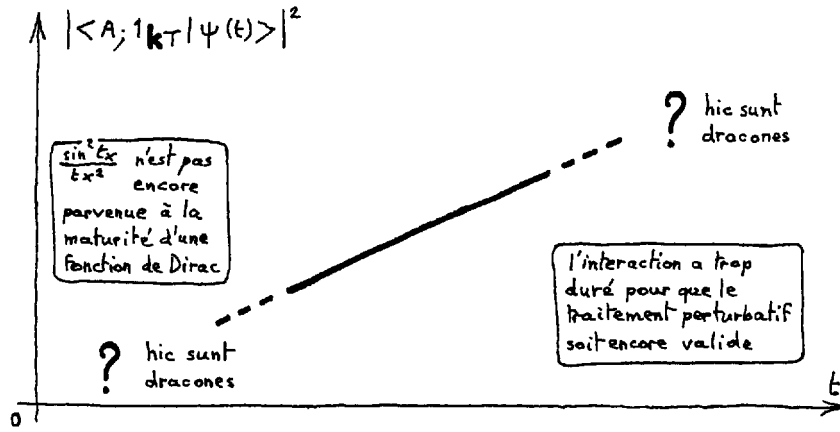


Figure VI.2: La période connue de la composante de  $|\psi(t)\rangle$  sur l'état final  $|A; 1_{\mathbf{k}_T}\rangle$ .

perturbatif ne nous permet d'estimer  $|\psi(t)\rangle$  que sur un laps de temps limité. La figure VI.2 nous rappelle que la probabilité de l'état  $|A; 1_{\mathbf{k}_T}\rangle$  au temps  $t$  n'est linéaire que dans une fenêtre de temps. Au début de l'évolution ( $t$  petit) la courbe de la figure II.4 n'est qu'une caricature bouffie de fonction de Dirac et doit céder au détecteur son rôle de discriminateur d'énergie; la variation de l'élément de matrice de transition à travers cette large fenêtre d'énergie peut ne pas être négligeable et il est impossible de trouver une loi d'évolution temporelle de la probabilité dont la forme soit indépendante de la nature du système et du détecteur. En revanche, après un certain temps ( $t$  grand) l'interaction, même faible, a pu se traduire par des effets importants qui invalident tout traitement perturbatif, fort heureusement d'ailleurs car celui-ci nous prédit une forme linéaire pour une loi qui empiriquement s'avère exponentielle.

### VI.5.1 Estimation

Notre connaissance d'une loi expérimentale quasi universelle nous permet quelques conjectures raisonnées concernant la durée de vie théorique d'un niveau. Dans un modèle à une seule transition (pour simplifier) issue de l'état  $|B\rangle$ , la *probabilité de survie* de l'état initial obtenue par le calcul perturbatif est de la forme:

$$\begin{aligned} |\langle B; 0 | \psi(t) \rangle|^2 &= 1 - \int d^2 \hat{\mathbf{k}} \sum_T |\langle A; 1_{\mathbf{k}_T} | \psi(t) \rangle|^2 \\ &\approx 1 - \Gamma t. \end{aligned}$$

En identifiant cette expression théorique — pour les temps faibles auxquels elle est valide — avec la loi de survie empirique,

$$|\langle B; 0 | \psi(t) \rangle|^2 = e^{-\frac{t}{\tau}} \approx 1 - \frac{t}{\tau},$$

on obtient une estimation semi phénoménologique de la durée de vie,  $\tau = 1/\Gamma$ , et si plusieurs possibilités de transition à des états finals s'offrent à l'état  $|B\rangle$ , la durée de vie de celui-ci s'obtient en additionnant les diverses probabilités de transition pour donner:

$$\tau = \frac{1}{\sum_A \Gamma_{B \rightarrow A}}.$$

Cette cuisine heuristique typique nous rassure en ce qui concerne la compatibilité de l'estimation théorique perturbative avec la loi exponentielle empirique, mais est loin de rassasier la *libido sciendi* de la lectrice avide. Ce sont Weisskopf et Wigner [80] qui ont les premiers montré — sans recourir aux approximations d'un calcul perturbatif — qu'une évolution de l'amplitude de survie de la forme

$$\langle B; 0 | \psi(t) \rangle = e^{-\frac{\Gamma}{2}t}, \quad (\text{VI.30})$$

était une solution viable des équations d'évolution d'un modèle à deux niveaux atomiques couplés par le rayonnement quantique.

Avant de passer à la description technique du modèle de Weisskopf et Wigner, il convient de remarquer que notre état “initial”,  $|B; 0\rangle$ , n'est pas un état stationnaire de la nature. Il n'est stationnaire que par rapport à un hamiltonien atomique idéal, sans interaction avec le rayonnement. Cette approximation — utilisée pour l'analyse de l'atome d'hydrogène dans les cours élémentaires de théorie quantique — est satisfaisante dans la mesure où la perturbation apportée par l'indébranchable interaction avec le rayonnement est faible (parce que la constante de structure fine  $\alpha$ , telle qu'elle apparaît par exemple dans le taux de transition (VI.19), est bien inférieure à l'unité), mais l'histoire ne peut se terminer là; les états excités, prétendument stationnaires, peuvent se désexciter.

La situation est en cela tout à fait similaire à celle du pendule oscillant classique. On commence par l'étude d'un pendule idéal sans freinage, dont les équations admettent des solutions stationnaires: les solutions oscillantes dont la pulsation est caractéristique du pendule idéal et dont les amplitudes et phases sont déterminées par les conditions initiales. Mais le pendule de la nature est toujours soumis à d'autres interactions (le freinage de l'air, les frottements au point de suspension). Chacun sait que les mouvements stationnaires précédents ne représentent que grossièrement les mouvements réels. C'est d'ailleurs la pertinence de ces solutions périodiques idéales, sans fin, qui est la plus difficile à faire admettre à un profane; en ce sens, le modèle le plus simple est souvent aussi le plus abstrait. Mais lorsque les

interactions de freinage sont faibles les mouvements stationnaires sont quand même proches des mouvements réels. On peut tenir compte phénoménologiquement de cette modification par un facteur exponentiel décroissant, ou analyser les interactions supplémentaires, écrire des équations du mouvement correspondantes, et en chercher effectivement les solutions.

Retenons de ces considérations que des états stationnaires pour l'hamiltonien idéal d'un premier modèle (pesanteur pour le pendule oscillant, interaction coulombienne pour l'atome d'hydrogène) ne le sont pas tout à fait pour la nature ou, plus modestement, pour l'hamiltonien d'un modèle amélioré (interaction avec l'air, ou suspension manquant de souplesse, pour le pendule, interaction avec le rayonnement électromagnétique pour l'électron de l'atome). Cette analogie du pendule oscillant a été proposée par Fermi, dans un article historique qui garde toute sa valeur pédagogique [27], pour illustrer concrètement l'évolution de la matière et du rayonnement couplés. Le pendule (la matière) et une corde vibrante (le rayonnement dans une boîte) indépendants ont leurs fréquences d'oscillations stationnaires (l'énergie de l'état excité de la matière et les fréquences modales du rayonnement). Qu'ils viennent à être couplés, faiblement, par un fil élastique (le couplage matière-rayonnement par la charge électrique) et le pendule peut se désexciter spontanément au profit d'une corde qui était initialement au repos. On peut retrouver aussi dans cet analogue mécanique les processus d'absorption ou d'émission stimulée.

Un produit de kets d'états stationnaires pour deux systèmes pris séparément n'est évidemment pas un ket d'état stationnaire pour ces deux systèmes lorsqu'ils interagissent. Dès lors que l'état initial (ou pris comme tel)  $|B;0\rangle$  n'est pas stationnaire pour  $H_{\text{él}} + H_{\text{int}}$  (en représentation d'interaction pour le rayonnement), ce n'est pas un état propre; il lui est associé une distribution des valeurs de l'énergie, caractérisée par une dispersion, appelée *largeur du niveau* initial, qui se retrouve dans la distribution de la probabilité d'émission non plus monochromatique mais elle-aussi caractérisée par cette largeur. La dispersion doit être d'autant plus grande que l'interaction est forte, c'est-à-dire que le taux de transition  $\Gamma$  est élevé. Pour de simples raisons dimensionnelles il faut donc s'attendre à ce que les diverses largeurs soient du même ordre:

$$(\Delta E)_{\text{él}} \propto (\Delta \hbar\omega)_{\text{ray}} \propto \hbar\Gamma.$$

La justification formelle de cette relation va nous être apportée par l'analyse de Weiskopf et Wigner qui, de l'hypothèse d'une amplitude de survie de la forme  $e^{-\Gamma t/2}$ , conduit à une forme de raie lorentzienne dont la largeur à mi-hauteur vaut, en fait, précisément  $\hbar\Gamma$ .

### VI.5.2 Le modèle de Weisskopf et Wigner

Considérons d'abord le cas, simplifié à l'extrême, d'un atome dont l'espace des états est réduit à être sous-tendu par deux états stationnaires: le fondamental  $A$  et un état excité  $B$ . Dans la mesure où la transition  $B \rightarrow A$  avec émission spontanée d'un photon est permise, nous nous permettons de négliger les processus d'émission de deux (ou plus) photons. Ainsi, partis de l'état excité de l'atome en absence de tout rayonnement, nous avons pour le ket d'état du système atome-rayonnement quantique, sous une présentation commode pour la suite des opérations,

$$|\psi(t)\rangle = c_{B0}(t) e^{-i\frac{E_B t}{\hbar}} |B; 0\rangle + \sum_{\mathbf{k}T} c_{A\mathbf{k}T}(t) e^{-i\frac{E_A t}{\hbar}} |A; 1_{\mathbf{k}T}\rangle. \quad (\text{VI.31})$$

L'équation d'évolution du ket d'état est, en représentation d'interaction,

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = (H_{\text{el}} + H_{\text{int}}(t)) |\psi(t)\rangle, \quad (\text{VI.32})$$

mais, pour des kets limités au type ci-dessus, la partie effective de l'hamiltonien d'interaction — celle dont les éléments de matrice entre états stationnaires figurant dans le développement (VI.31) sont non nuls — est, d'après (VI.2) et (III.9), de la forme

$$H_{\text{int. eff.}} = \sum_{\mathbf{k}T} (W_{\mathbf{k}T} e^{-i\omega t} + W_{\mathbf{k}T}^+ e^{i\omega t}) \quad (\text{VI.33})$$

avec par exemple, en cas d'approximation dipolaire électrique,

$$W_{\mathbf{k}T} = -\frac{q}{m} \sqrt{\frac{\hbar}{2\varepsilon_0 \omega \mathcal{V}}} a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \sum_{i=1}^Z \mathbf{P}_i e^{i\mathbf{k} \cdot \mathbf{R}_i}, \quad (\text{VI.34})$$

ou toute autre écriture loisible. Les opérateurs  $W$  et  $W^+$ , “à un photon”, n'ont pas d'éléments de matrice diagonaux dans les états stationnaires de base. On en déduit les équations couplées régissant l'évolution des amplitudes des états stationnaires dans l'état  $|\psi(t)\rangle$ :

$$\begin{aligned} i\hbar \dot{c}_{B0} &= \sum_{\mathbf{k}T} \langle B; 0 | W_{\mathbf{k}T} | A; 1_{\mathbf{k}T} \rangle e^{i(\omega_0 - \omega)t} c_{A1_{\mathbf{k}T}}, \\ i\hbar \dot{c}_{A1_{\mathbf{k}T}} &= \langle A; 1_{\mathbf{k}T} | W_{\mathbf{k}T}^+ | B; 0 \rangle e^{-i(\omega_0 - \omega)t} c_{B0}, \end{aligned}$$

après avoir posé  $\omega_0 \stackrel{\text{df}}{=} (E_B - E_A)/\hbar$ .

Empiriquement, la probabilité de survie de l'état initial,  $|\langle B; 0 | \psi(t) \rangle|^2$ , est une exponentielle décroissante, ce qui nous suggère d'essayer, en accord avec la condition initiale  $c_{B0}(0) = 1$ , une solution de la forme

$$c_{B0}(t) = e^{-\frac{\gamma}{2}t}, \quad (\text{VI.35})$$

où  $\gamma$  est une constante *a priori* complexe dont la partie réelle doit être positive, flanquée d'un facteur  $1/2$  dont la raison esthétique se dévoilera bientôt. Le rôle joué par cette forme exponentielle dans l'expression (VI.31) conduit parfois à définir la quantité

$$E'_B \stackrel{\text{df}}{=} E_B - i\hbar \frac{\gamma}{2},$$

surnommée *énergie complexe* de l'état  $B$  en raison du facteur temporel  $e^{-iE'_B t/\hbar}$  maintenant associé à l'état de base  $|B; 0\rangle$  du fait de l'interaction avec le rayonnement. Sa partie imaginaire,  $\Im E'_B = -(\hbar/2) \Re \gamma$  devra bien entendu être négative si l'on veut obtenir un comportement d'exponentielle décroissante lorsque  $t$  croît.

Par substitution dans la deuxième équation couplée, et après intégration définie par la condition initiale  $c_{A1_{\mathbf{k}_T}}(0) = 0$ , on obtient la solution:

$$i\hbar c_{A1_{\mathbf{k}_T}}(t) = \langle A; 1_{\mathbf{k}_T} | W_{\mathbf{k}_T}^+ | B; 0 \rangle \frac{e^{i(\omega - \omega_0 + i\frac{\gamma}{2})t} - 1}{i(\omega - \omega_0 + i\frac{\gamma}{2})}. \quad (\text{VI.36})$$

Nous avons maintenant la possibilité de nous affranchir de l'hypothèse qui restreignait le calcul perturbatif; après un temps infini, c'est-à-dire bien supérieur à  $1/\Re \gamma$ , on a:

$$|c_{A1_{\mathbf{k}_T}}(\infty)|^2 = \frac{1}{\hbar^2} \frac{|\langle A; 1_{\mathbf{k}_T} | W_{\mathbf{k}_T}^+ | B; 0 \rangle|^2}{(\omega - \omega_0 - \frac{1}{2} \Im \gamma)^2 + \frac{1}{4} (\Re \gamma)^2}. \quad (\text{VI.37})$$

Nous sommes ainsi en possession d'une expression pour la probabilité d'avoir un photon du mode  $\mathbf{k}_T$ , ou  $\omega$  (et non pas  $\omega_0$ ),  $\hat{\mathbf{k}}$ , et  $T$ , lorsque l'atome a certainement effectué la transition ( $c_{B0}(\infty) = 0$ ), autrement dit lorsqu'on a certainement un photon, quel qu'il soit.

Le report des solutions (VI.35) et (VI.36) dans la première équation couplée nous donne enfin une condition nécessaire pour la constante  $\gamma$  (jusqu'alors arbitraire),

$$-i\hbar \frac{\gamma}{2} e^{-\frac{\gamma}{2}t} = \sum_{\mathbf{k}_T} \langle B; 0 | W_{\mathbf{k}_T} | A; 1_{\mathbf{k}_T} \rangle \frac{1}{i\hbar} \langle A; 1_{\mathbf{k}_T} | W_{\mathbf{k}_T}^+ | B; 0 \rangle \frac{e^{-\frac{\gamma}{2}t} - e^{i(\omega_0 - \omega)t}}{i(\omega - \omega_0 + i\frac{\gamma}{2})},$$

ou, sous forme intégrale, lorsque le volume de la caisse à modes devient suffisant (souvenez vous de (II.50)):

$$i\hbar \frac{\gamma}{2} = \frac{1}{\hbar} \frac{\mathcal{V}}{(2\pi)^3} \int_0^\infty dk k^2 \int d^2 \hat{\mathbf{k}} \sum_T |\langle B; 0 | W_{\mathbf{k}_T} | A; 1_{\mathbf{k}_T} \rangle|^2 \frac{1 - e^{i(\omega_0 - \omega - i\frac{\gamma}{2})t}}{\omega - \omega_0 + i\frac{\gamma}{2}}.$$

Autrement dit, la constante  $\gamma$  doit être solution de l'équation intégrale:

$$\gamma = \frac{2i}{\hbar} \int_0^\infty d\omega f(\omega) \frac{1 - e^{i(\omega_0 - \omega - i\frac{\gamma}{2})t}}{\omega - \omega_0 + i\frac{\gamma}{2}}, \quad (\text{VI.38})$$

où

$$f(\omega) \stackrel{\text{df}}{=} \frac{\omega^2 \mathcal{V}}{(2\pi)^3 \hbar c^3} \int d^2 \mathbf{k} \sum_T |\langle B; 0 | W_{\mathbf{k}T} | A; 1_{\mathbf{k}T} \rangle|^2 \quad (\text{VI.39})$$

est une fonction réelle positive de l'énergie du photon émis, qui n'a aucune raison de présenter des singularités.

Dans tous les cas de désexcitation d'un niveau par interaction avec le rayonnement électromagnétique, la durée de vie  $1/\Re \gamma$  du niveau excité est en fait beaucoup plus longue que la période  $2\pi/\omega_0$  du rayonnement émis (une manifestation de la faiblesse de l'interaction électromagnétique ou, autrement dit, de  $\alpha \ll 1$ ). Pour un atome d'hydrogène par exemple, les durées de vie des niveaux sont de l'ordre de  $10^{-8}$  s, tandis qu'aux énergies de transition de l'ordre du Rydberg (13,6 eV) correspondent des périodes de rayonnement:

$$\frac{2\pi}{\omega_0} = 2\pi \frac{\hbar c}{\hbar \omega_0} \frac{1}{c} \approx 2\pi \frac{2 \times 10^8}{13,6} \frac{1}{3 \times 10^8 \times 10^{15}} \approx 10^{-16} \text{ s}.$$

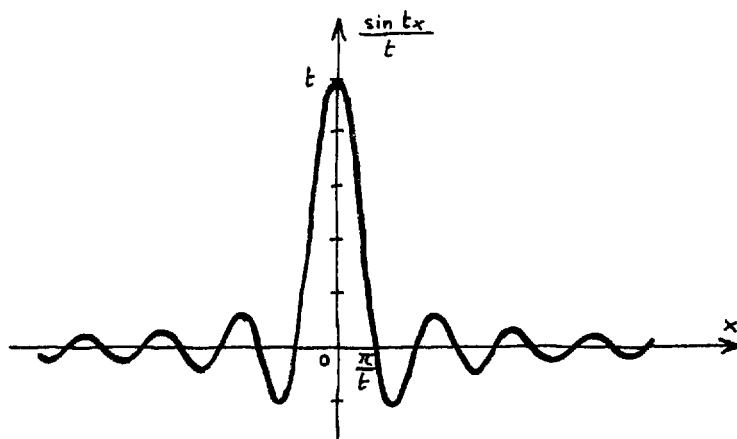
La partie imaginaire de  $\gamma$  peut être absorbée dans une redéfinition de la pulsation  $\omega_0$  de la transition, mais elle est, en tout état de cause, petite... généralement après soustraction d'un infini par une opération de renormalisation! Dans ces conditions, nous sommes à même d'obtenir une première approximation de la solution de l'équation intégrale (VI.38) en remplaçant, dans le second membre de celle-ci, la (petite) valeur inconnue  $\gamma$  par zéro, soit:

$$\gamma \approx \frac{2i}{\hbar} \int_0^\infty d\omega f(\omega) \frac{1 - e^{i(\omega_0 - \omega)t}}{\omega_0 - \omega}. \quad (\text{VI.40})$$

La validité du procédé repose sur la continuité du terme intégral de (VI.38) en  $\gamma = 0$ , continuité pas évidente mais qui sera confirmée *a posteriori* par le fait que nous trouverons effectivement, après quelques contorsions, une valeur pour l'intégrale (VI.40). Remarquons d'entrée que la décomposition

$$\frac{1 - e^{i(\omega_0 - \omega)t}}{\omega_0 - \omega} = \frac{1 - \cos(\omega_0 - \omega)t}{\omega_0 - \omega} - i \frac{\sin(\omega_0 - \omega)t}{\omega_0 - \omega}$$

met en évidence des intégrales de fonctions de  $\omega$  qui oscillent rapidement lorsque  $t \rightarrow \infty$ , et dont les contributions, sauf au point  $\omega = \omega_0$  singularisé par le dénominateur, sont alors négligeables. Ainsi, le deuxième terme n'est autre qu'une représentation de la "fonction" de Dirac à un facteur près, égal à l'aire "sous" la courbe  $\sin x/x$ ,

Figure VI.3: Le graphe de la fonction  $\frac{\sin tx}{x}$ .

tandis que le premier terme, nul en  $\omega = \omega_0$  grâce à son numérateur, a ailleurs un rôle équivalent à  $1/(\omega_0 - \omega)$  et va extraire une *partie principale*. La validité de notre tentative de solution (VI.35) exige que l'intégrale (VI.40) soit indépendante du temps, ce qui n'a donc de chance d'être le cas que lorsque  $t \rightarrow \infty$ , et est l'indice que la loi de désintégration n'est effectivement exponentielle que pour des valeurs de temps élevées.

Nous allons évaluer (VI.40) en calculant la limite, lorsque  $t$  croît, de l'intégrale apparentée:

$$\int_{-a}^b dx f(x) \frac{1 - e^{itx}}{x} = \int_{-a}^b dx f(x) \frac{1 - \cos tx}{x} + i \int_{-a}^b dx f(x) \frac{\sin tx}{x}, \quad (\text{VI.41})$$

sur un domaine  $[-a, b]$  qui comprend le point  $x = 0$ . Le calcul est fondé sur le fait que

$$\lim_{t \rightarrow \infty} \int_{-a}^b dx g(x) \sin tx = \lim_{t \rightarrow \infty} \int_{-a}^b dx g(x) \cos tx = 0,$$

pour toute fonction  $g(x)$ , continue bien sûr.<sup>5</sup>

La fonction  $\sin tx/x$  — analogue en cela de la fonction déjà rencontrée dans la figure II.4 — présente un maximum de plus en plus aigu lorsque  $t$  croît, tout en recouvrant une aire constante (voir figure VI.3). Ce facteur finit donc par sélectionner,

---

<sup>5</sup>Si la chose vous préoccupe, vous pourrez retrouver la démonstration de cette propriété, "physiquement évidente", dans les prémisses de n'importe quel cours sur les transformées de Fourier.



dans la partie imaginaire de l'intégrale (VI.41), la valeur de  $f(x)$  au point  $x = 0$ , et l'on a :

$$\int_{-a}^b dx f(x) \frac{\sin tx}{x} \underset{t \rightarrow \infty}{\sim} f(0) \int_{-a}^b dx \frac{\sin tx}{x},$$

avec

$$\int_{-a}^b dx \frac{\sin tx}{x} = \int_{-at}^{bt} dz \frac{\sin z}{z} \xrightarrow{t \rightarrow \infty} \int_{-\infty}^{\infty} dz \frac{\sin z}{z}.$$

On ne change pas la valeur de cette dernière intégrale en ajoutant à l'intégrand une fonction impaire et l'on a donc

$$\int_{-\infty}^{\infty} dz \frac{\sin z}{z} = \frac{1}{i} \int_{-\infty}^{\infty} dz \frac{e^{iz}}{z},$$

qui s'évalue sans difficulté en passant dans le plan complexe, le long du sentier rebattu,  $\Gamma$ , indiqué sur la figure II.5. (Le lemme de Jordan n'est pas applicable sur le demi cercle  $\mathcal{C}$ , mais le calcul explicite de l'intégrale permet de la borner par une intégrale qui tend vers zéro lorsque le rayon  $R$  croît.) Il vient

$$\int_{-\infty}^{\infty} dz \frac{\sin z}{z} = -\frac{1}{i} \lim_{\epsilon \rightarrow 0} \int_c dz \frac{e^{iz}}{z} = -\frac{1}{i} \lim_{\epsilon \rightarrow 0} \int_c dz \left( \frac{1}{z} + i + \dots \right) = \pi,$$

et donc :

$$\int_{-a}^b dx f(x) \frac{\sin tx}{x} \xrightarrow{t \rightarrow \infty} \pi f(0).$$

La partie réelle de l'intégrale (VI.41), à savoir

$$\int_{-a}^b dx f(x) \frac{1 - \cos tx}{x}, \quad (\text{VI.42})$$

s'évalue pour sa part en décomposant le domaine d'intégration en trois zones délimitées par  $-a$ ,  $-\epsilon$ ,  $\epsilon$  et  $b$ . Dans la zone centrale, on a :

$$\int_{-\epsilon}^{\epsilon} dx f(x) \frac{1 - \cos tx}{x} = \int_{-\epsilon t}^{\epsilon t} dz f\left(\frac{z}{t}\right) \frac{1 - \cos z}{z},$$

et, dans la mesure où  $\epsilon$  tend vers zéro plus rapidement que  $1/t$ , le recours aux développements limités de  $f(z/t)$  et  $\cos z$  au voisinage de zéro montre que cette intégrale est équivalente à  $\frac{1}{3} f'(0) \epsilon^3 t^2$ ; elle tend donc vers zéro. Ne reste à évaluer que la limite, lorsque  $\epsilon$  tend vers zéro, de

$$\int_{-a}^{-\epsilon} + \int_{\epsilon}^b dx f(x) \frac{1 - \cos tx}{x},$$

appelée *partie principale* ( $\mathcal{PP}$ ) de l'intégrale (VI.42). Nous avons, dans le cas  $a < b$ :

$$\begin{aligned} \mathcal{PP} \int_{-a}^b dx f(x) \frac{1 - \cos tx}{x} = \\ \mathcal{PP} \int_{-a}^a dx f(x) \frac{1 - \cos tx}{x} + \int_a^b dx \frac{f(x)}{x} - \int_a^b dx \frac{f(x)}{x} \cos tx. \end{aligned}$$

La fonction  $f(x)/x$  étant continue dans l'intervalle  $[a, b]$ , la dernière intégrale tend vers zéro lorsque  $t$  croît. Avec son domaine d'intégration symétrique, l'intégrale dont il reste à prendre la partie principale n'est plus sensible qu'à la composante impaire de la fonction  $f$ :

$$f_i(x) \stackrel{\text{df}}{=} \frac{1}{2}(f(x) - f(-x)).$$

En particulier:

$$\begin{aligned} \mathcal{PP} \int_{-a}^a dx f(x) \frac{\cos tx}{x} &= \mathcal{PP} \int_{-a}^a dx f_i(x) \frac{\cos tx}{x} = 2 \lim_{\epsilon \rightarrow 0} \int_{\epsilon}^a dx \frac{f_i(x)}{x} \cos tx \\ &= 2 \int_0^a dx \frac{f_i(x)}{x} \cos tx, \end{aligned}$$

quelle que soit la manière dont  $\epsilon$  tend vers zéro, puisque  $f_i(x)/x$  est régulière à l'origine. Cette dernière intégrale tend vers zéro lorsque  $t$  croît, et l'on a finalement

$$\int_{-a}^b dx f(x) \frac{1 - e^{itx}}{x} \xrightarrow{t \rightarrow \infty} \mathcal{PP} \int_{-a}^b dx \frac{f(x)}{x} - i\pi f(0),$$

ou, symboliquement,

$$\frac{1 - e^{itx}}{x} \xrightarrow{t \rightarrow \infty} \mathcal{PP} \frac{1}{x} - i\pi \delta(x).$$

Dans notre cas, admettant au bénéfice du doute que le bon chic bon genre de notre fonction (VI.39) s'étende jusqu'à l'infini, nous avons donc

$$\gamma \approx \frac{2\pi}{\hbar} f(\omega_0) + \frac{2i}{\hbar} \mathcal{PP} \int_0^\infty d\omega \frac{f(\omega)}{\omega_0 - \omega}. \quad (\text{VI.43})$$

Nous sommes (pour l'instant) au bout de nos peines avec cette expression de  $\gamma$  qui a été obtenue ici en prenant comme guide Heitler [35, § 18] dont les hypothèses sont moins nombreuses et plus claires que celles du calcul original de Weisskopf [80].

