
LP43 : EXEMPLES DE PHENOMENES QUANTIQUES

Stéphane Vittoz

Rapports de Jury

- [2012] Cette leçon ne doit pas se résumer à un catalogue d'expériences retraçant l'histoire de la mécanique quantique, ni à une juxtaposition d'exemples sans une logique didactique. On attend par suite du candidat un choix raisonné d'un petit nombre d'exemples, traités de manière suffisamment approfondie afin de faire ressortir quelques concepts propres à la mécanique quantique.
- [2010] On attend du candidat un choix raisonné d'un petit nombre d'exemples, traités de manière suffisamment approfondie, de manière à faire ressortir la logique propre de la mécanique quantique.

Prérequis

- Dualité onde-corpuscule
- Bases de la mécanique quantique
- Formalisme de Dirac

Bibliographie

- [1] **Basdevant**, Mécanique Quantique (et son CD)
- [2] **Aslangul**, Mécanique Quantique I (& II pour aller plus loin)

Plan de la leçon

INTRODUCTION.....	2
I QUANTIFICATION DES ÉCHANGES ÉNERGÉTIQUES : EFFET PHOTOELECTRIQUE	2
I.1 Effet photoélectrique : mise en évidence et caractéristiques	2
I.2 Limites de la vision classique	4
I.3 Interprétation quantique – Einstein, 1905.....	4
II DESCRIPTION PROBABILISTE : INTERFERENCES DE MATIERE	5
II.1 Dispositif expérimental	6
II.2 Prévision classique	6
II.3 Résultats expérimentaux	7
II.4 Interprétation quantique.....	7
III EXPERIENCE DE STERN ET GERLACH ET NOTION DE SPIN	8
III.1 Dispositif expérimental	8
III.2 Prévision classique	8
III.3 Résultats expérimentaux	9
III.4 Mise en évidence du spin	10
IV MESURE ET MÉCANIQUE QUANTIQUE	10
IV.1 Prévision semi-classique.....	11
IV.2 Observations expérimentales	11
IV.3 Interprétation quantique	11
CONCLUSION	12

Introduction

La mécanique quantique est une science fondamentalement ancrée dans l'expérimentation. Depuis sa naissance au début du 20^e siècle à travers l'étude du rayonnement du corps noir par Planck et celle de l'effet photoélectrique par Einstein, elle trouve aujourd'hui son expression la plus directe dans la vie quotidienne à travers les nanotechnologies. De nombreux phénomènes physiques ne peuvent dès lors être expliqués qu'à travers une approche quantique du problème. Invariablement, ces situations sont liées aux propriétés propres à la mécanique quantique à savoir la description probabiliste des grandeurs physiques et l'abandon de la notion de trajectoire d'une part, et de la quantification des états énergétiques et donc des échanges d'énergie d'autre part. Il est aussi possible de mettre en évidence des propriétés intrinsèquement quantiques telles le spin. Dans la suite, nous allons nous attacher à décrire quelques situations mettant en évidence ces différents aspects.

I Quantification des échanges énergétiques : effet photoélectrique

Jusqu'à l'intervention de Planck, L'approche classique du rayonnement du corps noir reposant sur les lois de Maxwell et la physique statistique aboutissait à un résultat aberrant alors connu sous le nom de catastrophe ultraviolette qui voulait que la densité d'énergie rayonnée tende vers l'infini. Afin de résoudre cette incohérence, l'idée absolument géniale de Planck a été de supposer que les échanges d'énergie liés au rayonnement thermique était nécessairement des multiples d'une « brique » énergétique alors appelée « quantum » dépendant de la fréquence dudit rayonnement soit :

$$E = h\nu = \hbar\omega$$
$$\hbar = \frac{h}{2\pi} = 1.0546 \cdot 10^{-34} \text{ Js}$$

Nous ne rentrerons pas aujourd'hui dans le détail de son approche mais retenons qu'il fut le premier à introduire la notion de quantification des échanges énergétiques. Ce qui paraissait alors être un simple artifice mathématique (aux yeux de Planck avant tout) va pourtant permettre au jeune Albert Einstein à proposer une explication à un autre phénomène sur lequel la communauté scientifique se casse les dents : l'effet photoélectrique, pourtant découvert par Hertz en 1887.

I.1 Effet photoélectrique : mise en évidence et caractéristiques

[2]

Afin de poser les bases de notre réflexion, nous allons mettre en évidence les particularités du problème à travers une expérience relativement simple utilisant un électroscope. Tout d'abord qu'est-ce qu'un électroscope ? Une plaque métallique (habituellement du zinc) est reliée par un axe conducteur à deux bras conducteurs. L'un de ces bras est mobile par rapport à l'autre selon un mouvement pendulaire. Les deux bras sont isolés électriquement de l'extérieur et peuvent être considérés comme uniquement reliés à la plaque métallique. On utilise ensuite une tige d'ébonite que l'on charge électrostatiquement en la frottant avec une peau de chat. On la met ensuite en contact avec la plaque métallique. Le transfert de charges de la tige vers l'électroscope s'équilibre entre la plaque et les deux bras. Quelque soit le signe de la charge transférée, la conservation de la charge impose que les deux bras sont chargés au même signe. Le bras mobile va donc s'écarter du bras fixe.

