



Atomes et Cavité: Complémentarité et Fonctions de Wigner

Patrice Bertet

► To cite this version:

Patrice Bertet. Atomes et Cavité: Complémentarité et Fonctions de Wigner. Physique Atomique [physics.atom-ph]. Université Pierre et Marie Curie - Paris VI, 2002. Français. NNT : . tel-00002496

HAL Id: tel-00002496

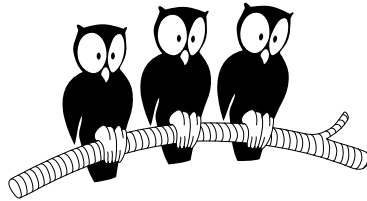
<https://theses.hal.science/tel-00002496>

Submitted on 28 Feb 2003

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.

**DÉPARTEMENT DE PHYSIQUE
DE L'ÉCOLE NORMALE SUPÉRIEURE
LABORATOIRE KASTLER-BROSSEL**



THÈSE DE DOCTORAT DE L'UNIVERSITÉ PARIS VI
spécialité : Physique Quantique

présentée par
Patrice BERTET

pour obtenir le titre de Docteur
de l'Université Paris VI

Sujet de la thèse :
**ATOMES ET CAVITE : COMPLEMENTARITE ET
FONCTIONS DE WIGNER**

Soutenue le 9 octobre 2002 devant le jury composé de :

Pr.	C.	FABRE	
Pr.	B.	GIRARD	Rapporteur
M.	P.	GRANGIER	Rapporteur
M.	K.	HARMANS	
Pr.	J.M.	RAIMOND	Directeur de thèse

Table des matières

Table des matières	i
Remerciements	v
Introduction	1
1 Complémentarité et intrication	9
1.1 Petit historique du concept de complémentarité	9
1.1.1 Les difficultés de la mécanique des quanta	9
1.1.2 L'interprétation de Schrödinger	10
1.1.3 L'interprétation de Heisenberg et le principe d'incertitude . . .	11
1.1.4 Le principe de complémentarité	12
1.1.4.a Enoncé	12
1.1.4.b La dualité onde/corpuscule	13
1.1.4.c La portée du principe de complémentarité	14
1.1.5 Complémentarité et intrication	15
1.2 Illustrations expérimentales du principe de complémentarité	17
1.2.1 Le paradigme du principe de complémentarité : les fentes d'Young	18
1.2.1.a Eclairer les fentes F1 et F2	20
1.2.1.b Le "diaphragme mobile"	21
1.2.2 Les "microscopes de Heisenberg"	22
1.2.3 Tester la complémentarité avec un "diaphragme mobile"	26
2 Le dispositif expérimental	31
2.1 Description générale	31
2.1.1 Motivations et ordres de grandeur	31
2.1.2 Schéma du dispositif expérimental	33
2.2 Les atomes de Rydberg circulaires	35
2.2.1 Propriétés	35
2.2.2 La préparation	39
2.2.3 La sélection de vitesse	42
2.2.4 La détection	48
2.3 La cavité supraconductrice	50

2.3.1	Propriétés générales	50
2.3.2	Le temps de vie de la cavité	53
3	Techniques d'électrodynamique quantique en cavité	57
3.1	Observer les oscillations de Rabi quantiques	58
3.1.1	Contrôler l'accord atome-cavité	58
3.1.1.a	Accorder la cavité sur la transition atomique	58
3.1.1.b	Contrôler le temps d'interaction	60
3.1.2	Refroidir la cavité avant une expérience	61
3.1.3	Oscillations de Rabi	65
3.2	L'interféromètre de Ramsey	66
3.2.1	Principe de fonctionnement	67
3.2.2	Les franges de Ramsey	68
3.3	Injecter un champ cohérent $ \alpha\rangle$	75
3.3.1	Propriétés d'un champ cohérent	75
3.3.2	Comment calibrer un champ cohérent ?	77
4	Réalisation d'un interféromètre à la frontière quantique/classique	83
4.1	Principe de l'expérience	83
4.1.1	La lame séparatrice quantique/classique	83
4.1.2	Le contraste des franges	85
4.2	Réalisation de l'expérience de complémentarité	87
4.2.1	L'accord de la cavité	87
4.2.2	Calibration du champ	87
4.2.2.a	Description générale de la méthode	88
4.2.2.b	Les résultats expérimentaux	88
4.2.2.c	Exploitation des franges de calibration	89
4.2.3	L'interféromètre à lame séparatrice quantique	93
4.2.4	Résultats expérimentaux	95
4.2.5	Contraste des franges et fluctuations quantiques	98
4.3	La gomme quantique	99
4.3.1	Principe de l'expérience	99
4.3.2	Réalisation	103
4.3.3	Résultats	103
4.4	Qu'est-ce qu'un champ classique ?	105
5	Détermination directe de la fonction de Wigner d'un état de Fock à un photon	107
5.1	La fonction de Wigner	108
5.1.1	Présentation générale	108
5.1.2	Quelques propriétés de la fonction de Wigner	109
5.1.3	Quelques fonctions de Wigner	112

5.1.4	Mesures de la fonction de Wigner déjà effectuées	114
5.1.4.a	Méthode tomographique	114
5.1.4.b	Méthode directe	116
5.2	Une méthode directe pour mesurer la fonction de Wigner d'un mode de la cavité	118
5.3	Expériences préliminaires	120
5.3.1	Le déphasage de π par photon	120
5.3.1.a	Comment l'obtenir	120
5.3.1.b	Réglage des franges e-g	121
5.3.2	La calibration du champ	126
5.3.3	Le temps de vie de la cavité	128
5.4	Fonction de Wigner du vide : démonstration de la méthode	129
5.4.1	Réalisation expérimentale	129
5.4.2	Discussion des résultats	131
5.4.2.a	La phase des franges	131
5.4.2.b	Résultats sur le contraste	133
5.4.2.c	Discussion du facteur de normalisation	134
5.5	Fonction de Wigner d'un état de Fock à un photon	135
5.5.1	Réalisation expérimentale	135
5.5.2	Discussion des résultats	137
5.5.2.a	Résultats sur la phase	137
5.5.2.b	Distribution de Wigner d'un état de Fock $ 1\rangle$	138
5.5.2.c	Discussion du facteur de normalisation	141
5.6	Discussion et perspectives	142
5.6.1	Intérêt des résultats précédents	142
5.6.2	Perspectives	143
5.6.2.a	Améliorations possibles	143
5.6.2.b	Extensions possibles	143
5.6.3	Perspectives à plus long terme	144
A	Eléments de théorie	151
A.1	Le Hamiltonien de Jaynes-Cummings	151
A.2	L'évolution temporelle du système	152
A.2.1	Les états habillés	152
A.2.2	L'interaction résonante	153
A.2.3	L'interaction dispersive	155
B	Discussion des facteurs de normalisation des mesures de fonctions de Wigner	157
B.1	Analyse détaillée des imperfections de l'expérience de mesure de la fonction de Wigner du vide	157
B.1.1	Contraste intrinsèque de l'interféromètre	157

B.1.2	Événements à deux atomes	158
B.1.3	Champ thermique dans le mode BF	161
B.2	Analyse détaillée du facteur de normalisation de la mesure de la fonction de Wigner de $ 1\rangle$	162
B.2.1	La première impulsion de Ramsey	162
B.2.2	Les imperfections de l'impulsion π dans la cavité	165
Bibliographie		167

Remerciements

J'ai effectué mon travail de thèse au laboratoire Kastler-Brossel. Je souhaite remercier ici ses directeurs successifs, Michèle Leduc, Elisabeth Giacobino, et Frank Laloë, d'avoir bien voulu m'y accueillir.

Jean-Michel Raimond a encadré ce travail de thèse. Je tiens à le remercier vivement pour sa constante disponibilité, l'attention et la bienveillance dont il a toujours fait preuve à mon égard. Sur un plan purement scientifique, je lui dois énormément. Ses explications physiques sont toujours extrêmement imagées et intuitives. Sa maîtrise technique, dans chacun des domaines pourtant variés que peut couvrir cette expérience, est proprement impressionnante.

Travailler dans l'équipe de Serge Haroche est une chance extraordinaire pour mieux comprendre la mécanique quantique. C'est un professeur exceptionnel, par la clarté et la profondeur de ses exposés. J'ai beaucoup appris en écoutant ses cours, aussi bien ceux du DEA de physique quantique que ceux du collège de France. Serge a en outre toujours suivi avec beaucoup d'attention les résultats que nous obtenions en salle de manip.

Michel Brune a assuré le suivi quotidien de mon travail expérimental. Il connaît les moindres recoins de l'expérience, et son aide a été précieuse. A chaque fois qu'une difficulté imprévue surgissait, et malgré ses nombreuses occupations extérieures, Michel était disponible et venait nous apprendre quelque nouvelle subtilité expérimentale qui nous avait échappé jusque-là. Il m'a impressionné notamment par son efficacité, ses qualités d'expérimentateur hors-pair, et ses fulgurantes intuitions physiques. Sa gentillesse est pour beaucoup dans l'excellente ambiance qui règne au sein du groupe.

Gilles Nogues et Arno Rauschenbeutel ont guidé mes premiers pas en S15, et devrais-je dire en physique expérimentale. Je tiens ici à les en remercier chaleureusement. Il est très agréable de travailler avec eux. Leur capacité de travail, leur niveau d'exigence, et leurs qualités d'expérimentateurs, ont été de constants stimulants pour moi. Je leur suis tout particulièrement redevable, puisque toutes les expériences présentées dans ce mémoire ont été possibles grâce à la qualité du montage effectué par Gilles et Arno avant que je n'arrive dans le groupe.

Stefano Osnaghi a été mon "co-manip" pendant toute la durée de ma thèse. Nous avons effectué ensemble tout le travail exposé ici. Nous avons souffert ensemble des problèmes posés par les lasers, la cryogénie, ou les micro-ondes. Nous avons passé des nuits à acquérir des données, des semaines à nous battre pour faire marcher un

asservissement, ou comprendre les données. Grâce à son humour, son ouverture d'esprit, sa curiosité sans cesse en éveil, je garde un excellent souvenir même des moments les plus éprouvants de notre travail d'équipe. Je le compte aujourd'hui parmi mes amis. Je tiens aussi à remercier les nombreux autres étudiants qui ont travaillé avec Stefano et moi, pendant des durées plus ou moins longues, pour les bons moments passés en leur compagnie. Il y a tout d'abord les étudiants en thèse qui ont pris la relève en S15 : Alexia Auffèves, Paolo Maioli, Tristan Meunier, Sébastien Gleyzes. Mathieu Perrin et Benoît Darquié ont passé quelques mois avec nous à l'occasion de divers stages. Merci enfin à Michel Gross et Philippe Goy, deux anciens membres du groupe, pour leur aide constante en électronique et en micro-onde.

Notre collaboration avec le groupe de théoriciens de Rio de Janeiro m'a permis de rencontrer un certain nombre de physiciens brésiliens. J'ai été ainsi très heureux de connaître Luiz Davidovich, dont les séminaires donnés dans le groupe m'ont marqué par leur qualité et leur clarté, Nicim Zagury, Enrique Solano. Perola Milman travaille depuis deux ans dans le groupe, et j'ai pu ainsi apprécier sa gentillesse, sa joie de vivre, ainsi que ses qualités de théoricienne, notamment lors de notre séjour aux Houches.

L'équipe animée par Valérie Lefèvre et Jean Hare travaille sur des microsphères en silice, qui constituent des résonateurs optiques. Je tiens à remercier Valérie, pour les discussions intéressantes et variées que nous avons pu avoir, Jean pour la patience dont il a toujours fait preuve lorsque nous faisons appel une nième fois à ses talents d'informaticien, ainsi que tous les étudiants de P10, Sébastien Steiner, Romain Long, Xavier Brockmann, pour les bons moments passés avec eux. Je voudrais enfin souhaiter à Gilles Nogues ainsi qu'à ses étudiants Philippe Hyafil et Jack Mozley, beaucoup de succès et de réussite dans le montage d'une nouvelle manip en R14.

Introduction

Un débat vieux de 80 ans

La mécanique quantique, telle qu'elle a été élaborée entre 1900 et 1930 par quelques physiciens de génie, est aujourd'hui la base de notre description physique de la réalité. Ses premiers succès historiques ont consisté à élucider des énigmes qui irritaient la physique classique de la fin du dix-neuvième siècle : calculer le spectre du rayonnement du corps noir, et comprendre les spectres des atomes. Puis la mécanique quantique a envahi progressivement tous les champs de la physique, de la physique du solide à l'astrophysique, faisant progresser considérablement tous les modèles, permettant des calculs plus précis. C'est aujourd'hui la théorie qui permet de rendre compte du plus vaste ensemble de phénomènes et de calculer certaines valeurs expérimentales avec une précision inégalée. L'impact de la mécanique quantique ne s'arrête pas aux livres de physique, elle a rendu possible deux des inventions les plus importantes du siècle, le laser et le transistor.

Si l'efficacité de la mécanique quantique est incontestable, paradoxalement, elle pose des difficultés d'interprétation qui ont été ressenties depuis l'origine par ses inventeurs, notamment par Einstein, et qui, encore aujourd'hui, peuvent être l'objet de débats animés entre physiciens.

Une question notamment a suscité beaucoup de controverses dans les toutes premières années qui ont suivi l'élaboration de la théorie. Alors que la théorie ondulatoire de la lumière est bien établie au début du vingtième siècle, certains faits expérimentaux nouveaux (effet Compton) suggèrent que la lumière peut avoir un comportement corpusculaire. La lumière est-elle finalement une onde ou un corpuscule ? Pour tenter de réconcilier ces deux théories, Niels Bohr proposa que les deux comportements en question étaient complémentaires, mais qu'ils ne pouvaient se manifester simultanément dans la même expérience. Dans certaines conditions expérimentales, la lumière se comporte comme une onde ; dans d'autres, elle se comporte comme une particule, et le physicien a besoin du recours à ces deux concepts contradictoires que sont l'onde et la particule pour décrire la totalité des faits expérimentaux. En d'autres termes, le comportement d'un objet quantique dépend du dispositif expérimental utilisé pour l'étudier.

Bohr a précisé son idée avec une expérience de pensée schématisée à la figure 1¹.

¹Sur cette figure, on a associé deux schémas dessinés séparément par Bohr [97].

Il s'agit d'un interféromètre d'Young dans lequel le diaphragme du dessus ($D1$) est suspendu par deux ressorts et peut donc se déplacer, tandis que le diaphragme du dessous ($D2$) est solidement fixé au bâti. Dans une expérience de fentes d'Young classique, l'observation de franges d'interférence révèle le caractère ondulatoire de la particule P (neutrons, électrons, photons ...). Dans le dispositif imaginé par Bohr, le passage de la particule par $D1$ lui transfère une certaine impulsion et le fait osciller. Si l'amplitude de cette oscillation est plus grande que la dispersion de position initiale du diaphragme, elle fournit un moyen de déterminer par quel diaphragme, $D1$ ou $D2$, P est passée. Cette information de type corpusculaire est incompatible avec un comportement ondulatoire. On ne devrait donc plus observer de franges d'interférence. C'est ce qui doit arriver lorsque $D1$, que l'on suppose initialement dans l'état fondamental de son mouvement, est suffisamment microscopique. Lorsqu'on augmente la masse du diaphragme $D1$ pour une raideur donnée, on perd l'information sur le chemin suivi, et l'on retrouve les franges d'interférence classiques et le comportement ondulatoire de la particule étudiée. En particulier, on peut étudier la transition entre les régimes de fonctionnement quantique et classique de l'interféromètre.