La probabilité de survie de l'état initial  $B$  correspondant à notre solution (VI.35) a pour expression  $|c_{B0}(t)|^2 = e^{-\Re \gamma t}$ , dans laquelle nous trouvons la durée de vie de l'état:  $\tau = 1/\Re \gamma$ . Quelle n'est pas alors notre joie, au vu de (VI.43) et (VI.39), de reconnaître précisément, dans cette partie réelle de  $\gamma$  qui nous intéresse

tant, la somme — sur toutes les directions et polarisations — des taux d'émissions spontanées issues de  $B$  précédemment calculés en (VI.10) dans l'approximation de la règle d'or de Fermi, c'est-à-dire durant un temps assez court pour légitimer un calcul perturbatif. Nous avons donc  $\Re \gamma = \Gamma$  et, pour ce qui est de la durée de vie du niveau  $B$ , le résultat subodoré au début de cette section:

$$\tau = \frac{1}{\Gamma}. \quad (\text{VI.44})$$

Revenons à la probabilité (VI.37) pour qu'un photon ayant été émis, il ait la pulsation  $\omega$ , la direction  $\hat{\mathbf{k}}$  et la polarisation  $T$ . La partie imaginaire de  $\gamma$  peut être absorbée dans une redéfinition de la pulsation  $\omega_0$  de la transition. Mais, grâce à l'absolue faiblesse de l'interaction électromagnétique (la constante de structure fine  $\alpha$  est inférieure à  $10^{-2}$ ), nous sommes maintenant dans la situation pratique où  $\Re \gamma \ll \omega_0$ , ce qui entraîne  $\Re \gamma = \Gamma$ . Cette faiblesse de l'interaction a aussi pour conséquence  $|\Im \gamma| \ll \omega_0$ , et même  $|\Im \gamma| \ll \Re \gamma$  (tout comme la correction, du second ordre, apportée à la fréquence propre d'un pendule par l'amortissement); nous pouvons pour cette raison nous dispenser, sauf préoccupation particulière, de la redéfinition de la pulsation de la transition. Ainsi, notre probabilité vaut:

$$\begin{aligned} \Pr(\mathbf{k}, T) &= |c_{A1_{\mathbf{k}T}}(\infty)|^2, \\ &\approx \frac{1}{\hbar^2} |\langle B; 0 | W_{\mathbf{k}T} | A; 1_{\mathbf{k}T} \rangle|^2 \frac{1}{(\omega - \omega_0)^2 + \frac{1}{4}\Gamma^2}. \end{aligned}$$

En tant que fonction de  $\omega$ , la *lorentzienne*  $((\omega - \omega_0)^2 + \frac{1}{4}\Gamma^2)^{-1}$  est quasi nulle partout, à l'exception d'un pic effilé dont la largeur à mi-hauteur est  $\Gamma$  (voir la figure VI.4). Dans l'étroit domaine de ce pic, le comportement de l'élément de matrice de transition en fonction de  $\omega$  n'a aucune raison de ne pas être débonnaire, et nous pouvons aussi bien remplacer ledit élément de matrice par sa valeur calculée pour la pulsation  $\omega_0$ . Passant à une boîte à modes assez grosse pour que la distribution des pulsations  $\omega$  ait effectivement la continuité implicitement supposée dans la représentation de la figure VI.4, la probabilité pour que la pulsation du photon soit dans l'intervalle  $[\omega, \omega + d\omega]$ , toutes directions et polarisations confondues, est:

$$d\Pr = \frac{1}{\hbar^2} \frac{\mathcal{V}}{(2\pi)^3} \frac{\omega^2 d\omega}{c^3} \int d^2\hat{\mathbf{k}} \sum_T |\langle B; 0 | W_{\mathbf{k}T} | A; 1_{\mathbf{k}T} \rangle|_{\omega_0}^2 \frac{1}{(\omega - \omega_0)^2 + \frac{1}{4}\Gamma^2}.$$

C'est cette somme sur les directions et polarisations qui nous intéresse si nous observons simplement l'intensité du rayonnement émis par un ensemble d'atomes identiques non orientés. Mais, pour obtenir l'*intensité spectrale*  $I(\omega)$ , encore faut-il multiplier la probabilité du photon par son énergie:  $I(\omega) d\omega \stackrel{\text{df}}{=} \hbar\omega d\Pr$ . Là aussi, le facteur  $\omega^3$  qui apparaît reste pratiquement égal à  $\omega_0^3$  sur le domaine du pic et nous

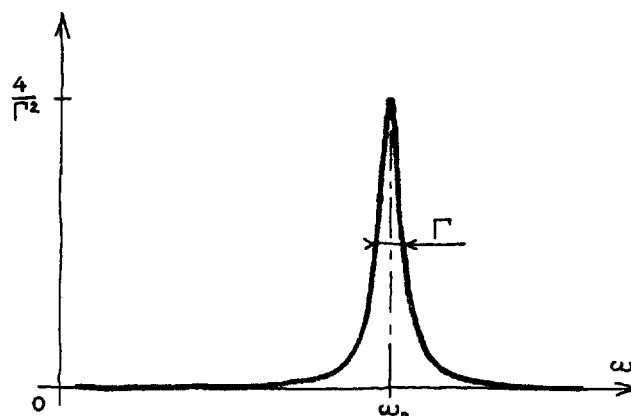


Figure VI.4: Le graphe de la lorentzienne  $\frac{1}{(\omega - \omega_0)^2 + \frac{1}{4}\Gamma^2}$ .

retrouvons, dans l'expression de l'intensité spectrale, notre vieille connaissance le taux total de transition  $\Gamma$ , en fonction duquel nous pouvons finalement écrire:

$$I(\omega) = \frac{\Gamma}{2\pi} \frac{\hbar\omega_0}{(\omega - \omega_0)^2 + \frac{1}{4}\Gamma^2}. \quad (\text{VI.45})$$

Ainsi, le spectre d'émission spontanée n'est plus monochromatique, la raie d'émission s'est épaissie et présente une largeur à mi hauteur égale au taux de transition. Le maximum de l'intensité se trouve à la pulsation  $\omega_0$ , pratiquement donnée par la différence d'énergie des deux états de l'atome éventuellement corrigée du petit décalage  $\Im m \gamma$ .

La démarche suivie pour aboutir finalement au profil de raie (VI.45) comporte une bonne dose d'empirisme. Partis de ce que nous voulions trouver, la loi exponentielle de désintégration (VI.35), nous avons dû, pour parvenir à nos fins, supposer le temps infini. Une analyse théorique plus rigoureuse[53, § XXI-13] montre effectivement la validité de la loi exponentielle et du profil lorentzien dans des conditions mieux cernées. Mais l'obligation de se placer dans des conditions plus restrictives pour traiter le problème mathématique de la désintégration est la rançon de notre désir d'une forme de loi universelle, largement indépendante de la nature du système, de son environnement et du temps. Nos résultats n'épuisent pas, loin s'en faut, toutes les possibilités. Des déviations théoriques à la loi exponentielle, pour les temps courts ou les temps infinis, sont inévitables et dépendent, elles, du système, de son mode de préparation (que j'ai passé sous silence en déclarant partir d'un état stationnaire... qui ne l'est pas!) et du procédé de détection (qui en plus nous replonge dans les vertiges pas seulement épistémologiques du problème de la mesure

en théorie quantique).

La largeur de la distribution d'énergie du rayonnement (VI.45) permet d'attribuer aussi une largeur à la distribution d'énergie associée à l'état atomique initial (qui, rappelons le, n'est pas stationnaire, et donc pas propre, vis-à-vis du rayonnement), disons  $(\Delta E)_B \stackrel{\text{df}}{=} \hbar\Gamma$ , appelée *largeur naturelle* du niveau. Il est sans importance que cette quantité ne soit sans doute pas précisément une dispersion au sens formel d'une moyenne quadratique; contrairement à ce que peut laisser entendre le discours quantique orthodoxe on est bien incapable d'observer directement l'énergie d'un niveau atomique, et ce sont de fait la position et la largeur de la distribution observée pour le rayonnement qui nous fournissent des définitions opérationnelles de l'énergie et de la largeur d'un état atomique (quasi) stationnaire.

### VI.5.3 Largeurs de raies, largeurs de niveaux

On peut observer et mesurer directement la loi exponentielle de décroissance d'un ensemble d'atomes ou de noyaux excités lorsque la durée de vie du système est suffisante,  $\tau \geq 10^{-8}$  s. Cette limite correspond à une largeur naturelle

$$\Delta E = \hbar\Gamma = \frac{\hbar c}{c\tau} \approx \frac{2 \times 10^2 \times 10^6}{3 \times 10^8 \times 10^{15} \times 10^{-8}} \approx 10^{-7} \text{ eV},$$

ou encore, en fréquence,  $\Gamma = 1/\tau \approx 100$  MHz. En pratique, on est capable de mesurer aussi la largeur du spectre d'énergie du rayonnement lorsqu'elle est supérieure à cette valeur. L'accord est confirmé expérimentalement dans les cas où les deux mesures sont possibles. Le succès de la relation, *a priori* saugrenue,  $\Delta E \tau \approx \hbar$ , entre deux grandeurs aussi étrangères que la vie moyenne d'un niveau et la largeur du spectre d'énergie du rayonnement qu'il émet, est à mettre à l'actif de la théorie quantique. L'origine de cette réussite quantique n'est pas à rechercher dans les arcanes du calcul qui mène au résultat (VI.44), car on trouve aussi la même relation  $\Delta\omega \tau = 1$  entre la constante de temps d'un oscillateur classique amorti par rayonnement classique et la largeur du spectre en fréquence du rayonnement émis.<sup>6</sup> En dernière analyse, ce succès repose donc sur l'équivalence  $E = \hbar\omega$  que la théorie quantique établit entre fréquence du rayonnement et énergie du photon.

Dans le cas (fort répandu!) d'un atome qui possède plus d'un niveau excité, les choses se compliquent. Il n'y a plus la relation directe entre largeur d'une raie de transition et vie moyenne du niveau de départ à laquelle on pouvait s'attendre classiquement. On peut évidemment encore calculer, dans le cadre de la règle d'or de Fermi, un taux de transition pour chaque couple de niveaux initial et final, puis se risquer à définir la largeur d'un niveau comme la somme des taux de transitions ayant ce niveau pour origine. Le traitement quantique non perturbatif analogue au cas à deux niveaux (voir par exemple [44, § 63]) donne un intérêt à cette définition en

---

<sup>6</sup>Voir par exemple [15, p. 80].

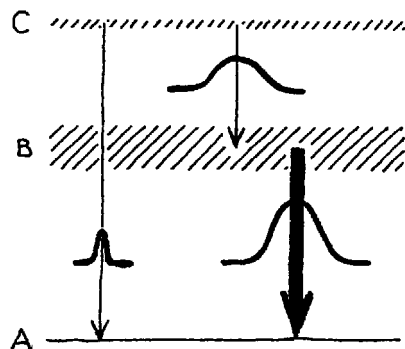


Figure VI.5: Un cas à trois niveaux, avec indications respectives des largeurs de niveaux, des largeurs de raies et de leurs intensités.

montrant qu'une raie de transition a encore une forme lorentzienne dont la largeur est égale — en toute logique — à la somme des largeurs des niveaux initial et final. Il n'y a plus la proportionnalité classique (VI.45) entre intensité et largeur de raie. Ainsi, quantiquement, une transition lente peut être large. La figure VI.5 représente un cas à trois niveaux où toutes les transitions issues de  $C$  (soit  $C \rightarrow A$  et  $C \rightarrow B$ ) sont lentes (faibles taux, faibles intensités) tandis que la transition  $B \rightarrow A$  est rapide (taux élevé, grande intensité); la raie  $C \rightarrow B$  présente alors la particularité d'être tout à la fois large et peu intense.

Le mécanisme d'élargissement d'une raie de transition que nous venons de trouver ne concerne encore qu'un atome (ou un ensemble d'atomes) isolé(s), immobile(s) et dont le seul mécanisme de désexcitation serait l'émission de rayonnement. Pour des atomes en situation plus réaliste, d'autres mécanismes interviennent, aussi bien pour élargir les niveaux (par diminution de leur espérance de vie) que les raies elle-mêmes:

- Dans un atome à plusieurs électrons, la désexcitation d'un électron vers une vacance dans une couche inférieure (en pratique de la couche  $L$  à la couche  $K$ ) peut fournir une énergie suffisante pour que l'interaction mutuelle coulombienne, à elle seule, éjecte un électron extérieur peu lié, avec un reliquat d'énergie cinétique (après déduction de l'énergie de liaison). De concert avec l'émission radiative, ce phénomène d'auto-ionisation, l'effet Auger, abrège la vie de l'état excité.
- Les collisions avec d'autres atomes, ou avec les électrons éjectés par l'effet Auger peuvent, toujours par interaction coulombienne, perturber l'atome excité pour le ramener à un état inférieur (ou supérieur), diminuant d'autant sa durée de vie moyenne.

- Les atomes d'un ensemble sont soumis à l'agitation thermique. Mesurée dans le laboratoire, l'énergie du photon émis par un atome mobile n'est plus donnée par la simple différence d'énergie  $E_B - E_A$  entre niveaux d'énergie dans le repère de l'atome, mais doit être calculée par une transformation de Lorentz de l'énergie  $\hbar\omega_0$  et de l'impulsion  $\hbar\mathbf{k}_0$  dans le repère de l'atome, en fonction de la vitesse de celui-ci. On obtient ainsi un résultat équivalent au décalage de fréquence (effet Doppler) calculé pour un rayonnement classique [27]. Pour un gaz peu dense (gaz parfait), la distribution maxwellienne des vitesses se traduit par une répartition des fréquences observées s'étendant sur une largeur de l'ordre de 1000 MHz, dans les conditions normales, qui excède amplement la largeur naturelle<sup>7</sup> et donne à la raie une forme gaussienne. Cet effet masque complètement le faible déplacement de la raie dû à la conservation de l'impulsion qui intervient dès que l'atome n'est plus infiniment lourd (il aurait fallu tenir compte dans notre analyse des degrés de liberté du centre de masse de l'atome); l'atome, même initialement immobile, emporte après émission une impulsion (et une énergie) de recul qui vient grever d'autant l'impulsion (et la pulsation) nominale  $\hbar\mathbf{k}_0$  du photon.
- Un atome non solitaire subit des interactions à distance avec les autres atomes, même dans un gaz (nul n'est parfait), par exemple l'interaction de Van der Waals entre moment dipolaire électrique de l'atome et champ électrique dipolaire créé par un autre atome. Lorsque l'atome ressent un champ magnétique des autres atomes, il peut se produire une levée de dégénérescence de ses sous niveaux magnétiques (effet Zeeman) qui donne lieu à un groupe de raies de transition voisines dont l'écart est inférieur à la résolution du détecteur. On observe alors une seule raie élargie.

Parmi tous ces mécanismes, seul l'effet Auger est indépendant de l'environnement de l'atome, en particulier de la température et de la densité lorsqu'il s'agit d'un gaz. La contribution de cet effet à la largeur du niveau est donc caractéristique de l'atome. Comme elle est inséparable de la contribution radiative, on est convenu d'entendre par largeur naturelle du niveau la somme des largeurs Auger et radiative.

Nous avons maintenant fait connaissance avec le concept de largeur de niveau, manifestation du couplage, toujours présent, avec le rayonnement quantique. Le profil d'une raie d'émission est une caractéristique des niveaux initial et final, et il n'y a donc pas à être surpris qu'une étude détaillée de l'absorption de rayonnement conduise à la même forme de raie d'absorption, avec la même largeur, pour un processus du genre  $A + n\gamma_\omega \rightarrow B + (n-1)\gamma_\omega$ , à cette nuance près qu'avec un faisceau électromagnétique intense et un grand nombre d'atomes dans leur état fondamental (cible solide) le processus d'absorption peut être sensible, même loin de la résonance ( $\omega \neq \omega_0$ ).

---

<sup>7</sup>Voir par exemple [15, p. 22].

## VI.6 Déplacement de niveau

Venons en maintenant au décalage de raie résultant de la partie imaginaire de la constante  $\gamma$  donnée par l'équation (VI.43). Nous avons déjà pris conscience du fait que ce sont les raies d'émission de rayonnement qui fournissent une définition opérationnelle des niveaux de l'atome, tant en largeur qu'en position. Historiquement, c'est d'ailleurs l'observation de ces raies qui a révélé l'existence des niveaux atomiques. Inspirés par la définition (VI.35) de  $\gamma$  et par les rôles que joue sa partie imaginaire dans l'évolution temporelle du ket d'état du système — équation (VI.31) — ou dans une redéfinition de la pulsation nominale  $\omega_0$  de la transition — évidente dans le second membre de (VI.38) —, associons à cette partie imaginaire un déplacement d'énergie du niveau initial:

$$\delta E_B \stackrel{\text{df}}{=} \hbar \Im \frac{\gamma}{2} = \mathcal{PP} \int_0^\infty d\omega \frac{f(\omega)}{\omega_0 - \omega}. \quad (\text{VI.46})$$

Contrairement à ce qui se passait pour la largeur de niveau (ou le taux de transition), donnée par la partie réelle de  $\gamma$  — encore l'équation (VI.43) — et qui n'impliquait que des photons de pulsation  $\omega_0$  conservant l'énergie, toutes les pulsations  $\omega$  (sauf  $\omega_0$ ) contribuent au déplacement du niveau.

Insérant l'opérateur de la transition (VI.34) dans la définition (VI.39) de  $f(\omega)$ , nous reste à évaluer l'expression:

$$f(\omega) = \frac{1}{(2\pi)^3} \frac{q_e^2}{2\varepsilon_0} \frac{\omega}{m^2 c^3} \int d^2 \hat{\mathbf{k}} \sum_T |\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} e^{i\mathbf{k} \cdot \mathbf{R}} | A \rangle|^2,$$

où, pour alléger l'écriture, nous supposons l'électron unique. Plus audacieusement, nous nous autorisons l'approximation dipolaire,  $e^{i\mathbf{k} \cdot \mathbf{R}} | A \rangle \approx | A \rangle$ , bien que l'intégration (VI.46) fasse intervenir des valeurs de pulsation élevées. La question referra surface. Quoi qu'il en soit, on peut ainsi calculer facilement la somme sur les polarisations et directions d'émission (par le procédé déjà utilisé pour obtenir le taux de transition (VI.25)). Il vient ainsi:

$$\int d^2 \hat{\mathbf{k}} \sum_T |\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle|^2 = \frac{8\pi}{3} |\langle B | \mathbf{P} | A \rangle|^2.$$

Enfin, nous intéressant au déplacement d'énergie du niveau, il est plus parlant de tout exprimer en fonction de l'énergie du photon,  $E \approx \hbar\omega$ , plutôt que sa pulsation, et de la différence d'énergie des niveaux,  $E_B - E_A = \hbar\omega_0$ . Ceci a pour mérite accessoire de faire apparaître la constante de structure fine  $\alpha$ . Alors:

$$f(E) = \frac{2}{3\pi} \alpha |\langle B | \mathbf{P} | A \rangle|^2 E,$$



et

$$\delta E_B = \frac{2}{3\pi} \alpha \frac{|\langle B|\mathbf{P}|A\rangle|^2}{(mc)^2} \mathcal{P}\mathcal{P} \int_0^\infty dE \frac{E}{E_B - E_A - E}. \quad (\text{VI.47})$$

La singularité de l'intégrant en  $E = E_B - E_A$  est certes évacuée par la partie principale, mais cela n'empêche pas l'intégrale de diverger linéairement:

$$\mathcal{P}\mathcal{P} \int_0^x dE \frac{E}{E_B - E_A - E} \underset{x \rightarrow \infty}{\sim} -x.$$

Le déplacement de niveau est infini!

En y regardant de plus près, la situation n'est peut être pas si désespérée: lorsque  $E$  croît, l'approximation dipolaire est illicite, il faut garder le facteur  $e^{i\mathbf{k}\cdot\mathbf{R}}$  associé à  $\mathbf{P}$  dans l'élément de matrice de transition qui dépend donc de  $E$  (par l'intermédiaire de  $|\mathbf{k}| = \omega/c = E/\hbar c$ ) et ne peut être factorisé hors de l'intégrale. En représentation de position de l'électron,  $e^{i\mathbf{k}\cdot\mathbf{r}}$  est une fonction de  $\mathbf{r}$  oscillant rapidement lorsque  $E$  croît. En conséquence de quoi l'élément de matrice, qui résulte d'une intégration sur toutes les valeurs  $\mathbf{r}$  de la position de l'électron, est une fonction décroissante de  $E$ . Nous nous sommes d'autre part cantonnés à l'équation de Schrödinger, relativiste galiléenne, pour décrire notre électron. Il est donc irréaliste de faire intervenir des énergies de photon élevées par rapport à l'énergie de repos de l'électron. Notre théorie exige certainement des raffinements einsteiniens pour  $E \geq mc^2$ , dont on peut espérer une action favorable sur la divergence en notant par exemple que le remplacement heuristique de la masse de l'électron — dans (VI.47) — par la “masse relativiste”  $m/\sqrt{1-(v/c)^2}$  vient heureusement affaiblir les contributions des énergies élevées. Nous pouvons pallier de façon phénoménologique les insuffisances de notre analyse théorique en supprimant purement et simplement les contributions à l'intégrale au delà d'une énergie de coupure  $E_c$ , disons de l'ordre de l'énergie de repos de l'électron. Ce faisant, nous obtenons une expression parfaitement calculable,

$$\delta E_B = \frac{2}{3\pi} \alpha \frac{|\langle B|\mathbf{P}|A\rangle|^2}{(mc)^2} \mathcal{P}\mathcal{P} \int_0^{E_c} dE \frac{E}{E_B - E_A - E},$$

qui a le mérite d'être finie, et dont nous avons même l'espoir qu'elle cerne d'un peu plus près la réalité. Malheureusement, sa sensibilité quasi linéaire au paramètre  $E_c$  par ailleurs indéterminé semble bien lui ôter toute valeur prédictive.