***NB :** ne pas trop frotter la tige d'ébonite et ne pas rester trop longtemps en contact avec l'électroscope. En effet, premièrement, si la charge transférée est trop importante, le bras mobile a tendance à rester bloqué en position haute. Deuxièmement, si la tige reste en contact trop longtemps avec l'électroscope, il semblerait que se dernier se décharge (peut être à travers la tige et notre bras puisque l'on représente un bon gros réservoir de charges). A noter que l'effet maximal sur l'électroscope s'obtient pour une très faible charge. Donc il faut ABSOLUMENT éviter de toucher l'appareil une fois chargé : un simple contact cutané avec la plaque suffit à le décharger.*

On peut constater que l'appareil ne se décharge pas sous l'influence de la lumière ambiante. Néanmoins, si l'on éclaire la plaque métallique avec une lampe spectrale au mercure Hg, on constate

un mouvement du bras mobile. Suivant le signe de la charge initial transférée, ce dernier va alors se rapprocher du bras fixe ou s'en éloigner. Dans les deux cas, ce comportement traduit la fuite de charges de l'électroscope, c'est-à-dire que des électrons sont arrachés de la plaque. Par conservation de la charge, les charges présentes dans le deux bras migrent vers la plaque. Si les bras étaient chargés négativement, la force de répulsion électrostatique diminue. Dans le cas positif, elle augmente.

L'appareil étant isolé électriquement de l'extérieur, on en déduit que le rayonnement de la lampe est responsable de la fuite des charges. Maintenant, si l'on interpose une plaque de verre entre la lampe et la plaque, la fuite semble s'arrêter. Or la lumière absorbe le rayonnement ultraviolet émis par la lampe Hg. On en déduit que c'est ce rayonnement qui doit être responsable du phénomène observé.

Afin de mieux mettre en avant les particularités de l'effet photoélectrique, nous allons nous intéresser à un dispositif permettant une étude quantitative simple de ses caractéristiques. Le schéma est disponible en Figure 1.

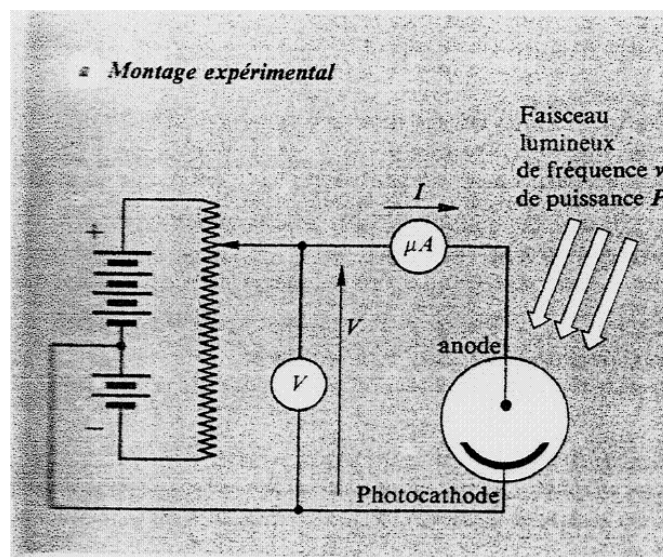


Figure 1 : Schéma du dispositif de Hertz

Une cellule photoélectrique composée d'une photocathode et d'une anode placée dans un vide considéré poussé. Elle est reliée à un potentiomètre permettant de faire varier la tension entre les deux électrodes entre des valeurs négatives et positives. Un voltmètre et un ampèremètre permettent de mesurer cette tension ainsi que le courant circulant entre l'anode et la photocathode. Comme dans le cas de l'électroscope, on soumet ensuite la cellule à un rayonnement spectral connu. On peut alors faire les observations suivantes :

- Suivant le **métal utilisé** pour la photocathode, il existe une **fréquence-seuil** ν_s du rayonnement en dessous duquel aucun effet ne se produit (cf Figure 2). Au-delà, l'effet est **instantané** au sens où aucun délai mesurable n'existe entre le début de l'illumination et le début de déviation du courant mesuré par l'ampèremètre.
- La caractéristique $I(V)$ démarre linéairement pour une tension V_0 **négative** pour saturer à $I = I_{dsat}$ lorsque la ddp devient élevée (cf Figure 2a). La valeur d' I_{dsat} dépend linéairement de l'intensité du flux incident mais pas la contre-tension V_0 .
- En revanche, V_0 varie **linéairement** avec la **fréquence** du rayonnement en s'annulant pour $\nu = \nu_s$, comme le montre la Figure 2b.