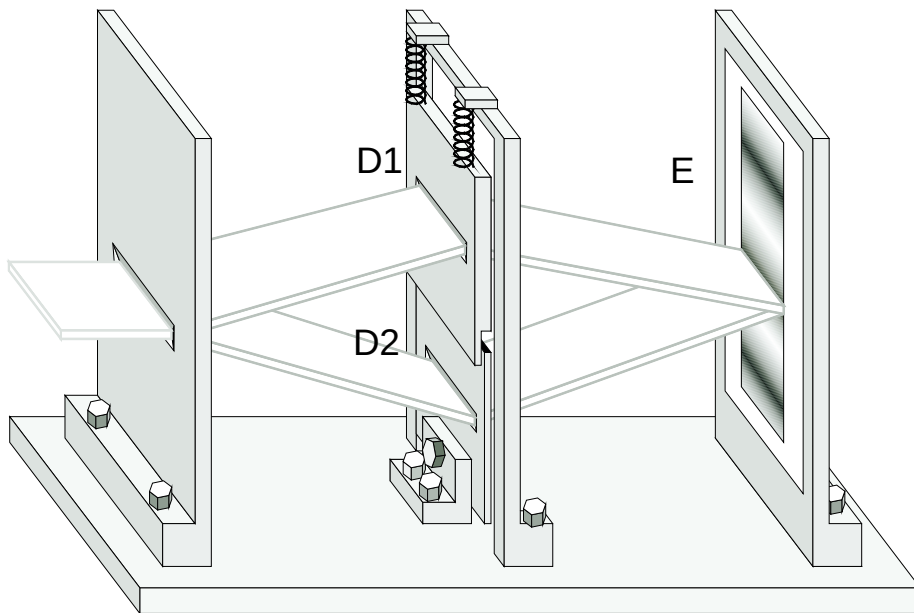


FIG. 1 – Dispositif imaginé par Bohr pour illustrer la notion de complémentarité.

Le concept d'“intrication”, discuté en 1935 par Einstein, Podolsky et Rosen dans leur fameux article [32], donne un éclairage décisif à cette expérience de pensée. Lorsque le diaphragme $D1$ est suffisamment léger, la position de la particule s'intrique avec l'impulsion du diaphragme. P et $D1$ sont alors les deux composantes d'une paire EPR. Observer l'impulsion de $D1$ lève l'ambiguïté sur la trajectoire de la particule et supprime l'interférence. Au contraire, lorsque $D1$ est suffisamment massif, son état final est le même, que la particule passe à-travers lui ou non. Il n'y a alors aucune intrication et

les franges sont visibles. Le contraste des franges fournit une mesure directe du degré d'intrication entre la particule et l'interféromètre.

Il serait techniquement très difficile de réaliser l'interféromètre de la figure 1. Nous avons réalisé une expérience analogue à celle imaginée par Bohr, mais nous utilisons un processus d'interférence qui a lieu dans l'espace de Hilbert à deux dimensions des états d'énergie d'un atome à deux niveaux $|a\rangle$ et $|b\rangle$, et non pas dans l'espace réel. Il s'agit d'appliquer deux impulsions R1 et R2 à l'atome, séparées par un délai T . La première impulsion crée une superposition cohérente des deux niveaux $|a\rangle$ et $|b\rangle$, et la deuxième l'analyse. Comme l'atome a deux “chemins” quantiques indiscernables pour finir dans l'état $|b\rangle$, la probabilité de détecter finalement l'atome dans le niveau $|b\rangle$ révèle des franges d'interférence, en fonction de la phase relative des deux impulsions R1 et R2. Un tel dispositif est appelé un interféromètre de Ramsey [77]. Les impulsions de Ramsey sont l'analogue des lames séparatrices dans un interféromètre de Mach-Zehnder.

Une impulsion de Ramsey classique se fait habituellement dans un champ important, contenant un nombre N de photons grand devant 1. Lorsque N devient petit ($N \sim 1$), l'atome s'intrique avec le champ, de la même manière qu'une particule s'intrique avec le diaphragme D1 lorsque celui-ci est microscopique dans le dispositif de Bohr. Ainsi, le nombre de photons contenus dans le champ d'une impulsion de Ramsey joue un rôle analogue à la masse de D1. Pour réaliser l'équivalent du diaphragme mobile de Bohr avec un interféromètre de Ramsey, il faut donc être capable d'effectuer une impulsion de Ramsey dans un champ microscopique. Le couplage atome-champ doit par conséquent être particulièrement important. Ces conditions de “couplage fort” sont réunies dans le dispositif d'électrodynamique quantique en cavité construit dans notre groupe.

Le dispositif d'électrodynamique quantique en cavité

Pour atteindre le régime de couplage fort entre un atome et le champ électromagnétique, nous étudions l'interaction d'atomes de Rydberg circulaires avec une cavité micro-onde supraconductrice de très haut facteur de qualité. Les états de Rydberg circulaires sont des états atomiques “exotiques” dans lesquels l'électron de valence d'un atome alcalin est à la fois loin du noyau et possède un moment angulaire maximal. Un atome excité dans un tel état possède une longue durée de vie à l'échelle atomique et un grand élément de matrice dipolaire, ce qui assure un couplage important au champ électromagnétique. Les états de Rydberg circulaires sont entièrement caractérisés par la donnée du nombre quantique principal n de l'électron de valence. Nous nous intéressons tout particulièrement aux deux niveaux circulaires $|n = 50\rangle$ (que nous appellerons dans la suite $|g\rangle$) et $|n = 51\rangle$ (que nous appellerons $|e\rangle$). La transition entre les niveaux $|e\rangle$ et $|g\rangle$ a une fréquence de 51.1GHz , dans le domaine des ondes millimétriques, résonante avec la fréquence de l'un des modes de notre cavité micro-onde. Celui-ci est particulièrement bien isolé de l'environnement extérieur, puisque le temps caractéristique de dissipation de l'énergie qui y est stockée est de $T_{cav} = 1\text{ms}$.

Le régime de couplage fort entre l'atome et la cavité se manifeste par le phénomène d'“oscillations de Rabi quantiques”. On envoie un atome, initialement préparé dans l'état $|e\rangle$, dans la cavité vide (état $|0\rangle$). Si la fréquence de la cavité est égale à la fréquence de la transition atomique, les niveaux $|e, 0\rangle$ et $|g, 1\rangle$ ont même énergie. Ils sont en outre fortement couplés. Par conséquent, on assiste à l'échange périodique d'un quantum d'excitation entre l'atome et la cavité, à la fréquence de Rabi qui caractérise la force du couplage. Cela se traduit expérimentalement par une modulation sinusoïdale de la probabilité de détecter l'atome dans le niveau $|e\rangle$ en fonction de la durée de son interaction avec la cavité, comme on peut le constater sur la figure 2 [21].

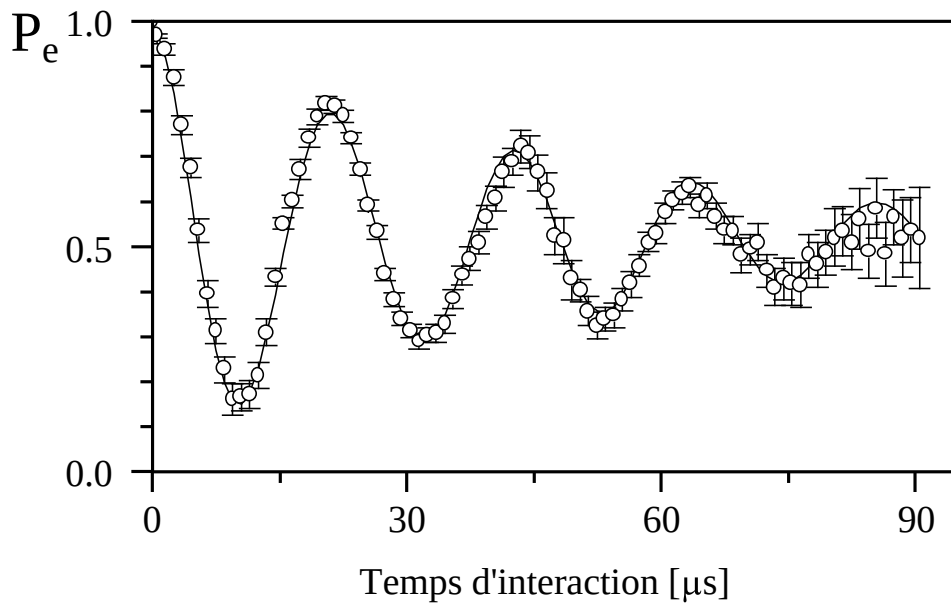


FIG. 2 – Oscillations de Rabi quantiques

Dans la suite de ce mémoire, nous décrirons plus en détail l'appareillage. Soulignons tout de suite que ce dispositif nous a déjà permis de réaliser plusieurs expériences sur les fondements de la mécanique quantique. Parmi celles-ci, nous citerons la réalisation d'une paire maximalelement intriquée d'atomes (paire EPR) [63], la préparation d'un état type “chat de Schrödinger” [16], la détection non-destructive d'un seul photon [70], et la préparation d'un état intriqué tripartite [80]. Nos expériences s'inscrivent en outre dans un effort plus général pour contrôler et manipuler des systèmes quantiques individuels, dans un but soit fondamental (meilleure compréhension de la mécanique quantique, métrologie [90], ...) soit plus appliqué (ordinateurs quantiques [26, 33]). De nombreux autres groupes, travaillant sur des systèmes physiques pourtant extrêmement différents, étudient ainsi le comportement quantique d'ions piégés [64], de photons émis lors de cascades radiatives dans les atomes [2] ou par conversion paramétrique dans un cristal [53], de boîtes quantiques [86, 48], de jonctions Josephson [67, 93, 91, 62], ...

L'expérience du “diaphragme mobile” avec un interféromètre de Ramsey

Grâce au phénomène d'oscillations de Rabi quantiques, il devient possible de réaliser une impulsion de Ramsey dans la cavité même lorsque celle-ci ne contient aucun photon. A fortiori, lorsque la cavité contient un champ plus important, le couplage atome-champ est aussi suffisant. Nous avons donc pu réaliser un interféromètre de Ramsey à la frontière quantique/classique en appliquant la première impulsion R1 à l'atome dans un champ stocké dans le mode de la cavité qui a un très long temps de relaxation, contenant un petit nombre de photons N . La deuxième impulsion de Ramsey est classique, et correspond sur la figure de Bohr au diaphragme fixe $D2$.

Nous avons étudié la transition quantique/classique en faisant varier le nombre de photons N . Comme le suggérait le principe de complémentarité, lorsque $N = 0$, les franges d'interférence disparaissent complètement, alors que lorsque N est grand on retrouve les franges observées dans un interféromètre de Ramsey habituel. Dans le régime intermédiaire, les résultats sont quantitativement prédits par un modèle extrêmement simple dans lequel nous considérons uniquement l'intrication de l'atome avec la cavité. En outre, nous avons pu montrer que lorsque $N = 0$, bien que la cohérence atomique soit en apparence perdue puisque le contraste des franges de Ramsey est alors nul, nous avons la possibilité de retrouver cette cohérence en effaçant l'information contenue dans la cavité. Ainsi, comme le suggérait Bohr, nous avons le *choix* d'observer un comportement ondulatoire ou corpusculaire, selon l'appareil de mesure que nous utilisons.

Décohérence et fonctions de Wigner

Ces notions éclairent considérablement un autre problème fondamental de la théorie quantique : l'inexistence de cohérences quantiques dans le monde macroscopique. La question a été posée dès 1935 par Schrödinger [82] sous la forme paradoxale que l'on connaît : si l'on croit à la mécanique quantique, pourquoi n'observe-t-on jamais un chat, par exemple, dans une superposition d'états “mort” et “vivant” ? Un tel état serait fort différent de la situation que nous connaissons habituellement, d'un chat “mort” ou “vivant”, puisqu'on devrait pouvoir observer des phénomènes d'interférence entre ces deux états, similaires à ceux qui nous sont familiers dans le domaine microscopique. Pourquoi n'observe-t-on couramment que l'aspect corpusculaire des objets macroscopiques, alors que la dualité onde/corpuscule devrait les concerner aussi ?

Cette question joue un rôle crucial dans la théorie de la mesure en mécanique quantique. Imaginons qu'un système microscopique, par exemple un atome à deux niveaux notés $|+\rangle$ et $|-\rangle$, interagisse avec un appareil macroscopique qui mesure son état interne. L'aiguille de cet appareil, initialement dans l'état $|\uparrow\rangle$ finira donc dans deux états macroscopiquement distincts $|\rightarrow\rangle$ ou $|\leftarrow\rangle$ selon l'état initial de l'atome :

$$\begin{aligned} |+\rangle|\uparrow\rangle &\longrightarrow |+\rangle'|\rightarrow\rangle \\ |-\rangle|\uparrow\rangle &\longrightarrow |-\rangle'|\leftarrow\rangle \end{aligned} \tag{1}$$

Maintenant l'atome peut être placé dans une superposition linéaire des états $|+\rangle$ et $|-\rangle$. Dans ce cas, si l'on suppose que l'évolution du système est bien décrite par l'équation de Schrödinger, l'état final sera le suivant :

$$\frac{1}{\sqrt{2}}(|+\rangle + |-\rangle)|\uparrow\rangle \longrightarrow \frac{1}{\sqrt{2}}(|+\rangle'|\rightarrow\rangle + |-\rangle'|\leftarrow\rangle) \quad (2)$$

On devrait alors observer des effets d'interférence quantique entre les deux états macroscopiques de l'appareil de mesure, or il est clair que personne n'a jamais observé cela dans la réalité ! C'est ainsi que Von Neumann a été amené à distinguer deux types d'évolution : l'évolution unitaire des systèmes microscopiques en interaction, décrite par l'équation de Schrödinger d'une part, et une évolution probabiliste associée à tout processus de mesure. Après cette évolution, le système microscopique finit dans l'un des états propres de l'observable correspondant à l'appareil de mesure, avec une probabilité égale au poids de cet état propre dans la fonction d'onde initiale du système (postulat de réduction du paquet d'onde).

Une telle description pose beaucoup de questions. Pour quelle raison l'équation de Schrödinger, qui décrit si bien tous les phénomènes microscopiques, cesserait brutalement d'être applicable dans le domaine macroscopique ? Existe-t-il une "taille critique" pour un système à partir duquel cette équation cesserait d'être valable ?

Pour répondre à ces questions, certains physiciens ont mis en avant le rôle prépondérant joué par le couplage à l'environnement dans la disparition des cohérences quantiques entre états macroscopiques. Cette approche a notamment été développée par Zurek ([100],[101]), et est qualifiée de "théorie de la décohérence". L'environnement joue vis-à-vis des systèmes macroscopiques le rôle d'un détecteur mesurant en permanence l'état du système, similaire en cela au diaphragme mobile de Bohr, à ceci près que dans l'expérience du diaphragme mobile, nous pouvions choisir quel type d'effet nous voulions observer, corpusculaire ou ondulatoire, alors que l'environnement extérieur est par définition ce qui échappe à notre contrôle. C'est le type d'environnement et de couplage à l'environnement qui détermine la possibilité d'observer des franges d'interférence, quelle que soit la taille du système. La compréhension de ce point a suscité beaucoup d'efforts expérimentaux pour voir des effets d'interférence quantique dans des systèmes de plus en plus macroscopiques, en les isolant le mieux possible de leurs environnements naturels. En quelque sorte, on essaye de pousser le plus loin possible la frontière quantique/classique [56].

D'autre part, pour un système donné, la théorie de la décohérence prévoit que les superpositions linéaires d'états sont d'autant plus fragiles que les états sont "distants" l'un de l'autre. En effet, la mesure permanente effectuée par l'environnement sur l'état du système discrimine d'autant plus facilement entre les deux états que ceux-ci sont macroscopiquement distincts. La superposition linéaire se transforme alors en mélange statistique. Certains modèles ont permis de transformer ce raisonnement qualitatif en une évaluation quantitative du "temps de décohérence". Le paradigme de ces modèles est celui d'un oscillateur harmonique couplé linéairement à un bain thermique d'oscillateurs

harmoniques, placé dans un état du type de celui décrit par l'équation 2 [18].

La seule manière de valider expérimentalement ces modèles est d'isoler suffisamment bien un système pour se placer dans une situation où le couplage à l'environnement n'empêche pas l'observation de ces cohérences. En 1996, deux équipes (dont la nôtre) ont réussi à observer la cohérence d'états qualifiés de “chats de Schrödinger” dans des systèmes physiques très différents (ions piégés [66] et électrodynamique quantique en cavité [16]). Par ailleurs, notre équipe a pu étudier la relaxation de cette cohérence et constater qu'elle avait lieu d'autant plus rapidement que la séparation entre les états était plus importante, ce qui montre bien la fragilité extrême des superpositions d'états dans le domaine macroscopique.