Avant de montrer que nous pouvons quand même tirer quelque chose de cette expression, il nous faut la généraliser à des cas plus réalistes que celui d'un atome à deux niveaux. Pour cela, on ajoute les contributions de tous les états stationnaires de l'atome, sans considération d'énergie. On obtient alors le déplacement du niveau  $B$ :

$$\delta E_B = \frac{2}{3\pi} \frac{\alpha}{(mc)^2} \sum_I |\langle B|\mathbf{P}|I\rangle|^2 \mathcal{P}\mathcal{P} \int_0^{E_c} dE \frac{E}{E_B - E_I - E}, \quad (\text{VI.48})$$

avec des parties principales à prendre autour de chaque singularité  $E_B - E_I$ . Bien entendu, les états  $I$  pour lesquels l'élément de matrice de transition est nul n'apporteront aucune contribution.

### VI.6.1 Renormalisation de la masse de l'électron

Je n'ai jusqu'ici discuté du déplacement de l'énergie que pour des niveaux d'électron lié dans un atome. Qu'en est-il pour un électron libre, ou prétendu tel puisque lui aussi toujours en interaction avec le rayonnement quantique? Plus précisément, essayons d'appliquer l'analyse précédente à un électron libre, d'impulsion  $\mathbf{p}$  dans l'état initial, qui peut effectuer une transition vers un état d'impulsion finale  $\mathbf{p}'$ , accompagné d'un photon  $\mathbf{k}T$ . La lectrice alerte aura sursauté à l'évocation d'une telle transition qui ne conserve pas l'énergie, mais nous pouvons tenter notre chance, confortés par l'idée que le déplacement de niveau implique de toute façon une intégrale sur les énergies de photon nonobstant toute contingence conservatrice (équation (VI.46)). En retournant dans notre boîte à modes — protection contre les pathologies du carré d'une "fonction" de Dirac — et en sommant les contributions de tous les états finals  $\mathbf{p}'$ , nous avons (voir page 181):

$$i\hbar\frac{\gamma}{2} = \frac{1}{\hbar} \sum_{\mathbf{p}', \mathbf{k}, T} |\langle \mathbf{p}; 0 | W_{\mathbf{k}T} | \mathbf{p}'; 1_{\mathbf{k}T} \rangle|^2 \frac{1 - e^{i(\omega_0 - \omega - i\frac{\gamma}{2})t}}{\omega - \omega_0 + i\frac{\gamma}{2}},$$

où  $\omega_0$  est la variation d'énergie de l'électron:  $\omega_0 \stackrel{\text{df}}{=} (E_{\mathbf{p}} - E_{\mathbf{p}'})/\hbar$ . Supposant encore  $|\gamma| \ll \omega_0$ , et tenant compte de l'expression (VI.34) de  $W_{\mathbf{k}T}$ , nous obtenons:

$$\gamma \approx \frac{2i}{\hbar^2} \frac{q_e^2}{m^2} \frac{\hbar}{2\varepsilon_0 \mathcal{V}} \sum_{\mathbf{p}', \mathbf{k}, T} \frac{1}{\omega} |\langle \mathbf{p} | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} e^{i\mathbf{k} \cdot \mathbf{R}} | \mathbf{p}' \rangle|^2 \frac{1 - e^{i(\omega_0 - \omega)t}}{\omega_0 - \omega}.$$

La fonction d'onde normée d'un électron d'impulsion  $\mathbf{p}$  dans notre boîte-laboratoire de volume  $\mathcal{V}$  est  $\langle \mathbf{r} | \mathbf{p} \rangle = \mathcal{V}^{-1/2} e^{i\mathbf{p} \cdot \mathbf{r}/\hbar}$ , d'où l'élément de matrice

$$\langle \mathbf{p} | \mathbf{P} e^{i\mathbf{k} \cdot \mathbf{R}} | \mathbf{p}' \rangle = \int_{\mathcal{V}} d^3\mathbf{r} \frac{e^{-i\frac{\mathbf{p} \cdot \mathbf{r}}{\hbar}}}{\sqrt{\mathcal{V}}} \frac{\hbar}{i} \nabla \frac{e^{i\frac{(\hbar\mathbf{k} + \mathbf{p}') \cdot \mathbf{r}}{\hbar}}}{\sqrt{\mathcal{V}}} = \mathbf{p} \delta_{\mathbf{p}, \mathbf{p}' + \hbar\mathbf{k}},$$

et

$$\gamma = \frac{i}{\hbar} \frac{q_e^2}{m^2 \varepsilon_0 \mathcal{V}} \sum_{\mathbf{k}T} \frac{1}{\omega} (\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{p})^2 \frac{1 - e^{i(\omega_0 - \omega)t}}{\omega_0 - \omega}.$$

Dans un grand laboratoire, cette expression devient:

$$\gamma = \frac{i}{\hbar} \frac{q_e^2}{m^2 \varepsilon_0} \frac{1}{(2\pi)^3} \int_0^{\omega_c/c} d\left(\frac{\omega}{c}\right) \left(\frac{\omega}{c}\right)^2 \int d^2\hat{\mathbf{k}} \frac{1}{\omega} \sum_T (\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{p})^2 \frac{1 - e^{i(\omega_0 - \omega)t}}{\omega_0 - \omega},$$

l'intégration étant coupée à la pulsation  $\omega_c \stackrel{\text{df}}{=} E_c/\hbar$ . Le déplacement d'énergie de l'électron libre d'impulsion  $\mathbf{p}$  est donc donné par

$$\begin{aligned}\delta E_{\mathbf{p}} &= \hbar \Im \frac{\gamma}{2} \\ &= \frac{q_e^2}{2\varepsilon_0 m^2 c^3} \frac{1}{(2\pi)^3} \int_0^{\omega_c} d\omega \omega \frac{1 - \cos(\omega_0 - \omega)t}{\omega_0 - \omega} \int d^2\mathbf{k} \sum_T (\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{p})^2.\end{aligned}$$

Le temps  $t$  tendant vers l'infini, la contribution oscillante du  $\cos(\omega_0 - \omega)t$  dans l'intégrale devient négligeable; l'intégrale sur les directions est triviale (voir p. 174), et l'on a:

$$\delta E_{\mathbf{p}} = \frac{2}{3\pi} \alpha \frac{\mathbf{p}^2}{(mc)^2} \int_0^{E_c} dE \frac{E}{E_{\mathbf{p}} - E_{\mathbf{p}-\hbar\mathbf{k}} - E}.$$

Pour un électron “non relativiste” — et qui le reste —, le dénominateur de l'intégrand se simplifie:

$$\begin{aligned}E_{\mathbf{p}} - E_{\mathbf{p}-\hbar\mathbf{k}} - E &= \frac{1}{2m} (\mathbf{p}^2 - (\mathbf{p} - \hbar\mathbf{k})^2) - E, \\ &= \frac{\hbar\mathbf{k} \cdot \mathbf{p}}{m} - \frac{\hbar^2\mathbf{k}^2}{2m} - E = -E \left( 1 - \hat{\mathbf{k}} \cdot \frac{\mathbf{p}}{mc} + \frac{E}{2mc^2} \right), \\ &\approx -E.\end{aligned}$$

L'intégration sur  $E$  est alors triviale, et en posant

$$C \stackrel{\text{df}}{=} -\frac{2}{3\pi} \alpha \frac{1}{(mc)^2} E_c, \quad (\text{VI.49})$$

on a finalement un déplacement d'énergie de l'électron libre  $\delta E_{\mathbf{p}} = C\mathbf{p}^2$ , proportionnel à  $\mathbf{p}^2$ . L'énergie corrigée,  $\mathbf{p}^2/2m + C\mathbf{p}^2$ , est donc elle aussi proportionnelle à  $\mathbf{p}^2$ .

C'est cette énergie corrigée, la seule effectivement observée puisqu'il n'existe pas d'interrupteur pour l'interaction avec le rayonnement, qui nous montre que cette interaction se traduit, en ce qui concerne l'électron libre par une simple renormalisation de la *masse nue*,  $m$ , ou *masse mécanique*, d'un électron fictif sans interaction avec le rayonnement, pour donner la *masse physique*  $m_\varphi$ , ou *masse habillée*, ou *masse renormalisée*, que l'on trouve dans les tables: c'est celle qui est mesurée. La relation entre masse physique et masse nue est:

$$\frac{\mathbf{p}^2}{2m_\varphi} = \frac{\mathbf{p}^2}{2m} + C\mathbf{p}^2,$$

soit

$$m_\varphi = \frac{m}{1 + 2mC} \approx (1 - 2mC)m,$$

dans l'hypothèse — que nous allons vérifier immédiatement — où la correction de déplacement est faible, soit encore

$$m_\varphi = \left(1 + \frac{4}{3\pi} \alpha \frac{E_c}{mc^2}\right) m.$$

Lorsque l'énergie de coupure  $E_c$  est égale à  $mc^2$ , la correction relative vaut  $(4/3\pi)\alpha$ , de l'ordre de 0,4 %. Evidemment, ce genre d'évaluation est affecté par l'arbitraire du paramètre de coupure, sans conséquence tangible puisque la masse ne saurait se laisser observer nue.

### VI.6.2 Déplacement Lamb

Avant de continuer, il est bon de mettre un peu d'ordre dans nos notations des diverses masses. Pour ne pas alourdir en permanence l'écriture, et pour respecter l'usage, désignons par  $m$  la masse physique, celle que l'on trouve dans les tables empiriques et avec laquelle nous avons d'ailleurs effectué tous nos calculs numériques jusqu'à présent. Lorsque nécessaire, la masse nue sera notée  $m_0$ .

Il n'existe pas d'interrupteur de l'interaction avec le rayonnement, mais certains niveaux atomiques ne sont pas branchés! Dans l'atome d'hydrogène en particulier, des états dégénérés dans la seule interaction coulombienne sont diversement couplés au rayonnement; les déplacements de niveaux différent d'un état à l'autre et on peut observer — et mesurer — une levée de dégénérescence directement causée par l'interaction avec le rayonnement. Mais l'expression (VI.48) que nous allons utiliser ne reflète pas tout à fait les différences entre des mondes avec ou sans interaction avec le rayonnement. En toute innocence nous avons déjà pris en compte une partie de cette interaction en résolvant l'équation de Schrödinger pour les fonctions-d'onde des états  $B$  et  $I$  avec la valeur de la masse physique de l'électron. Le déplacement d'énergie de l'électron libre qui conduit à cette masse est donc compté en trop si nous voulons effectivement comparer l'énergie d'un état en présence d'interaction avec le rayonnement avec son énergie en absence d'interaction. Au déplacement calculé par (VI.48), il faut retrancher la contribution déjà incluse dans la masse physique pour obtenir le déplacement effectif:

$$\delta E_B^{\text{eff}} \stackrel{\text{df}}{=} \delta E_B - C \langle B | \mathbf{P}^2 | B \rangle.$$

En remplaçant  $C$  par sa définition, et en commettant la peccadille d'y substituer la masse physique à la masse nue, on obtient:

$$\delta E_B^{\text{eff}} = \frac{2}{3\pi} \frac{\alpha}{(mc)^2} \mathcal{P}\mathcal{P} \int_0^{E_c} dE \left\{ \sum_I |\langle B | \mathbf{P} | I \rangle|^2 \frac{E}{E_B - E_I - E} + |\langle B | \mathbf{P} | B \rangle|^2 \right\},$$

car la partie principale coïncide avec l'intégrale au sens habituel lorsque l'intégrant est régulier. Par une décomposition de l'unité sur les kets  $|I\rangle$  supposés complets, et

une réduction au même dénominateur, il vient:

$$\Delta E_B^{\text{eff}} = \frac{2}{3\pi} \frac{\alpha}{(mc)^2} \sum_I |\langle B|\mathbf{P}|I\rangle|^2 (E_B - E_I) \mathcal{PP} \int_0^{E_c} dE \frac{E}{E_B - E_I - E}.$$

Fait remarquable, la divergence de l'intégrale n'est plus que logarithmique,

$$\mathcal{PP} \int_0^x dE \frac{E}{E_B - E_I - E} \underset{x \rightarrow \infty}{\sim} -\ln x,$$

illustration du comportement général de la soustraction de deux divergences du même degré dont résulte une divergence de degré inférieur. Il ne faut pas sous estimer l'importance pratique de cette comptabilité; tout l'intérêt du procédé de renormalisation en découle. Notre résultat ne dépendra plus que du logarithme du paramètre de coupure  $E_c$  et sera donc assez peu sensible à sa valeur, fort heureusement car nous n'avons aucune condition pour le déterminer précisément.<sup>8</sup>

La partie principale s'évalue sans peine:

$$\begin{aligned} & \mathcal{PP} \int_0^{E_c} dE \frac{E}{E_B - E_I - E} \\ &= \lim_{\epsilon \rightarrow 0} \left( -\ln |E_B - E_I - E| \Big|_0^{E_B - E_I - \epsilon} - \ln |E_B - E_I - E| \Big|_{E_B - E_I + \epsilon}^{E_c} \right), \\ &= -\ln \left| \frac{E_B - E_I - E_c}{E_B - E_I} \right|, \\ &\approx -\ln \left| \frac{E_c}{E_B - E_I} \right|, \end{aligned}$$

car l'énergie de coupure  $E_c$  est de l'ordre de la masse de l'électron (0,5 MeV), tandis que les écarts entre niveaux,  $|E_B - E_I|$ , sont de l'ordre du Rydberg (10 eV), sauf pour les états  $I$  du continuum, mais la contribution de ceux-ci est réduite par défaut de recouvrement des fonctions-d'onde dans  $|\langle B|\mathbf{P}|I\rangle|^2$ . Comme annoncé, l'expression renormalisée du déplacement de niveau ne dépend plus que du logarithme du paramètre de coupure:

$$\delta E_B^{\text{eff}} = \frac{2}{3\pi} \frac{\alpha}{(mc)^2} \sum_I |\langle B|\mathbf{P}|I\rangle|^2 (E_I - E_B) \ln \left| \frac{E_c}{E_I - E_B} \right|.$$

Comme  $E_c/|E_I - E_B| \gg 1$ , nous sommes dans la partie aplatie du logarithme, la valeur de celui-ci doit peu dépendre de  $E_I$  et nous pouvons sans risque le remplacer

---

<sup>8</sup>En théorie relativiste einsteinienne, comme nous nous y attendions (page 192), les intégrands sont effectivement décroissants, les intégrales divergent logarithmiquement et leur soustraction converge; ainsi le résultat physique ne dépend plus du tout de la valeur du paramètre de coupure.

par une moyenne, constante, judicieusement évaluée:

$$\delta E_B^{\text{eff}} = \frac{2}{3\pi} \frac{\alpha}{(mc)^2} \ln \frac{E_c}{|E_{\text{moy}} - E_B|} \sum_I |\langle B|\mathbf{P}|I\rangle|^2 (E_I - E_B). \quad (\text{VI.50})$$

La somme restante, sur les états  $I$ , peut se calculer simplement grâce, justement, à une *règle de somme*. Nous évaluons d'abord le commutateur de l'opérateur impulsion de l'électron et de son hamiltonien atomique,

$$[\mathbf{P}, H_{\text{at}}] = \left[ \mathbf{P}, \frac{\mathbf{P}^2}{2m} + V(\mathbf{R}) \right] = \frac{\hbar}{i} (\nabla V),$$

fonction de l'opérateur  $\mathbf{R}$  définie à partir de la fonction gradient du potentiel  $V(\mathbf{r})$ , en pratique le potentiel coulombien. De cette identité entre opérateurs, on déduit l'identité entre éléments de matrice:

$$(E_I - E_B) \langle B|\mathbf{P}|I\rangle = \frac{\hbar}{i} \langle B|(\nabla V)|I\rangle.$$

En multipliant scalairement par  $\langle I|\mathbf{P}|B\rangle$ , puis en sommant sur les états  $I$ , complets, on a

$$\begin{aligned} \sum_I (E_I - E_B) |\langle B|\mathbf{P}|I\rangle|^2 &= \frac{\hbar}{i} \sum_I \langle B|(\nabla V)|I\rangle \cdot \langle I|\mathbf{P}|B\rangle, \\ &= -\frac{\hbar}{i} \sum_I \langle I|(\nabla V)|B\rangle \cdot \langle B|\mathbf{P}|I\rangle, \end{aligned}$$

puisque cette quantité est réelle, d'où

$$\begin{aligned} \sum_I (E_I - E_B) |\langle B|\mathbf{P}|I\rangle|^2 &= \frac{\hbar}{i} \langle B|(\nabla V) \cdot \mathbf{P}|B\rangle = -\frac{\hbar}{i} \langle B|\mathbf{P} \cdot (\nabla V)|B\rangle \\ &= -\frac{\hbar}{2i} \langle B|[P_j, (\partial_j V)]|B\rangle \\ &= \frac{\hbar^2}{2} \langle B|(\Delta V)|B\rangle, \end{aligned}$$

et

$$\delta E_B^{\text{eff}} = \frac{\alpha}{3\pi} \ln \frac{E_c}{|E_{\text{moy}} - E_B|} \left( \frac{\hbar}{mc} \right)^2 \langle B|(\Delta V)|B\rangle (E_I - E_B). \quad (\text{VI.51})$$

En représentation de position, le potentiel de l'électron dans l'atome d'hydrogène est  $V(\mathbf{r}) = q_e \phi(\mathbf{r})$ , où  $\phi$  est le potentiel électrostatique créé par le noyau (supposé fournir, de par sa masse, l'origine d'un référentiel inertiel), solution de l'équation de Poisson  $\Delta\phi = -\rho(\mathbf{r})/\varepsilon_0$ . D'où la valeur moyenne de  $\Delta V$ :

$$\langle B|(\Delta V)|B\rangle = -\frac{q_e}{\varepsilon_0} \int d^3\mathbf{r} |\psi_B(\mathbf{r})|^2 \rho(\mathbf{r}) \approx \frac{q_e^2}{\varepsilon_0} |\psi_B(0)|^2, \quad (\text{VI.52})$$

dans la mesure où la distribution de charge du proton,  $\rho(\mathbf{r})$ , apparaît ponctuelle en regard de l'étendue de la distribution  $|\psi_B(\mathbf{r})|^2$  des positions de l'électron (Fermi contre Ångström).

Dans cette approximation, seuls les états  $B$  dont la fonction-d'onde n'est pas nulle à l'origine sont effectivement déplacés par l'interaction avec le rayonnement. Ce sont les états  $s$  ( $l = 0$ ), la probabilité de position étant repoussée du centre par le potentiel centrifuge  $(\hbar^2/2m)l(l+1)/r^2$  dès que l'on a affaire à un état de moment orbital non nul. L'interaction coulombienne entre l'électron et le proton d'un atome d'hydrogène présente encore une autre particularité, bien connue de la lectrice, à savoir la dégénérescence des états  $ns$ ,  $np$ ,  $nd \dots$  jusqu'à  $l = n - 1$ . Grâce à ces merveilleuses propriétés nous disposons avec l'atome d'hydrogène d'un système où des états voient leur dégénérescence en partie levée par l'interaction avec le rayonnement, mettant ainsi en évidence le déplacement de niveau.

La fonction-d'onde d'un état  $ns$  de l'atome d'hydrogène vaut, à l'emplacement du proton, dans notre approximation  $m_p \gg m$  [53, app. B-3]:

$$\psi_{ns}(0) = \frac{1}{\sqrt{4\pi}} \frac{2}{(na_\infty)^{3/2}},$$

où  $a_\infty \stackrel{\text{df}}{=} (4\pi\epsilon_0/e^2)(\hbar^2/m)$  est le rayon de Bohr, tandis que pour tout autre état,  $l \neq 0$ , on a  $\psi_{nl}(0) = 0$ . Par reports dans les équations (VI.52) et (VI.51), on obtient le déplacement d'un état  $ns$ :

$$\delta E_B^{\text{eff}} = \frac{8}{3\pi} \frac{\alpha^3}{n^3} R_\infty \ln \frac{E_c}{|E_{\text{moy}} - E_B|}, \quad (\text{VI.53})$$

en fonction de la constante de Rydberg  $R_\infty \stackrel{\text{df}}{=} \frac{1}{2}(e^2/4\pi\epsilon_0 a_\infty)$ .

Un argument qualitatif [54, p. 80] permettait de prévoir le sens de ce déplacement: le couplage avec le rayonnement entraîne des fluctuations de position de l'électron qui se traduisent par un étalement de sa distribution de charge électrique, d'où une diminution (en valeur absolue) de son énergie potentielle d'interaction attractive (négative) avec le noyau, et une élévation de l'énergie totale. La variation d'énergie potentielle moyenne due à ces fluctuations est nulle au premier ordre (la valeur moyenne de la fluctuation est nulle dans l'état fondamental du rayonnement), tandis que la contribution du second ordre est proportionnelle à  $\langle \Delta V \rangle$ , c'est-à-dire au recouvrement de la densité de charge du noyau et de la probabilité de position de l'électron, avec pour conséquence la nullité du déplacement lorsque l'électron a un moment cinétique orbital.

Ce déplacement de niveau, relativement anodin puisque les descriptions élémentaires — expérimentales ou théoriques — du spectre de l'atome d'hydrogène ne le mentionnent pas, constitue en fait un jalon historique fondamental dans le développement de la théorie quantique du rayonnement. Jusqu'en 1937, les corrections relativistes de Dirac rendaient parfaitement compte de la structure fine du

spectre de l'atome d'hydrogène: la dégénérescence des états  $nl$ ,  $l = 0, 1, \dots, n-1$  était en partie levée, mais subsistait encore une dégénérescence des couples d'états de mêmes  $nl$ , avec  $j = |l \pm \frac{1}{2}|$ .<sup>9</sup> Une situation de doute théorique s'instaura à la suite de mesures spectroscopiques qui ne semblaient s'expliquer que par une montée du niveau  $2s\frac{1}{2}$  par rapport au  $2p\frac{1}{2}$ . Cet écart fut définitivement confirmé par l'expérience de Lamb et Retherford [42] permise par l'avancée des techniques de micro-ondes de longueur 3 cm.<sup>10</sup> La longueur d'onde de la transition  $2p\frac{3}{2} \rightarrow 2p\frac{1}{2}$  utilisée pour l'expérience (bien décrite, par exemple, dans [15, p. 455]) est de 2,74 cm. Le résultat de la mesure de Lamb et Retherford, maintenant connu sous le nom de *déplacement Lamb*,

$$\frac{E_{2s\frac{1}{2}} - E_{2p\frac{1}{2}}}{2\pi\hbar} \approx 1000 \text{ MHz}, \quad (\text{VI.54})$$

fut présenté à la Conférence de Physique Théorique qui se tenait à Shelter Island du 2 au 4 juin 1947. La lectrice friande d'anecdotes pourra lire et comparer les récits qu'ont donnés Weinberg [77], Dyson [23] et Weisskopf [79] de cette première réunion des pères fondateurs de l'électrodynamique quantique, à la fleur de l'âge après un long interlude turbulent. Bethe résolut le problème dans le train qui le ramenait à Philadelphie, et le 27 juin, la *Physical Review* recevait son article [10] dans lequel il établissait la formule (VI.53). La valeur moyenne des énergies de transitions dipolaires électriques virtuelles depuis le niveau  $2s$ , déterminée numériquement, donnait  $|E_{\text{moy}} - E_{2s}| \approx 20 R_{\infty}$ .<sup>11</sup> En adoptant la coupure  $E_c = mc^2 = 0,5 \text{ MeV}$ , la lectrice trouvera, comme Bethe, le déplacement de l'état  $2s$  dû au couplage avec le rayonnement quantique,

$$\frac{\delta E_{2s}^{\text{eff}}}{2\pi\hbar} = 1040 \text{ MHz},$$

en parfait accord avec l'expérience, puisque la même analyse théorique prédit un déplacement nul pour l'état  $2p\frac{1}{2}$  et que la quantité mesurée (VI.54) ne dépend pas du déplacement de l'état fondamental  $1s$ .