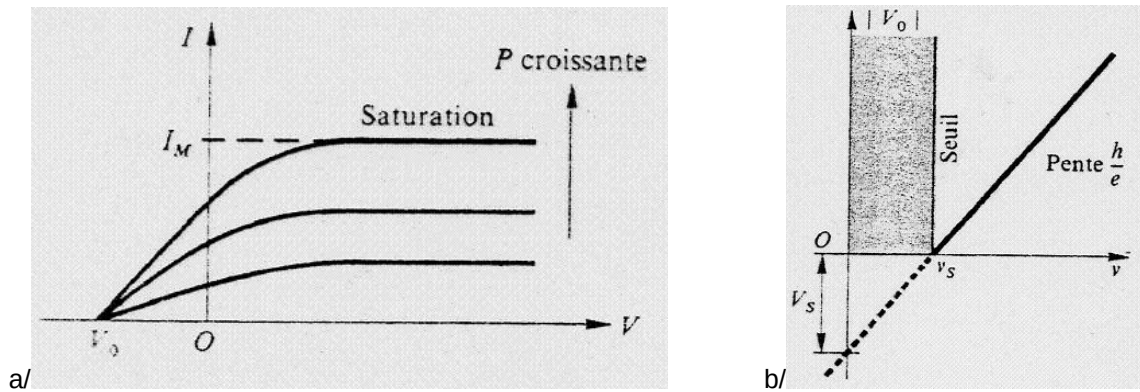


Figure 2 : a) Caractéristique $I(V)$ de la cellule ; b) Variations de la contre-tension V_0

I.2 Limites de la vision classique

La vision classique du rayonnement ne permet d'expliquer **ni l'instantanéité du phénomène, ni l'existence d'une fréquence-seuil ni la négativité de la contre-tension**. La saturation peut être en revanche expliquée par la formation d'une charge d'espace entravant le mouvement des électrons. En se basant sur une vision continue des échanges énergétiques, en attendant suffisamment longtemps et quelque soit la fréquence du rayonnement, on s'attendrait à ce qu'un électron acquiert l'énergie nécessaire à la traversée de la cellule. Cela s'oppose à l'existence d'un phénomène de type « Tout-ou-rien » dépendant qui plus est de la fréquence. Il est nécessaire de se baser sur la quantification des échanges pour expliquer ces caractéristiques.

I.3 Interprétation quantique – Einstein, 1905

La vision d'Einstein sur l'effet photoélectrique repose tout d'abord sur une vision du gaz d'électron dans les métaux conducteurs qui ne sera pas convenablement établie avant le travail de Fermi dans les années 20 et l'élaboration encore plus tardive de la théorie des bandes. Nous avons donc aujourd'hui que dans les matériaux conducteurs, le niveau de Fermi se situe au milieu d'une bande appelée « bande de conduction ». Les électrons qui occupent les états d'énergie présents dans cette bande sont ceux qui participent aux phénomènes de conduction. Se basant sur une approche similaire mais moins poussée, Einstein postule qu'il existe une énergie limite pour laquelle les électrons de cette couche peuvent être arrachés à l'influence du métal.

En termes modernes, cette énergie limite correspond la différence entre l'énergie de Fermi dans le métal et un niveau d'énergie fictif appelé « énergie du vide » correspondant à l'énergie que doit atteindre l'électron pour s'échapper du gaz. L'énergie limite est communément appelée **travail de sortie** du métal noté W_s .

Ceci étant posé, Einstein utilise alors le quantum d'énergie de Planck et suppose que l'échange d'énergie au niveau de l'électron se fait de manière quantifié. Ainsi, pour qu'un électron s'arrache du métal, l'énergie amenée par le quanta doit être égale ou supérieure à W_s . Par conservation de l'énergie avant et après le transfert d'énergie on obtient :

$$h\nu = W_s + \frac{1}{2}mv^2 \geq W_s$$

On observe ainsi un cas limite : celui où le quanta amène très exactement W_s à un électron au niveau de Fermi. Ce dernier est alors arraché sans vitesse initiale. Si le quanta correspond à une énergie plus faible, rien ne se passera ! La **fréquence-seuil** ainsi obtenue ne dépend effectivement que de la **nature du métal** comme observé :

$$\boxed{\nu_s = \frac{W_s}{h}}$$

Pour ce qui est de la **contre-tension négative** $V_0(\nu)$, elle s'explique alors en considérant l'énergie cinétique après arrachement des électrons du niveau de Fermi E_{c0} :

$$E_{c0} = h(\nu - \nu_s) \quad (\text{pour } \nu \geq \nu_s)$$

Une tension négative sur la cellule s'oppose au mouvement des électrons de la cathode à l'anode. Pour une fréquence du rayonnement donnée, V_0 correspond alors à la tension de freinage permettant de ralentir les électrons les plus énergétiques pour une fréquence donnée ie ceux arrachés depuis le niveau de Fermi. On obtient alors l'égalité suivante entre l'énergie cinétique initiale de ces derniers et l'énergie prélevée par la contre-tension :

$$E_{c0} = h(\nu - \nu_s) = eV_0$$

$$\boxed{V_0 = -\frac{h}{e}(\nu - \nu_s)}$$

On remarquera au passage que l'étude la courbe $V_0(\nu)$ permet une mesure direct de la constante de Planck. Notons de plus l'ordre de grandeur du travail de sortie d'un métal qui varie de 2 eV à quelques eV ce qui correspond à des rayonnements de fréquence limite située dans l'ultraviolet. Dans le cas du zinc, $W_s = 3.4$ eV soit $\lambda_s = 365$ nm.

Revenons sur l'expérience de l'électroscope. Nous comprenons désormais que si la plaque ne se décharge avec la lumière ambiante c'est parce que la densité spectrale de ce rayonnement est très atténuée dans l'ultraviolet. Lorsque l'on allume la lampe Hg qui émet dans ces fréquences, la plaque se décharge. L'interposition d'une plaque de verre, absorbant l'ultraviolet, met ainsi logiquement fin au phénomène puisque le rayonnement incident est composé de fréquences trop faibles.