Cependant, cette première étude expérimentale de la décohérence était *partielle* puisque nous n'avions étudié à cette occasion qu'un signal de corrélation qui révélait la cohérence de l'état préparé. Ce paramètre est une fonction complexe de tous les coefficients de la matrice densité, et il est clair qu'une grande partie de l'information contenue dans cet état nous échappait. Il serait extrêmement intéressant de reprendre cette étude de la décohérence en caractérisant *complètement* l'état quantique du champ au cours du temps. En particulier, on pourrait ainsi voir directement la transformation d'une superposition cohérente en un mélange statistique.

La suite de mon travail de thèse a consisté à démontrer qu'il était possible de caractériser complètement l'état quantique de la cavité en faisant une succession de mesures par des atomes de Rydberg circulaires. Au lieu de mesurer la matrice densité du champ, nous mesurons la fonction de Wigner du champ, qui est l'analogue quantique d'une distribution de probabilité dans l'espace des phases. La fonction de Wigner, au même titre que la matrice densité, contient toutes les informations sur l'état physique d'un système. Elle est très utile précisément pour étudier la transition quantique/classique car elle est positive pour tout état classique alors qu'elle prend des valeurs négatives pour beaucoup d'états qui n'ont pas d'équivalent classique.

En outre il est possible de mesurer directement la fonction de Wigner en tout point α de l'espace des phases dans notre dispositif. Il suffit de mesurer la valeur moyenne de l'opérateur $\hat{P} = (-1)^{\hat{N}}$ qui mesure la parité du nombre de photons contenus dans la cavité dans l'état déplacé de $-\alpha$. Une méthode très élégante de mesure de $\langle \hat{P} \rangle$ a été proposée par Lutterbach et Davidovich [58] et par Englert [36] indépendamment. Elle consiste à utiliser l'interaction dispersive entre un atome et la cavité dans le régime où la présence d'un seul photon suffit à déphaser la fonction d'onde atomique d'un angle π .

Nous avons pu effectuer la mesure de la fonction de Wigner du vide, ainsi que d'un état de Fock à un photon $|1\rangle$ en utilisant cette méthode. Ceci constitue un premier pas vers la reconstitution d'états quantiques plus complexes, par exemple un “chat de Schrödinger”, et vers l'observation du “film de la décohérence” dont il a déjà été question plus haut.

Le premier chapitre de ce mémoire est dévolu à une étude détaillée du principe de complémentarité. Je m'attacherai à détailler les problèmes que Bohr a cherché à

résoudre en énonçant ce principe, à-travers une courte analyse historique. J'essayerai ensuite de souligner son importance en discutant certaines implications. Enfin pour illustrer ce concept un peu abstrait, je ferai une revue des différentes expériences qui ont été déjà réalisées sur des systèmes différents du nôtre, et qui montrent clairement toute la généralité du concept de complémentarité. Je montrerai l'originalité de notre expérience par rapport à celles qui l'ont précédée.

Dans le deuxième chapitre je décris brièvement le dispositif expérimental d'électrodynamique en cavité que nous utilisons. La construction de ce dispositif a été commencée bien avant mon arrivée dans le groupe, et depuis lors il est en perpétuelle évolution et progression. Je ne chercherai pas à faire une description exhaustive du montage, qui pourra être trouvée dans les thèses des étudiants qui m'ont précédé, ainsi que dans la référence [46]. Je vais plutôt chercher à faire ressortir les différents points qui rendent ce dispositif unique, ainsi que ma contribution à son développement.

Le troisième chapitre décrira les techniques que nous employons couramment pour les expériences d'électrodynamique quantique en cavité. Nous présentons dans ce chapitre les "briques" élémentaires que nous utiliserons dans les expériences présentées dans les chapitres 4 et 5.

Dans le quatrième chapitre je présenterai en détail l'expérience de complémentarité réalisée à l'aide de ce dispositif. Je décrirai notamment le principe du fonctionnement de l'interféromètre qui est le coeur de cette expérience, ainsi que sa réalisation. Puis je présenterai les résultats obtenus.

Enfin dans le cinquième et dernier chapitre je présenterai les résultats obtenus sur la mesure de la fonction de Wigner du champ. Je commencerai par détailler le principe de la mesure, en suivant pour cela la référence [58]. Puis je montrerai comment nous avons mis en oeuvre expérimentalement la méthode proposée, en insistant notamment sur l'obtention pour la première fois d'un déphasage de π par photon dans le régime dispersif. Je présenterai ensuite les résultats obtenus pour un état de Fock à un photon, en les comparant aux autres mesures existantes. Enfin j'évoquerai les perspectives ouvertes par ces résultats.

Chapitre 1

Complémentarité et intrication

Le but de ce chapitre est d'expliciter la notion de complémentarité, afin de mieux saisir les enjeux de l'expérience qui sera présentée par la suite.

1.1 Petit historique du concept de complémentarité

Dans cette section, on commence par évoquer les problèmes d'interprétation qui ont conduit à l'élaboration du principe de complémentarité. Le but n'est bien sûr pas une étude historique exhaustive de la genèse de la mécanique quantique, mais plutôt de mieux saisir la portée de la réponse apportée par Bohr à ces problèmes. Cette section est basée sur plusieurs références dont [97], [14], et les articles originaux de Bohr [12], [13].

1.1.1 Les difficultés de la mécanique des quanta

En 1923 la “théorie quantique” repose sur un certain nombre de règles et de postulats appliqués aux systèmes microscopiques. Le postulat originel a été proposé par Planck en 1900 [76]. Il dit que l'énergie rayonnée par un oscillateur matériel est émise par paquets discontinus d'énergie : $E = h\nu$. Bohr et Sommerfeld en 1913 [11] construisirent un modèle planétaire de l'atome d'hydrogène dans lequel l'électron occupe certaines orbites, celles pour lesquelles le moment cinétique de l'électron est un multiple de $\hbar = \frac{h}{2\pi}$. Une autre règle énoncée par Bohr est le “principe de correspondance”, qui postule que, dans la limite des grands nombres quantiques, l'on doit retrouver les résultats de la théorie classique. Tous ces postulats tirent leur légitimité du fait qu'ils expliquent avec une grande précision un grand nombre de faits expérimentaux restés mystérieux auparavant. Le postulat de Planck permet le calcul du spectre de rayonnement du corps noir à l'équilibre thermique [76], et permet aussi l'élaboration d'une théorie de l'effet photoélectrique [31]. La théorie de Bohr-Sommerfeld donne un début d'explication microscopique de la table périodique des éléments.

Tous ces succès rendaient de plus en plus claire la nécessité d'un saut conceptuel de grande envergure par rapport à la mécanique classique pour décrire les phénomènes microscopiques. Il fallait inventer une théorie qui puisse donner un cadre général aux règles un peu disparates inventées par Planck et Bohr, et qui puisse établir un lien clair entre la physique quantique et la physique classique, de manière plus précise que le principe de correspondance. Notamment, il fallait réconcilier la continuité inhérente à la mécanique classique avec la discontinuité des "sauts quantiques" d'une orbite à l'autre, tels qu'ils sont décrits dans le modèle de Bohr de l'atome d'hydrogène.

Le formalisme mathématique complet de la mécanique quantique telle qu'elle est enseignée aujourd'hui encore fut inventé dans les années 1925-26. Passons rapidement en revue les principaux résultats de cette époque. Heisenberg inventa la mécanique des matrices en 1925 [41] ; Pauli introduisit son principe d'exclusion en 1925 ; la découverte du spin date de cette même époque. Enfin Schrödinger en 1926, approfondissant l'hypothèse de de Broglie sur le comportement ondulatoire des corpuscules, proposa une équation pour l'onde de matière d'une particule de masse m dans un potentiel $V(\mathbf{r}, t)$:

$$i\hbar \frac{\partial \psi(\mathbf{r}, t)}{\partial t} = \left(-\frac{\hbar^2}{2m} \nabla^2 + V(\mathbf{r}, t) \right) \psi(\mathbf{r}, t) \quad (1.1)$$

Schrödinger montra par la suite l'équivalence mathématique de la mécanique des matrices de Heisenberg avec sa mécanique ondulatoire, renforçant par là-même la crédibilité de chacune des deux approches. Mais si le formalisme quantique est parvenu à maturité dès 1926, subsistait le problème fondamental de l'interprétation de ce formalisme. Soulignons tout de suite que par "interprétation" il n'est question que d'interprétation physique et non pas d'interprétation philosophique. La question est : quelle est la relation entre le formalisme et la réalité physique ? quel sens exact faut-il donner aux diverses quantités symboliques introduites tant dans la théorie de Heisenberg que celle de Schrödinger ?

1.1.2 L'interprétation de Schrödinger

Dans la mécanique ondulatoire de Schrödinger, les conditions de quantification découlent de conditions de résonance pour l'onde de matière ; on peut retrouver par un classique calcul d'équation aux dérivées partielles l'énergie d'un oscillateur harmonique quantique (postulat de Planck) et de l'électron dans l'atome d'hydrogène. Ces succès firent croire à Schrödinger que la mécanique ondulatoire était une extension de la mécanique classique, de la même manière que l'optique ondulatoire est une extension de l'optique géométrique. Il essaya donc de décrire tous les phénomènes quantiques comme des phénomènes purement ondulatoires. Il interpréta la fonction d'onde $\psi(\mathbf{r}, t)$ comme une sorte de vibration électromagnétique ou acoustique représentant directement la densité électronique. Il supposa que les corpuscules, supposés ponctuels, étaient en fait des paquets d'onde très piqués spatialement. Il essaya de décrire les transitions radiatives d'un atome entre deux états comme des phénomènes de battement entre modes

propres.

L'avantage de cette interprétation était de se débarrasser des mystérieuses discontinuités, les “sauts quantiques”, introduits par Bohr, si radicalement éloignés de toute intuition classique. Mais un certain nombre de difficultés furent très vite soulignées par Heisenberg et Bohr. Citons-en deux. Tout d'abord, si l'on applique l'équation 1.1 à un paquet d'onde, il est bien connu que celui-ci se disperse en un temps très bref (sauf dans le cas très particulier où le potentiel est harmonique). Comment, dans ce cadre, expliquer qu'il soit possible de visualiser la trajectoire d'un électron sur un écran, alors que le paquet d'ondes imaginé par Schrödinger aurait tendance à se délocaliser ? D'autre part Bohr montra que, sans les “sauts quantiques”, il devenait impossible de retrouver la loi de Planck du rayonnement du corps noir.

1.1.3 L'interprétation de Heisenberg et le principe d'incertitude

Heisenberg refusait l'interprétation de Schrödinger et essaya de rendre compte autrement du comportement ondulatoire des corpuscules. Pour bien saisir la difficulté du problème, citons Heisenberg lui-même : “Neither of [Bohr nor I] could tell how so simple a phenomenon as the trajectory of an electron in a cloud chamber could be reconciled with the mathematical formulations of quantum or wave mechanics” []. Autrement dit : un électron a bien la trajectoire qu'une particule de masse m aurait en mécanique classique puisqu'on peut la voir directement dans une chambre à bulles ; or ce simple fait semble incompatible avec les équations de la mécanique quantique !

Heisenberg se persuada qu'il s'agissait bien d'un problème d'*interprétation* : il était convaincu que le formalisme avait déjà une forme définitive, mais comment faire le lien avec la réalité ? Pour surmonter cette difficulté, Heisenberg se demanda ce que l'on observait réellement dans une chambre à bulles. La “trajectoire” visualisée n'est après tout rien d'autre qu'une succession discontinue de mesures de position et d'impulsion faites **avec une certaine imprécision** Δq et Δp . Rien dans le formalisme quantique n'interdit cela, pourvu que Δq et Δp vérifient l'inégalité :

$$\Delta p \Delta q \geq \hbar \quad (1.2)$$

Cette inégalité étant une conséquence directe de la non-commutation des deux observables p et q [42].

Heisenberg chercha par la suite, dans des expériences de pensée, des situations où il serait possible de mesurer à la fois la position et l'impulsion avec une précision illimitée. La plus célèbre de ces Gedankenexperiment est celle du “microscope d'Heisenberg”. Il essaya d'imaginer ce qui se passerait si l'on essayait de regarder un électron avec un microscope suffisamment puissant pour résoudre les dimensions atomiques. Pour cela il faut utiliser des rayons gamma ; mais alors le recul de l'électron lors de la diffusion d'un seul photon γ fait que l'électron est éjecté de son orbite. On ne peut donc pas observer la *trajectoire* de l'électron sur son orbite autour du noyau puisqu'on ne peut voir qu'un

seul point de cette orbite. On peut par contre avoir une distribution statistique de la position de l'électron sur son orbite, en refaisant la même expérience avec une collection d'atomes dans le même état. Cette distribution est donnée par le module au carré de la fonction d'onde ψ selon le précepte énoncé par Born.

Notons que la relation d'incertitude n'est pas propre au couple "position-impulsion", tout couple de variables dont les observables ne commutent pas vérifient une telle relation d'incertitude ; on parle alors de variables conjuguées.

Heisenberg fut persuadé que le principe d'incertitude, qui était inscrit au coeur du formalisme, fournissait la clé des difficultés d'interprétation de la mécanique quantique. Résumons-le une dernière fois en disant qu'il est impossible de connaître avec une infinie précision la valeur de deux variables conjuguées A et B, et il est donc absurde d'attribuer à un système, à un instant donné, des valeurs bien définies pour A et B. Notamment la notion de "trajectoire" n'est qu'une approximation, qui est justifiée dans le cas d'une chambre à bulles mais n'a aucun sens pour une orbitale atomique.

1.1.4 Le principe de complémentarité

1.1.4.a Enoncé

La notion de complémentarité a été inventée par Bohr en 1927, au moment même où Heisenberg développait ses idées sur le principe d'incertitude. Bohr l'a présentée à une conférence à Côme, puis dans un article dans *Nature* en 1928 [12]. Il tente d'y répondre aux mêmes problèmes fondamentaux d'interprétation du formalisme quantique que Heisenberg, et sa réponse est proche de la sienne. Néanmoins le principe de complémentarité a un statut différent des relations d'incertitude, pour plusieurs raisons. D'une part il n'existe pas un énoncé précis, défini, et encore moins mathématique du principe de complémentarité. D'autre part son contenu, quoique plus vague que le principe d'incertitude, a une portée plus générale. On peut dire que Bohr, avec sa notion de complémentarité, propose **un cadre conceptuel et philosophique cohérent** au formalisme quantique. Soulignons que le texte le plus clair de Bohr sur la complémentarité et le plus profond est à notre avis son article de 1935 en réponse à Einstein, Podolsky et Rosen [13]. C'est sur ce texte que nous nous appuyons essentiellement pour le raisonnement qui suit.

1) Bohr remarque que toute opération de "mesure", d'"observation", est quelque chose de classique en ce sens que, lorsqu'on parle de mesure, on présuppose toujours une dissymétrie fondamentale entre l'objet observé (l'atome, l'électron, ...) d'une part et l'appareil de mesure d'autre part (photodiodes, oscilloscopes,...), qui est censé perturber *aussi peu que possible* ce dernier. Pour rendre compte de toute expérience, même quand le système étudié est microscopique, le scientifique est amené à dire : "Utilisant tel *appareillage* M, j'ai observé telle *propriété* P de tel *système* S". Il s'agit d'une contrainte qui est presque d'ordre linguistique.

2) or lorsque le système S étudié est microscopique, toute interaction avec un appa-

reil de mesure macroscopique M le perturbera notablement. Et à cause de l'existence du quantum d'action $\hbar$, il est impossible, *même en principe*, de concevoir un appareil de mesure qui perturberait aussi peu que possible le système S . En effet, si l'on essayait de réduire autant que possible les dimensions de M pour perturber le moins possible S , on arriverait à la situation où l'appareil de mesure M doit lui aussi être décrit par la mécanique quantique. A ce moment-là, la distinction entre le système S et l'appareil de mesure M n'a plus de sens. On ne peut plus savoir "qui mesure qui", l'état de S et de M est inséparable.

Par conséquent il faut bien reconnaître que, à cause de l'existence du quantum d'action, la *propriété* P évoquée ci-dessus n'appartient plus intrinsèquement au système S , mais dépend de M . C'est l'énoncé le plus général du principe de complémentarité. "Evidence obtained under different experimental conditions cannot be comprehended within a single picture, but must be regarded as *complementary* in the sense that only the totality of the phenomena exhausts the possible information about the objects" ("Les résultats obtenus dans différentes conditions expérimentales ne peuvent pas être rassemblés en une seule image, mais doivent être vues comme *complémentaires* dans le sens que seule la totalité des phénomènes observés contient toutes les propriétés d'une objet") ([97] p.18).