La mariée n'est-elle pas trop belle? Un tel accord ne serait-il pas dû à une chance adroitement aidée par le choix de la valeur  $E_c = mc^2$  pour un paramètre de coupure par ailleurs peu déterminé? Vraisemblablement non. Le facteur logarithmique dans l'expression (VI.53), qui vaut 7,6 pour les valeurs adoptées par Bethe, ne varie, par exemple, que de 5,3 pour  $E_c = 0,1 mc^2$  à 9,9 pour  $E_c = 10 mc^2$ . Si effectivement l'accord quantitatif est quelque peu aidé par le choix de la valeur de  $E_c$ , ce choix particulier n'est pour rien dans la qualité de l'ordre de grandeur

<sup>9</sup>Vous trouverez une discussion et de belles illustrations du spectre de l'hydrogène dans [34].

<sup>10</sup>Ce progrès était consécutif à l'effort de guerre dans la mise au point du radar. Un cas de transfert de technologie du domaine militaire vers la physique (ai-je bien dit qu'elle était pure?).

<sup>11</sup>Cette valeur était assez élevée, par rapport à l'énergie d'ionisation  $R_{\infty}$ , pour surprendre son auteur. Il s'est ensuite attaché à l'expliquer qualitativement dans [11].



d'un résultat obtenu, rappelons-le, par différence entre deux quantités infinies (!), qualité dont le mérite revient plutôt au mécanisme de soustraction de divergences.

A propos du déplacement des niveaux par interaction avec le rayonnement, Bethe écrivait en introduction, non sans goût du paradoxe ironique: «Dans toutes les théories existantes on trouve que ce déplacement est infini, on l'a donc toujours négligé». Le calcul du déplacement Lamb marque pour la théorie quantique du rayonnement le passage d'une adolescence illustrée par des brillants succès de débutante à la plénitude d'une majorité débarassée de contradictions sans avoir perdu aucun des charmes troublants de sa jeunesse. N'ayons plus peur, menons une vie dangereuse. N'ignorons plus les infinis, soustrayons les! Le paramètre de masse initialement assigné à l'électron dans la formalisation du système quantique électron-rayonnement n'est pas la quantité mesurée dans les circonstances ordinaires. L'électron charrie son interaction avec le rayonnement électromagnétique (même dans le vide, c'est-à-dire si le rayonnement est dans son fondamental) et la contribution de celui-ci à l'inertie du système. Il faut donc se livrer à une opération de renormalisation dans laquelle les paramètres de la théorie (ici la masse nue) cèdent la place aux paramètres qui ont une signification physique immédiate (la masse habillée). La mesure et le calcul du déplacement Lamb inauguraient tout un champ nouveau pour la théorie quantique du rayonnement. Notons que la correction au facteur  $g$  de l'électron (voir page 27), mesurée à la même époque et grâce aux mêmes techniques, dut, pour cause de difficultés techniques de calcul, attendre un peu plus longtemps son explication, bien que l'origine en résidât aussi dans l'interaction de l'électron avec le rayonnement. Nous avons déjà étudié un autre phénomène trahissant la présence permanente du rayonnement quantique et exigeant une soustraction, l'effet Casimir, mais il n'était encore, à cette époque, soupçonné ni théoriquement, ni expérimentalement.

C'est la mesure du déplacement de niveau par Lamb et Retherford qui a stimulé les calcul d'effets d'interaction avec le rayonnement — les *corrections radiatives* — par soustraction des divergences plutôt qu'oubli pur et simple des infinis. L'interprétation correcte de ces effets confirmait la robustesse de la théorie quantique du rayonnement en dépit des infinis qui l'affectaient depuis sa naissance. Ce succès donnait l'élan pour une formulation relativiste complète de la théorie, l'*électrodynamique quantique*, et la démonstration de sa renormalisabilité: tous les infinis que l'on rencontre dans la théorie peuvent être absorbés dans la redéfinition d'un nombre *fini* de paramètres physiques: la masse de l'électron — ce que nous venons de faire — et sa charge.

## VI.7 Rayonnement multipolaire

Pour étudier l'émission spontanée à partir d'un état initial  $B$  vers un état  $A$ , nous avons vu (§ VI.4) que l'on peut en général utiliser l'approximation  $e^{-i\mathbf{k}\cdot\mathbf{R}}|B\rangle \approx |B\rangle$ ,

laquelle permet d'exprimer la probabilité de transition en fonction du simple élément de matrice  $\langle A|\mathbf{R}|B\rangle$ , analogue d'un moment dipolaire électrique. Il peut se produire, pour certains couples d'états  $A$  et  $B$  — ceux qui ne satisfont pas les règles de sélection dipolaires électriques de la section VI.4.2 —, que cet élément de matrice soit nul. Bien que la transition  $B \rightarrow A$  soit alors parfois qualifiée d'*interdite*, il s'agit d'une interdiction toute relative. La transition  $B \rightarrow A$  n'est pas impossible, mais elle doit passer par une approximation de  $e^{-i\mathbf{k}\cdot\mathbf{R}}$  d'ordre plus élevé et va donc avoir une probabilité plus petite. S'il faut aller jusqu'au terme  $(\mathbf{k}\cdot\mathbf{R})^n$  dans le développement de l'exponentielle pour trouver un élément de matrice de transition non nul, il est clair que l'ordre de grandeur de la probabilité de transition, proportionnelle au carré de l'élément de matrice, sera réduit par un facteur  $(r/\lambda)^{2n}$ , où  $r$  est l'extension spatiale du système et  $\lambda \propto k^{-1}$  la longueur d'onde de la transition. Nous avons vu (page 166) que  $r/\lambda$  est de l'ordre de  $10^{-3}$  pour les électrons dans un atome et de  $10^{-2}$  pour les protons dans un noyau.

Revenons à l'expression générale (VI.11), spécialisée au cas d'un seul quanton pour alléger l'écriture. Dans le cas où l'élément de matrice de l'approximation dipolaire électrique est nul, il nous faut recourir au terme suivant dans l'approximation de l'exponentielle et examiner la possible action de l'opérateur  $(\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R})$ , pour lequel on va trouver d'autres règles de sélection. Si celles-ci ne sont pas satisfaites par le couple d'états  $A$  et  $B$ , il faudra pousser le développement plus loin. Ce développement et la mise en évidence des règles de sélection à chaque ordre peuvent être réalisés systématiquement grâce à l'artillerie formelle de l'algèbre du moment angulaire. Pour notre discussion — qui se bornera au successeur immédiat du terme dipolaire électrique — nous allons procéder de manière artisanale en décomposant l'opérateur (allégé de ses indices, il suffit de se souvenir que  $\hat{\mathbf{e}} \cdot \mathbf{k} = 0$ ):

$$\begin{aligned} (\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) &= \frac{1}{2} \left\{ (\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) + (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P}) \right\} \\ &\quad + \frac{1}{2} \left\{ (\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) - (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P}) \right\}. \end{aligned} \quad (\text{VI.55})$$

Inspirés par l'astuce qui a déjà si bien servi page 168, commençons par calculer le commutateur de l'hamiltonien atomique avec le produit de deux composantes de  $\mathbf{R}$ :

$$[H_{\text{at}}, R_i R_j] = \left[ \frac{\mathbf{P}^2}{2m}, R_i R_j \right] = -i \frac{\hbar}{m} (R_i P_j + P_i R_j),$$

expression grâce à laquelle le contenu de la première accolade dans (VI.55) peut se réécrire

$$\varepsilon_i k_j (R_i P_j + P_i R_j) = -\frac{m}{i\hbar} \varepsilon_i k_j [H_{\text{at}}, R_i R_j],$$

et son élément de matrice,

$$\langle A | \left\{ (\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) + (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P}) \right\} | B \rangle$$

$$\begin{aligned}
&= -\frac{m}{i\hbar} (E_A - E_B) \varepsilon_i k_j \langle A | R_i R_j | B \rangle \\
&= -im\omega_0 \varepsilon_i k_j \langle A | (R_i R_j - \frac{1}{3} \delta_{ij} \mathbf{R}^2) | B \rangle, \quad (\text{VI.56})
\end{aligned}$$

puisque  $\varepsilon_i k_j \delta_{ij} = 0$ . Ainsi écrits, les opérateurs qui interviennent dans la transition sont les composantes d'un tenseur constitué avec les composantes de  $\mathbf{R}$ , car ils engendrent une représentation linéaire des rotations (ils se transforment évidemment comme des produits de composantes de vecteurs). Cette représentation est de rang cinq car il n'y a que cinq composantes indépendantes (le tenseur est symétrique et a une trace nulle). Ces composantes sont en fait inséparables vis-à-vis de l'ensemble des rotations. La représentation engendrée est donc irréductible, traduction de la possibilité d'écrire ces composantes sous forme de combinaisons linéaires des harmoniques sphériques d'ordre deux:

$$r_i r_j - \frac{1}{3} \delta_{ij} \mathbf{r}^2 = \mathbf{r}^2 C_{ijm} Y_m^{(2)}(\hat{\mathbf{r}}),$$

en complète analogie avec le cas d'un opérateur vectoriel (identité (VI.20)). La lectrice qui connaît son théorème de Wigner-Eckart — que les autres me croient — en déduira immédiatement que l'élément de matrice de cet opérateur entre deux états propres du moment angulaire total est nul si la condition

$$|j_A - j_B| \leq 2 \leq j_A + j_B$$

n'est pas satisfaite. Les transitions induites par l'opérateur  $(\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) + (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P})$  sont donc du type

$$j_B \rightarrow j_A = j_B, |j_B \pm 1|, |j_B \pm 2|$$

sauf  $0 \rightarrow 0$ ,  $0 \rightarrow 1$ , et  $\frac{1}{2} \rightarrow \frac{1}{2}$ . L'opérateur de spin du quanton,  $\mathbf{S}$ , n'intervient pas dans cet opérateur. Dans le cas où plusieurs quantons participent, et pour des états  $A$  et  $B$  états propres de  $\mathbf{S}_{\text{tot}}^2$ , on doit donc avoir  $s_A = s_B$ , ce qui interdit les transitions entre multiplets d'ordres différents (la même règle valait bien sûr pour les transitions dipolaires électriques). Pour des états propres de  $\mathbf{L}_{\text{tot}}^2$ , les nombres quantiques  $l_A$  et  $l_B$  satisfont les mêmes règles de sélection que  $j_A$  et  $j_B$ . Mais les harmoniques sphériques  $Y_m^{(2)}$  sont paires et, à toutes fins pratiques, les états  $A$  et  $B$  ne peuvent être que pairs ou impairs (voir page 172). Ainsi, pour cet opérateur, les parités de  $A$  et  $B$  doivent obéir à la règle de sélection  $\pi_A = \pi_B$ . Dans le cas d'un quanton unique, cela implique les possibilités  $l_B = |l_A \pm 2|$  et (contrairement au cas dipolaire électrique)  $l_B = l_A$ , mais l'interdiction de  $l_B = |l_A \pm 1|$ .

Les opérateurs de transition dans (VI.56) sont formellement identiques aux composantes du moment quadripolaire électrique d'une distribution de charge classique [15, p. 518]. Pour cette raison, une transition médiée par l'opérateur  $(\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) + (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P})$  est dite *quadripolaire électrique*.

Examinons maintenant le deuxième terme de la décomposition (VI.55). Nous avons:

$$(\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) - (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P}) = \varepsilon_i k_j (P_i R_j - R_i P_j)$$

$$\begin{aligned}
&= \varepsilon_j k_i (P_j R_i - R_j P_i) \\
&= \varepsilon_j k_i (R_i P_j - P_i R_j - 2i\hbar \delta_{ij}) \\
&= -\varepsilon_j k_i (P_i R_j - R_i P_j),
\end{aligned}$$

grâce, une fois de plus, à la propriété de transversalité  $\hat{\mathbf{e}} \cdot \mathbf{k} = 0$ . On peut donc écrire:

$$\begin{aligned}
(\hat{\mathbf{e}} \cdot \mathbf{P})(\mathbf{k} \cdot \mathbf{R}) - (\hat{\mathbf{e}} \cdot \mathbf{R})(\mathbf{k} \cdot \mathbf{P}) &= \frac{1}{2}(\varepsilon_i k_j - \varepsilon_j k_i)(P_i R_j - R_i P_j) \\
&= (\mathbf{k} \wedge \hat{\mathbf{e}}) \cdot (\mathbf{R} \wedge \mathbf{P}).
\end{aligned}$$

L'apparition de l'opérateur moment cinétique orbital,  $\mathbf{L} = \mathbf{R} \wedge \mathbf{P}$ , nous signifie de ne peut plus nous contenter de négliger la contribution du spin du quanton qui figurait initialement dans l'hamiltonien d'interaction (VI.2). Le rotationnel de l'opérateur de champ fait encore apparaître un facteur  $\mathbf{k} \wedge \hat{\mathbf{e}}$ , et nous pouvons finalement — à l'ordre zéro du développement de l'exponentielle dans ce terme — tenir compte sans surprise de cette interaction en ajoutant la contribution  $g\mathbf{S}$  à l'opérateur  $\mathbf{L}$ . L'élément de matrice de transition correspondant au deuxième terme dûment complété est donc:

$$\frac{1}{2} (\mathbf{k} \wedge \hat{\mathbf{e}}) \cdot \langle A | (\mathbf{L} + g\mathbf{S}) | B \rangle, \quad (\text{VI.57})$$

manifestation, au facteur  $q/2m$  près, de l'opérateur moment magnétique quantique, qui s'était révélé dans (II.21), et ainsi nommé pour sa ressemblance formelle avec le moment magnétique classique tel qu'il intervient dans le premier terme du développement de l'énergie d'interaction d'un système de particules et d'un champ magnétique. Le caractère vectoriel de cet opérateur lui dicte le même comportement dans les rotations que le moment dipolaire électrique. Une telle transition est dite *dipolaire magnétique*. Ses règles de sélection sont, puisqu'il s'agit d'un opérateur vectoriel pour les rotations, mais pseudo pour les réflexions (il se réfléchit comme un produit vectoriel d'authentiques vecteurs):

$$\begin{cases} j_B \rightarrow j_A = j_B, & |j_B \pm 1|, \quad \text{sauf } 0 \rightarrow 0, \\ \pi_A = \pi_B. \end{cases}$$

Remarquez que lorsqu'il y a plusieurs quantons, l'expression de l'opérateur moment magnétique est  $(q/2m) \sum_{i=1}^Z (\mathbf{L}_i + g\mathbf{S}_i)$ , et que leur moment angulaire total,  $j$ , est entier s'ils sont en nombre  $Z$  pair.

Ce mécanisme de transition dipolaire magnétique agit généralement en collaboration avec le mécanisme quadripolaire électrique sauf lorsque leurs règles de sélection ne sont pas satisfaites simultanément, par exemple une transition  $\frac{1}{2}- \rightarrow \frac{1}{2}-$ , une transition  $2+ \rightarrow 0+$ , ou une transition entre états propres de  $\mathbf{S}_{\text{tot}}^2$  satisfaisant toutes les règles de sélection communes sur  $j_A$ ,  $j_B$ ,  $\pi_A$ , et  $\pi_B$ , mais avec  $s_A \neq s_B$ .

Pour sa part, le terme dépendant du spin dans l'hamiltonien d'interaction (VI.2) permet aux neutrons de participer à l'émission spontanée de rayonnement électromagnétique par un noyau. Toujours pour simplifier le récit, dans le cas d'une transition impliquant un seul neutron — dont la charge électrique  $q_n$  est nulle — il ne peut y avoir de contributions électrique ou magnétique orbitale et la première contribution, si elle n'est pas nulle, est du type magnétique dipolaire avec un élément de matrice de transition effectif, déduit de (VI.57):

$$\frac{1}{2} (\mathbf{k} \wedge \hat{\mathbf{e}}) \cdot \langle A | g_n \mathbf{S} | B \rangle,$$

où, empiriquement,  $g_n \approx -3,8$  (voir page 27).

Plus généralement, la technique d'analyse en opérateurs tensoriels irréductibles des produits de  $\mathbf{P}$  et  $\mathbf{S}$  avec les termes  $(\mathbf{k} \cdot \mathbf{R})^n$  correspondants aux différents ordres du développement de  $e^{-i\mathbf{k} \cdot \mathbf{R}}$  permet d'obtenir les règles de sélection pour tous les ordres. A une transition  $B \rightarrow A$ , sont susceptibles de contribuer tous les opérateurs  $2l$ -polaires tels que

$$|j_A - j_B| \leq l \leq j_A + j_B,$$

avec

$$\pi_A = \pm (-)^l \pi_B \quad \text{pour les opérateurs} \quad \begin{cases} \text{électriques} \\ \text{magnétiques} \end{cases}, \text{ dits } \begin{cases} E_l \\ M_l \end{cases}.$$

La contribution dominante vient de l'ordre le plus bas non nul. En passant d'un ordre multipolaire au suivant, ainsi qu'en passant, à ordre égal, d'une transition électrique à une transition magnétique, l'ordre de grandeur de la probabilité de transition est réduit par un facteur  $(r/\lambda)^2$ .

Pourquoi chez les atomistes va-t-il le plus souvent sans dire que les transitions sont dipolaires électriques, tandis que les divers ordres multipolaires font les choux gras des spectroscopistes nucléaires? Nous nous attendons à une réduction de la probabilité de transition lorsque la transition dipolaire électrique est interdite par une règle de sélection. Effectivement, l'expérience montre que lorsque la transition est interdite, ou plutôt empêchée, la durée de vie de l'état atomique initial est très grande, de  $10^{-3}$  à 1 s. Paradoxalement, à cette stabilité de l'état correspond une plus grande difficulté d'observation, puisque la raie émise est peu intense, et ces états *métastables* ont peu d'occasions de se manifester en laboratoire. Pour les noyaux, les durées de vie dans les transitions dipolaires électriques sont, d'emblée, plus courtes que pour les atomes. De plus, les facteurs d'empêchement  $(r/\lambda)^2$  sont moins prohibitifs:  $r$  passe de 1 Å à 1 fm =  $10^{-5}$  Å, tandis que l'énergie de transition passe de la dizaine d'électron-volts à la dizaine de Mev. Enfin, deux causes différentes peuvent venir favoriser les transitions multipolaires dans les noyaux, en dotant les nucléons de charges et facteurs  $g$  effectifs plus grands que pour un nucléon isolé:

- Dans un noyau les distances mutuelles entre nucléons sont de l'ordre de leur rayon, rayon qui marque les limites de pertinence du modèle du quanton. A ces

distances, les interactions (par exemple l'échange de pion) entre les structures complexes que sont chacun des nucléons peuvent modifier les valeurs effectives des paramètres supposés résumer leurs propriétés.

- Dans certains états nucléaires, les nucléons peuvent se livrer à des comportements collectifs qui apparentent le noyau plus à une goutte liquide qu'à un sac de billes. La transition est alors mieux décrite en considérant l'évolution globale de la goutte, avec sa densité de charge et de courant, plutôt que la somme des contributions de nucléons indépendants, représentation que l'on peut conserver en modifiant phénoménologiquement les caractéristiques des nucléons participant à la transition.

Pour toutes ces raisons, les probabilités de transitions multipolaires au delà de  $E_1$  restent beaucoup plus grandes dans les noyaux que dans les atomes. A l'intensité accrue des raies s'ajoute la facilité technique de détecter des photons d'origine nucléaire qui passeront d'autant moins inaperçus que l'énergie qu'ils sont à même d'abandonner dans le détecteur est plus élevée.

En revanche, un atome étant électriquement neutre, il n'est pas, contrairement au noyau, protégé des influences extérieures par une barrière coulombienne repoussant toute tentative d'approche d'un système chargé dont l'énergie serait insuffisante. Ainsi les comportements d'ensemble d'atomes peuvent être fort différents des comportements individuels et, dans certaines circonstances, des raies interdites atteignent des intensités beaucoup plus grandes que les raies permises. Par exemple, des raies interdites se manifestent de façon particulièrement spectaculaire dans les aurores boréales ou le rayonnement des nébuleuses [18, § 109].

Enfin, des transitions peuvent ne satisfaire aucune des règles de sélection des transitions multipolaires. C'est le cas par exemple pour la transition  $2s\frac{1}{2} \rightarrow 1s\frac{1}{2}$  de l'atome d'hydrogène. Mais nous n'avons jusqu'à présent fait appel qu'à la règle d'or de Fermi, c'est-à-dire au premier ordre des perturbations. En passant au second ordre, toujours avec l'hamiltonien d'interaction (VI.2), on rend compte de telles transitions (très lentes mais néanmoins existantes) par l'émission spontanée de deux photons (peu probable mais possible), laquelle prédit par exemple pour l'état  $2s\frac{1}{2}$  de l'hydrogène une durée de vie de 0,14 s, en regard de  $1,6 \times 10^{-9}$  s pour les états  $2p$ .

## Exercices



## Chapitre VII

# Histoires de photons

La théorie quantique du rayonnement n'est pas limitée à la description des traditionnels processus d'émission ou d'absorption lors d'une transition atomique entre deux niveaux, et nous allons en voir quelques applications. L'expérimentatrice peut légitimement se demander quel rapport il y a entre son photon, bien concret (tel qu'elle le détecte avec un scintillateur et un photomultiplicateur), et celui du théoricien (paquet d'énergie et d'impulsion associées à un état stationnaire du système dit "rayonnement quantique"). L'analyse théorique du mécanisme de la détection, sur un modèle de détecteur quelque peu idéalisé il est vrai, permettra d'éclairer, sinon éclaircir, la question. Nous allons également étudier l'équilibre thermodynamique du rayonnement et quelques processus de diffusion de la lumière dont l'analyse se révèle particulièrement directe dans le cadre de la théorie quantique. Tous les avatars classiques du rayonnement, la propagation, la réflexion, la diffusion, les interférences et l'effet Doppler peuvent être retrouvées dans la théorie quantique. Une telle conformité des résultats des théories classique et quantique nous conduira finalement à aborder, à l'autre extrémité du spectre des préoccupations de la lectrice, la question de la nécessité de la théorie quantique du rayonnement. Le photon, au-delà de la cohérence intellectuelle et de la simplification technique qu'il apporte à notre description de la nature, est-il absolument nécessaire?

### VII.1 La détection des photons

Les exemples abondent de dispositifs les plus divers — des plus courants (littéralement?) aux plus élaborés — qui sont plus ou moins explicitement considérés comme des détecteurs de photons. Citons par exemple:

- l'atome de la photocathode d'un photomultiplicateur, efficace dans le domaine des  $X$  (pour des photons d'énergie supérieure, la photocathode doit



- être précédée d'un scintillateur qui joue le rôle de convertisseur),
- le cristal de bromure d'argent de l'émulsion d'une pellicule photographique,
- un bâtonnet de la rétine de l'œil,
- les systèmes à photo-ionisation, du type chambre à fils, *etc.*

Dans tous ces systèmes, la détection d'un photon est en fait la matérialisation macroscopique, au moyen d'un amplificateur convenable — électronique ou chimique — de l'excitation d'un système matériel quantique (atome ou molécule généralement) bien localisé. Le détecteur est en pratique constitué d'un grand nombre de systèmes quantiques élémentaires identiques. L'information sur la "position" du photon, autrement (et bien mieux) dit la position du système excité, est, selon les types de détecteurs, transmise avec plus ou moins de précision au cours de l'amplification. Dans le cas de la photocathode, la précision s'étale sur toute la surface de celle-ci, tandis qu'avec la pellicule photo, elle reste de la taille du grain de bromure de l'émulsion. Dans une chambre à bulles, ce n'est que s'il a une énergie suffisante pour créer une paire électron-positron — ou, tout au moins, éjecter un électron d'une molécule — que le photon se signale; dans le cas contraire, l'unique bulle qui dénote la simple excitation d'un atome se confond avec tous les germes d'ébullition parasites.

Idéalisons le détecteur de photons en le réduisant à un atome qui attend, dans son état fondamental  $A$ , d'être excité, à un état  $B$ , par absorption dipolaire électrique d'un photon. Pour simplifier encore l'écriture, supposons que les éléments de matrice des composantes de l'opérateur dipolaire électrique entre les états  $A$  et  $B$  sont tous nuls sauf un. La nature vectorielle du champ électrique n'intervient plus et l'hamiltonien effectif du système constitué par les électrons de l'atome et le rayonnement quantique est, en représentation d'interaction, de la forme

$$H_{\text{el}} - \mathcal{P} E(\mathbf{r}, t), \quad (\text{VII.1})$$

pour un atome fixé au point  $\mathbf{r}$  et dont l'opérateur moment dipolaire électrique est  $\mathcal{P}$ . Un système amplificateur que nous n'étudierons pas ici — ce qui ne signifie nullement qu'il ne pose aucun problème fondamental quant à l'interprétation de la mesure en théorie quantique — a pour fonction de manifester macroscopiquement et irréversiblement que l'atome placé en  $\mathbf{r}$  a effectué la transition  $|A\rangle \rightarrow |B\rangle$ .