Pour conclure, notons que l'instantanéité du phénomène suggère a priori un transfert par collision entre l'électron et « quelque chose » amenant le quanta d'énergie. Cette observation convaincra Einstein que le quanta d'énergie se comporte comme une particule (qui deviendra 20 plus tard le photon) responsable du rayonnement électromagnétique.

Nous allons maintenant nous intéresser à un autre aspect de la mécanique quantique : la description probabiliste du comportement des particules à travers la notion de dualité onde-corpuscule.

II Description probabiliste : interférences de matière

La dualité onde-corpuscule est une propriété intrinsèque à toute particule mise en avant par De Broglie en 1923 qui établit la relation entre impulsion et longueur d'onde pour ces ondes de matière. Dès lors, on peut supposer que les phénomènes propres aux ondes peuvent être observés pour les ondes de matière. La preuve fondamentale de la dualité onde-corpuscule est le fait de Davisson et Germer qui, en 1927, observèrent des phénomènes de diffraction d'un faisceau d'électron à travers un cristal de nickel. Nous n'aborderons pas ici cette expérience mais nous allons plutôt présenter les phénomènes d'interférences de matière. Longtemps resté une expérience de pensée, l'obtention rigoureuse de ce type de phénomènes pour des électrons date des années 60 et pour des atomes des années 80.

II.1 Dispositif expérimental

Nous allons ici nous intéresser à la description de l'interférence entre atomes à travers un dispositif à fentes d'Young (cf Figure 3).

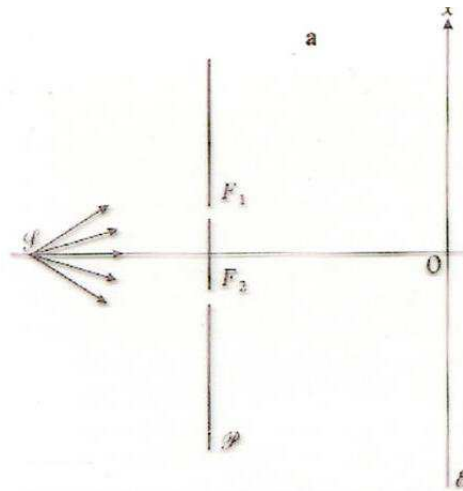


Figure 3 : Schéma du dispositif à fentes d'Young

On considère une source d'atomes produisant un flux incident sur deux fentes suffisamment fines et on observe l'impact des atomes sur un écran de détection situé à une distance grande devant la distance entre les deux fentes.

II.2 Prédiction classique

Dans une approche classique corpusculaire, on raisonne comme suit. En bouchant l'une ou l'autre des fentes, on obtient une répartition des impacts comportant un maximum face à la fente puis décroissant symétriquement de part et d'autre de ce maximum. Lorsque l'on ouvre les deux fentes, on s'attend donc à ce que les électrons passent aléatoirement mais avec la même probabilité par l'une ou l'autre fente. Auquel cas, le profil d'impact attendu correspondrait directement à la somme des profils obtenus pour chaque fente comme présenté sur la Figure 4.

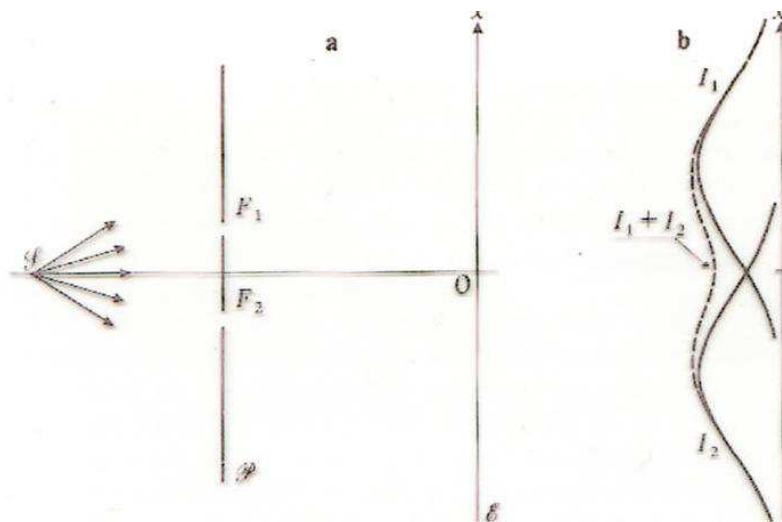


Figure 4 : Résultats de la prédiction classique

Néanmoins, les résultats expérimentaux montrent un résultat tout à fait différent qui remet la description classique en cause et plus encore.

II.3 Résultats expérimentaux

[Appli CD Basdevant 1.1]

Notons que si l'une ou l'autre des fentes on obtient un résultat en accord avec la prévision classique. Néanmoins, lorsque les deux fentes sont ouvertes, le résultat est profondément différent. L'application utilisée permet de retranscrire les observations réellement obtenues par l'expérience.