1.1.4.b La dualité onde/corpuscule

On voit bien en quoi le raisonnement précédent donne une manière de comprendre les paradoxes concernant la nature des systèmes microscopiques. L'électron par exemple est-il une onde (comme le supposait Schrödinger) ou une particule ? La réponse de Bohr commence par une critique presque sémantique de la proposition précédente. Pour lui, dire que l'électron *est* une onde (ou une particule) présuppose que l'on serait capable, à l'aide d'un dispositif expérimental adéquat, d'observer une propriété *intrinsèque* de l'électron. Or le principe de complémentarité nous l'interdit précisément. Il faut admettre que, dans différents dispositifs expérimentaux, un objet microscopique peut se comporter soit comme une onde, soit comme un corpuscule, *sans que l'on puisse réduire l'un à l'autre*, car chaque dispositif le perturbe de manière différente. C'est la fameuse "dualité onde/corpuscule".

L'expérimentateur peut *choisir* librement quel aspect il veut mettre en évidence, en utilisant le dispositif expérimental adéquat. Le choix peut même être effectué après que la particule ait été détectée, puisque le temps n'intervient à aucun endroit dans le raisonnement de Bohr [96].

Par conséquent, Bohr prédit que toute tentative pour observer à la fois un comportement ondulatoire et corpusculaire est vouée à l'échec, car les dispositifs expérimentaux nécessaires sont mutuellement exclusifs. Il est impossible d'observer des franges d'interférence *et* de savoir par où la "particule" est passée (information de type "which-path" que nous noterons WP dans la suite).

Il est souvent possible de trouver l'origine de la disparition des franges d'interfé-

rence dans une perturbation de type mécanique induite par l'action du détecteur WP sur la particule. La prédiction du principe de complémentarité se confond alors avec les conséquences des relations d'incertitude de Heisenberg pour le couple des observables conjuguées "position-impulsion" (on les qualifie parfois également de "variables complémentaires"). Mais ce n'est pas toujours le cas. On pourrait imaginer des détecteurs WP qui n'aient aucune action mécanique sur la particule (nous en donnerons des exemples dans la suite). Le principe de complémentarité prédit que, même dans ce cas, les franges d'interférence disparaîtront, bien qu'aucune relation d'incertitude ne s'applique. En ce sens, on peut considérer que le principe de complémentarité a une portée plus générale que les relations d'incertitude de Heisenberg.

1.1.4.c La portée du principe de complémentarité

Les notions introduites par Heisenberg et Bohr d'incertitude minimale sur la connaissance d'un système physique et de complémentarité ont une grande importance pour la physique. Elles ont eu aussi des répercussions en épistémologie. Les relations d'incertitude de Heisenberg ne sont rien d'autre qu'une conséquence mathématique interne du formalisme quantique. Mais leur statut est fondamental car elles sont l'un des garants de la cohérence de ce formalisme. Dès lors que quelqu'un aura trouvé une seule situation où ce principe ne s'applique pas, toute la mécanique quantique s'effondrera. Einstein avait bien compris cet état de fait, qui s'efforça pendant de longues années de découvrir une telle situation. Pour l'instant bien sûr, 80 ans plus tard, malgré toutes les tentatives expérimentales, ce principe tient toujours.

Dans ses écrits sur la complémentarité, Bohr nous explique en des termes simples, sans une seule formule, en quoi réside la spécificité ultime des phénomènes quantiques. Pour Bohr, si la mécanique quantique nous paraît contre-intuitive, c'est à cause de l'utilisation irraisonnée de termes du langage courant tel que "l'électron est ceci", ou "nous mesurons telle propriété". Avant d'énoncer ce genre de propositions, il faudrait effectuer une espèce de "critique de l'observation pure", à laquelle Bohr se livre dans ses écrits sur la notion de complémentarité. C'est uniquement après avoir analysé quelles étaient vraiment nos possibilités de définition et de mesure que nous pouvons réutiliser le langage courant, en prenant conscience de toute son ambiguïté, pour tenter de décrire la réalité. En fait, seul le langage formel de la mécanique quantique fournit une description physique non-ambigüe.

L'originalité de la position de Bohr est bien dans cette irruption de la philosophie dans la physique et réciproquement, avec des conséquences profondes pour l'une comme pour l'autre. Il est probable que l'intuition physique de Bohr a été guidée par des convictions philosophiques kantienne. (Cette démarche a été moins féconde lorsque Bohr a tenté d'exporter le principe de complémentarité dans des domaines comme la biologie et les sciences humaines). Le grand nombre d'expériences qui se sont attachées à tester ou illustrer les principes de Heisenberg et de Bohr, et que je vais présenter par la suite, manifestent bien cette importance cruciale qu'ils ont tous deux dans la

physique moderne.

1.1.5 Complémentarité et intrication

Une leçon essentielle que l'on peut tirer du principe de complémentarité est que l'on doit nécessairement tenir compte du dispositif expérimental utilisé pour décrire un système microscopique. Pour se conformer à cette prescription, il faut considérer que le système physique étudié n'est plus seulement la particule microscopique P, mais l'ensemble de la particule et de l'appareillage expérimental "P+M".

Dans le cas particulier où cet appareillage est dans un état quantique bien défini, le système joint "P+M" évolue selon l'équation de Schrödinger vers un état qui fait en général apparaître des corrélations très fortes entre l'état de P et celui de M. On dit que l'état final est *intriqué*. Nous allons voir comment la prise en compte de l'intrication entre P et M permet de retrouver la complémentarité entre les comportements ondulatoire et corpusculaire.

Le fonctionnement d'un interféromètre classique est décrit schématiquement à la figure 1.1 a). La particule P est initialement dans un état $|\Psi_{in}\rangle$. Son évolution future est gouvernée par l'opérateur d'évolution $\hat{U}(t, 0)$ spécifique à l'interféromètre considéré. A l'instant final t_{fin} , l'état de la particule sera $|\Psi_{fin}\rangle = \hat{U}(t_{fin}, 0)|\Psi_{in}\rangle$. Mais l'on peut également considérer l'état de la particule à un instant intermédiaire t_1 , que l'on peut en toute généralité décomposer sur une base d'états orthogonaux $|a_i\rangle$:

$$|\Psi_{in}\rangle \longrightarrow |\Psi(t_1)\rangle = \hat{U}(t_1, 0)|\Psi_{in}\rangle = \sum_i |a_i\rangle \langle a_i| \hat{U}(t_1, 0)|\Psi_{in}\rangle \quad (1.3)$$

La probabilité P_f que P finisse dans l'état $|\Psi_0\rangle$ à un instant t_{fin} est donnée par le module au carré du produit scalaire de $|\Psi_0\rangle$ avec l'état final $|\Psi_{fin}\rangle$ que l'on peut écrire en fonction de l'état de la particule à l'instant t_1 $|\Psi_{fin}\rangle = \hat{U}(t, t_1)|\Psi_{t_1}\rangle$:

$$P_f = |\langle \Psi_0 | \Psi_{fin} \rangle|^2 = \left| \sum_i \langle \Psi_0 | \hat{U}(t, t_1) | a_i \rangle \langle a_i | \hat{U}(t_1, 0) | \Psi_{in} \rangle \right|^2 \quad (1.4)$$

Les interférences sont dûes au développement du module carré, dans cette expression, qui donne des termes "rectangles" $i \neq j$ non nuls.

Imaginons maintenant que l'on cherche à mesurer l'état de la particule à l'instant t_1 en la couplant à un détecteur WP initialement dans un état quantique bien défini $|W_0\rangle$ (figure 1.1 b)). Sans connaître le détail microscopique de l'interaction entre P et le détecteur, on peut dire en toute généralité que chaque état $|a_i\rangle$ sera couplé à un état $|W_i\rangle$ du détecteur. L'état joint du système "P+M" à l'instant t_1 s'écrit alors

$$\Psi^{(P+M)}(t_1) = \sum_i |W_i\rangle |a_i\rangle \langle a_i | \hat{U}(t_1, 0) | \Psi_{in} \rangle \quad (1.5)$$

Il s'agit en général d'un état intriqué. L'état de P pris individuellement n'est alors plus un état pur ; il est décrit par sa matrice densité obtenue en calculant la trace partielle

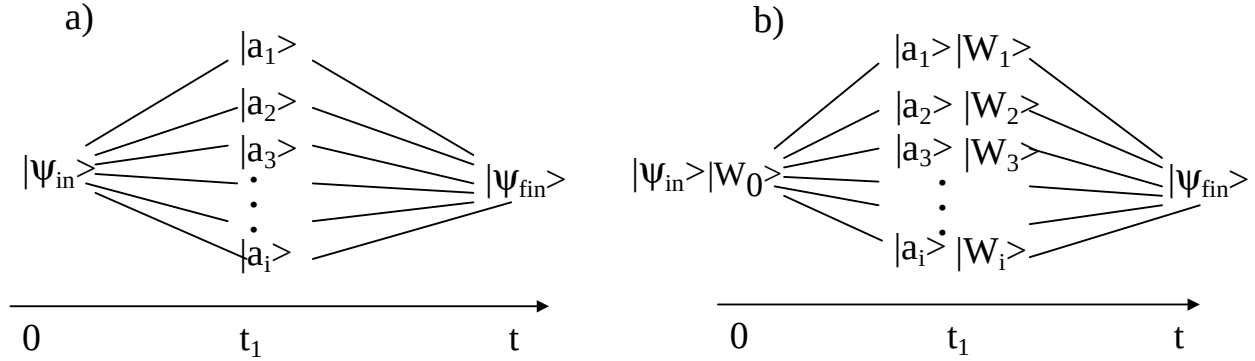


FIG. 1.1 – (a) Schéma de principe d'un interféromètre classique. La particule est initialement dans un état $|\Psi_{in}\rangle$, elle finit dans un état $|\Psi_f\rangle$ en passant par les états intermédiaires $|a_i\rangle$ à l'instant t_1 . (b) Si l'on ajoute un détecteur WP, l'état du détecteur $|W_i\rangle$ est corrélé à l'état de la particule $|a_i\rangle$

de la matrice densité du système global par rapport aux variables de M. On trouve

$$\rho_P = \sum_{i,j} |a_i\rangle\langle a_i| \hat{U}(t_1, 0) |\Psi_{in}\rangle\langle W_j| W_i\rangle \langle \Psi_{in}| \hat{U}^\dagger(t_1, 0) |a_j\rangle\langle a_j| \quad (1.6)$$

La visibilité des franges d'interférence est directement proportionnelle aux termes non-diagonaux de ρ . Or on voit intervenir dans l'expression de $\rho_{i,j}$, pour $i \neq j$, des termes du type $\langle W_j| W_i\rangle$ en facteur des termes déjà présents dans l'expression 1.4, dont l'importance dépend du fonctionnement du détecteur WP :

- Dès que le détecteur WP est capable de distinguer si la particule était dans l'état $|a_i\rangle$ ou $|a_j\rangle$ à l'instant t_1 , les états $|W_i\rangle$ et $|W_j\rangle$ ont nécessairement un recouvrement nul et leur produit scalaire vaut donc zéro. Dans ce cas, il est impossible d'observer de franges d'interférence, la particule P a un comportement du type corpusculaire.

- au contraire, lorsque les états $|W_i\rangle$ et $|W_j\rangle$ sont indiscernables, leur produit scalaire vaut 1 et l'on retrouve le comportement ondulatoire de l'équation 1.4.

Ainsi, en considérant uniquement les corrélations quantiques qui apparaissent ou non entre P et M, il a été possible de montrer en toute généralité que l'observation de franges d'interférence et la mesure du chemin suivi par P dans l'interféromètre nécessitaient des dispositifs expérimentaux contradictoires. Soulignons que nous n'avons pas utilisé de relation d'incertitude pour parvenir à ce résultat.

Jusqu'ici nous n'avons considéré que les deux cas limites : 1) les franges ont un contraste maximal (et il est alors impossible de distinguer les chemins suivis par la particule dans l'interféromètre) 2) les franges ont un contraste nul, et l'on peut alors avoir une information WP certaine. Il est possible de quantifier précisément la relation entre contraste ou visibilité des franges (V) et "distinguaibilité" du chemin suivi (D) dans le régime intermédiaire où les franges d'interférence ont un contraste limité. Pour cela, nous suivons un raisonnement dû à Englert [34], que nous présentons brièvement.

Englert considère la situation où une particule est envoyée dans un interféromètre à deux bras symétriques que l'on note $+$ et $-$. Un détecteur WP est placé dans l'interféromètre, sur lequel on peut mesurer une observable A après passage de la particule. La quantité D que nous allons définir précisément mesure la quantité d'information "stockée" dans le détecteur WP après passage de la particule. Elle ne dépend que de l'interféromètre étudié et du détecteur WP considéré.

On commence par définir la vraisemblance pour deviner le bon chemin pour une observable A donnée, dont les valeurs propres sont notées A_i , à partir des probabilités jointes de détecter la particule dans la voie $+$ ou dans la voie $-$ et de mesurer la valeur propre A_i :

$$L_A = \sum_i \max\{p(A_i, +), p(A_i, -)\} \quad (1.7)$$

Cette quantité vaut entre $1/2$ (aucune information sur le chemin suivi) et 1 (on connaît avec certitude le chemin suivi). Mais elle dépend de l'observable A que l'on a choisi de mesurer : même dans le cas où le détecteur WP a parfaitement enregistré le chemin suivi, un choix malheureux de l'observable A pourrait faire en sorte que l'on ne puisse pas la lire. Pour que la "distinguableté" ne dépende que du détecteur WP utilisé et plus de A , on introduit la quantité

$$D = -1 + 2\max_A(L_A) \quad (1.8)$$

qui vaut entre 0 (chemins parfaitement indistingables) et 1 (information complète sur le chemin suivi), tout comme la visibilité des franges. On peut alors montrer [34] que

$$V^2 + D^2 \leq 1 \quad (1.9)$$

Les quantités V et D peuvent être mesurées expérimentalement de manière indépendante, ce qui donne tout son intérêt à cette relation. Soulignons que pour l'établir, aucune supposition n'a été faite sur l'état initial de P ou du détecteur WP. En particulier, 1.9 reste valable même lorsque P ou le détecteur sont initialement dans des mélanges statistiques d'états (auquel cas on a $V^2 + D^2 < 1$). Lorsque tous les systèmes sont initialement dans des états purs, $V^2 + D^2 = 1$.

1.2 Illustrations expérimentales du principe de complémentarité

Nous allons maintenant procéder à une revue de différentes expériences qui ont été réalisées pour illustrer ou tester la complémentarité. Précisons tout de suite que cette revue ne prétend pas être exhaustive. En effet le nombre d'expériences réalisées dans ce domaine est très important. Nous choisirons certaines expériences qui nous semblent représentatives, et nous les discuterons plus précisément, puis nous insisterons sur les nouveautés apportées par notre expérience. Mais avant cela nous présenterons

les expériences de pensée proposées par Bohr et Einstein, qui ont servi de guide à toutes ces réalisations expérimentales ainsi qu'à la nôtre.