Un tel modèle met bien en évidence que  $\mathbf{r}$  n'est rien d'autre que l'endroit où l'on absorbe éventuellement un photon ou, pour reprendre notre métaphore musicale, l'emplacement du chevalet (l'atome) qui se laisse exciter en absorbant une partie de l'état de vibration, non localisé, de la corde de guitare (le champ, évidemment!). Le photon, pas plus que l'état de vibration de la corde ne possède de variable dynamique du type position spatiale. A un contradicteur qui, par exemple, s'obstinerait à voir des positions de photons révélées dans une émulsion placée dans la zone d'interférence entre une lumière monochromatique incidente normalement à un miroir et la lumière réfléchie (les franges de Lippman), la lectrice pourra

rétorquer que dans la même expérience avec une émulsion dont la sensibilité serait fondée sur une transition dipolaire magnétique, les plans de prétendue présence des photons seraient les plans ventraux du champ magnétique, c'est-à-dire décalés de  $\lambda/4$  par rapport aux précédents!

Reprenant le développement modal (III.12), nous pouvons écrire notre opérateur champ électrique effectif:

$$E(\mathbf{r}, t) = E^{(+)}(\mathbf{r}, t) + E^{(-)}(\mathbf{r}, t),$$

avec

$$\begin{aligned} E^{(+)}(\mathbf{r}, t) &\stackrel{\text{df}}{=} i \sum_{\mathbf{k}T} \sqrt{\frac{\hbar\omega}{2\varepsilon_0\mathcal{V}}} a_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)}, \\ E^{(-)}(\mathbf{r}, t) &\stackrel{\text{df}}{=} (E^{(+)}(\mathbf{r}, t))^+. \end{aligned}$$

En raison de leurs dépendances temporelles canoniques, en  $e^{\mp i\omega t}$ , les contributions  $E^{(+)}$  et  $E^{(-)}$  sont dites respectivement à *fréquences positives* et à *fréquences négatives*.<sup>1</sup> L'intérêt d'une telle décomposition est que, au premier ordre des perturbations et dans l'approximation de l'onde tournante (page 44), seules les fréquences positives, celles qui sont associées à l'annihilation d'un photon, interviennent sur le fondamental de l'atome. Vis à vis d'un hamiltonien tel que (VII.1), dit en représentation d'interaction, tout état du rayonnement est un état stationnaire de l'hamiltonien non perturbé  $H_{\text{el}}$ , puisqu'indifférent à celui-ci. On peut ainsi écrire le développement d'un état général du système atome-rayonnement sur les états stationnaires du système non perturbé,

$$\sum_{I,\psi} c_{I\psi}(t) e^{-i\frac{E_I t}{\hbar}} |I\rangle|\psi\rangle,$$

sans avoir à préciser autrement les états  $|\psi\rangle$  du rayonnement. Partis de l'état initial  $|A\rangle|\psi_i\rangle$ , nous obtenons alors, en guise d'amplitude de l'état  $|B\rangle|\psi_f\rangle$  à l'instant  $t$ , dans les susdites approximations,

$$i\hbar c_{Bf}(t) \approx - \int_0^t dt' e^{i\omega_0 t'} \langle B|\mathcal{P}|A\rangle \langle \psi_f|E^{(+)}(\mathbf{r}, t')|\psi_i\rangle,$$

avec  $\omega_0 \stackrel{\text{df}}{=} (E_B - E_A)/\hbar$ , et la probabilité de transition correspondante:

$$\begin{aligned} \text{Pr}_{Bf}(t) &= \frac{1}{\hbar^2} |\langle B|\mathcal{P}|A\rangle|^2 \\ &\times \int_0^t dt' \int_0^t dt'' e^{i\omega_0(t''-t')} \langle \psi_i|E^{(-)}(\mathbf{r}, t')|\psi_f\rangle \langle \psi_f|E^{(+)}(\mathbf{r}, t'')|\psi_i\rangle. \end{aligned}$$

---

<sup>1</sup>Attention: par analogie avec la dépendance temporelle conventionnellement adoptée pour un état stationnaire en mécanique quantique, ce sont bien les contributions en  $e^{-i\omega t}$  qui sont qualifiées de fréquence positive!

Dans le détecteur, seul est amplifié le signal de l'excitation de l'atome; l'état final du rayonnement reste généralement inaperçu. Après sommation sur l'ensemble complet de ces états on obtient, grâce à la relation de fermeture, la probabilité d'excitation du détecteur à l'état  $B$ :

$$\Pr_B(t) = \frac{1}{\hbar^2} |\langle B|\mathcal{P}|A\rangle|^2 \int_0^t dt' \int_0^t dt'' e^{i\omega_0(t''-t')} \langle \psi_i | E^{(-)}(\mathbf{r}, t') E^{(+)}(\mathbf{r}, t'') | \psi_i \rangle.$$

En pratique, le détecteur est à *bande large*, c'est-à-dire que ses états excités accessibles se répartissent quasi continument dans un domaine plus large que la fenêtre de fréquences du rayonnement pour lequel il est prévu. Dans ces conditions, la probabilité de réponse du détecteur est

$$\Pr(t) = \int dE_B \rho(E_B) \Pr_B(t),$$

en fonction de la densité  $\rho(E)$  de niveaux de l'atome autour de l'énergie  $E$ . En admettant que cette densité, ainsi que les éléments de matrice du moment dipolaire électrique, dépendent peu de l'énergie finale de l'atome, comparés à l'exponentielle  $e^{i\omega_0(t''-t')}$ , et grâce à la relation

$$\int dE_B e^{\frac{i}{\hbar}(E_B-E_A)(t''-t')} = 2\pi\hbar \delta(t''-t'),$$

on a:

$$\Pr(t) = \frac{2\pi}{\hbar} |\langle B|\mathcal{P}|A\rangle|^2 \rho(E_B) \int_0^t dt' \langle \psi_i | E^{(-)}(\mathbf{r}, t') E^{(+)}(\mathbf{r}, t') | \psi_i \rangle.$$

Le taux de transition du détecteur,  $w \stackrel{\text{df}}{=} (d/dt) \Pr$ , est donc finalement de la forme:

$$w(t) = S \langle \psi_i | E^{(-)}(\mathbf{r}, t) E^{(+)}(\mathbf{r}, t) | \psi_i \rangle, \quad (\text{VII.2})$$

où j'ai factorisé dans le coefficient  $S$  tout ce qui est caractéristique du seul détecteur.

Ce n'est qu'en passant sur des difficultés d'ordres fondamental comme technique — déjà mentionnées à propos de l'émission spontanée — que l'on peut assimiler d'emblée  $w$  au taux de comptage. Vouloir s'interroger sur le mécanisme d'un détecteur nous fait retomber dans les affres de l'énigme de la mesure en théorie quantique. La notion de probabilité, telle qu'elle intervient dans les règles de base de la théorie pour donner une signification au module carré du produit scalaire des kets d'état, est à vrai dire mal définie. Vouloir lui donner une définition opérationnelle en assimilant implicitement probabilité et fréquence relative des résultats d'un échantillon statistique de mesures ne fait que repousser le problème. Le concept de mesure n'est pas mieux défini dans la théorie, sauf à faire intervenir dans celle-ci l'acte volontaire d'une observatrice tout aussi peu définie, promouvant par là même

la théorie quantique à la dignité de science humaine! Outre cet abîme de perplexité où nous plonge la réflexion sur un banal taux de comptage, subsiste la question du lien entre la probabilité  $\text{Pr}(t)$ , calculée dans les conditions de la règle d'or de Fermi — c'est-à-dire pour une durée intermédiaire —, et la probabilité d'assister à une transition entre les instants  $t$  et  $t + dt$  dont on dérive directement le taux de comptage. L'analyse du processus de détection des photons amène donc des complications, étudiées par Glauber [31], mais qui ne semblent fort heureusement que théoriques.

Dans la pratique, l'expression (VII.2) semble bien donner toute satisfaction. C'est la composante à fréquence positive,  $E^{(+)}$ , qui y joue le rôle de champ électrique dont le détecteur mesure la valeur moyenne. Cette composante est une combinaison linéaire d'opérateurs d'annihilation; il est ainsi tout à fait compréhensible que les états cohérents — qui définissent au mieux le champ électrique — aient pour propriété caractéristique d'être états propres des opérateurs d'annihilation.

Un détecteur un peu plus réaliste va répondre à toutes les composantes du champ électrique, et son taux de transition sera donné par une expression tensorielle analogue:

$$w(t) = S_{pq} \langle \psi_i | E_p^{(-)}(\mathbf{r}, t) E_q^{(+)}(\mathbf{r}, t) | \psi_i \rangle, \quad (\text{VII.3})$$

où

$$S_{pq} = \frac{2\pi}{\hbar} \rho(E_B) \langle A | \mathcal{P}_p | B \rangle \langle B | \mathcal{P}_q | A \rangle.$$

Enfin, plus généralement, l'état initial du rayonnement peut ne pas être pur. Dans un état impur, le rayonnement est représenté par un mélange statistique de kets  $|\psi_i\rangle$  avec les probabilités  $\text{Pr}_{\psi_i}$ . Le taux de transition du détecteur exposé à cette population hétérogène s'écrit alors

$$w(t) = S_{pq} G_{pq}^{(1)}(\mathbf{r}, t, \mathbf{r}, t),$$

avec

$$G_{pq}^{(1)}(\mathbf{r}, t, \mathbf{r}', t') \stackrel{\text{df}}{=} \text{Tr} \{ \rho E_p^{(-)}(\mathbf{r}, t) E_q^{(+)}(\mathbf{r}', t') \}.$$

Dans la trace à calculer apparaît la notion commode d'*opérateur densité* du rayonnement,  $\rho \stackrel{\text{df}}{=} \sum_{\psi_i} |\psi_i\rangle \text{Pr}_{\psi_i} \langle \psi_i|$ , à ne pas confondre, nonobstant son symbole, avec la densité de niveaux du détecteur. La lectrice qui rencontre pour la première fois cet opérateur densité pourra s'y initier dans [53, § VIII-21] et, à tout le moins, vérifier que dans le cas pur,  $\rho = |\psi_i\rangle \langle \psi_i|$ , elle retrouve bien l'expression (VII.3) et, dans le cas impur, une moyenne de celle-ci pondérée par les probabilités  $\text{Pr}_{\psi_i}$ .

Les caractéristiques du rayonnement, ou tout au moins celles qui importent au détecteur que nous étudions, sont entièrement contenues dans la *fonction de cohérence du premier ordre* du rayonnement mesuré,  $G^{(1)}(\mathbf{r}, t, \mathbf{r}', t)$ . Les fonctions de cohérence d'ordre supérieur joueront un rôle dans la comparaison entre théories classique et quantique du rayonnement (§ ??).

section

## VII.2 L'équilibre thermodynamique matière-rayonnement

Comment notre théorie quantique du rayonnement accomode-t-elle le vieux problème du spectre du corps noir? Apporte-t-elle un éclairage microscopique aux considérations thermodynamiques globales d'Einstein (voir par exemple [15, p. 83], [24], [59, p. 406]) qui, de la théorie classique de la densité spectrale d'énergie du rayonnement à basse fréquence (la loi de Rayleigh-Jeans) et de l'hypothèse d'égalité des probabilités d'absorption et d'émission stimulée entre deux niveaux atomiques donnés (le principe du bilan détaillé), déduisait tout à la fois:

- une relation entre coefficient  $\mathcal{A}$  d'émission spontanée et coefficient  $\mathcal{B}$  d'absorption/émission stimulée,
- et la densité spectrale du rayonnement pour toute fréquence, ou loi de Planck.

Encore autrement dit, le schéma:

$$\left\{ \begin{array}{ll} \mathcal{B}_{\text{ém}} = \mathcal{B}_{\text{abs}} & \text{(bilan détaillé)} \\ u(\omega) = \frac{1}{\pi^2} \frac{\omega^2}{c^3} k_B T & \text{(Rayleigh-Jeans)} \end{array} \right. \xrightarrow[\text{à haute température}]{\text{Einstein, et équilibre}} \left\{ \begin{array}{ll} \mathcal{A} = \frac{\hbar \omega^3}{\pi^2 c^3} \mathcal{B} & \text{(Einstein)} \\ u(\omega) = \frac{\hbar}{\pi^2 c^3} \frac{\omega^3}{e^{\frac{\hbar \omega}{k_B T}} - 1} & \text{(Planck)} \end{array} \right. , \quad (\text{VII.4})$$

peut-il être un résultat de notre théorie?<sup>2</sup>

La théorie quantique du rayonnement ne distingue plus les phénomènes d'émissions stimulée et spontanée: un seul mécanisme les préside (voir § VI.2). D'une certaine façon, elle nous fournit comme point de départ les coefficients d'Einstein, et nous allons voir que l'on peut en déduire très simplement la densité spectrale du rayonnement à l'équilibre thermodynamique; autrement dit:

$$\left\{ \begin{array}{l} \mathcal{B}_{\text{ém}} \\ \mathcal{B}_{\text{abs}} \end{array} \right. \xrightarrow[\text{équilibre}]{} u(\omega) \quad (\text{Planck}). \quad (\text{VII.5})$$

Pour un atome dans l'état initial  $A$  (stationnaire si le rayonnement n'existait pas), la probabilité de l'état  $B$  après l'écoulement du temps  $t$  vaut  $\text{Pr}_{A \rightarrow B}(t) = \Gamma_{AB} t$ , avec les réserves encore mentionnées dans la section précédente à propos des scrupules légitimes, et pourtant sans conséquences, que peuvent faire naître cette utilisation heuristique...

- de la règle d'or de Fermi, étendue des temps intermédiaires à des temps quelconques,

---

<sup>2</sup>Ne pas confondre, dans cette section: la constante de Boltzmann  $k_B$  et le module  $k$  du vecteur de mode, la température  $T$  et l'indice de mode de polarisation  $T$ .

— et de la notion de probabilité qui, mal définie dans la théorie quantique, acquiert subrepticement, par son utilisation, la signification d'une fréquence relative.

On a de même  $\text{Pr}_{B \rightarrow A}(t) = \Gamma_{BA} t$ . Imaginons alors un ensemble d'atomes, identiques, en interaction avec le rayonnement, ces atomes étant suffisamment dilués dans l'espace pour être dispensés d'interactions mutuelles directes. Sauf repos éternel dans l'état de zéro absolu, les grandeurs physiques de ce système sont en perpétuelle évolution. En particulier, du fait des transitions d'émission et d'absorption, les nombres  $n_A$  et  $n_B$  d'atomes dans les états  $A$  et  $B$  et le nombre  $n_{\mathbf{k}T}$  de photons de l'état du rayonnement dans le mode  $\mathbf{k}T$  ne cessent de fluctuer. Dans la mesure où le nombre d'atomes est élevé, ces fluctuations relatives aux valeurs moyennes doivent être assez faibles et on peut admettre que, sur une durée  $\Delta t$ , la variation du nombre moyen d'atomes dans l'état  $A$  est donnée par

$$\Delta \bar{n}_A = -\bar{n}_A \bar{\Gamma}_{AB} \Delta t + \bar{n}_B \bar{\Gamma}_{BA} \Delta t.$$

J'ai pris soin de rappeler qu'il faut considérer des taux de transition moyens; en effet, les taux d'absorption et d'émission (équations (VI.3) et (VI.5)) fluctuent car ils dépendent du nombre de photons  $n_{\mathbf{k}T}$ , lui-même fluctuant.

A l'équilibre thermodynamique, par définition, les grandeurs macroscopiques que sont les nombres moyens n'évoluent plus,  $\Delta \bar{n}_A = 0$ . Dans ces conditions, on a  $\bar{n}_A \bar{\Gamma}_{AB} = \bar{n}_B \bar{\Gamma}_{BA}$ , et donc:

$$\frac{\bar{n}_B}{\bar{n}_A} = \frac{\bar{\Gamma}_{AB}}{\bar{\Gamma}_{BA}} = \frac{\bar{n}_{\mathbf{k}T}}{\bar{n}_{\mathbf{k}T} + 1}. \quad (\text{VII.6})$$

La théorie quantique du rayonnement, en nous faisant cadeau des expressions (VI.3) et (VI.5), nous a permis de faire l'économie du principe du bilan détaillé et du traitement séparé de l'émission spontanée pour évaluer  $\bar{\Gamma}_{AB}/\bar{\Gamma}_{BA}$ . Nous n'avons plus qu'un seul mécanisme d'émission, le 1 qui s'additionne au dénominateur représentant la contribution dite spontanée lorsqu'il n'y a aucun photon (rayonnement dans son état fondamental).

Nous considérons un système atomes-rayonnement macroscopique; l'enceinte qui contient atomes et rayonnement, quand bien même serait elle isolée, est assez grande pour constituer elle-même un thermostat vis à vis des atomes. A l'équilibre avec ce thermostat, les nombres moyens d'atomes dans les divers états suivent la distribution canonique, soit, pour des niveaux non dégénérés,

$$\frac{\bar{n}_B}{\bar{n}_A} = \frac{e^{-\beta E_B}}{e^{-\beta E_A}} = e^{-\beta \hbar \omega}, \quad (\text{VII.7})$$

avec  $\beta \stackrel{\text{df}}{=} 1/k_B T$ , où  $T$  est la température canonique du système. En identifiant cette expression avec la condition d'équilibre (VII.6), on obtient le nombre moyen

de photons dans le mode  $\mathbf{k}T$ :

$$\bar{n}_{\mathbf{k}T} = \frac{1}{e^{\beta\hbar\omega} - 1}. \quad (\text{VII.8})$$

Rappelons pour la lectrice férue de physique statistique qu'elle avait obtenu ce résultat...

- en prenant la distribution de Bose des nombres d'occupation moyens des niveaux d'énergie  $\varepsilon$  pour un gaz parfait,  $\bar{n}_\varepsilon = 1/(e^{\beta(\varepsilon - \mu)} - 1)$ ,
- et en appliquant la loi d'action de masse à l'équilibre photons-paroi ( $\gamma \leftrightarrow 0$ ), pour obtenir le potentiel chimique des photons,  $\mu_\gamma = 0$ .

Du nombre moyen de photons on déduit la densité d'énergie moyenne (énergie moyenne par unité de volume) du rayonnement dont la pulsation est comprise entre  $\omega$  et  $\omega + d\omega$ :

$$u(\omega) d\omega = \frac{1}{\mathcal{V}} \sum_{\mathbf{k}, T} \bar{n}_{\mathbf{k}T} \varepsilon(\mathbf{k}),$$

où la sommation sur  $\mathbf{k}$  est limitée aux vecteurs d'onde dont l'extrémité est dans la coquille de rayon  $k = \omega/c$  et d'épaisseur  $dk = d(\omega/c)$ . Soit, selon l'habituel comptage de modes,  $4\pi k^2 dk / (2\pi/L)^3$  vecteurs  $\mathbf{k}$  correspondant, à toute fin pratique, à la même pulsation  $\omega$ . On a donc

$$\begin{aligned} u(\omega) d\omega &= \frac{1}{\mathcal{V}} \sum_{T=1}^2 \frac{1}{e^{\beta\hbar\omega} - 1} \hbar\omega \frac{4\pi k^2 dk}{(2\pi/L)^3}, \\ &= \frac{1}{\mathcal{V}} \frac{2}{e^{\beta\hbar\omega} - 1} \hbar\omega \frac{\mathcal{V}}{2\pi^2} \left(\frac{\omega}{c}\right)^2 d\left(\frac{\omega}{c}\right), \end{aligned}$$

d'où la loi de Planck annoncée pour la densité spectrale:

$$u(\omega) = \frac{\hbar}{\pi^2 c^3} \frac{\omega^3}{e^{\beta\hbar\omega} - 1}. \quad (\text{VII.9})$$

Cette expression régurgite la loi de Rayleigh-Jeans à la limite des hautes températures (ou des basses fréquences, puisque la température n'intervient qu'en rapport avec la pulsation):

$$u(\omega) \underset{\frac{\hbar\omega}{k_B T} \ll 1}{\sim} \frac{1}{\pi^2} \frac{\omega^2}{c^3} k_B T. \quad (\text{VII.10})$$

La simple application du principe zéro de la physique au rayonnement donnait (presque) aussi bien cette densité spectrale (Exercice 10), au facteur  $\pi^2 \approx 10$  près! Les physiciens classiques, quelle que soit leur analyse, ne pouvaient trouver d'autre loi, et pour échapper à la divergence de ce spectre à haute fréquence (la catastrophe ultraviolette) il fallait l'invention d'une nouvelle donnée du problème, la constante

de Planck. Dès qu'intervient la nouvelle constante fondamentale  $\hbar$ , tous les espoirs sont permis puisque, au moins, rien n'interdit plus l'existence d'une fonction de la grandeur sans dimension  $\hbar\omega/k_B T$ , intervenant éventuellement en facteur de la loi de Rayleigh-Jeans et la régularisant. La théorie quantique (la constante  $\hbar$ ) n'existait pas, il fallait l'inventer, c'est ce que fit Planck. Mais évidemment, à part en autoriser l'existence, le principe zéro est incapable de toute prédiction concernant une fonction sans dimension.

N'oublions pas que la densité spectrale donnée par la loi de Planck (VII.9) ne concerne qu'un rayonnement en équilibre thermodynamique. Cette restriction s'est introduite dès l'instant que nous avons assigné la distribution canonique (VII.7) aux nombres moyens d'occupation des états atomiques. C'est précisément cette répartition qui est différente dans un laser (le premier problème dans la réalisation d'un laser est d'assurer le déséquilibre thermodynamique qui conduira à l'inversion de population  $\bar{n}_B > \bar{n}_A$ ), et la densité spectrale du rayonnement laser n'a rien à voir avec l'expression (VII.9), dans laquelle on serait d'ailleurs bien en peine d'insérer une valeur de température puisque, hors équilibre, cette dernière n'est même pas définie. Dans le même ordre d'idées, notons que l'existence du corps noir, avec tout son spectre (VII.9) — pour toute pulsation  $\omega$  —, nécessite implicitement une boîte à modes — une cavité — assez grande pour que le spectre soit pratiquement continu et, surtout, dont les parois jouent le rôle d'atomes supplémentaires dont les niveaux permettent de thermaliser toutes les pulsations du rayonnement autres que  $(E_B - E_A)/\hbar$ .

Le taux d'absorption  $\bar{\Gamma}_{AB}$  du niveau fondamental est proportionnel au nombre de photons  $\bar{n}_{\mathbf{k}T}$  dans le mode, nombre lui-même lié à la densité spectrale par la relation

$$\bar{n}_{\mathbf{k}T} = \frac{\pi^2 c^3}{\hbar \omega^3} u(\omega).$$

Le taux d'absorption est donc proportionnel à la densité spectrale, avec un facteur de proportionnalité  $\mathcal{B}$ , caractéristique propre à l'atome, indépendant du rayonnement, appelé *coefficient  $\mathcal{B}$  d'Einstein*:

$$\bar{\Gamma}_{AB} = \mathcal{B} u(\omega). \quad (\text{VII.11})$$

Le taux d'émission du niveau  $B$  est, quant à lui, proportionnel à  $\bar{n}_{\mathbf{k}T} + 1$ , avec le même facteur de proportionnalité, comme nous le rappelle la relation (VII.6). On en déduit donc ce taux,

$$\begin{aligned} \bar{\Gamma}_{BA} &= \frac{\bar{n}_{\mathbf{k}T} + 1}{\bar{n}_{\mathbf{k}T}} \mathcal{B} u(\omega) \\ &= \mathcal{B} u(\omega) + \frac{\hbar \omega^3}{\pi^2 c^3} \mathcal{B}, \end{aligned} \quad (\text{VII.12})$$

soit deux contributions:



- la première, égale au taux d'absorption du niveau  $A$ ,
- et la deuxième, rigoureusement indépendante de l'intensité du rayonnement, caractéristique de la paire de niveaux atomiques considérés, qui a reçu le nom de *coefficient  $\mathcal{A}$  d'Einstein*:

$$\mathcal{A} \stackrel{\text{df}}{=} \frac{\hbar\omega^3}{\pi^2 c^3} \mathcal{B}. \quad (\text{VII.13})$$

Cette dernière contribution, à l'œuvre même en absence de tout rayonnement, était distinguée sous le nom d'émission spontanée avant que la théorie quantique du rayonnement ne vienne unifier les descriptions des processus d'émissions spontanée et stimulée. Les coefficients  $\mathcal{A}$  et  $\mathcal{B}$  sont propres à l'atome, et la relation (VII.13) établie par Einstein sur des considérations d'équilibre thermodynamique avec le rayonnement du corps noir était donc valide pour l'interaction de l'atome avec tout type de rayonnement, les taux de transition dépendant, eux, du rayonnement suivant les relations (VII.11) et (VII.12). J'ajoute, pour rassurer l'éventuelle lectrice étonnée par la relation (VII.13), que la définition (VII.11) du coefficient  $\mathcal{B}$  dépend de la définition de la densité spectrale: si l'on utilise la densité en fréquence — plus traditionnelle que la densité en pulsation —, le coefficient  $\mathcal{B}$  se trouve modifié en sorte que la relation d'Einstein s'écrit  $\mathcal{A} = (2\hbar\omega^3/\pi c^3)\mathcal{B}$ .