Tout d'abord notons que pour un flux relativement faible, il est possible d'observer individuellement l'impact des atomes. Cela confirme la nature corpusculaire des atomes étudiés. Néanmoins, ces impacts se font de manière **aléatoire** ce qui traduit un comportement **probabiliste** sans pour autant entrer en conflit avec la mécanique classique. En effet, en appliquant le même raisonnement utilisé pour la description de phénomènes aléatoires classiques (ex : pile ou face) et en acquérant suffisamment, on devrait a priori être capable de décrire ce phénomène.

Mais les observations ne s'arrêtent pas là. En effet, après un grand nombre d'impacts, le profil observé ne correspond pas à la somme des profils obtenus pour chaque fente. La répartition des impacts présentent **des franges de densité maximales et minimales** très similaires aux interférences lumineuses. Il devient dès lors impossible de déterminer par quelles fentes chaque atome est passé. On aboutit à un constat : chaque atome est passé par « les deux fentes à la fois ». Une simple description probabiliste en termes de **trajectoire** d'un **corpuscule** se révèle **impossible** : nous avons atteint les limites de la description classique. Notons de plus que l'observations des franges n'est possible que dans la limite d'un grand nombre d'impacts et ne semble pas prévisible si l'on s'intéresse à un seul atome.

II.4 Interprétation quantique

Pour expliquer ce phénomène d'interférence, il faut avoir recours à la notion de fonction d'onde ou amplitude de probabilité. On attache à chaque particule une amplitude de probabilité $\psi(\vec{r}, t)$ dont la norme au carré fournit la probabilité de présence de la particule en un point donné de l'espace :

$$dP(\vec{r}, t) = |\psi(\vec{r}, t)|^2 \cdot d^3\vec{r}$$

Cette fonction contient alors toutes les informations sur la nature ondulatoire de la particule. Ainsi, pour les cas où l'une ou l'autre des fentes est bouchée, on peut définir pour chaque atome une fonction $\psi_1(\vec{r}, t)$ resp. $\psi_2(\vec{r}, t)$ représentant l'amplitude de probabilité relative au comportement de l'atome dans chaque cas. Puisque la fonction d'onde suit le principe de superposition, la situation où les deux fentes sont ouvertes peut être représentée par $\psi_{12}(\vec{r}, t) = \psi_1(\vec{r}, t) + \psi_2(\vec{r}, t)$. C'est cette particularité même qui permet d'expliquer le phénomène d'interférence !

En effet, la densité de probabilité liée à cette situation devient :

$$|\psi_{12}(\vec{r}, t)|^2 = |\psi_1(\vec{r}, t) + \psi_2(\vec{r}, t)|^2$$

On remarque donc que la densité de probabilité de la situation où l'on observe les interférences ne correspond pas la simple somme des densités de probabilité obtenues pour chaque fente séparément. Sa forme est d'ailleurs formellement similaire à celle de l'intensité lumineuse dans le cas...des interférences. On explique ainsi le phénomène observé.

Cette description tout en appuyant la nature ondulatoire des particules met surtout en exergue l'abandon de la notion de trajectoire en termes classiques : le comportement d'une particule n'est plus décrit qu'à travers sa fonction d'onde, une grandeur intrinsèquement probabiliste liés à des comportements variables de la particule.

Les nombreuses propriétés de la fonction d'onde ont permis d'aboutir à de nouvelles interprétations de phénomènes jusqu'alors difficiles à expliquer. Elle intervient par exemple en chimie

dans la description orbitale des liaisons chimiques ou encore dans la répartition énergétique des électrons dans un cristal à travers la théorie des bandes.

Enfin, l'équation de Schrödinger qui régit le comportement dynamique de la fonction d'onde a permis de mettre en avant la notion de confinement qui a de nombreuses applications théoriques mais aussi pratiques notamment dans les nanotechnologies.

Nous avons pu observer deux phénomènes représentatifs des caractéristiques de la description quantique mais possédant un lien fort avec des comportements classiques. Intéressons-nous désormais à une propriété intrinsèquement quantique de la matière : le spin.

III Expérience de Stern et Gerlach et notion de spin

[1] L'expérience de Stern et Gerlach réalisée en 1921 a abouti à travers son analyse par la mécanique quantique à la démonstration de l'existence d'un moment cinétique interne à chaque particule : le spin. Cette partie a pour but de décrire le cheminement de pensée opéré depuis l'expérience jusqu'à la notion même de spin.

III.1 Dispositif expérimental

Un faisceau d'atomes d'argent produit à l'aide d'un four est dirigé vers un aimant produisant un champ magnétique transversal au faisceau. On dispose un écran de détection derrière l'aimant afin de relever les modifications apportées à la trajectoire par le champ magnétique.