1.2.1 Le paradigme du principe de complémentarité : les fentes d'Young

La complémentarité entre comportement ondulatoire et corpusculaire est manifeste dans le dispositif des fentes d'Young, tel qu'il est représenté sur la figure 1.2. Feynman soulignait que ce dispositif contenait "le seul mystère de la mécanique quantique" [37]. Des particules P (photons, électrons, atomes ...) sont diffractées par un écran percé de deux diaphragmes $D1$ et $D2$, et on observe des franges d'interférence sur un écran E . Rappelons que, si la distance entre les diaphragmes et E est notée L , que l'écart entre $D1$ et $D2$ est noté d et la longueur d'onde de la particule λ , l'interfrange observé sera

$$i = \lambda L / d \quad (1.10)$$

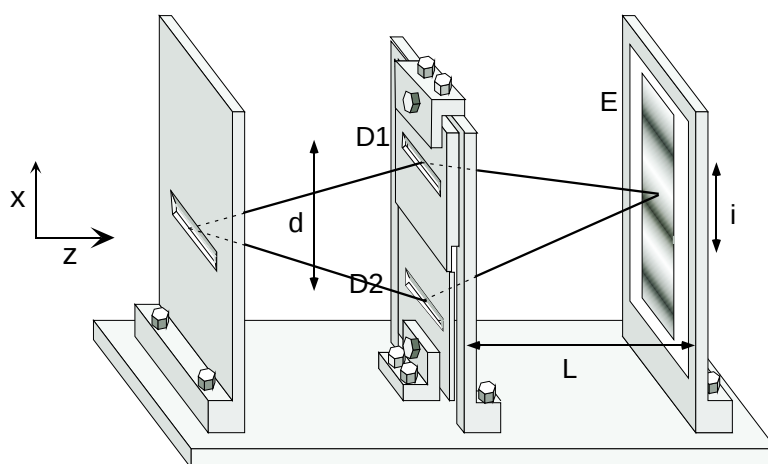


FIG. 1.2 – Schéma d'une expérience de fentes d'Young. La particule P est diffractée par un écran percé de deux diaphragmes $D1$ et $D2$. On observe sur l'écran E la probabilité P d'avoir un impact en fonction de x .

Mais pour peu que l'on tente d'observer par quelle fente est passée P (information "which-path" que l'on notera WP), on perd les franges d'interférence. Plusieurs variantes ont été proposées pour essayer d'obtenir l'information WP de la manière la moins invasive possible pour la particule P . On peut les répartir en deux classes, selon qu'on utilise un dispositif extérieur à l'interféromètre pour acquérir l'information WP , ou bien que l'on cherche à enregistrer cette information dans la structure même de l'interféromètre, sans aucune manipulation externe. (voir la figure 1.3)

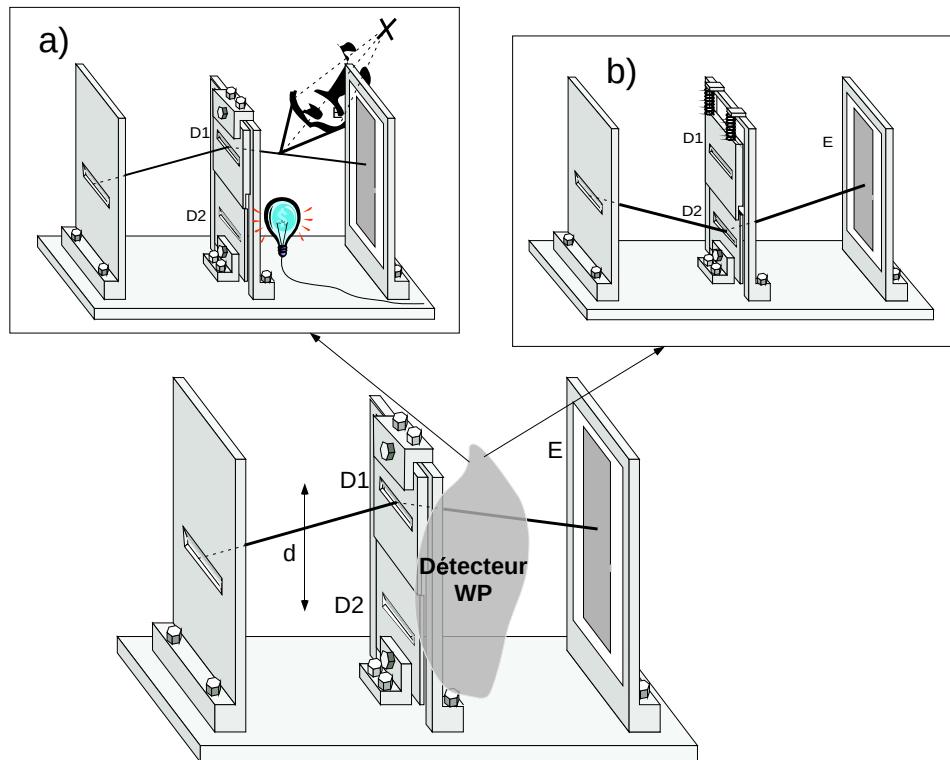


FIG. 1.3 – Deux manières de détruire les franges en observant le chemin par où la particule est passée. (a) Le "microscope de Heisenberg" : on place une source lumineuse derrière les diaphragmes. La particule diffuse un photon, que l'on peut détecter grâce à l'objectif d'un microscope. On peut ainsi distinguer si la particule est passée par $D1$ ou $D2$. (b) "diaphragme mobile" : l'un des deux diaphragmes (ici $D1$) est suspendu au bâti par des ressorts au lieu d'en être directement solidaire comme dans l'interféromètre classique. Le mouvement final de ce diaphragme dépend alors du passage de la particule à-travers lui ou non, ce qui fournit une information WP et brouille les franges.

1.2.1.a Eclairer les fentes F1 et F2

Le prototype du détecteur WP extérieur à l'interféromètre est le “microscope de Heisenberg”. Il consiste à avoir une source lumineuse derrière le diaphragme et à éclairer les deux fentes, comme représenté sur la figure 1.3a). Si les particules P sont des électrons par exemple, ils vont diffuser les photons de longueur d'onde λ' émis par la source lumineuse, et, pour peu que $\lambda' < d$, il devient possible, à l'aide d'un objectif de microscope, de voir par quel diaphragme est passée la particule.

Reprenons l'analyse faite par Bohr de cette situation. Il montre que les franges d'interférence sont inévitablement brouillées à cause de la relation d'incertitude position-impulsion [12]. En effet, si l'on arrive à connaître la position selon x des électrons avec une précision $\delta x < d$, l'imprécision sur la composante de l'impulsion selon x , déterminée par le principe d'incertitude, vaut $\delta p_x = h/\delta x > h/d$. Or cette imprécision sur l'impulsion se traduit par une imprécision sur la position sur l'écran qui vaut

$$\Delta x = \frac{\delta p_x}{p} L = \frac{\delta p_x \lambda L}{h} > \lambda L/d = i \quad (1.11)$$

d'après l'équation 1.10. On conçoit donc bien que les franges disparaissent.

Cette analyse n'est pourtant pas pleinement satisfaisante. D'abord elle ne s'applique qu'au dispositif considéré sur la figure 1.3a). Or on pourrait imaginer d'autres types d'interféromètres où elle ne serait pas valable (nous en verrons un exemple au paragraphe 1.2.2). D'autre part, elle laisse de côté un aspect du problème. Après la diffusion des photon par l'atome, celui-ci est simplement intriqué avec le photon diffusé. Certes, si l'on détecte les photons au microscope de la manière qui est indiquée sur la figure 1.3a), les franges seront brouillées car la détection d'un photon nous donne une information non-ambigüe sur le diaphragme par lequel l'atome est passé, laissant par contre l'impulsion transverse des atomes complètement indéterminée. Mais cette perte de cohérence *n'est pas* due à une perturbation incontrôlée de l'atome par la diffusion du photon. En effet, nous pourrions choisir de détecter les photons, non pas au microscope mais au télescope. Nous retrouverions alors des franges d'interférence en corrélant la détection des atomes à la détection d'un photon au télescope ! Car le télescope mesure la direction d'émission du photon diffusé, et donc l'impulsion transverse de l'atome après passage par les diaphragmes. Il respecte par contre toute l'ambigüité initiale sur le trajet des atomes à-travers les diaphragmes. En d'autres termes, nous avons le *choix* d'observer deux comportements de l'atome complémentaires mais incompatibles, selon le dispositif expérimental choisi. Le brouillage des franges est *réversible*, pour peu que l'on efface de manière cohérente l'information sur la position. On parle de “gomme quantique” [84]. Nous verrons dans la suite une démonstration expérimentale de cette possibilité.

Insistons encore sur le fait que la disparition des franges, due à l'intrication avec tout détecteur WP, et les relations d'incertitude de Heisenberg sont deux choses indépendantes. Il est en effet possible d'imaginer des détecteurs WP plus “raffinés” que

le microscope de Heisenberg, et qui ne tombent pas sous le coup de l'analyse précédente basée sur le principe d'incertitude. Une proposition intéressante de “détecteur WP non-destructif” a été faite par Scully et al dans un article de 1991 [83]. Au lieu de faire interférer des électrons, les auteurs imaginent que l'on puisse faire interférer dans un dispositif de franges d'Young des atomes de Rydberg (même si expérimentalement, cela semble extrêmement difficile). Nous verrons en détail au prochain chapitre les propriétés de ces atomes puisque nous les utilisons dans notre dispositif. Pour l'instant contentons-nous de savoir qu'ils ont des transitions dans le domaine micro-onde. Scully imagine qu'on a pu placer, derrière chaque fente D1 et D2, deux cavité micro-onde C1 et C2 accordées sur une transition atomique. Sous certaines conditions, on peut forcer l'atome à émettre un photon dans la cavité; ainsi la présence du photon dans une cavité ou l'autre indique par quel chemin est passé l'atome. Notons que l'atome ne subit aucun transfert d'impulsion lors de l'interaction avec les cavités C1 et C2; on ne peut donc pas accuser cette interaction de brouiller les franges par une action mécanique sur l'état spatial de l'atome [35, 87]. Le détecteur WP formé des deux cavités permet effectivement de savoir par quel chemin est passé l'atome sans le perturber mécaniquement.

Pourtant les franges sont bel et bien supprimées, car l'atome s'intrique avec les cavités C1 et C2. En résumé, même une mesure non-destructive du chemin suivi par l'atome dans l'interféromètre suffit à brouiller les franges.

1.2.1.b Le “diaphragme mobile”

Einstein avait imaginé un autre dispositif permettant d'enregistrer l'information WP dans l'interféromètre, représenté schématiquement sur la figure 1.3 dans le cadre b. Imaginons que le diaphragme D1 soit mobile, et qu'il soit suffisamment léger et bien isolé de l'environnement pour que l'on puisse mesurer son impulsion avant et après le passage de P, et ainsi connaître l'impulsion communiquée au diaphragme par la particule. On pourrait peut-être, alors, dans la même expérience, observer des franges d'interférence (fonctionnement normal de l'interféromètre), et faire une mesure précise de l'impulsion du diaphragme qui indiquerait le chemin suivi par l'atome. Analysons en détail ce dispositif.

Les ondes partielles venant des deux fentes D1 et D2 ont des vecteurs d'onde différents k_1 et k_2 dont la différence a une composante selon x qui vaut

$$\Delta k_x \sim 2\pi/\lambda \quad (1.12)$$

Pour pouvoir déterminer si le photon est passé par D1 ou D2, il faut donc être capable de mesurer l'impulsion de D1 avant et après le passage de P, avec une précision $\Delta p_x = \hbar \Delta k_x \sim \hbar/\lambda$. Or l'incertitude minimale δp_x sur l'impulsion de D1 est reliée par le principe d'incertitude à l'imprécision minimale sur sa position par rapport au bâti δx : $\delta p_x = \hbar/\delta x$. Donc, si l'on veut $\Delta p_x \gg \delta p_x$, alors nécessairement aussi $\Delta k_x \gg \delta k_x = \hbar/\delta x$ ce qui d'après 1.12 implique que $\delta x \gg \lambda$. Si l'incertitude sur la

position du diaphragme D1 est plus grande que l'interfrange, les franges d'interférence sont brouillées comme le prédit le principe de complémentarité.

Encore une fois, cette analyse, quoique juste, n'est pas aussi profonde que celle que donne le principe de complémentarité. Bohr l'a faite explicitement dans son article de 1935 en réponse à Einstein, Podolsky et Rosen [13]. Si l'on considère l'ensemble "P+D1" comme un seul système à observer avec des appareils macroscopiques M, on a le *choix* d'observer soit un comportement ondulatoire (pour cela il suffit de mesurer finalement la position du diaphragme), soit un comportement corpusculaire (mesure de l'impulsion du diaphragme).

Voyons maintenant comment ces expériences de pensée ont été réalisées avec des dispositifs expérimentaux. L'extrême sensibilité exigée par les expériences que je viens de présenter fait que l'optique quantique a toujours été un domaine privilégié pour les tests du principe de complémentarité.

1.2.2 Les "microscopes de Heisenberg"

De nombreux groupes ont déjà réalisé des dispositifs ressemblant au microscope de Heisenberg présenté au paragraphe 1.2.1.a. Nous allons présenter en détail deux de ces expériences qui nous semblent représentatives du domaine, puis nous passerons rapidement en revue les autres systèmes qui ont été utilisés.

L'expérience du MIT Une réalisation du microscope de Heisenberg a été effectuée dans le groupe de D. Pritchard à MIT [24]. Le dispositif expérimental est schématisé sur la figure 1.4. Le coeur de ce dispositif est un interféromètre atomique constitué de 3 réseaux d'amplitude matériels obtenus par nanofabrication. Un jet d'atomes de sodium collimaté traverse l'interféromètre. On éclaire les atomes à la sortie du réseau R1 avec un laser à résonance avec la transition $3S - 3P$ du sodium. Les atomes, à résonance, diffusent un photon laser (le temps d'interaction avec le faisceau est choisi de manière à ce que la probabilité de diffuser deux photons soit faible). On peut faire varier l'endroit du dispositif où a lieu la diffusion lumineuse, et ainsi choisir la séparation d entre les deux bras de l'interféromètre.

Les résultats obtenus sont montrés à la figure 1.5. La partie supérieure de la figure 1.5 a) montre le rapport entre le contraste des franges d'interférence observées en fonction de la séparation d , et le contraste mesuré lorsque le laser est éteint. Lorsque $d \ll \lambda$, ce rapport vaut 1 : le contraste des franges d'interférence est le même que quand le laser L est éteint. Pourtant l'effet mécanique du recul de l'atome au moment de l'absorption ou de l'émission d'un photon correspondrait à un déplacement Δx d'une centaine d'interfranges ; comme ce recul s'effectue dans une direction qui est tout-à-fait aléatoire on pourrait s'attendre à ce que toute interférence doive être supprimée, en faisant une analyse du type de celle présentée dans le paragraphe 1.2.1.a. Ce n'est pas le cas, grâce au type d'interféromètre à trois réseaux utilisé par les physiciens du MIT, qui n'est sensible ni à une dispersion d'angle d'incidence ni à une dispersion de vitesse

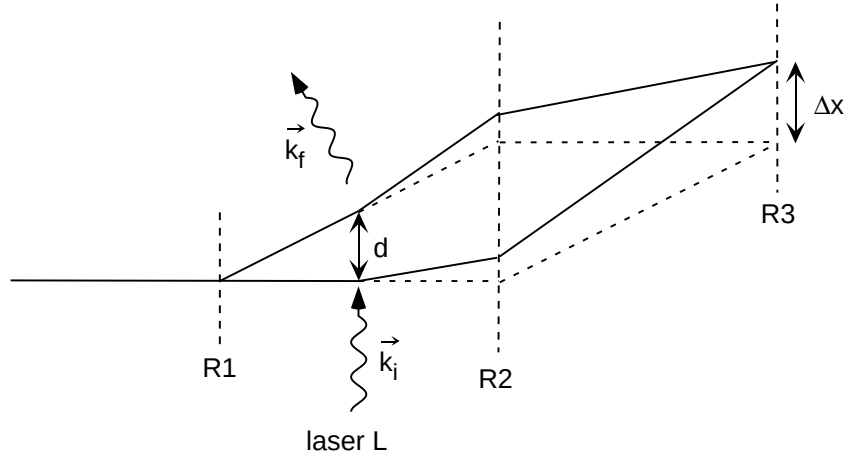


FIG. 1.4 – Schéma de l'expérience réalisée au MIT. Les trois réseaux R1, R2 et R3 sont des réseaux matériels. Les atomes diffusent les photons issus du laser L et leur point d'impact sur le réseau R3 est translaté d'une quantité Δx

des particules. Ainsi ils démontrent que les effets purement mécaniques causés par la diffusion n'affectent pas le fonctionnement de l'interféromètre. Dans ce cas $d \ll \lambda$, on ne serait pas capable, *même en principe*, d'observer par quel chemin l'atome est passé puisque la résolution spatiale donnée par l'observation du photon serait limitée par λ . Le principe de complémentarité n'interdit donc pas l'observation de franges d'interférence.

Par contre, lorsque $d \approx \lambda$, le contraste des franges diminue fortement. En effet, dans ce régime, on serait capable, en ajoutant juste un microscope comme représenté sur la figure 1.3, de voir par quel chemin l'atome est passé ; or cela contredirait le principe de complémentarité.