*exercice?*

On a, en ordre de grandeur,  $\mathcal{A}/\mathcal{B} \approx \hbar/\lambda^3$ . En passant d'une longueur d'onde de transition,  $\lambda$ , de 1 m à 1  $\mu\text{m}$ , le rapport  $\mathcal{A}/\mathcal{B}$  se trouve multiplié par un facteur  $10^{18}$  ce qui, toutes choses égales par ailleurs, explique que l'on néglige ordinairement la contribution de l'émission spontanée dans l'étude des transitions hertziennes et, inversement, la contribution de l'émission stimulée dans les transitions nucléaires. C'est d'ailleurs l'émission spontanée qui nous éclaire à longueur de journée. En effet, les deux contributions au taux d'émission sont, en cas d'interaction avec un rayonnement de Planck, dans le rapport

$$\frac{\mathcal{A}}{\mathcal{B}u(\omega)} = e^{\beta\hbar\omega} - 1.$$

Au maximum de densité spectrale du rayonnement du Soleil, assimilable tout au moins extérieurement à un corps noir de température  $T \approx 6000 \text{ K}$ , autrement dit à la longueur d'onde typique  $\lambda \approx 0,5 \mu\text{m}$ , la contribution spontanée est environ cent fois plus grande que la contribution stimulée.<sup>3</sup>

Terminons ces considérations un tantinet passéistes sur l'émission stimulée et l'émission spontanée. On lit parfois que cette dernière est un phénomène quantique. Si l'on entend par là un phénomène qui ne peut se produire dans le cadre de la physique classique, c'est une assertion fautive. Un modèle classique d'atome

<sup>3</sup>Au sortir de l'eau, nos lointains ancêtres disposaient déjà d'un détecteur efficace pour cette longueur d'onde car, pur et merveilleux hasard, elle coïncide avec l'étroite fenêtre de transparence de l'eau liquide au rayonnement électromagnétique, fenêtre pour laquelle l'œil avait inévitablement été sélectionné.

constitué par un électron orbitant autour d'un noyau massif peut parfaitement rayonner... Il ne fait même que cela, et il a fallu inventer la mécanique quantique pour l'en empêcher. C'est bien plutôt l'absence de rayonnement spontané du fondamental des atomes qui constitue un phénomène proprement quantique. Dès l'avènement du modèle proto-quantique de l'atome de Bohr qui reprenait à son compte cette constatation, la question était donc posée de trouver un moyen pour que les états excités de l'atome, eux, puissent rayonner. L'étape intermédiaire de la théorie semi-classique du rayonnement — atome quantique de Schrödinger en interaction avec le rayonnement classique — y apportait une réponse globale en permettant le calcul du coefficient  $\mathcal{B}$  d'absorption et, par la relation d'Einstein (VII.13), du coefficient  $\mathcal{A}$  d'émission spontanée.

Mais c'est la théorie quantique du rayonnement qui unifie les concepts d'émission et surtout rend compte de l'évolution d'un état initial constitué par l'atome dans un état excité et le rayonnement dans un état à  $n$  photons (du mode correspondant à la transition) vers un état final du rayonnement à  $n + 1$  photons *du même mode*. Ainsi se trouve expliquée la propriété quantique — dite de *cohérence* — de l'émission stimulée, initialement proposée par Einstein et qui devait, bien plus tard, aboutir à la mise au point du laser.

En raison même de la contradiction fondamentale qu'il y a à considérer un atome quantique comme source de rayonnement classique, il ne peut y avoir, dans le cadre de la théorie semi-classique, d'explication à la cohérence du rayonnement stimulé émis par *un* atome. Sous l'influence d'un rayonnement classique incident — en onde plane par exemple —, l'atome subit une transition et émet un photon, état de rayonnement sans rapport avec une onde plane et concept d'ailleurs totalement étranger à la théorie classique du rayonnement. La question même de la cohérence entre rayonnement incident et rayonnement émis est donc sans objet. *A fortiori*, dans une théorie intégralement classique, *un* dipôle oscillant excité par une onde plane émet son rayonnement sous forme d'onde sphérique — ce qui ne signifie nullement que l'amplitude en soit isotrope — de même fréquence que l'onde incidente bien sûr, mais sans autre caractéristique commune qui permette de leur attribuer une quelconque cohérence. Dès lors que le rayonnement est classique, le rayonnement stimulé par une onde plane devrait lui-même être émis en onde plane pour que la question de la cohérence puisse commencer à se poser. Mais nous avons vu (§ V.5.3) qu'un *état cohérent* du rayonnement quantique — à ne pas confondre avec photons cohérents ou ondes cohérentes —, dans la mesure où il correspond à un nombre moyen de photons élevé, a pratiquement un comportement de rayonnement classique en onde plane. La cohérence classique exige un grand nombre de sources.<sup>4</sup>

---

<sup>4</sup>L'article pédagogique [22], reproduit dans [41], montre comment un plan d'atomes quantiques, ou de dipôles classiques, stimulés par une onde plane, émet une onde plane cohérente, c'est-à-dire de même vecteur d'onde et phase que l'onde incidente.

### VII.3 Diffusion d'un photon par un atome

Sous ce titre, nous allons étudier les processus où, dans l'état initial comme dans l'état final, n'intervient qu'un seul photon, avec les mises en garde maintenant habituelles sur l'assimilation abusive à une bille classique — ou même à un objet quantique doué d'une quelconque position — à laquelle ce vocable lapidaire pourrait induire. Etant donné généralement le processus

$$A + \gamma_{\mathbf{k}T} \rightarrow B + \gamma_{\mathbf{k}'T'}$$

pour lequel, partis d'un état initial avec un atome dans l'état  $A$  (stationnaire en absence d'interaction avec le rayonnement) et le rayonnement quantique dans un état à un photon du mode  $\mathbf{k}T$ , on se demande quelle est la probabilité de trouver, dans l'état final, l'atome dans l'état  $B$  et le rayonnement dans un état à un photon du mode  $\mathbf{k}'T'$ .

Pour alléger l'écriture, je nous restreins, comme d'habitude, au cas d'un seul électron. Rappelons que l'hamiltonien d'interaction atome-rayonnement s'écrit

$$H_{\text{int}} \approx -\frac{q}{m} \mathbf{A}(\mathbf{R}, t) \cdot \mathbf{P} + \frac{q^2}{2m} \mathbf{A}^2(\mathbf{R}, t), \quad (\text{VII.14})$$

où l'opérateur de champ a pour développement modal:

$$\mathbf{A}(\mathbf{r}, t) \stackrel{\text{df}}{=} \sum_{\mathbf{k}T} \sqrt{\frac{\hbar}{2\varepsilon_0\omega\mathcal{V}}} \left\{ a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)} + a_{\mathbf{k}T}^+ \hat{\mathbf{e}}_{\mathbf{k}T}^* e^{-i(\mathbf{k}\cdot\mathbf{r}-\omega t)} \right\}.$$

Seuls les modes  $\mathbf{k}T$  et  $\mathbf{k}'T'$  interviennent dans la définition du problème. Nous négligeons donc — à un infini soustrayable près! — les contributions de tous les autres modes. En écrivant l'hamiltonien d'interaction ci-dessus, nous négligeons l'interaction entre spin et champ magnétique, dont les contributions à l'élément de matrice de transition sont plus faibles... dans la mesure où celles que nous envisageons ne seront pas nulles. Le passage de l'état initial à l'état final exige l'annihilation d'un photon du mode  $\mathbf{k}T$  et la création d'un photon du mode  $\mathbf{k}'T'$ . A l'ordre de perturbation le plus bas, seuls des monômes d'opérateurs  $a_{\mathbf{k}'T'}^+ a_{\mathbf{k}T}$  et  $a_{\mathbf{k}T} a_{\mathbf{k}'T'}^+$  pourront contribuer à la transition, soit:

- au premier ordre, le terme en  $\mathbf{A}^2$ ,
- et au second ordre, le terme en  $\mathbf{A} \cdot \mathbf{P}$ . (Le terme de spin ne contribuerait qu'au second ordre, plus faiblement que celui-ci.)

#### VII.3.1 La contribution diamagnétique

La contribution du terme en  $\mathbf{A}^2$  n'exige qu'un calcul de perturbation dépendant du temps, au premier ordre, analogue à la démonstration de la règle d'or de Fermi.

Dans des notations que j'espère maintenant évidentes, on a l'élément de matrice de transition correspondant:

$$\begin{aligned} \langle B; 0_{\mathbf{k}T}, 1_{\mathbf{k}'T'} | \frac{q^2}{2m} \mathbf{A}^2(\mathbf{R}, t) | A; 1_{\mathbf{k}T}, 0_{\mathbf{k}'T'} \rangle = \\ \frac{q^2}{2m} \frac{\hbar}{2\varepsilon_0 \mathcal{V} \sqrt{\omega\omega'}} \langle B; 0_{\mathbf{k}T}, 1_{\mathbf{k}'T'} | (a_{\mathbf{k}T} a_{\mathbf{k}'T'} + a_{\mathbf{k}'T'}^\dagger a_{\mathbf{k}T}) | A; 1_{\mathbf{k}T}, 0_{\mathbf{k}'T'} \rangle \\ \times \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* e^{i(\mathbf{k}-\mathbf{k}') \cdot \mathbf{R} - (\omega-\omega')t}. \end{aligned}$$

Une fois encore, nous faisons l'approximation de grande longueur d'onde,

$$e^{i(\mathbf{k}-\mathbf{k}') \cdot \mathbf{R}} |A\rangle \approx |A\rangle,$$

justifiée dans la mesure où les longueurs d'onde des photons sont beaucoup plus grandes que la dimension du système lié, ce qui est certainement le cas avec un rayonnement dans le domaine optique,  $\lambda \approx 5000 \text{ \AA}$ , et un atome d'extension  $r \approx 1 \text{ \AA}$ . Moyennant quoi l'élément de matrice vaut

$$\frac{q^2}{2m} \frac{\hbar}{2\varepsilon_0 \mathcal{V} \sqrt{\omega\omega'}} 2 \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* e^{-i(\omega-\omega')t} \langle B|A \rangle.$$

Ainsi, la partie diamagnétique de l'interaction donne — revoir (II.55) — une contribution à l'amplitude de l'état  $|B; 0_{\mathbf{k}T}, 1_{\mathbf{k}'T'}\rangle$  dans l'état du système au temps  $t$  après départ de l'état  $|A; 1_{\mathbf{k}T}, 0_{\mathbf{k}'T'}\rangle$ :

$$\begin{aligned} c_{\text{dia}}(t) \approx \\ \frac{1}{i\hbar} \frac{q^2}{2m} \frac{\hbar}{2\varepsilon_0 \mathcal{V} \sqrt{\omega\omega'}} 2 \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \delta_{AB} \int_0^t dt_1 e^{-\frac{i}{\hbar}(E_A - E_B + \hbar(\omega - \omega'))t_1}. \quad (\text{VII.15}) \end{aligned}$$

### VII.3.2 La contribution paramagnétique

L'opérateur de champ étant un opérateur “à un photon”, le terme en  $\mathbf{A} \cdot \mathbf{P}$  ne peut modifier le nombre de photons dans un mode que d'une unité, raison pour laquelle l'élément de matrice de la partie diamagnétique de l'hamiltonien d'interaction est nul entre l'état initial et l'état final précités. Ainsi, sa contribution à l'amplitude de transition est nulle au premier ordre de perturbation, et il nous faut reprendre l'équation générale (II.54) d'évolution de cette amplitude qui, intégrée de 0 à  $t$ , est équivalente au système d'équations intégrales couplées:

$$c_B(t) = c_B(0) + \frac{1}{i\hbar} \int_0^t dt_2 e^{\frac{i}{\hbar}(E_B - E_I)t_2} \langle B | H_{\text{int}}(t_2) | I \rangle c_I(t_2),$$

pour un système quantique quelconque dont nous notons  $A, B, I \dots$  les états stationnaires non perturbés. Au premier ordre, avec la condition initiale  $c_I(0) = \delta_{IA}$ , nous savons déjà que, de façon générale,

$$c_I^{(1)}(t) = \frac{1}{i\hbar} \int_0^t dt_1 e^{i\frac{E_I - E_A}{\hbar}t_1} \langle I | H_{\text{int}}(t_1) | A \rangle,$$

et que si l'élément de matrice correspondant à l'état  $I = B$  est nul, il nous faut évaluer le deuxième ordre de perturbation pour y déceler une éventuelle contribution qui puisse fournir une meilleure approximation de  $c_B(t)$ . Pour cela, substituons l'approximation du premier ordre dans le deuxième membre des équations intégrales (tout comme nous avons obtenu, par substitution de la condition initiale, l'approximation du premier ordre). On obtient ainsi:

$$c_B^{(2)}(t) = \frac{1}{(i\hbar)^2} \sum_I \int_0^t dt_2 e^{i\frac{E_B - E_I}{\hbar}t_2} \langle B | H_{\text{int}}(t_2) | I \rangle \\ \times \int_0^{t_2} dt_1 e^{i\frac{E_I - E_A}{\hbar}t_1} \langle I | H_{\text{int}}(t_1) | A \rangle,$$

expression qui a plus de chances d'être non nulle puisqu'il pourrait bien exister, parmi le spectre, quelques états  $I$  tels que les éléments de matrice entre  $A$  et  $I$  d'une part,  $I$  et  $B$  d'autre part, ne soient pas nuls.

Appliquant ce résultat à notre système atome-rayonnement pour calculer la contribution du terme diamagnétique, on trouve, après approximation de grande longueur d'onde:

$$c_{\text{para}}(t) \approx \frac{1}{(i\hbar)^2} \left(-\frac{q}{m}\right)^2 \frac{\hbar}{2\varepsilon_0 \mathcal{V} \sqrt{\omega\omega'}} \int_0^t dt_2 \int_0^{t_2} dt_1 \sum_I e^{i\frac{E_B - E_I}{\hbar}t_2} e^{i\frac{E_I - E_A}{\hbar}t_1} \\ \times \left\{ \langle B; 1_{\mathbf{k}'T'}, 0_{\mathbf{k}T} | a_{\mathbf{k}'T'}^+ \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} e^{i\omega't_2} | I; 0_{\mathbf{k}'T'}, 0_{\mathbf{k}T} \rangle \right. \\ \times \langle I; 0_{\mathbf{k}'T'}, 0_{\mathbf{k}T} | a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} e^{-i\omega t_1} | A; 0_{\mathbf{k}'T'}, 1_{\mathbf{k}T} \rangle \\ + \langle B; 1_{\mathbf{k}'T'}, 0_{\mathbf{k}T} | a_{\mathbf{k}T} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} e^{-i\omega t_2} | I; 1_{\mathbf{k}'T'}, 1_{\mathbf{k}T} \rangle \\ \left. \times \langle I; 1_{\mathbf{k}'T'}, 1_{\mathbf{k}T} | a_{\mathbf{k}'T'}^+ \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} e^{i\omega't_1} | A; 0_{\mathbf{k}'T'}, 1_{\mathbf{k}T} \rangle \right\}.$$

A cause des opérateurs de création et d'annihilation, seuls deux états intermédiaires du rayonnement contribuent. Ne restent que des éléments de matrice de ces opérateurs qui sont tout simplement égaux à un. On a donc, en fonction des seuls éléments de matrice atomiques, et après intégration sur  $t_1$ :

$$c_{\text{para}}(t) = -\frac{1}{\hbar} \left(\frac{q}{m}\right)^2 \frac{1}{2\varepsilon_0 \mathcal{V}} \frac{1}{\sqrt{\omega\omega'}}$$

$$\begin{aligned}
& \times \sum_I \left\{ \langle B | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle \right. \\
& \quad \times \int_0^t dt_2 e^{\frac{i}{\hbar}(E_B - E_I + \hbar\omega')t_2} \frac{e^{\frac{i}{\hbar}(E_I - E_A - \hbar\omega)t_2} - 1}{\frac{i}{\hbar}(E_I - E_A - \hbar\omega)} \\
& \quad + \langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | A \rangle \\
& \quad \times \left. \int_0^t dt_2 e^{\frac{i}{\hbar}(E_B - E_I - \hbar\omega)t_2} \frac{e^{\frac{i}{\hbar}(E_I - E_A + \hbar\omega')t_2} - 1}{\frac{i}{\hbar}(E_I - E_A + \hbar\omega')} \right\}.
\end{aligned}$$

Une erreur naïve à ne pas commettre: Croire applicable ici la relation de fermeture,  $\sum_I |I\rangle\langle I| = 1$ , en oubliant que subsistent, en facteur, des exponentielles qui dépendent de  $I$  !

Restent à effectuer trois types d'intégrales, triviales, sur  $t_2$ . Mais toute intégration d'exponentielle fait apparaître un dénominateur, et ces trois intégrales donnent des contributions respectivement proportionnelles à

$$\frac{1}{E_B + \hbar\omega' - E_A - \hbar\omega}, \quad \frac{1}{E_B - E_I + \hbar\omega'} \quad \text{et} \quad \frac{1}{E_B - E_I - \hbar\omega}.$$

Nous allons bientôt retrouver que le processus étudié n'a de chance notable de se produire que si  $E_B + \hbar\omega' = E_A + \hbar\omega$ . (Heureusement: c'est la conservation de l'énergie, conservation qui justement nous avait suggéré de baptiser "énergie du photon" la quantité  $\hbar\omega$ .) Comparée aux deux autres, la contribution du premier type est donc singulière (sauf en cas d'existence d'un état intermédiaire résonant); elle est dominante, et c'est la seule que nous gardons (approximation de l'onde tournante, page 44) pour obtenir enfin:

$$\begin{aligned}
c_{\text{para}}(t) & \approx -\frac{1}{i} \left( \frac{q}{m} \right)^2 \frac{1}{2\varepsilon_0 \mathcal{V}} \frac{1}{\sqrt{\omega\omega'}} \\
& \times \sum_I \left\{ \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle}{E_I - E_A - \hbar\omega} + \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | A \rangle}{E_I - E_A + \hbar\omega'} \right\} \\
& \times \int_0^t dt_2 e^{\frac{i}{\hbar}(E_B + \hbar\omega' - E_A - \hbar\omega)t_2}.
\end{aligned} \tag{VII.16}$$

### VII.3.3 Graphes préhistoriques

Il ne nous reste plus qu'à additionner les contributions diamagnétique (VII.15) et paramagnétique (VII.16) pour obtenir, après intégration, l'évaluation de l'amplitude cherchée:

$$c(t) \approx \frac{1}{i} \frac{q^2}{m} \frac{1}{2\varepsilon_0 \mathcal{V}} \frac{1}{\sqrt{\omega\omega'}} \frac{e^{\frac{i}{\hbar}(E_B + \hbar\omega' - E_A - \hbar\omega)t} - 1}{\frac{i}{\hbar}(E_B + \hbar\omega' - E_A - \hbar\omega)} \left\{ \delta_{AB} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \right.$$

$$+ \frac{1}{m} \sum_I \left( \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle}{E_A + \hbar\omega - E_I} + \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | A \rangle}{E_A - \hbar\omega' - E_I} \right) \Bigg\}.$$

C'est cette expression que nous utiliserons pour le calcul de la section efficace du processus.

Mais on peut auparavant s'amuser à en réarranger les termes pour l'écrire:

$$\begin{aligned} c(t) = & - \frac{e^{\frac{i}{\hbar}(E_B + \hbar\omega' - E_A - \hbar\omega)t} - 1}{\frac{E_B + \hbar\omega' - E_A - \hbar\omega}{\hbar}} \left\{ \langle B | \frac{\hat{\mathbf{e}}_{\mathbf{k}'T'}^*}{\sqrt{2\varepsilon_0\omega'\mathcal{V}}} \cdot \frac{q^2}{m} \frac{\hat{\mathbf{e}}_{\mathbf{k}T}}{\sqrt{2\varepsilon_0\omega\mathcal{V}}} | A \rangle \right. \\ & + \langle B | \frac{\hat{\mathbf{e}}_{\mathbf{k}'T'}^*}{\sqrt{2\varepsilon_0\omega'\mathcal{V}}} \cdot \frac{q}{m} \mathbf{P} \sum_I \frac{|I\rangle\langle I|}{E_A - E_I + \hbar\omega} \frac{q}{m} \mathbf{P} \cdot \frac{\hat{\mathbf{e}}_{\mathbf{k}T}}{\sqrt{2\varepsilon_0\omega\mathcal{V}}} | A \rangle \\ & \left. + \langle B | \frac{\hat{\mathbf{e}}_{\mathbf{k}T}}{\sqrt{2\varepsilon_0\omega\mathcal{V}}} \cdot \frac{q}{m} \mathbf{P} \sum_I \frac{|I\rangle\langle I|}{E_A - E_I - \hbar\omega'} \frac{q}{m} \mathbf{P} \cdot \frac{\hat{\mathbf{e}}_{\mathbf{k}'T'}^*}{\sqrt{2\varepsilon_0\omega'\mathcal{V}}} | A \rangle \right\}, \end{aligned}$$

c'est-à-dire sous forme...

- d'un facteur universel, dépendant du temps, qui assure la conservation de l'énergie au cours du processus en favorisant singulièrement les valeurs  $E_B + \hbar\omega' = E_A + \hbar\omega$ ,
- et d'un facteur spécifique de l'interaction et des états du système électron-rayonnement.

Nous pouvons alors mettre au point un moyen mnémorique permettant de retrouver facilement cette expression:

- Au premier élément de matrice, d'origine diamagnétique, on associe le dessin de la figure VII.1. Au bas de la figure l'état initial, en haut l'état final, pour rappeler la disposition ordinaire des graphes d'espace-temps (axe du temps vers le haut).<sup>5</sup> Chaque état d'électron est représenté par une ligne en trait plein, munie d'une flèche vers l'avenir, chaque état de photon par une ligne ondulée. L'interaction diamagnétique est figurée par la jonction des lignes en un *vertex* "à quatre pattes" (deux lignes d'électron, deux lignes de photon). A chaque composant (lignes et vertex) de ce graphe on associe (voir les phylactères) un facteur. Remarquez que le facteur associé à une ligne sortante est simplement le conjugué du facteur associé à la même ligne entrante. Enfin, le vertex implique une conservation de l'énergie totale de ses lignes:  $E_A + \hbar\omega = E_B + \hbar\omega'$ . L'expression de l'élément de matrice s'obtient en écrivant les facteurs associés aux composants, au fur et à mesure de leur rencontre au cours d'une lecture du graphe *en sens inverse* des flèches d'électron, et en effectuant les produits scalaires (on dit "en contractant") les facteurs

<sup>5</sup>Vous verrez aussi très souvent adoptée la convention de la bande dessinée où le temps progresse de la gauche vers la droite.