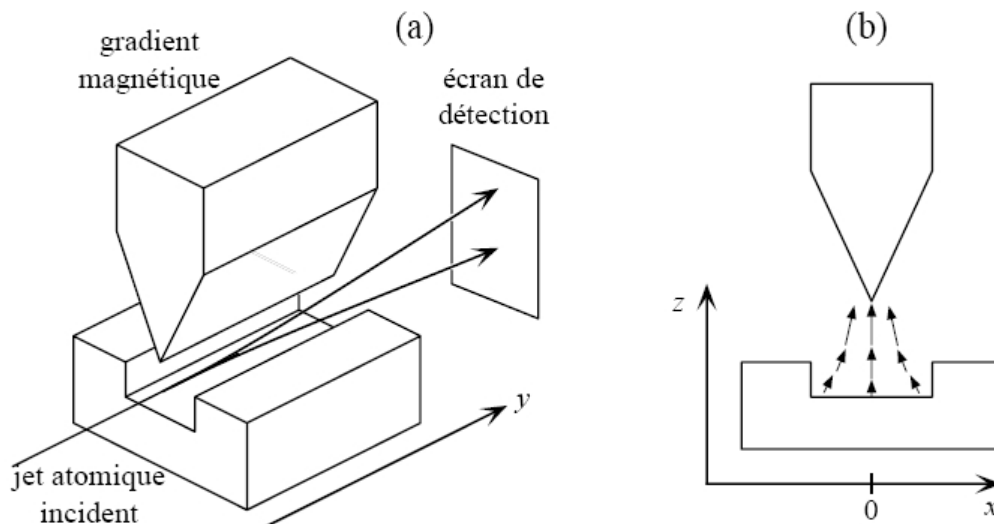


Figure 5 : Schéma du dispositif de Stern et Gerlach et allure du champ magnétique

III.2 Prévision classique

Les atomes d'argent possèdent un moment magnétique $\vec{\mu}$ qui interagit avec le champ $\vec{B}$ à la traversée de l'aimant. L'énergie d'interaction E et le couple Γ s'exerçant sur le moment magnétique correspondants s'écrivent :

$$E = -\vec{\mu} \cdot \vec{B} \quad \Gamma = \vec{\mu} \wedge \vec{B}$$

La force exercée sur la particule se déduit de E et peut s'écrire :

$$\vec{F} = \vec{\nabla}(\vec{\mu} \cdot \vec{B}) = \mu_x \vec{\nabla} B_x + \mu_y \vec{\nabla} B_y + \mu_z \vec{\nabla} B_z$$

En considérant un modèle classique de l'atome d'hydrogène, il est possible d'établir et de généraliser à tout atome la relation de proportionnalité qui lie le moment magnétique au moment cinétique $\vec{L}$ de l'atome :

$$\boxed{\vec{\mu} = \gamma_0 \vec{L}} \quad \gamma_0 = -\frac{q}{2m_e}$$

Le coefficient γ_0 est appelé *rapport gyromagnétique*. Notons que le moment magnétique suit dès lors un mouvement de précession autour de l'axe du champ magnétique à la pulsation ω_0 dite de Larmor :

$$\omega_0 = -\gamma_0 \cdot B$$

Si l'on s'intéresse au mouvement d'un atome dans le plan de symétrie $x = 0$, et en supposant les variations du champ faible devant la pulsation de Larmor, la force exercée sur l'atome ne conserve que sa composante selon z :

$$\boxed{\vec{F} = \mu_z \frac{\partial B_z}{\partial z} \vec{e}_z}$$

Dans une approche statistique classique où l'on suppose que le moment magnétique de chaque atome est de norme μ_0 et direction aléatoire, on s'attend à observer sur l'écran un segment distribué selon z dont les extrema correspondent à des valeurs de $\mu_z = \pm\mu_0$. Du fait de la largeur du faisceau et de la forme du champ imposée, on peut aussi s'attendre à observer un léger étalement de ce segment selon l'axe x . En champ nul, on s'attend à obtenir une tâche située en $z = 0$. Ces prévisions classiques sont résumées en Figure 6.

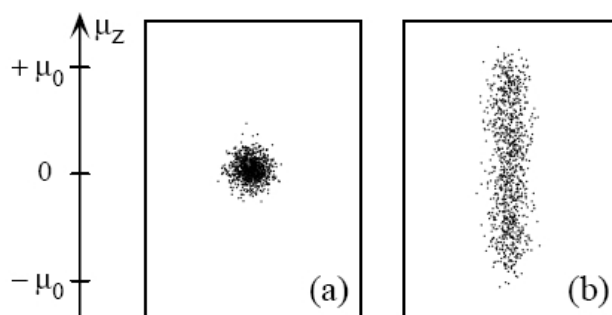


Figure 6 : Prévision classique a) en champ nul b) avec un champ selon z

Et personne ne s'étonnera de constater que les résultats obtenus sont radicalement différents.

III.3 Résultats expérimentaux

Si en champ nul, le résultat obtenu suit bien la prévision classique, lorsqu'un champ magnétique est appliqué on obtient le résultat de la Figure 7c.

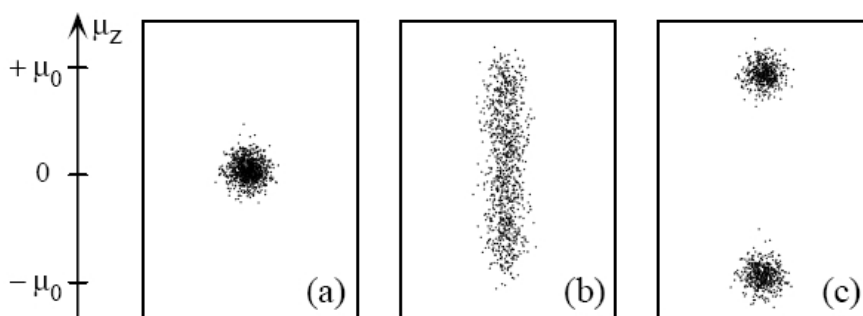


Figure 7 : Prévisions classiques et résultats expérimentaux (c)

On constate que les atomes ont suivi deux déviations privilégiées correspondants aux moments magnétiques extrémaux prévus par l'approche classique. La valeur correspondante de μ_0 pour l'atome d'argent est de l'ordre de $9,27 \cdot 10^{-24} \text{ J.T}^{-1}$. Notons que l'utilisation d'autres types d'atomes, on observe plus de deux tâches, parfois avec une tâche centrale et dans tous les cas symétriquement réparties de part et d'autre de $z = 0$. L'approche classique est ainsi mise en défaut : le moment magnétique et donc le moment cinétique de l'atome (on peut généraliser à toute particule) est quantifié spatialement puisque son influence sur l'atome (la déviation) ne prend qu'un nombre fini de valeurs discrètes.