Cette expérience est donc une très belle réalisation de l'expérience de pensée décrite dans le paragraphe 1.2.1.a. Insistons encore sur le fait que la perte de contraste observée n'est pas le résultat d'un effet purement mécanique de dispersion de l'impulsion au moment de la diffusion, mais bien d'une dispersion de la phase relative entre les amplitudes associées à chaque bras de l'interféromètre. Les résultats obtenus montrent par ailleurs qu'il n'est bien sûr pas nécessaire de réellement *observer* l'atome avec un vrai microscope ayant la bonne ouverture angulaire, qui serait de toutes façons difficile à installer. La seule intrication des atomes avec les photons diffusés suffit à induire un comportement corpusculaire.

Les physiciens du MIT ont enfin ajouté un diaphragme très étroit devant leur détecteur, de manière à ne détecter que des atomes ayant une impulsion transverse donnée. Cela revient à corréler la détection des atomes avec la diffusion du photon laser dans un certain angle solide, déterminé par la position et la forme précise du diaphragme (ce qui équivaut à une détection des photons au télescope, que nous avons analysée au paragraphe 1.2.1.a).

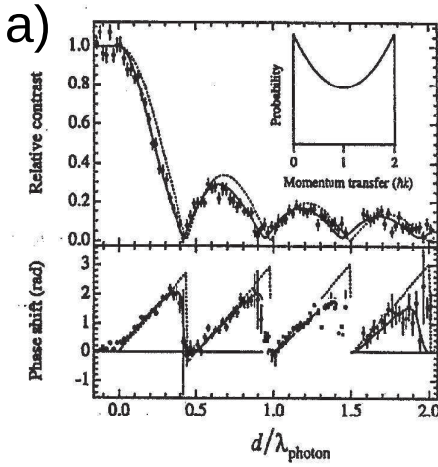


FIG. 2. Relative contrast $C(\text{laser on})/C(\text{laser off})$ and phase shift of the interferometer as a function of d , where $d = z\lambda_{\text{dB}}/\lambda_g$ and $d < 0$ indicates scattering before the first grating. The dashed curve corresponds to purely single photon scattering, and the solid curve is a best fit that includes contributions from atoms that scattered 0 photons (5%) and 2 photons (18%). The inset shows $P(\Delta k_x)$, the distribution of the net transverse momentum transfer.

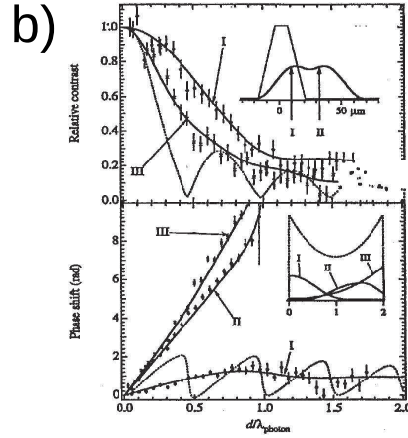


FIG. 3. Relative contrast and phase shift of the interferometer as a function of d for the cases in which atoms are correlated with photons scattered into a limited range of directions. The solid curves are calculated using the known collimator geometry, beam velocity, and momentum recoil distribution and are compared with the uncorrelated case (dashed curves). We do not attribute statistical significance to small deviations between these data and fits. The upper inset shows atomic beam profiles at the third grating when the laser is off (thin line) and when the laser is on (thick line). The arrows indicate the third grating positions for cases I and II. The lower inset shows the acceptance of the detector for each case, $P'_i(\Delta k_x)$, compared to the original distribution (dotted line).

FIG. 1.5 – Résultats de l'expérience de Chapman et al.

Le contraste des franges est indiqué à la figure 1.5 b) (partie supérieure), normalisé au contraste maximal observé (qui n'est cependant pas le même que lorsque le laser n'est pas allumé). On constate que le contraste diminue moins rapidement lorsqu'on corréle la détection des atomes avec l'émission des photons dans une certaine direction. En effet, sélectionner certaines directions d'émission est équivalent à détecter les photons dans la base des ondes planes, et donc à remélanger les deux ondes sphériques émises par l'atome lors de la diffusion, au lieu de les observer dans la base des ondes sphériques, c'est-à-dire au microscope. La courbe 1.5 b) est donc une réalisation expérimentale de l'idée de "gomme quantique" discutée plus haut.

Signalons qu'une expérience similaire a été réalisée dans le contexte pourtant très différent de la physique mésoscopique par Buks et al. [22]. Ils ont réalisé un interféromètre électronique du type Aharonov-Bohm dans un gaz d'électrons 2D à haute mobilité. Sur l'une des deux voies de l'interféromètre, les physiciens israéliens ont placé un "détecteur" dont l'efficacité est contrôlée par une électrode. Ils ont observé une réduction du contraste des franges lorsque le détecteur était capable de distinguer le passage de l'électron par le bras de l'interféromètre dans lequel il était. Il s'agit de la première illustration du principe de complémentarité en physique du solide, et même devrait-on dire en-dehors du champ de la physique atomique et de l'optique quantique.

Les photons jumeaux D'autres expériences utilisent un degré de liberté interne de la particule P pour enregistrer l'information WP. Nous allons discuter plus en détail

une de ces expériences, qui met en jeu des “photons jumeaux”, c’est-à-dire des paires de photons émis par conversion paramétrique, réalisée à Innsbruck dans le groupe d’Anton Zeilinger [43].

Le système étudié est décrit sur la figure 1.6. Une impulsion laser traverse un cristal non-linéaire ($LiIO_3$) puis est rétro-réfléchiée dans le cristal. A chaque passage, le laser a pu créer une paire de photons “idler” et “signal” (conversion paramétrique). Appelons D (pour Direct) les photons créés au premier passage et R (pour Réfléchi) ceux créés au deuxième passage des photons. Comme on le voit sur la figure 1.6, les photons D sont eux aussi rétro-réfléchis dans le cristal, et ils sont alignés et diaphragmés de telle sorte que les modes spatio-temporels des photons D rétro-réfléchis et R coïncident. Alors, selon la phase relative des deux paires R et D (qui peut être modifiée en déplaçant n’importe lequel des trois miroirs M1, M2 et M3), la création de photons dans le mode “signal” et “idler” sera soit amplifiée, soit supprimée : il y a interférence destructive ou constructive entre les deux amplitudes de probabilité pour que les photons soient émis en D ou en R. Notons que ce signal d’interférence résulte bien d’une interférence à deux photons, car si l’on supprime l’un des trois miroirs M1, M2 ou M3 les franges d’interférence disparaissent, aussi bien sur le détecteur D1 que D2.

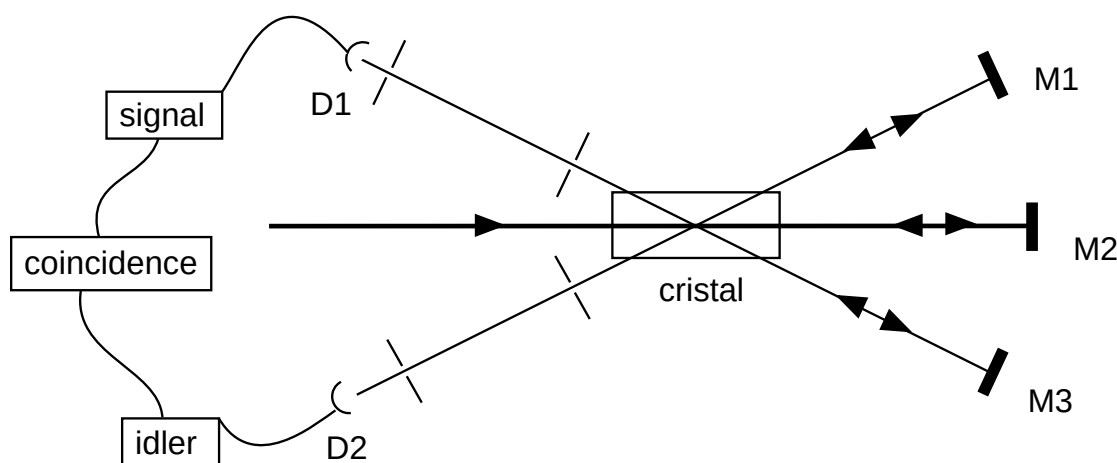


FIG. 1.6 – Schéma de l’expérience réalisée à Innsbruck. Le laser de pompe produit une paire de photons “signal” et “idler” soit directement (D) soit après une réflexion (R). Les deux modes R et D sont rendus indiscernables par les miroirs M1 et M3.

Cette interférence entre deux chemins quantiques est sujette au principe de complémentarité. Si le dispositif expérimental nous laisse la moindre possibilité de savoir si les photons détectés sont D ou R, les franges d’interférence doivent disparaître. C’est ce qu’ont montré les physiciens autrichiens. En introduisant une lame $\frac{\lambda}{2}$ sur le chemin du photon D “idler” qui transforme sa polarisation initialement verticale en polarisation horizontale, ils ont observé la suppression de l’interférence sur le photon “signal”. Avec d’autres variantes de ce dispositif, ils ont pu montrer que les franges pouvaient être récupérées si l’on plaçait avant les détecteurs des polariseurs à 45° permettant d’effacer

de manière cohérente l'information contenue dans la polarisation en la détectant dans une autre base.

Mentionnons pour finir que d'autres expériences, révélant différentes manifestations de la complémentarité ont été réalisées avec des dispositifs expérimentaux très différents [3, 75, 54]. Le groupe de Dave Wineland a étudié les franges d'interférence causées par la diffusion d'un laser par deux ions piégés, et a vu que le signal d'interférence dépendait de la polarisation émise, ce qui est une manifestation de la complémentarité [30]. Le groupe de G. Rempe a utilisé la diffraction de Bragg causée par deux ondes lumineuses hors-résonante pour réaliser un interféromètre atomique avec des atomes froids. Lorsque l'on applique des impulsions radiofréquence qui font que le chemin suivi par l'atome dans l'interféromètre est corrélé à son état hyperfin, le principe de complémentarité s'applique et l'on constate effectivement que les franges disparaissent [29]. Les physiciens allemands ont pu en outre tester la relation de dualité 1.9 pour la première fois [28].

1.2.3 Tester la complémentarité avec un “diaphragme mobile”

Tous ces résultats expérimentaux, obtenus avec des systèmes si différents, montrent clairement toute la généralité du concept dégagé par Bohr. Nous avons voulu le tester dans le régime évoqué dans le paragraphe 1.2.1.b où c'est l'interféromètre lui-même qui peut enregistrer une information WP, ce qui à notre connaissance n'avait jamais été réalisé. On a rappelé à la figure 1.3 b) le dispositif imaginé par Bohr et qui a servi d'inspiration à l'expérience.

Interférométrie de Ramsey Au lieu de réaliser un interféromètre avec diaphragme mobile dans l'espace réel, ce qui serait extrêmement difficile, nous utilisons un interféromètre qui agit dans l'espace abstrait de l'état interne d'atomes à deux niveaux. Nous allons commencer par détailler le principe de l'interférométrie de Ramsey [77] pour un atome à deux niveaux $|e\rangle$ et $|g\rangle$.

Le principe de fonctionnement d'un interféromètre de Ramsey est décrit sur la figure 1.7. Une source externe S_{Ram} , accordée à la fréquence d'une transition atomique entre deux niveaux $|e\rangle$ et $|g\rangle$, est allumée brièvement à deux instants distincts t_1 et t_2 . (L'endroit où se trouve l'atome lors des impulsions est appelé une zone de Ramsey). L'impulsion R1 place l'atome dans une superposition des états $|e\rangle$ et $|g\rangle$, et l'impulsion R2 les remélange. La probabilité de détecter finalement l'atome dans $|g\rangle$ va donc exhiber des franges d'interférence, dues à l'existence de deux chemins quantiques indiscernables allant de $|e\rangle$ à $|g\rangle$:

$$P(g) = \frac{1}{2}(1 + \cos(\Delta\Phi)) \quad (1.13)$$

La réalisation expérimentale d'un interféromètre de Ramsey “classique” sera décrite en détail au chapitre 3.

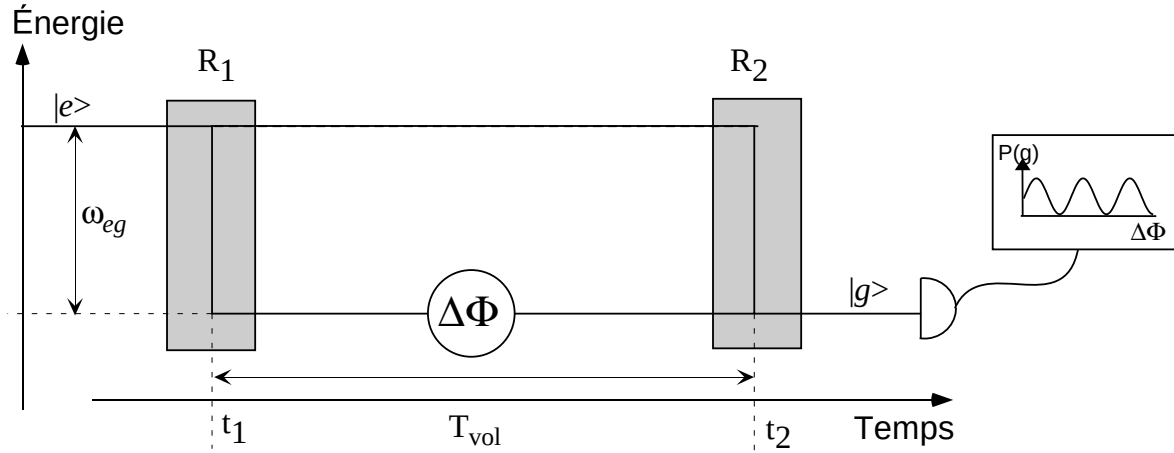


FIG. 1.7 – Principe de fonctionnement d'un interféromètre de Ramsey. Les niveaux d'énergie $|e\rangle$ et $|g\rangle$ sont mélangés en $R1$ et recombinaés en $R2$. La probabilité de détecter l'atome dans $|g\rangle$ révèle des interférences lorsqu'on fait varier le déphasage $\Delta\Phi$

Un interféromètre de Ramsey avec un “diaphragme mobile” Les impulsions de Ramsey classiques sont généralement effectuées dans des champs contenant un grand nombre de photons. Nous avons cherché à réaliser un interféromètre de Ramsey dans lequel l'une des deux impulsions a lieu dans un champ microscopique, contenant un nombre de photons $N \sim 1$. Pour cela, nous réalisons le dispositif montré à la figure 1.8 b. L'atome effectue la première impulsion de Ramsey dans le mode d'une cavité contenant un champ classique à faible nombre de photons. La deuxième impulsion a lieu dans un champ classique. On peut comprendre qualitativement le comportement de l'atome dans un tel interféromètre en discutant le fonctionnement d'un interféromètre de Mach-Zehnder dont il est l'analogue formel dans l'espace des énergies internes. On peut comparer les deux dispositifs à la figure 1.8.

Sur la figure 1.8 a), un faisceau lumineux est divisé en deux faisceaux a et b par une lame séparatrice $L1$ et est recombinaé sur la lame $L2$. Les amplitudes quantiques associées à chacun des deux chemins a et b présentent une différence de phase $\Delta\Phi$ que l'on peut faire varier. Si $L1$ et $L2$ sont macroscopiques, la probabilité de détecter le photon au niveau du détecteur D présente des franges d'interférence parfaitement contrastées. Supposons maintenant que $L2$ soit toujours massive mais $L1$ soit une lame très légère tournant autour d'un axe perpendiculaire au plan de l'interféromètre. Quand la particule interagit avec $L1$, elle est, avec une probabilité 50% soit transmise dans la voie a (la lame ne bouge pas), soit réfléchiée dans la voie b (la lame $L1$ reçoit une impulsion et finit dans un état cohérent du mouvement). L'état final de $L1$ mesure le chemin choisi par le photon dans l'interféromètre, ce qui brouille les franges d'interférence habituelles.