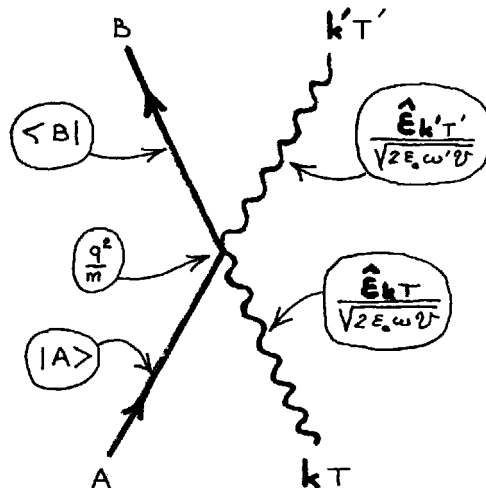


Figure VII.1: La contribution diamagnétique à l'amplitude de transition.

vectoriels dans l'ordre de leur apparition.<sup>6</sup>

- Le deuxième élément de matrice est représenté sur la figure VII.2. Aux composants qui apparaissaient déjà dans le dessin précédent (les quatre lignes représentant les états initial et final, dites *externes*), il va sans dire que l'on affecte encore les mêmes facteurs. Restent les nouveaux composants: deux vertex à trois pattes (deux lignes d'électron, une ligne de photon), et une ligne d'électron *interne* (entre deux vertex), avec sa flèche, auxquels on affecte les nouveaux facteurs indiqués dans les phylactères. L'opérateur associé à la ligne interne est le *propagateur* de l'électron (dans cette théorie "non relativiste"); son dénominateur est l'excédent (algébrique) d'énergie aux vertex extrémités de la ligne: énergie externe, ici  $E_A + \hbar\omega$  ou  $E_B + \hbar\omega'$ , moins énergie de l'état intermédiaire,  $E_I$ . Là encore, l'élément de matrice s'obtient par transcription, et contraction, des termes rencontrés au cours d'une lecture en sens inverse des flèches.
- Mais il n'y avait aucune raison d'attacher la ligne de photon entrant au premier vertex de la ligne d'électron interne et la ligne de photon sortant au deuxième vertex. L'inverse ne s'en distingue pas extérieurement et correspond au graphe de la figure VII.3, topologiquement distinct du graphe précédent. L'un et l'autre se disent respectivement *graphe direct* et *graphe d'échange* (à votre choix). Toutes les conventions sont maintenant fixées, le troisième élément de

<sup>6</sup>Par construction, notre théorie jouit de l'invariance par rotation, une probabilité (un module carré d'amplitude) qui dépendrait d'un choix d'axe particulier serait signe d'une erreur de calcul.



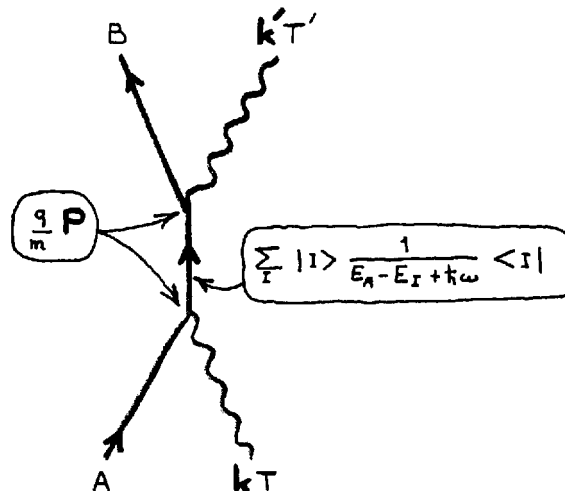


Figure VII.2: Première contribution paramagnétique.

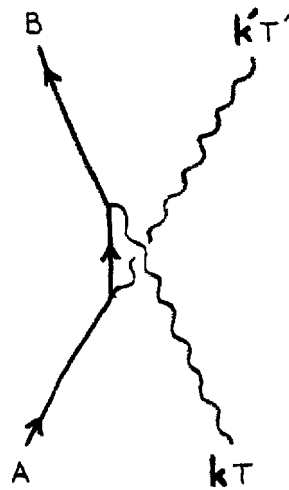


Figure VII.3: Contribution paramagnétique d'échange.

matrice se lit sur ce nouveau graphe, sans autre phylactère, mais en faisant attention à la valeur que prend le dénominateur du propagateur lorsque l'on échange les lignes de photon entrant et sortant.

Comme toujours en théorie quantique, l'amplitude de probabilité  $c(t)$  s'obtient en additionnant les amplitudes de toutes les modalités selon lesquelles peut se dérouler le processus défini par ses états initial et final. Ceux-ci donnés (les pattes externes), on trace tous les graphes topologiquement distincts que l'on peut construire à l'ordre  $q^2$  avec les ingrédients dont on dispose (deux types de vertex et ligne interne d'électron). On transcrit les éléments de matrice correspondant à ces diverses modalités, puis on les additionne (y compris les diverses modalités représentées par chaque état intermédiaire  $I$ ). Ne reste plus qu'à adjoindre le facteur temporel universel pour disposer de l'expression de l'amplitude  $c(t)$ .

Ces dessins, et les règles afférentes, sont des ancêtres non relativistes des graphes de Feynman. Ils ont l'avantage d'en constituer une introduction dépouillée de la quincaillerie géométrique des spineurs de Dirac, mais d'utilité pratique beaucoup plus restreinte: pas d'électrons relativistes, pas de positrons, pas de propagateur du photon, et d'inextricables complications, liées à la non invariance relativiste de la théorie, lorsqu'on désire étendre cette représentation au calcul perturbatif à des ordres plus élevés. Quoi qu'il en soit, l'éloquence de ces graphes leur confère une très réelle valeur heuristique, témoins les termes en lesquels ils sont commentés dans le langage courant: par exemple, sur la figure VII.1, un électron dans l'état  $A$  "entre", "absorbe" un photon  $\mathbf{k}T$  "puis émet" un photon  $\mathbf{k}'T'$ , et "ressort" enfin dans l'état  $B$ . En dépit de cette impression d'histoire vécue, il faut toutefois se garder d'y voir une description temporelle d'événements réels; nous ne faisons appel qu'à des états stationnaires, et surtout ces dessins ne font qu'illustrer les divers termes d'un calcul perturbatif. Aucun de ces graphes, pris individuellement, ne peut prétendre représenter une réalité.

### VII.3.4 La formule de Kramers et Heisenberg

En calculant la probabilité de transition  $|c(t)|^2$ , la représentation de la fonction de Dirac (II.57) se manifeste à nouveau:

$$\left| \frac{e^{\frac{i}{\hbar}Et} - 1}{\frac{i}{\hbar}E} \right|^2 = \frac{4\hbar^2}{E^2} \sin^2 \frac{Et}{2\hbar} \underset{t \rightarrow \infty}{\sim} 2\pi\hbar t \delta(E).$$

Additionnant, comme d'habitude, les probabilités de toutes les transitions qui peuvent se produire sous couvert de la fonction de Dirac — ou de ce qui en tient lieu —, il nous faut prendre en compte le nombre de modes du rayonnement, de polarisation

donnée, dans un domaine donné  $d^3\mathbf{k}$ :

$$\frac{k^2 dk d^2\hat{\mathbf{k}}}{(2\pi/L)^3} = \frac{\mathcal{V}}{(2\pi)^3} \frac{E^2 dE}{(\hbar c)^3} d^2\hat{\mathbf{k}},$$

en terme d'énergie du photon  $E = \hbar\omega = \hbar ck$ . La probabilité d'avoir, après le temps  $t$ , l'état  $B$  pour l'électron et l'état à un photon dont la direction soit dans le domaine  $d^2\hat{\mathbf{k}}'$  autour de  $\hat{\mathbf{k}}'$  et la polarisation  $\hat{\mathbf{e}}_{\mathbf{k}'T'}$  (orthogonale à  $\hat{\mathbf{k}}'$ ), est alors, en reprenant la première écriture de  $c(t)$ :

$$\begin{aligned} \text{Pr}(t) &= d^2\hat{\mathbf{k}}' \int dE' \frac{\mathcal{V}}{(2\pi)^3} \frac{E'^2}{(\hbar c)^3} |c(t)|^2, \\ &\underset{t \rightarrow \infty}{\sim} d^2\hat{\mathbf{k}}' \int dE' \frac{\mathcal{V}}{(2\pi)^3} \frac{E'^2}{(\hbar c)^3} \left( \frac{q^2}{m} \frac{1}{2\varepsilon_0 \mathcal{V}} \right)^2 \frac{1}{\omega\omega'} \\ &\quad \times 2\pi\hbar t \delta(E_B + E' - E_A - \hbar\omega) \left| \delta_{AB} \dots + \frac{1}{m} \sum_I \dots \right|^2. \end{aligned}$$

Comme annoncé, seule la valeur de la pulsation du photon sortant qui conserve l'énergie intervient, et dorénavant:

$$\hbar\omega' = E' = E_A + \hbar\omega - E_B. \quad (\text{VII.17})$$

Le taux de transition,  $\Gamma \stackrel{\text{df}}{=} d/dt \text{Pr}(t)$ , du processus  $A + \gamma_{\mathbf{k}T} \rightarrow B + \gamma_{(\hat{\mathbf{k}}', d^2\hat{\mathbf{k}}')T'}$ , vaut donc

$$\Gamma = d^2\hat{\mathbf{k}}' \frac{1}{(2\pi)^2 c^3 \mathcal{V}} \left( \frac{q^2}{2\varepsilon_0 m} \right)^2 \frac{\omega'}{\omega} \left| \delta_{AB} \dots + \frac{1}{m} \sum_I \dots \right|^2. \quad (\text{VII.18})$$

De la manière la plus économique, définissons la section efficace du processus par

$$d\sigma \stackrel{\text{df}}{=} \frac{-\delta E / \Delta t}{E / \Delta t \Delta \mathcal{S}}, \quad (\text{VII.19})$$

où:

- $\delta E$  est l'énergie soustraite au faisceau de photons  $\mathbf{k}T$  incidents, pendant  $\Delta t$ , pour cause de photons finals  $E', (\hat{\mathbf{k}}', d^2\hat{\mathbf{k}}'), T'$ ,
- $E$  est l'énergie charriée par le faisceau incident pendant  $\Delta t$ ,
- $\Delta \mathcal{S}$  est la section du faisceau.

On a évidemment  $\delta E = -\Gamma \Delta t \hbar\omega$ . L'état initial — un photon d'énergie  $\hbar\omega$  dans le volume  $\mathcal{V}$  de la boîte — correspond à une densité d'énergie  $\hbar\omega/\mathcal{V}$  et, pour une tranche de faisceau de section  $\Delta \mathcal{S}$  et longueur  $c\Delta t$ , on a  $E = (c\Delta t \Delta \mathcal{S}/\mathcal{V}) \hbar\omega$ . D'où la section efficace en fonction du taux de transition:

$$d\sigma = \frac{\mathcal{V}}{c} \Gamma. \quad (\text{VII.20})$$

Il y a bien sûr quelque incohérence dans le mélange désinvolte des notions de photon et de faisceau, auxquelles j'ai fait appel pour ce calcul de la section efficace. C'est l'un des multiples cas où l'assimilation implicite du photon à une bille permet des raccourcis saisissants: il suffit d'imaginer un faisceau incident constitué (dans l'espace!) d'une suite de ces billes indépendantes, et le tour est joué. Assimiler un seul photon initial à un faisceau se révèle beaucoup moins audacieux lorsqu'on remarque que l'on trouve le même taux de transition pour une réaction avec, dans l'état initial  $n$  photons  $\mathbf{k}T$ , et dans l'état final  $n - 1$  photons  $\mathbf{k}T$  et un photon  $\mathbf{k}'T'$ . En pratique, les états quasi-classiques du rayonnement sont des états cohérents à nombre de photons élevé en moyenne et relativement bien défini. On conçoit alors que parler d'état initial à un ou à  $n$  photons, ou d'une onde électromagnétique incidente, ne change en rien le résultat et que l'on puisse donc obtenir impunément celui-ci à partir d'une définition (VII.19) de la section efficace, héritée de la diffusion du rayonnement classique.

Les relations (VII.20) et (VII.18) nous permettent de calculer la section efficace différentielle  $d\sigma$  du processus. Celle-ci a la dimension d'une aire, raison pour laquelle elle s'exprime au mieux en fonction de la longueur

$$r \stackrel{\text{df}}{=} \frac{q^2/4\pi\epsilon_0}{mc^2} = \frac{q^2/4\pi\epsilon_0}{\hbar c} \frac{\hbar}{mc}. \quad (\text{VII.21})$$

Dans le cas d'un électron ( $q_e = -|e|$ ), cette longueur est donc le produit de la constante de structure fine et de la longueur d'onde Compton dudit électron,  $r_e = \alpha \lambda_e \approx 2,82 \text{ Fm}$ , que l'on continue d'appeler *rayon classique de l'électron* parce que, à l'évidence, la constante  $\hbar$  n'intervient pas dans sa définition. Historiquement, ce rayon s'est introduit lorsqu'on tentait d'édifier un modèle de structure de l'électron en forme de boule, de rayon  $r_e$ , plus ou moins uniformément chargée électriquement. Il était alors tentant d'attribuer à l'énergie du champ électrique inséparable de l'électron, la totalité de l'inertie, ou masse, ou énergie de repos, observée pour cet électron. A un facteur numérique sans dimension près — dépendant du modèle de distribution radiale de charge dans la boule et, naturellement, voisin de un —, on avait donc  $m_e c^2 \approx q_e^2/(4\pi\epsilon_0 r_e)$ , conduisant à la définition (VII.21). Toujours est-il que nous obtenons ainsi la distribution angulaire (formule de Kramers et Heisenberg),

*Exercice, énergie électromagnétique par principe zéro, calcul divers modèles.*

$$\frac{d\sigma}{d^2\hat{\mathbf{k}}'} = r_e^2 \frac{\omega'}{\omega} \left| \delta_{AB} \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* + \frac{1}{m} \sum_I \left\{ \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle}{E_A + \hbar\omega - E_I} + \frac{\langle B | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | A \rangle}{E_A - \hbar\omega' - E_I} \right\} \right|^2, \quad (\text{VII.22})$$

où  $d\sigma$  est la section efficace différentielle du processus de diffusion photon-atome:

$$A, \omega, \hat{\mathbf{k}}, \hat{\mathbf{e}}_{\mathbf{k}T} \rightarrow B, (\hat{\mathbf{k}}', d^2\hat{\mathbf{k}}'), \hat{\mathbf{e}}_{\mathbf{k}'T'}.$$

Rappelons que la conservation de l'énergie, retrouvée, est sous entendue:  $\hbar\omega' = E_A + \hbar\omega - E_B$ . La somme sur les états stationnaires intermédiaires  $I$  de l'atome doit comprendre en particulier les états du *continuum*, mais il est des situations où, grâce aux dénominateurs, la contribution de ces états peut être négligée. Dans le cas d'un atome qui comporte plusieurs électrons, il suffit de remplacer l'opérateur  $\mathbf{P}$  par la somme de leurs impulsions,  $\sum_i \mathbf{P}_i$ .

### VII.3.5 Diffusion élastique

Etudions d'abord, à l'aide de la formule de Kramers et Heisenberg, un processus où l'atome se retrouve dans son état initial:

$$A, \omega, \hat{\mathbf{k}}, \hat{\mathbf{e}}_{\mathbf{k}T} \rightarrow A, (\hat{\mathbf{k}}', d^2\hat{\mathbf{k}}'), \hat{\mathbf{e}}_{\mathbf{k}'T'}.$$

On a alors  $\omega' = \omega$ , et cette réaction est dite de diffusion élastique, un peu abusivement il est vrai, car la quantité de mouvement totale n'est pas conservée: la quantité de mouvement de l'atome reste nulle, et celle du photon change de direction. Mais c'est là un comportement parfaitement en accord avec notre hypothèse initiale d'un atome infiniment lourd, sur lequel un photon peut diffuser tel le cochonnet sur une boule de pétanque.<sup>7</sup> Dans le cas élastique, la formule de Kramers et Heisenberg devient:

$$\frac{d\sigma}{d^2\hat{\mathbf{k}}'} = r_e^2 \left| \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}^* + \frac{1}{m} \sum_I \left\{ \frac{\langle A | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | A \rangle}{E_A + \hbar\omega - E_I} + \frac{\langle A | \hat{\mathbf{e}}_{\mathbf{k}T} \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}}_{\mathbf{k}'T'}^* \cdot \mathbf{P} | A \rangle}{E_A - \hbar\omega - E_I} \right\} \right|^2. \quad (\text{VII.23})$$

Toute l'astuce consiste maintenant à mettre le premier terme — la contribution diamagnétique — sous une forme analogue au second! En effet, et en abrégant quelque peu les notations, on a:

$$\begin{aligned} \hat{\mathbf{e}} \cdot \hat{\mathbf{e}}'^* &= \varepsilon_i \varepsilon_j'^* \delta_{ij} = \varepsilon_i \varepsilon_j'^* \frac{1}{i\hbar} [R_i, P_j], \\ &= \frac{1}{i\hbar} \left\{ (\hat{\mathbf{e}} \cdot \mathbf{R})(\hat{\mathbf{e}}'^* \cdot \mathbf{P}) - (\hat{\mathbf{e}}'^* \cdot \mathbf{P})(\hat{\mathbf{e}} \cdot \mathbf{R}) \right\}, \end{aligned}$$

et, comme le premier membre est un nombre,

$$\begin{aligned} \hat{\mathbf{e}} \cdot \hat{\mathbf{e}}'^* &= \langle A | \hat{\mathbf{e}} \cdot \hat{\mathbf{e}}'^* | A \rangle, \\ &= \frac{1}{i\hbar} \sum_I \left\{ \langle A | \hat{\mathbf{e}} \cdot \mathbf{R} | I \rangle \langle I | \hat{\mathbf{e}}'^* \cdot \mathbf{P} | A \rangle - \langle A | \hat{\mathbf{e}}'^* \cdot \mathbf{P} | I \rangle \langle I | \hat{\mathbf{e}} \cdot \mathbf{R} | A \rangle \right\}. \end{aligned}$$

<sup>7</sup>Que cette hypothèse se révèle insuffisante il faudrait alors inclure les degrés de liberté du centre de masse de l'atome dans la description quantique de celui-ci (en particulier dans son hamiltonien non perturbé); ainsi la quantité de mouvement des états initial et final de l'atome interviendrait dans leur caractérisation — en plus de leur énergie interne — et on en trouverait, au même titre, une loi de conservation.

En exprimant — voir (VI.17) — les éléments de matrice (entre états propres non perturbés) de la position de l'électron en fonction des éléments de matrice de son impulsion,

$$\langle A|\mathbf{R}|I\rangle = -\frac{i\hbar}{m} \frac{1}{E_A - E_I} \langle A|\mathbf{P}|I\rangle,$$

on a:

$$\hat{\mathbf{e}} \cdot \hat{\mathbf{e}}'^* = \frac{1}{m} \sum_I \frac{\langle A|\hat{\mathbf{e}}'^* \cdot \mathbf{P}|I\rangle \langle I|\hat{\mathbf{e}} \cdot \mathbf{P}|A\rangle + \langle A|\hat{\mathbf{e}} \cdot \mathbf{P}|I\rangle \langle I|\hat{\mathbf{e}}'^* \cdot \mathbf{P}|A\rangle}{E_I - E_A},$$

sans oublier que le dénominateur dépendant de  $I$  empêche toute fermeture sur les états  $I$ . En posant  $\hbar\omega_{IA} \stackrel{\text{df}}{=} E_I - E_A$ , on peut écrire finalement:

$$\begin{aligned} \frac{d\sigma}{d^2\hat{\mathbf{k}}'} = r_e^2 \left| \frac{1}{m\hbar} \sum_I \left\{ \left( \frac{1}{\omega_{IA}} - \frac{1}{\omega_{IA} - \omega} \right) \langle A|\hat{\mathbf{e}}'^* \cdot \mathbf{P}|I\rangle \langle I|\hat{\mathbf{e}} \cdot \mathbf{P}|A\rangle \right. \right. \\ \left. \left. + \left( \frac{1}{\omega_{IA}} - \frac{1}{\omega_{IA} + \omega} \right) \langle A|\hat{\mathbf{e}} \cdot \mathbf{P}|I\rangle \langle I|\hat{\mathbf{e}}'^* \cdot \mathbf{P}|A\rangle \right\} \right|^2. \end{aligned}$$

Le cas d'un photon incident à basse énergie, autrement dit d'énergie faible par rapport aux excitations de l'atome ( $\omega \ll \omega_{IA}$  pour tout  $I$ ), est traditionnellement appelé *diffusion de Rayleigh*. On a alors,

$$\frac{1}{\omega_{IA}} - \frac{1}{\omega_{IA} \pm \omega} = \pm \frac{\omega}{\omega_{IA}^2} \frac{1}{1 \pm \omega/\omega_{IA}} \approx \pm \frac{\omega}{\omega_{IA}^2} \left( 1 \mp \frac{\omega}{\omega_{IA}} \right),$$

et la distribution angulaire, réexprimée en fonction des seuls éléments de matrice de l'opérateur position de l'électron, devient:

$$\begin{aligned} \frac{d\sigma}{d^2\hat{\mathbf{k}}'} \underset{\omega \ll \omega_{IA}}{\sim} \alpha^2 \frac{\omega^4}{c^2} \\ \times \left| \sum_I \frac{\langle A|\hat{\mathbf{e}}'^* \cdot \mathbf{R}|I\rangle \langle I|\hat{\mathbf{e}} \cdot \mathbf{R}|A\rangle + \langle A|\hat{\mathbf{e}} \cdot \mathbf{R}|I\rangle \langle I|\hat{\mathbf{e}}'^* \cdot \mathbf{R}|A\rangle}{\omega_{IA}} \right|^2. \quad (\text{VII.24}) \end{aligned}$$

On voit que la section efficace différentielle de diffusion élastique dans une direction donnée est, à basse énergie, proportionnelle à  $\omega^4$ . C'est la seule dépendance en  $\omega$  que l'on trouve dans cette formule, les autres facteurs ne faisant intervenir que des constantes fondamentales et des éléments de matrice et énergies propres à l'atome. Remarquons que, dans ce cas, il ne peut se produire de toute façon aucun processus de diffusion inélastique, car la situation  $\omega \ll \omega_{IA}$  pour le photon incident implique, en particulier, que  $E_A + \hbar\omega$  est beaucoup plus petit que l'énergie du premier niveau excité. Il est alors impossible d'émettre un photon correspondant à cette excitation et satisfaisant la condition de conservation (VII.17).

Le fait qu'un gaz puisse paraître incolore sous un éclairage normal (spectre du Soleil filtré par la ionosphère, ou spectre d'une lampe à incandescence), autrement dit l'absence de raie d'émission du gaz, signifie qu'aucun niveau excité des atomes, ou molécules, du gaz n'est atteint. Les pulsations de transition  $\omega_{IA}$  sont donc toutes supérieures à l'ultraviolet qui constitue pratiquement la limite du spectre incident. La formule de la diffusion Rayleigh convient alors pour décrire un rayonnement incident visible qui diffuse sur ce gaz. C'est le cas de la diffusion de la lumière solaire dans l'atmosphère. Le facteur  $\omega^4$  explique le bleu du ciel: lorsqu'on ne regarde pas le Soleil directement (en particulier si on lui tourne le dos), seule parvient dans nos yeux la lumière diffusée, dont l'intensité est renforcée à haute fréquence, c'est-à-dire dans l'extrémité bleue du spectre visible. Inversement, en vision directe du Soleil, cette intensité diffusée, de dominante bleue, est soustraite du faisceau incident; ne s'offre à nos regards que ce qui reste, c'est-à-dire l'autre extrémité, rouge, du spectre, surtout si l'épaisseur d'air traversée multiplie les chances de diffusion. Ainsi, le Soleil matinal ou crépusculaire nous apparaît rouge. Ajoutons que lors de la diffusion d'un photon incident par un ensemble de molécules — comme cela se produit dans l'atmosphère —, il faut *a priori* ajouter de façon cohérente les amplitudes de diffusion par chacune des molécules, le module carré donnant la probabilité du processus. Ce sont finalement les fluctuations de position des molécules qui estompent, en moyenne temporelle, les termes d'interférence, et ramènent la résultante à une somme incohérente des sections efficaces individuelles. Mais surtout, le gommage de ces termes débarasse la section efficace résultante de leur dépendance angulaire caractéristique et assure au ciel la quasi uniformité de sa coloration bleue. Il n'en serait pas de même avec des molécules disposées bien régulièrement. Un cristal éclairé ne diffuse de lumière que vers l'avant et dans quelques directions bien déterminées (les pics de Bragg); par comparaison, observé depuis une direction quelconque, le cristal reste quasi invisible dans le domaine des rayons  $X$ . Enfin, il ne faut pas oublier les objets macroscopiques abondants dans l'atmosphère (gouttelettes d'eau, cristaux de glace, poussières, aérosols) et qui, eux-aussi, diffusent la lumière par un processus d'ailleurs beaucoup plus complexe à analyser: la *diffusion de Mie* (voir, par exemple [36]).