III.4 Mise en évidence du spin

[1, chap.10]

Pris telle quelle, cette quantification du moment cinétique est une manifestation quantique mais ne permet pas directement de mettre en évidence le spin de l'atome. Pour cela il faut s'attarder sur la nature de l'observable liée au moment cinétique orbital $\hat{L}$:

$$\hat{L} = \hat{r} \wedge \hat{p} \quad \hat{L} \wedge \hat{L} = i\hbar \hat{L}$$

$\hat{r}$ et $\hat{p}$ correspondant aux observables position et impulsion. On définit de manière générale une observable vectorielle moment cinétique $\hat{J}$ à partir des relations particulières du moment cinétique orbital, notamment la seconde. Auquel cas une étude approfondie de cette observable permet d'établir les résultats suivants :

- L'observable $\hat{J}^2 = \hat{J}_x^2 + \hat{J}_y^2 + \hat{J}_z^2$ a des valeurs propres de la forme $j(j+1)\hbar^2$ où j est un entier ou demi-entier et $j \geq 0$.
- Auquel cas, les valeurs propres de l'opérateur scalaire $\hat{J}_i$ ($i = x, y \text{ ou } z$) a des valeurs propres de la forme $m\hbar$, m étant un entier ou demi-entier tel que $m = \{-j, -j+1, \dots, j-1, j\}$

En appliquant ces résultats au cas particulier du moment cinétique orbital, il est possible de démontrer que j et m sont nécessairement entiers. Notamment $m = 0$ est une possibilité. Cela signifie que dans le cas d'un atome possédant un moment cinétique de type orbital, la probabilité d'une déviation nulle par l'expérience de Stern et Gerlach est non nulle. Or dans le cas des atomes d'argent, cette possibilité n'est pas observée. On en déduit que le comportement du moment magnétique est dans ce cas régit par un **moment cinétique non orbital et interne** à la particule : le spin. Les valeurs propres de l'observable qui lui est liée correspond à des valeurs de j et m demi-entières : on parle de **spin** $\frac{1}{2}$. Cette propriété des particules est intrinsèque à la mécanique quantique et ne trouve pas d'équivalent en mécanique classique. On peut grossièrement l'interpréter comme une rotation de la particule sur elle-même. Notons néanmoins que le spin peut aussi être entier comme dans le cas du photon (spin 1) ou nul comme dans le cas de l'atome de carbone...du boson de Higgs.

Intéressons-nous maintenant à une dernière caractéristique de la mécanique quantique : la notion quantique de la mesure.

IV Mesure et mécanique quantique

Pour aborder le grand thème de la mesure en mécanique quantique et la notion de réduction du paquet d'onde, nous allons réutiliser l'expérience de Stern et Gerlach. Nous allons cette fois-ci mettre plusieurs de ces dispositifs « en série » comme le montre la Figure 8.

Un premier dispositif mesurant le moment selon z est muni d'un bloqueur sélectionnant uniquement les atomes de moment positif. Ce faisceau est alors dirigé à travers un dispositif mesurant le moment selon x . On observe une quantification du moment selon x symétriquement équivalente à celle selon z . Imaginons maintenant que l'on bloque à nouveau le faisceau correspondant au moment négatif. Puis réinjectons-le dans un dispositif mesurant le moment selon z . Que va-t-on observer ?

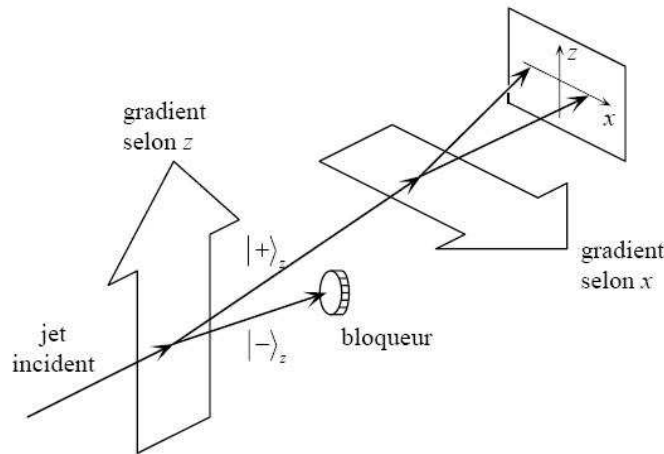


Figure 8 : Dispositifs de Stern et Gerlach « en série »

IV.1 Pr vision semi-classique

Selon une approche semi-classique tenant compte de la quantification du moment magn tique mais abordant classiquement la mesure, on pr suppose l'ind pendance de la mesure des diff rentes composantes du moment sous r serve de poss der un contr le suffisant sur le syst me pour que le champ magn tique du dispositif de mesure d'une composante n'influence pas les autres composantes non mesur es.