L'interféromètre de Ramsey montré à la figure 1.8 b) est formellement analogue à l'interféromètre de Mach-Zehnder précédent. Lors de la première impulsion de Ramsey, l'atome est transmis soit dans l'état $|e\rangle$ soit dans l'état $|g\rangle$ avec une probabilité de

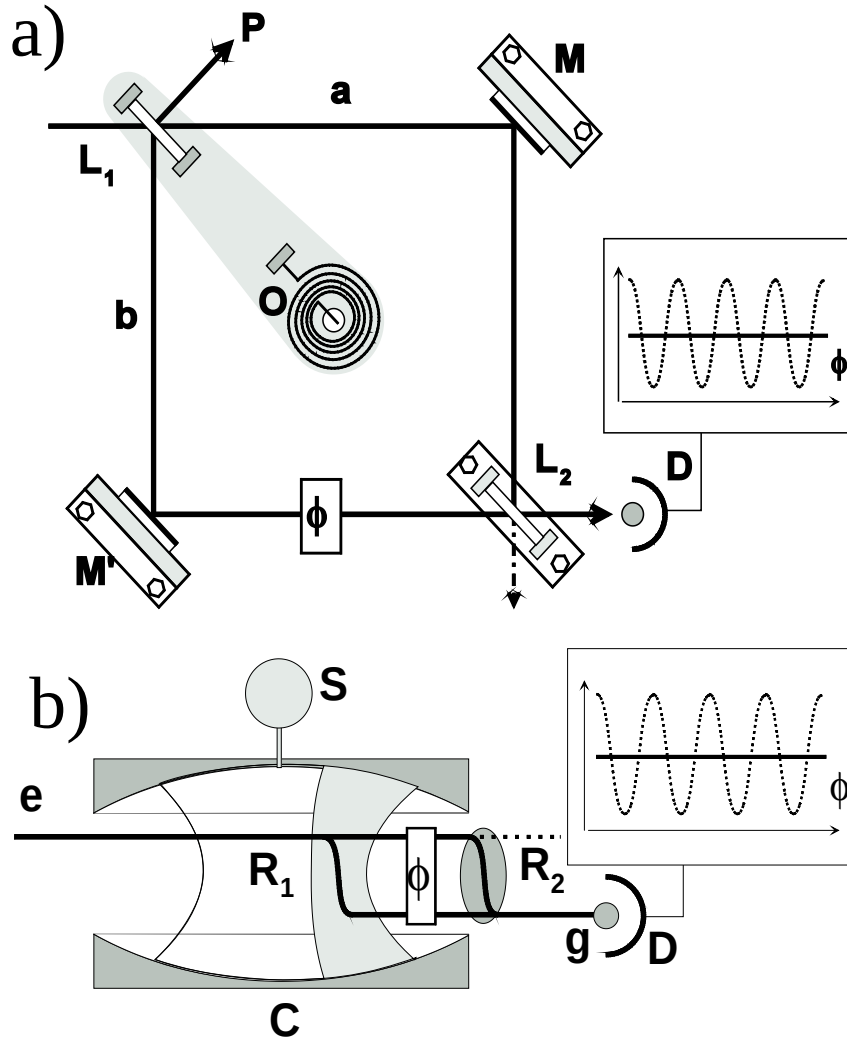


FIG. 1.8 – Analogie entre l'interféromètre de Mach-Zehnder à lame séparatrice mobile et l'interféromètre de Ramsey quantique. Figure (a) L'interféromètre de Mach-Zehnder. La lame séparatrice L_1 tourne autour d'un axe perpendiculaire au plan de l'interféromètre, fixé en O . Un ressort exerce une force de rappel. L'ensemble "ressort+ L_1 " forme un oscillateur harmonique Ω . (b) Un atome subit une première impulsion de Ramsey dans le mode d'une cavité à très haut facteur de qualité, et une deuxième impulsion de Ramsey dans un champ classique.

50%. Après le passage de l'atome, le champ dans la cavité comporte donc un photon en plus. Lorsque la cavité contient initialement un champ classique intense, l'ajout d'un photon ne modifie presque pas son état, et l'on devrait observer des franges parfaitement contrastées. Mais lorsque le nombre de photons initial du champ classique est faible, le champ est substantiellement modifié par l'ajout d'un photon. L'état de la cavité mesure l'état de l'atome entre les deux impulsions de Ramsey, ce qui brouille les franges.

Pour réaliser cette expérience, il est crucial d'avoir un couplage atom-champ qui permette d'effectuer une impulsion de Ramsey dans un champ de quelques photons. Nous allons voir en détail dans le prochain chapitre comment le dispositif d'électrodynamique quantique en cavité permet d'atteindre le régime du "couplage fort" entre l'atome et la cavité. Dans ce régime, le couplage atome-champ est si intense que même lorsque la cavité est vide, il est possible d'y effectuer une impulsion de Ramsey. On pourra ainsi étudier la transition quantique/classique au chapitre 4 sur toute la gamme des champs possibles.

Chapitre 2

Le dispositif expérimental

Le dispositif expérimental que j'ai utilisé dans ce travail de thèse est construit et amélioré depuis 1990. Il est le fruit du travail de nombreux étudiants qui m'ont précédé, ainsi que des permanents de l'équipe. Lorsque je suis arrivé dans l'équipe, en janvier 1999 pour mon stage de DEA, l'enceinte à vide venait d'être fermée et l'expérience était tout juste refroidie par Gilles Nogues et Arno Rauschenbeutel. Elle n'a pas revu le jour avant juillet 2001, soit 28 mois de fonctionnement ininterrompu. Par conséquent, pour réaliser les expériences présentées aux chapitres 4 et 5, je n'ai pas eu à travailler directement sur le coeur du dispositif. J'ai néanmoins pu pleinement goûter aux joies de la physique expérimentale en travaillant sur la partie externe de l'expérience, qui comprend en particulier les lasers et la micro-onde. En outre, j'ai pu voir tous les détails de l'expérience sur laquelle j'avais travaillé pendant deux ans et prendre part à la réparation et l'amélioration du dispositif avec les nouveaux étudiants du groupe de septembre 2001 à mars 2002. On trouvera la description de ces améliorations dans les thèses d'Alexia Auffèves, Paolo Maioli et Tristan Meunier.

Dans ce chapitre je vais décrire les grandes lignes de ce dispositif. On trouvera beaucoup plus de détails sur certains points précis dans les thèses des étudiants qui m'ont précédé ([6, 71, 61, 27, 63, 68, 78, 73]) ainsi que dans la référence [46].

2.1 Description générale

2.1.1 Motivations et ordres de grandeur

Rappelons brièvement les motivations qui ont présidé à la réalisation du dispositif que je vais détailler dans la suite de ce chapitre. Le modèle de Jaynes-Cummings décrit l'interaction d'un système à deux niveaux avec un champ quantifié sans relaxation par un terme de couplage linéaire. Il est suffisamment simple pour être résolu analytiquement, tout en conduisant à une dynamique qui manifeste clairement la nature quantique des systèmes de départ (oscillations de Rabi quantiques, intrication, ...).

Le but de cette expérience est de produire un système qui obéisse à une dynamique

du type Jaynes-Cummings, dans l'espoir de tester les principes les plus fondamentaux de la mécanique quantique. Dans notre dispositif, des atomes à deux niveaux interagissent avec un mode d'une cavité résonant avec la transition atomique. Nous avons cherché à atteindre le régime où la dissipation de l'énergie se produit à un rythme beaucoup plus lent que le temps caractéristique du couplage atome-champ (régime de *couplage fort*). Pour un atome, la dissipation de l'énergie se manifeste par l'émission spontanée, dont l'échelle de temps est donnée par la durée de vie radiative de l'état considéré. Pour la cavité, elle est due à la diffusion du champ stocké dans le mode de la cavité vers les modes du continuum, ainsi qu'à l'absorption des miroirs. La dissipation dans la cavité sous toutes ses formes est caractérisée par un seul paramètre, le taux de décroissance de l'intensité d'un champ stocké dans la cavité, T_{cav}^{-1} .

Il est très difficile expérimentalement d'atteindre le régime de couplage fort. Pour l'instant, seules deux classes d'expérience utilisant des techniques très différentes ont atteint ce régime :

- des expériences d'électrodynamique quantique en cavité dans le domaine optique étudiant l'interaction d'un atome dans l'état fondamental avec une cavité du type Fabry-Pérot constituée de deux miroirs diélectriques se faisant face ([88, 89]).

- des expériences d'électrodynamique quantique en cavité dans le domaine micro-onde, étudiant l'interaction d'atomes de Rydberg avec une cavité supraconductrice ([65, 20]). Dans notre groupe, nous utilisons des atomes de Rydberg circulaires en interaction avec une cavité micro-onde supraconductrice. Cela nous impose de travailler à basse température (environ 1 K) pour que le facteur de qualité de la cavité soit le plus important possible, et pour que le rayonnement du corps noir à la fréquence de la cavité (51.1GHz) qui perturbe les expériences soit le moins important possible.

Les atomes de Rydberg circulaires sont des états électroniques très excités, proches de la limite d'ionisation, pour lesquels le moment angulaire orbital de l'électron est maximal. Ils ont une grande durée de vie (typiquement 30ms). Or les atomes sont initialement dans l'état fondamental ; il est donc nécessaire de les exciter dans les états de Rydberg circulaires. Il faut aussi les détecter à la fin de l'expérience. Les techniques que nous utilisons seront décrites en détail dans la section 2.2.

La cavité supraconductrice a été réalisée par Gilles Nogues. Il s'agit d'une cavité du type Fabry-Pérot ouverte constituée de deux miroirs sphériques se faisant face en Niobium massif. Elle sera décrite plus en détail au paragraphe 2.3. Contentons-nous de dire que son temps de vie est d'environ 1 ms quand elle est refroidie à 1K.

Comme le temps caractéristique d'évolution du système couplé atome-champ est d'environ $20\mu s$ (période de Rabi), on voit qu'il y a plus d'un ordre de grandeur d'écart entre le couplage atome-champ et la dissipation. Notre système permet donc d'atteindre le couplage fort.

2.1.2 Schéma du dispositif expérimental

Le dispositif expérimental est schématisé sur la figure 2.1. Les atomes de Rubidium (^{85}Rb) sont issus d'un four en un jet horizontal qui traverse l'expérience. Le four, placé dans une enceinte à vide séparée de l'enceinte principale de l'expérience, est chauffé à 190° . Des diaphragmes successifs limitent l'extension transverse du jet à moins d'un millimètre.

Les atomes croisent en premier lieu la zone de sélection de vitesse. Un système de pompage optique sélectif en vitesse permet de n'exciter dans le niveau circulaire que les atomes qui ont la vitesse souhaitée avec une sélectivité $\Delta v/v \sim 1\%$. Il sera décrit en détail au paragraphe 2.2.3.

Puis le jet pénètre dans l'enceinte à vide principale qui contient le coeur de l'expérience, refroidi à 1.3K . Les atomes sont excités dans le niveau circulaire correspondant à $|g\rangle$ ou $|e\rangle$ dans une région d'environ 0.5mm de rayon définie par l'intersection des trois faisceaux laser d'excitation. La suite du processus d'excitation dans les niveaux circulaires est complexe et sera décrite en détail au paragraphe 2.2.2. Retenons seulement pour l'instant que l'ensemble du processus de circularisation a lieu à une position et un instant bien définis puisque le premier laser d'excitation est pulsé avec une fenêtre temporelle d'environ $2\mu\text{s}$. Comme en plus on connaît la vitesse de chaque atome, on peut suivre la position de l'atome à-travers l'expérience avec une résolution d'environ 1mm .

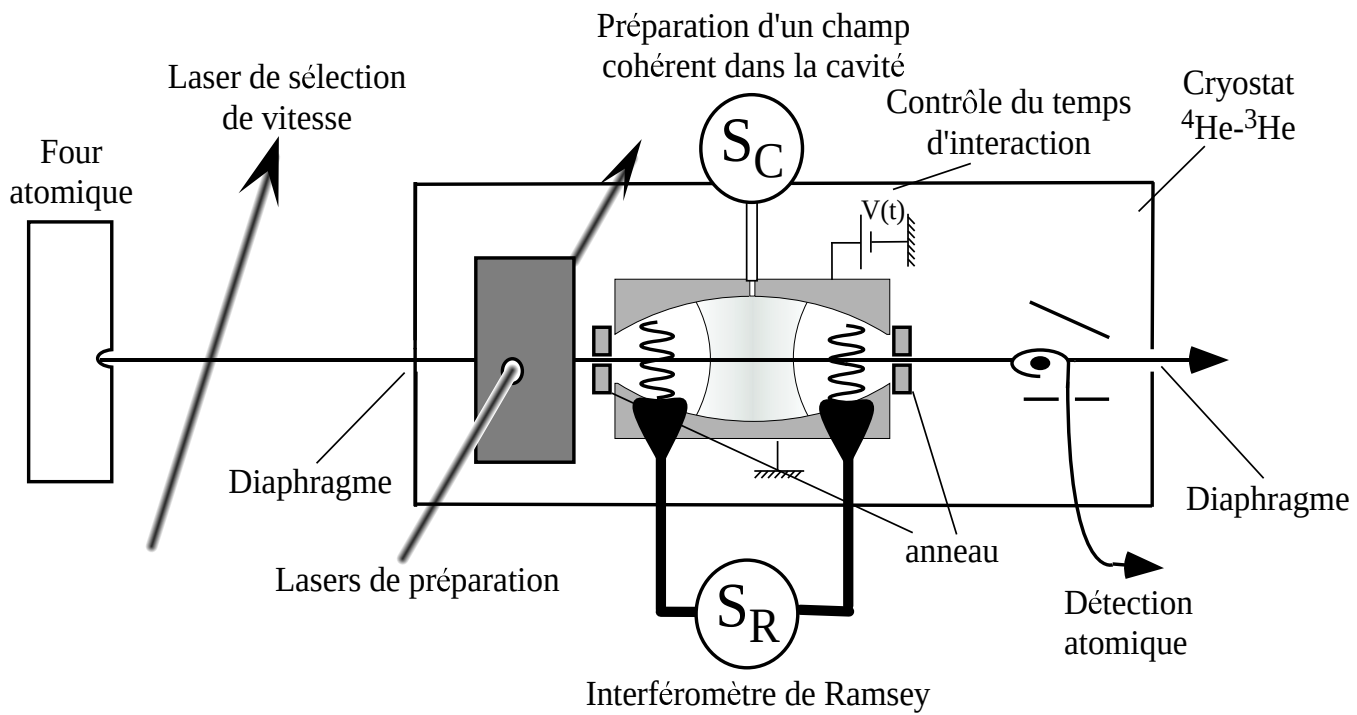
Les atomes pénètrent ensuite dans la zone où sont effectuées les opérations de contrôle cohérent de l'état atomique, entre les miroirs de la cavité. Cette dernière sera décrite en détail dans la partie 2.3. On peut effectuer deux types d'opérations sur l'atome :

- en premier lieu, l'atome peut interagir avec le mode de la cavité, et y effectuer des oscillations de Rabi quantiques (car nous sommes dans les conditions de couplage fort). Un champ électrique que l'on applique entre les miroirs permet alors de contrôler en temps réel la fréquence atomique, et donc notamment de choisir avec une résolution de 100ns le temps d'interaction de chaque atome avec la cavité. Notons, comme on peut le constater sur la figure 2.1, que l'on peut injecter un petit champ cohérent dans la cavité, par un guide d'onde qui permet de coupler la cavité à une source micro-onde extérieure.

- d'autre part, on peut appliquer des impulsions micro-onde aux atomes lorsqu'ils sont entre les miroirs de la cavité. En effet, un autre guide d'onde qui débouche sur le côté de la cavité permet de coupler les atomes à une source extérieure. Grâce à ce guide, nous pouvons réaliser un interféromètre de Ramsey, dont le fonctionnement sera décrit au chapitre suivant.

Enfin les atomes arrivent au niveau du détecteur. Celui-ci ionise sélectivement les atomes de Rydberg circulaires en fonction de leur état interne. Pour obtenir des renseignements statistiques sur l'état final de l'atome, il suffit alors de recommencer la séquence expérimentale un grand nombre de fois et de reconstituer les probabilités

FIG. 2.1 – Schéma global du dispositif expérimental



finales de détecter l'atome dans le niveau $|e\rangle$ ou dans le niveau $|g\rangle$.