Après cette analyse de la section efficace de diffusion élastique en fonction de la fréquence, voyons de quelle manière elle dépend des polarisations. Considérons un photon initial de vecteur d'onde  $\mathbf{k}$  et vecteur de polarisation rectiligne  $\hat{\mathbf{e}}$ . Ces caractéristiques données permettent de définir un repère dont l'axe  $\hat{\mathbf{z}}$  est choisi suivant  $\mathbf{k}$ , et l'axe  $\hat{\mathbf{x}}$  suivant  $\hat{\mathbf{e}}$  (figure VII.4). Nous nous intéressons à la probabilité de diffusion d'un photon dans la direction  $\hat{\mathbf{k}}'$  avec une polarisation rectiligne donnée  $\hat{\mathbf{e}}'$  (orthogonale à  $\hat{\mathbf{k}}'$  bien sûr). C'est ce que nous faisons par exemple en regardant le ciel à travers des lunettes de soleil du type dit "polaroid". L'analyse est plus simple pour des vecteurs polarisation particuliers  $\hat{\mathbf{e}}'_1$  et  $\hat{\mathbf{e}}'_2$ , orthogonaux, en terme desquels on peut écrire une polarisation quelconque  $\hat{\mathbf{e}}' = \cos \psi \hat{\mathbf{e}}'_1 + \sin \psi \hat{\mathbf{e}}'_2$ . Choisissons, en guise de vecteurs polarisation privilégiés,  $\hat{\mathbf{e}}'_1$  dans le plan de diffusion  $(\mathbf{k}, \mathbf{k}')$ , orienté

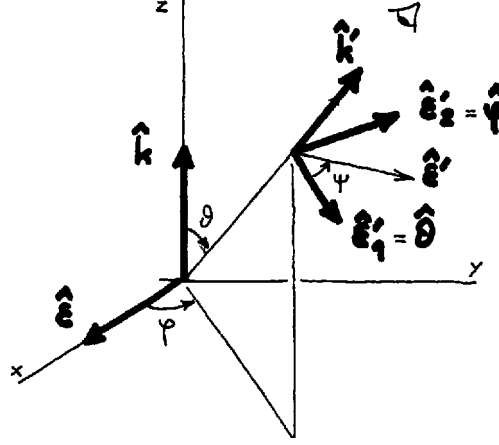


Figure VII.4: Disposition relative des caractéristiques du photon initial (vecteur d'onde  $\hat{k}$  et vecteur polarisation  $\hat{e}$ ) et des grandeurs finales observées (direction de diffusion  $\hat{k}'$  et vecteurs de polarisation de base  $\hat{e}'_1$  et  $\hat{e}'_2$ ).

dans le sens de la colatitude  $\theta$  croissante, et  $\hat{e}'_2$  de façon à constituer un trièdre direct avec  $\hat{k}'$  et  $\hat{e}'_1$  (voir la figure VII.4). Autrement dit, en fonction des vecteurs de la base des coordonnées sphériques, normés:  $\hat{e}'_1 \stackrel{\text{df}}{=} \hat{\theta}$ , et  $\hat{e}'_2 \stackrel{\text{df}}{=} \hat{\varphi}$ . Nos divers vecteurs polarisation ont pour composantes cartésiennes:

$$\hat{e} = \begin{Bmatrix} 1 \\ 0 \\ 0 \end{Bmatrix}, \quad \hat{e}'_1 = \begin{Bmatrix} \cos \theta \cos \varphi \\ \cos \theta \sin \varphi \\ -\sin \theta \end{Bmatrix} \quad \text{et} \quad \hat{e}'_2 = \begin{Bmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{Bmatrix}.$$

Le report de ces composantes dans la formule (VII.24) donne les sections efficaces cherchées en fonction des éléments de matrice des opérateurs position  $X$ ,  $Y$  et  $Z$  entre les états atomiques  $A$  et  $I$ . En particulier, dans le cas de la “diffusion à 90 degrés”,  $\theta = \pi/2$ , le terme général de la sommation sur les états intermédiaires se réduit à

$$\langle A j_A m_A | (-Z) | I j_I m_I \rangle \langle I j_I m_I | X | A j_A m_A \rangle,$$

lorsque la polarisation finale est  $\hat{e}'_1 = -\hat{z}$ , et

$$\langle A j_A m_A | (-\sin \varphi X + \cos \varphi Y) | I j_I m_I \rangle \langle I j_I m_I | X | A j_A m_A \rangle,$$

lorsque la polarisation finale est  $\hat{e}'_2$ . Dans ces expressions, j'ai rappelé de façon explicite que les états stationnaires de l'atome isolé sont des états propres du carré du moment angulaire total et de l'une de ses composantes,  $J_z$  si l'on sacrifie au conformisme. L'instant est à présent particulièrement bien venu de se souvenir des



composantes sphériques de l'opérateur position, définies en (VI.20) et en termes desquelles nous pouvons écrire:

$$\begin{cases} X = -\frac{1}{\sqrt{2}}(R_1^{(1)} - R_{-1}^{(1)}) \\ Y = \frac{i}{\sqrt{2}}(R_1^{(1)} + R_{-1}^{(1)}) \\ Z = R_0^{(1)}. \end{cases}$$

Les règles de sélection fondées sur la parité et le moment angulaire, déjà utilisées dans la section VI.4.2, nous assurent alors que les éléments de matrice des opérateurs  $X$  et  $Y$  ne sont éventuellement non nuls que si  $m_I = m_A \pm 1$ , tandis que pour l'opérateur  $Z$  il faut  $m_I = m_A$ . Le produit d'éléments de matrice qui contribue à la section efficace de diffusion dans un état final de polarisation  $\hat{\epsilon}'_1 = -\hat{z}$  est donc toujours nul, mais il n'y a pas de raison qu'il en soit de même pour la polarisation  $\hat{\epsilon}'_2$ . Ce résultat ne dépend pas de l'angle  $\varphi$  et l'on peut donc en conclure que le rayonnement diffusé à  $90^\circ$  est polarisé rectiligne dans le plan perpendiculaire au faisceau incident, quelle que soit la polarisation de celui-ci.

Ce phénomène de polarisation du rayonnement diffusé à  $90^\circ$  s'observe facilement en contemplant un ciel ensoleillé à travers une paire de lunettes polaroïd. Il s'explique aussi en théorie classique du rayonnement et était même, semble-t-il, connu des navigateurs vikings. C'est en tout cas la seule raison que l'on ait trouvée à la présence de morceaux de cordiérite — un cristal dichroïque naturel, qui peut servir de filtre polariseur — dans les tombeaux des chefs vikings. Aux latitudes élevées fréquentées par ces marins peu frileux, le Soleil est — de par sa faible hauteur — fréquemment invisible, qu'il soit caché par des nuages bas, ou des brumes, ou même par l'horizon pendant la nuit polaire. Il suffisait qu'une portion de ciel soit dégagée pour que le navigateur, en cherchant une direction où son filtre assombrissait le ciel au maximum, puisse estimer une direction du Soleil (figure VII.5), bien que le rayonnement émis par celui-ci ne soit pas polarisé (on y trouve tous les vecteurs polarisation possibles). Le sens marin du navigateur — l'observation des vagues, du vent, du ciel, — devait lui permettre ensuite de déterminer de quel côté se trouvait le Soleil sur la direction trouvée. Le procédé nécessitait toutefois une éclaircie car le mécanisme de diffusion du rayonnement par les gouttes d'eau (diffusion multiple, voir [12]) étant tout autre que celui de la diffusion par les molécules ou atomes d'un gaz, la lumière diffusée par les nuages est généralement blanche et non polarisée. Bien d'autres effets, comme la réfraction et la dispersion chromatique, interviennent encore sur la propagation de la lumière dans l'atmosphère pure ou encombrée, et occasionnent des manifestations spectaculaires telles que les arcs en ciel ou le rayon vert<sup>8</sup>, mais la théorie quantique n'éprouve, en tout cas, aucune difficulté à expliquer aussi bien l'évident azur que sa plus discrète polarisation.

---

<sup>8</sup>Lire, contempler, ou voir, par exemple [56],[58],[73],[75],[62].

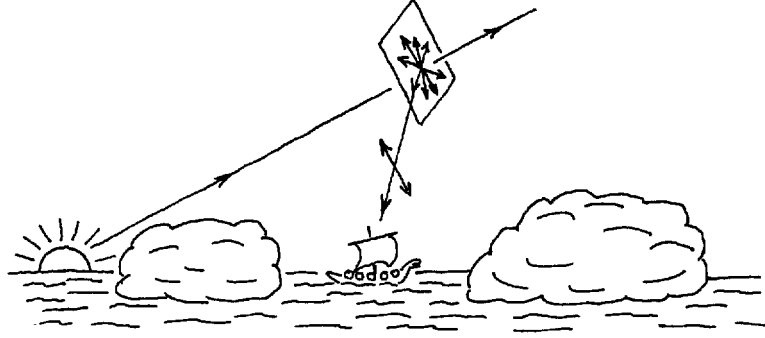


Figure VII.5: Navigateur viking dans le brouillard, trouvant quand même la direction du Soleil à l'aide d'un filtre polariseur.

Reprenons maintenant la formule de la section efficace élastique (VII.23) et passons au cas d'un photon incident d'énergie élevée, appelé *diffusion Thomson*. Par énergie élevée, on entend que la pulsation  $\omega$  est assez grande pour que les contributions paramagnétiques, dans l'expression (VII.23) soient rendues négligeables par leurs dénominateurs. On obtient tout simplement:

$$\frac{d\sigma}{d^2\mathbf{k}'} \sim r_e^2 |\hat{\mathbf{e}}_{\mathbf{k}T} \cdot \hat{\mathbf{e}}_{\mathbf{k}'T'}|^2. \quad (\text{VII.25})$$

Mais, ce faisant, on omet une sommation sur tous les états stationnaires  $I$  de l'électron dans l'atome, y compris — et même si la notation ne le montre pas — ceux du continu lorsque l'environnement moyen de l'électron est représenté par un potentiel un peu plus libéral qu'un puits infini ou un oscillateur harmonique. La condition  $\hbar\omega \gg |E_I - E_A|$  est bien entendu irréalisable pour tous les états  $I$ , mais les éléments de matrice au numérateur viennent prendre le relais: des états  $A$  et  $I$  d'énergie très différente ont de faibles recouvrements, leurs fonctions-d'onde se mettent en valeur dans des régions de l'espace différentes, et quelque traitement que leur inflige un opérateur (ici  $\mathbf{P} = -i\hbar\nabla$ ) il ne peut donner de valeur à une fonction là où elle est restée quasi nulle. On peut donc négliger les contributions paramagnétiques parce que leur dénominateur est grand pour les bas niveaux  $I$ , et parce que leur numérateur est petit pour les niveaux  $I$  élevés. En outre, le défaut de recouvrement permet, en venant tronquer effectivement la somme sur les états  $I$ , d'arrêter celle-ci bien avant d'atteindre des énergies  $E_I$  de l'ordre de la masse au repos de l'électron, sauvant ainsi la cohérence de ce calcul non relativiste.

Dans l'expression (VII.25), obtenue en cherchant à décrire la diffusion d'un photon de haute énergie par l'électron d'un atome, cet atome n'intervient finalement pas. Cette expression décrit donc aussi bien la diffusion par un électron libre,

mais n'oublions pas qu'elle a été obtenue par approximation dipolaire électrique. L'énergie du photon a donc beau être grande par rapport aux premières excitations de l'atome (dont les longueurs d'onde typiques correspondent au visible), il faut quand même que sa longueur d'onde reste supérieure à la dimension de l'atome (de l'ordre de l'Ångström). Dans le domaine de validité  $1 \text{ Å} < \lambda < 5000 \text{ Å}$ , la longueur d'onde reste bien supérieure à la longueur d'onde Compton de l'électron,

$$\lambda_e \stackrel{\text{df}}{=} \frac{\hbar}{m_e c} = \frac{\hbar c}{m_e c^2} \approx 4 \times 10^2 \text{ fm},$$

longueur d'onde à laquelle commencent à se manifester des changements de fréquence, d'apparence inélastique, mais dus en fait au recul de l'électron sous l'effet d'une impulsion incidente  $\hbar k = \hbar/\lambda$ , ou d'une énergie

$$\hbar\omega = \hbar kc = \frac{\hbar c}{\lambda} = mc^2,$$

*Diffusion inélastique, effet Raman, raies Stokes et anti Stokes. Corrélations.*

qui nécessite donc de toute façon un traitement relativiste einsteinien du problème.

## Exercices

**10.** Les grandeurs dimensionnées caractérisant entièrement chaque fréquence d'un rayonnement électromagnétique en équilibre thermodynamique sont cette fréquence,  $\nu$  ou (encore plus fondamental)  $\omega$ , et la température d'équilibre  $T$ .

*i)* Le monde classique repose sur les deux constantes fondamentales que sont la vitesse  $c$  et la constante de Boltzmann  $k_B$ . L'énergie thermodynamique étant une grandeur extensive, en déduire, à l'aide du "principe zéro" de la physique, la forme de la densité spectrale d'énergie (une énergie sur un volume et une fréquence). Comparez cette prédiction avec la loi de Rayleigh-Jeans (VII.10).

*ii)* L'univers quantique est édifié sur une constante fondamentale supplémentaire,  $\hbar$ . Par quel facteur, fonction (sans dimension) de  $\hbar\omega/k_B T$ , la loi (quantique) de Planck (VII.9) diffère-t-elle de la loi (classique) de Rayleigh-Jeans?

# Bibliographie

- [1] M. ABRAMOWITZ and I.A. STEGUN, *Handbook of mathematical functions*, Dover (New York 1972).
- [2] A.I. AKHIEZER and V.B. BERESTETSKII, *Quantum Electrodynamics*, Interscience (New York 1965).
- [3] H.L. ARMSTRONG, *No place for a photon?*, Am. J. Phys. **51** (1983) 103.
- [4] Y. AYANT et M. BORG, *Fonctions spéciales*, Dunod “Université” (Paris 1971).
- [5] F. BALIBAR, *Le mètre au fil du temps*, La Recherche **15** (février 1984) 263.
- [6] R. BALLIAN and C. BLOCH, *Distribution of Eigenfrequencies for the Wave Equation in a Finite Domain II. Electromagnetic Field. Riemannian Spaces*, Ann. Phys. (N.Y.) **64** (1971) 271.
- [7] S.M. BARNETT and C.R. GILSON, *Manipulating the vacuum: squeezed states of light*, Eur. J. Phys **9** (1988) 257.
- [8] J. BASS, *Cours de mathématiques*, Masson (Paris 1977).
- [9] G. BAYM, *Lectures on quantum mechanics*, W.A. Benjamin (New York 1969).
- [10] H.A. BETHE, *The Electromagnetic Shift of Energy Levels*, Phys. Rev. **72** (1947) 339.
- [11] H.A. BETHE, L.M. BROWN and J.R. STEHN, *Numerical Values of the Lamb Shift*, Phys. Rev. **77** (1950) 370.
- [12] C.J. BOHREN, *Multiple scattering of light and some its observable consequences*, Am. J. Phys. **55** (1987) 524.
- [13] T.H. BOYER, *Quantum Zero-Point Energy and Long-Range Forces*, Ann. Phys. (N.Y.) **56** (1970) 474.

- [14] M. BUNGE, *Philosophie de la physique*, Seuil “Science ouverte” (Paris 1975).
- [15] B. CAGNAC et J.-C. PEBAY-PEROULA, *Physique atomique*, t.1, *Expériences et principes fondamentaux*, Dunod “Université” (Paris 1975).
- [16] B. CAGNAC et J.-C. PEBAY-PEROULA, *Physique atomique*, t.2, *Applications de la mécanique quantique*, Dunod “Université” (Paris 1971).
- [17] P. CARRUTHERS and M.M. NIETO, *Coherent states and the number phase uncertainty relation*, Phys. Rev. Lett. **14** (1965) 387.
- [18] E. CHPOLSKI, *Physique atomique*, t.II, *Fondements de la mécanique quantique et structure de l’enveloppe électronique de l’atome*, Mir (Moscou 1978).
- [19] C. COHEN-TANNOUDJI, B. DIU et F. LALÖE, *Mécanique quantique*, 2 t., Hermann (Paris 1977).
- [20] C. COHEN-TANNOUDJI, J. DUPONT-ROC et G. GRYNBERG, *Photons et atomes, Introduction à l’électrodynamique quantique*, InterEditions/CNRS “Savoirs actuels” (Paris 1987).
- [21] C. COHEN-TANNOUDJI, J. DUPONT-ROC et G. GRYNBERG, *Processus d’interaction entre photons et atomes*, InterEditions/CNRS “Savoirs actuels” (Paris 1988).
- [22] A.V. DURRANT, *Some basic properties of stimulated and spontaneous emission: a semiclassical approach*, Am. J. Phys. **44** (1976) 630.
- [23] F.J. DYSON, *Les dérangeurs de l’univers*, Payot “Espace des sciences” (Paris 1986).
- [24] A. EINSTEIN, *Théorie quantique du rayonnement*, in *Albert Einstein, Œuvres choisies, t. I, Quanta*, Seuil/CNRS “Sources du savoir” (Paris 1989), p. 134.
- [25] E. ELIZALDE and A. ROMEO, *Essentials of the Casimir effect and its computation*, Am. J. Phys. **59** (1991) 711.
- [26] H. EVERETT, III, *The theory of the universal wave function*, in *The Many-Worlds Interpretation of Quantum Mechanics*, ed. B.S. De Witt and N. Graham, Princeton University Press (Princeton 1973), Introduction p. 3.
- [27] E. FERMI, *Quantum theory of radiation*, Rev. Mod. Phys. **4** (1932) 87.
- [28] R.P. FEYNMAN, *La nature de la physique*, Seuil “Point Sciences” (Paris 1980).
- [29] R.P. FEYNMAN, R.B. LEIGHTON et M. SANDS, *Le Cours de physique de Feynman*, InterEditions (Paris 1979).

- [30] V. GINZBURG, *Physique théorique et astrophysique*, Ed. Mir (Moscou 1978).
- [31] R.J. GLAUBER, *Optical coherence and photon statistics*, in *Optique et électronique quantique*, eds C. de Witt *et al.*, Gordon and Breach (New York 1965).
- [32] W. GOUGH, *Poynting in the wrong direction?*, Eur. J. Phys. **3** (1982) 83.
- [33] M. HAMERMESH, *Group theory and its application to physical problems*, Addison-Wesley (Reading 1962).
- [34] T. HÄNSCH, A. SCHAWLOW et G. SERIES, *Le spectre de l'hydrogène atomique*, Pour la Science **19** (1979) 46.
- [35] W. HEITLER, *The Quantum theory of radiation*, Dover (New York 1954).
- [36] H.C. VAN DE HULST, *Light Scattering by Small Particles*, Dover (New York 1981).
- [37] M. IONA, *The photon*, The Physics Teacher (October 1983) 480.
- [38] J.M. JAUCH, *Gauge Invariance as a Consequence of Galilei-Invariance for Elementary Particles*, Helv. Phys. Acta **37** (1964) 284.
- [39] J. KALCKAR, *On the Measurability of the Spin and Magnetic Moment of the Free Electron*, Nuov. Cim. **8A** (1972) 759.
- [40] J.R. KLAUDER and B. SKAGERSTAM, *Coherent states; Applications in Physics and Mathematical Physics*, World Scientific (Singapore 1985).
- [41] P.L. KNIGHT and L. ALLEN, *Concepts of Quantum Optics*, Pergamon (Oxford 1983).
- [42] W.E. LAMB, Jr. and R.C. RETHERFORD, *Fine Structure of the Hydrogen Atom IV*, Phys. Rev. **86** (1952) 1014.
- [43] R.H. LAMBERT, *Density of States in a Sphere and Cylinder*, Am. J. Phys. **36** (1968) 417, erratum p. 1169.
- [44] L. LANDAU et E. LIFCHITZ, *Physique théorique, t. IV, Théorie quantique relativiste, première partie*, Mir (Moscou 1972).
- [45] J.-M. LÉVY-LEBLOND, *The Pedagogical Role and Epistemological Significance of Group Theory in Quantum Mechanics*, Riv. Nuov. Cim. **4** (1974) 99.
- [46] J.-M. LÉVY-LEBLOND, *Who is Afraid of Non Hermitian Operators? A quantum description of angle and phase*, Ann. Phys. (N.Y.) **101** (1976) 319.

- [47] J.-M. LÉVY-LEBLOND, *Quantum fact and classical fiction: Clarifying Landé's pseudo-paradox*, Am. J. Phys. **44** (1976) 1130.
- [48] J.-M. LÉVY-LEBLOND, *On the Conceptual Nature of the Physical Constants*, Riv. Nuov. Cim. **7** (1977) 187.
- [49] J.-M. LÉVY-LEBLOND, *Towards a Proper Quantum Theory*, in *Quantum Mechanics, a half century later*, ed. J. Leite Lopes and M. Paty, Reidel (Dordrecht 1977).
- [50] J.-M. LÉVY-LEBLOND et F. BALIBAR, *Quantique, Rudiments*, InterEditions (Paris 1984).
- [51] J.-M. LÉVY-LEBLOND et F. BALIBAR, *Quantique, Eléments* (à paraître).
- [52] R. LOUDON, *The quantum theory of light*, Clarendon Press (Oxford 1983).
- [53] A. MESSIAH, *Mécanique quantique*, 2 vol., Dunod (Paris 1962).
- [54] A.B. MIGDAL, *Qualitative Methods in Quantum Theory*, W.A. Benjamin (Reading 1977).
- [55] P.W. MILONNI, *Why spontaneous emission?*, Am. J. Phys. **52** (1984) 340.
- [56] M. MINNAERT, *The nature of light & colour in the open air*, Dover (New York 1954).
- [57] C.W. MISNER, K.S. THORNE et J.A. WHEELER, *Gravitation*, Freeman (New York 1970).
- [58] D.J.K. O'CONNELL, *The Green Flash*, Scientific American **202** (Jan. 1960) 112.
- [59] A. PAIS, *Albert Einstein, La vie et l'œuvre*, InterEditions (Paris 1993).
- [60] R. PEIERLS, *Surprises in Theoretical Physics*, Princeton University Press (Princeton 1979).
- [61] H. POINCARÉ, *Les méthodes nouvelles de la mécanique céleste*, Gauthier-Villars (Paris 1893).
- [62] E. ROHMER, *Le rayon vert*, avec Marie Rivière, (Biarritz 1986).
- [63] J.J. SAKURAI, *Advanced Quantum Mechanics*, Addison-Wesley (Reading 1973).
- [64] M. SARGENT III, M.O. SCULLY and W.E. LAMB JR., *Laser physics*, Addison-Wesley (Reading 1974).

- [65] E. SCHRÖDINGER, *Der stetige Übergang von der Mikro- zur Makromechanik*, Naturw. **14** (1926) 664.
- [66] G.C. SCORGIE, *Simple discovery of the coherent states of the quantised electromagnetic field*, Eur. J. Phys. **2** (1981) 114.
- [67] M.O. SCULLY and M. SARGENT III, *The concept of the photon*, Physics Today (March 1972) 38.
- [68] I. STEWART, *Beating out the shape of a drum*, New Scientist **1825** (1992) 26.
- [69] J. STRNAD, *Photons in introductory quantum physics*, Am. J. Phys. **54** (1986) 650.
- [70] D. TABOR and R.H.S. WINTERTON, *The direct measurement of normal and retarded van der Waals forces*, Proc. Roy. Soc. A. **312** (1969) 435.
- [71] E.F. TAYLOR and J.A. WHEELER, *A la découverte de l'espace-temps*, Dunod (Paris 1970).
- [72] R. THOM, *Le Monde Dimanche*, 30 octobre 1983.
- [73] J. WALKER, *The flying circus of physics with answers*, Wiley (New York 1977).
- [74] J. WALKER, *Expériences d'amateur, Ou les mille et unes façons de rater une béarnaise*, Pour la Science **28** (février 1980) 118.
- [75] J. WALKER, *Expériences d'amateur, Le mystère des arcs-en-ciel et des arcs surnuméraires, si peu observés*, Pour la Science **34** (août 1980) 109.
- [76] D.F. WALLS, *Squeezed states of light*, Nature **306** (1983) 141.
- [77] S. WEINBERG, *The Forces of Nature*, American Scientist (March-April 1977).
- [78] G. WEINREICH, *Comment vibrent les cordes jumelées d'un piano*, Pour la Science **17** (mars 1979) 84.
- [79] V. WEISSKOPF, *The development of field theory in the last 50 years*, Physics Today (Nov. 1981).
- [80] V. WEISSKOPF und E. WIGNER, *Berechnung der natürlichen Linienbreite auf Grund der Diracschen Lichttheorie*, Zs. f. Phys. **63** (1930) 54.
- [81] H. WEYL, *Espace, Temps, Matière, Leçons sur la théorie de la relativité générale*, Blanchard (Paris 1958).