Dans ce cas, on s'attendrait   ce que la mesure de μ_z en bout de cha ne ne fournisse que le moment positif puisque dans cette approche la mesure de μ_x ne saurait influencer la pr paration du faisceau selon z op r e en d but de cha ne.

IV.2 Observations exp rimentales

La r alisation d'un tel dispositif permet de montrer qu'en bout de cha ne, on observe non pas une possibilit  de μ_z mais bien les deux possibilit s $\mu_z = \pm\mu_0$ que l'on aurait obtenues par une mesure directe sans bloqueur. L'approche semi-classique de la mesure est mise en d faut et   raison puisque l'hypoth se d'ind pendance des mesures de chaque composante est erron e au sens quantique.

En effet, dans l'approche quantique, la mesure selon x a d truit la pr paration des atomes dans l' tat μ_z positif. Pour comprendre de quelle mani re, mod lisons la situation dans le formalisme de Dirac.

IV.3 Interpr tation quantique

Nous nous int ressons ici   la description d'un ph nom ne li  au degr  de libert  interne de l'atome (son spin). On note $\hat{\mu}_i$ l'observable du moment magn tique de la composante selon i . Au vu des r sultats exp rimentaux, dans le cas des atomes d'argent, ces observables ne poss dent que deux valeurs propres $\pm\mu_0$. Du fait de la nature du spin $1/2$, on admettra qu'un espace de dimension est suffisant   la description du probl me. Auquel cas, on note $|+\rangle_i$ et $|-\rangle_i$ les  tats propres de l'observable $\hat{\mu}_i$ fournissant $+\mu_0$ resp. $-\mu_0$. Ces  tats propres forment donc une base de l'espace de Hilbert dans lequel on travaille. Dans le cas de la composante selon z , un  tat quelconque de moment magn tique d'un atome peut donc s' crire :

$$|\mu\rangle = a|+\rangle_z + b|-\rangle_z \quad |a|^2 + |b|^2 = 1$$

La repr sentation matricielle de $\hat{\mu}_z$ est une matrice 2×2 qui va comme suit :

$$\hat{\mu}_z = \mu_0 \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}$$

En utilisant les propriétés des observables, les caractéristiques du moment magnétique et les résultats expérimentaux, il est alors possible de déduire la forme des observables $\hat{\mu}_x$ et $\hat{\mu}_y$:

$$\hat{\mu}_x = \mu_0 \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad \hat{\mu}_y = \mu_0 \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}$$

On en déduit la forme des états propres de ces deux observables :

$$|\pm\rangle_x = \frac{1}{\sqrt{2}}(|+\rangle_z \pm |-\rangle_z) \quad |\pm\rangle_y = \frac{1}{\sqrt{2}}(|+\rangle_z \pm i|-\rangle_z)$$

On remarque alors une fait important : ces observables ne commutent pas. Notamment :

$$[\hat{\mu}_z, \hat{\mu}_x] = 2i\mu_0\hat{\mu}_y$$

Or le principe d'incertitude de Heisenberg nous fournit l'inégalité liant l'indétermination des deux observables selon x et selon z pour état quelconque $|\mu\rangle$:

$$\Delta\hat{\mu}_z \cdot \Delta\hat{\mu}_x \geq \frac{\langle \mu | [\hat{\mu}_z, \hat{\mu}_x] | \mu \rangle}{2}$$

Il est donc impossible de connaître précisément deux composantes du moment magnétique au cours de la même mesure. En préparant le système dans l'état $|+\rangle_x$ en milieu de chaîne, nous avons rétabli l'indétermination sur μ_z . C'est pourquoi l'on observe en fin de chaîne les deux états possibles pour cette composante.

Cette analyse permet de se rendre compte de l'importance de l'étude des observables liées à une expérience donnée car suivant les cas, la mesure simultanée de plusieurs grandeurs avec une précision arbitraire est impossible. Il s'agit d'une illustration de la problématique de la mesure en mécanique quantique reposant sur la non commutativité des observables. Elle est subtilement différente de la réduction du paquet d'onde.

Conclusion

L'exploitation du formalisme de la mécanique quantique permet d'expliquer de nombreux phénomènes qui mettent les approches classiques en échec. Néanmoins, il est important de retenir que ces phénomènes ont quelques particularités communes qui rendent nécessaire l'approche quantique.

Ainsi, les interactions énergétiques de particules à particules demandent une approche tenant compte de la quantification énergétique. Les phénomènes faisant intervenir la nature ondulatoire de la matière sont nécessairement décrits à l'aide de fonctions d'ondes si l'on veut rendre compte des résultats expérimentaux. Enfin, les phénomènes faisant intervenir le spin, intrinsèquement quantique.

Il faut retenir que la mécanique loin d'être une simple assise théorique doté d'un formalisme abstrait est profondément ancrée dans l'expérimental. Et si nous n'avons vu ici que ses vertus explicatives, sa nature prédictive est un sujet de recherche de pointe et a déjà abouti à de nombreuses applications dans les hautes technologies.