2.2 Les atomes de Rydberg circulaires

Les états dits de Rydberg sont des états où l'électron de valence d'un atome alcalin est sur une orbite très éloignée du noyau (nombre quantique $n \gg 1$), proche de la limite d'ionisation. On qualifie l'orbite occupée par l'électron de "circulaire" si en plus le moment angulaire orbital de l'électron y est maximal, ce qui correspond à des nombres quantiques l et m vérifiant $|m| = l = n - 1$. Il y a donc deux états circulaires pour un nombre quantique principal n donné, de nombres quantiques m opposés.¹ Au signe de m près, un état circulaire est donc complètement caractérisé par la donnée de n . Il s'agit d'états quasi-classiques hydrogénoïdes où l'électron évolue sur une orbite circulaire autour du coeur ionique de l'atome. Les structures fine et hyperfine de ces états sont négligeables comparées aux fréquences caractéristiques de l'expérience ([6]).

2.2.1 Propriétés

Les états utilisés et les fréquences de transition

Au cours des expériences que nous allons présenter aux chapitres suivants nous utiliserons trois niveaux atomiques différents de nombres quantiques $n=51, 50, 49$, que nous noterons dans toute la suite les niveaux $|e\rangle$, $|g\rangle$, et $|i\rangle$.

On peut avoir une idée du domaine du spectre correspondant aux transitions entre états circulaires en faisant l'approximation de considérer le Rubidium comme un atome hydrogénoïde dont l'énergie varie en n^{-2} . La fréquence de transition ω_n varie donc en n^{-3} . Elle se trouve dans le domaine des micro-ondes, autour de $50GHz$. On trouvera le schéma des niveaux ainsi que les fréquences exactes des transitions à la figure 2.3.

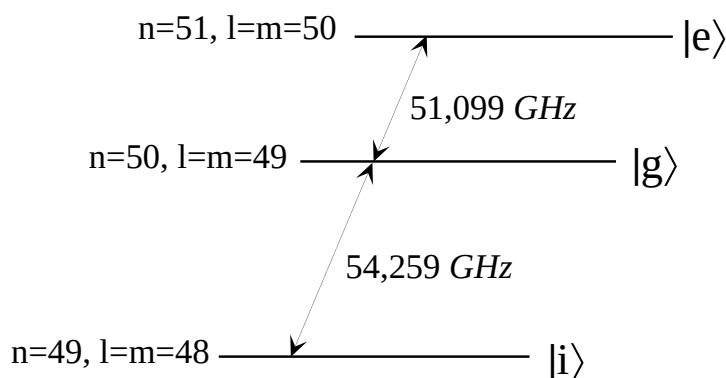


FIG. 2.2 – Schéma des niveaux atomiques utilisés dans l'expérience

¹Le signe de m dépend du sens du champ magnétique lors de la procédure de préparation, comme nous le verrons.

Un grand dipôle

Le rayon d'un atome de Rydberg circulaire vaut environ $n^2 a_0$ si $a_0 = 0.53 \cdot 10^{-10} m$ est le rayon de Bohr de l'atome d'hydrogène. Par conséquent l'élément de matrice du dipôle devrait varier en n^2 . Le calcul peut être fait exactement [10, 38] et l'on trouve en effet que

$$\begin{aligned} |\mathbf{d}_n| &= |q| |\langle n, l=n-1, m=n-1 | \mathbf{r} | n-1, l=n-2, m=n-2 \rangle| \\ |\mathbf{d}_n| &= \sqrt{2} \frac{n(n-1)^2}{2n-1} \left(1 - \frac{1}{(2n-1)^2}\right)^n |q| a_0 \\ |\mathbf{d}_n| &\simeq n^2 \frac{|q| a_0}{\sqrt{2}} \quad \text{si } n \gg 1, \end{aligned} \tag{2.1}$$

où q représente la charge électrique de l'électron. Pour $n=51$, $|\mathbf{d}_{51}| \simeq 1776 \cdot |q| a_0$ et pour $n=50$, $|\mathbf{d}_{50}| \simeq 1680 \cdot |q| a_0$. Ces valeurs de dipôle sont particulièrement grandes à l'échelle atomique. Elles reflètent la taille de l'atome de Rydberg et permettent un couplage fort avec le champ. Nous pouvons voir l'atome de Rydberg comme une antenne atomique de dimensions exceptionnelles, comme une sonde de champ électrique micro-onde très sensible.

Une grande durée de vie

Une propriété cruciale pour notre expérience concerne la durée de vie de ces atomes qui est exceptionnellement longue malgré leur fort couplage au champ électromagnétique. Cela est dû au fait que les circulaires ne peuvent se désexciter par une transition dipolaire électrique que vers le niveau circulaire immédiatement inférieur à cause de la règle de sélection $|\Delta m| = 1$. Ainsi, l'état $|e\rangle$ se désexcite vers $|g\rangle$ qui à son tour se désexcite vers $|i\rangle$. Dans l'espace libre, le taux d'émission spontanée d'un état circulaire à l'autre est [1] :

$$\tau_{at}^{-1} = \gamma_{at} = \frac{\omega_n^3 |\mathbf{d}_n|^2}{3\pi\epsilon_0 \hbar c^3}. \tag{2.2}$$

Dans l'équation 2.2, l'augmentation du dipôle avec n ($\mathbf{d}_n$ varie en n^2) est plus que compensée par le terme en ω_n^3 qui diminue en n^{-9} . Finalement, on constate que $\tau_{at} \propto n^5$. Par exemple, $\tau_{at} = 30 ms$ pour $n = 50$.

Atomes de Rydberg circulaires en champs électriques et magnétiques

Vu les grandes dimensions des atomes de Rydberg circulaires, on peut prévoir qu'ils seront très sensibles aux champs électriques.

En présence d'un champ électrique $\mathbf{E}$, la symétrie sphérique est brisée et l n'est plus un bon nombre quantique. Par contre m le demeure car la symétrie par rotation autour de l'axe du champ directeur est toujours respectée. La présence du champ électrique

couple les états de l différents qui ont le même nombre quantique magnétique m [99]. Le nouveau nombre quantique qui remplace l en présence d'un champ électrique $\mathbf{E}$ est noté n_1 et s'appelle le "nombre quantique parabolique". Il prend ses valeurs entre 0 et $n - |m| - 1$. Pour les états circulaires, $n_1 = 0$ et $|m| = n - 1$. On peut calculer exactement les déplacements Stark induits par un champ $\mathbf{E}$ par un développement en perturbation [71]. On calcule ainsi le diagramme Stark des états de Rydberg, qui est représenté à la figure 2.3.

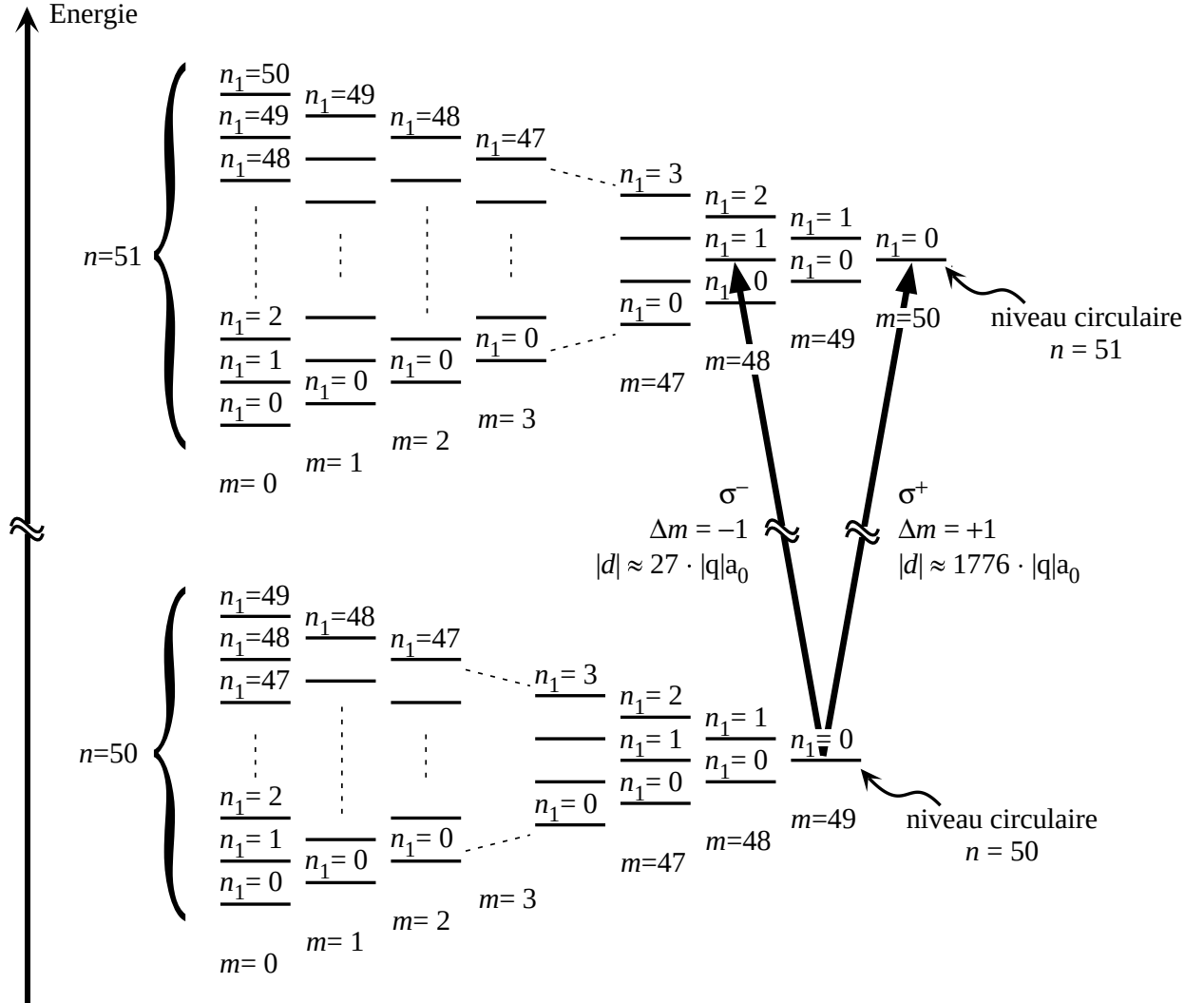


FIG. 2.3 – Diagramme Stark des deux niveaux de Rydberg $n=50$ et $n=51$ de l'atome de Rubidium pour $m > 0$.

Les atomes de Rydberg circulaires n'ont pas de dipôle électrique permanent à cause de leurs propriétés de symétrie (le nombre quantique l est bien défini pour eux). Par conséquent, ils ont un effet Stark quadratique et sont beaucoup moins sensibles aux

champs électriques que les états elliptiques de même nombre quantique principal qui eux ont un effet Stark linéaire. Cependant, cet effet Stark quadratique reste important car l'électron de valence étant loin du noyau, les atomes dans des états de Rydberg circulaires sont très polarisables. Pour la transition qui nous intéresse, entre les niveaux circulaires $n = 50$ et $n = 51$, on calcule que l'effet Stark quadratique est de $-255kHz/(V/cm)^2$.

Notons qu'en l'absence de champ électrique, tous les 2500 états de la multiplicité $n = 50$ ont la même énergie. Un atome initialement dans un état de Rydberg circulaire n'y resterait pas. La moindre perturbation (un champ électrique parasite par exemple) couplerait en effet cet état à beaucoup d'autres et l'orbite atomique serait très rapidement décircularisée. Comme nous souhaitons conserver des atomes dans des états circulaires dans toute l'expérience, on en déduit que la présence d'un champ électrique directeur est nécessaire tout le long du dispositif. Ce champ s'ajoute aux champs inhomogènes, de sorte qu'un atome traversant l'expérience voit un champ variable en amplitude et en direction. Il est crucial que les variations de champ soient suffisamment lentes pour que les niveaux circulaires suivent adiabatiquement ces variations. Toute variation "rapide" du champ électrique provoquerait un couplage de l'état circulaire à une orbite elliptique voisine. L'échelle de temps est bien sûr donnée par la fréquence de Bohr de la transition vers le niveau elliptique le plus proche. Or, plus le champ directeur est important, plus la dégénérescence entre les niveaux circulaire et elliptique est levée, donc plus la fréquence de Bohr de la transition est importante. On conçoit donc que plus le champ directeur est important, plus l'orbite circulaire est stable.

La nécessité d'appliquer un champ électrique statique sur tout le trajet des atomes est une contrainte expérimentale importante. Notamment, cela nous oblige à utiliser une cavité micro-onde *ouverte* pour laquelle il est très difficile d'avoir un bon facteur de qualité. Nous reviendrons longuement sur ce point au paragraphe 2.3.

Des atomes à deux niveaux

Nous venons de voir que "vers le bas", un niveau circulaire n'est couplé qu'au niveau circulaire de nombre quantique principal inférieur. Par contre, le niveau circulaire $n = 50$ par exemple est couplé à plusieurs états de la multiplicité $n = 51$: il peut en effet absorber un photon polarisé σ^- , π ou σ^+ . En fait, tout le long de l'expérience, le champ électrique directeur que nous venons d'évoquer lève la dégénérescence entre les transitions σ et π (comme on pourra le constater sur la figure 2.3). Restent deux transitions σ^+ et σ^- de même fréquence (à l'effet Stark quadratique différentiel près) qui sont représentées à la figure 2.3.

Ces transitions couplent le niveau circulaire $n = 50$ soit au niveau circulaire supérieur $|n = 51, l = m = 50\rangle$, soit au niveau $|n = 51, n_1 = 1, m = 48\rangle$ qui est elliptique. Le couplage au niveau elliptique peut être calculé exactement ; on trouve qu'il est 70 fois moindre qu'au niveau circulaire. Comme dans toutes nos expériences, le vecteur représentant l'état de l'atome sur la sphère de Bloch ne fait en général qu'un tour au

plus, on voit qu'on peut négliger le couplage au niveau elliptique dans l'étude de l'évolution atomique (en effet, le transfert de population associé à la transition σ^- serait négligeable).

Finalement, chaque niveau circulaire n'est couplé qu'à d'autres niveaux circulaires, chacun d'eux étant stables à l'échelle de temps d'une séquence expérimentale. Les états $|n = 51\rangle$ ($|e\rangle$), $|n = 50\rangle$ ($|g\rangle$) et $|n = 49\rangle$ ($|i\rangle$) forment un système à trois niveaux quasiment fermé.

Choix des niveaux

On peut se demander pourquoi avoir choisi de travailler avec les niveaux circulaires $n=50$ et 51 . Il s'agit d'un compromis entre deux contraintes contradictoires :

- lorsque n augmente, le temps de vie des atomes augmente ainsi que leur dipôle. Il semble donc souhaitable de s'approcher autant que possible de la limite d'ionisation.
- mais d'autre part la longueur d'onde augmente avec n . La taille du résonateur finirait donc par poser un problème. En outre, le nombre de photons thermiques émis par le corps noir environnant à la fréquence de la transition atomique augmenterait rapidement, ce qui nous gênerait considérablement pour les expériences envisagées.

2.2.2 La préparation

La préparation des états de Rydberg circulaires avec une bonne efficacité est un point difficile de l'expérience. Il est nécessaire de communiquer aux atomes de Rubidium initialement dans l'état fondamental à la fois de l'énergie (nombre quantique principal n) et du moment angulaire [45]. La mise au point du processus que je vais décrire brièvement dans la suite a requis beaucoup de travail il y a quelques années [71, 72]. C'est devenu aujourd'hui une opération de routine que nous nommons "circularisation".

Elle a lieu dans un champ magnétique $\mathbf{B}$ d'environ 18 Gauss, qui est généré par des bobines de Helmholtz supraconductrices installées au niveau de la zone de circularisation. Ce champ doit impérativement être écranté, sans quoi toute l'expérience serait envahie par un champ magnétique très inhomogène. C'est pourquoi la zone de circularisation est entourée d'un cylindre creux de Niobium de 5 cm de diamètre, qui est déjà supraconducteur lorsqu'on fait passer le courant dans les bobines pour la première fois, et qui confine le champ magnétique produit par effet Meissner. Nous appelons ce cylindre la "boîte à circulariser". Soulignons d'autre part que toute l'expérience est protégée par une série de blindages contre les champs magnétiques inhomogènes. Des bobines de compensation annulent le champ magnétique terrestre. Un blindage en μ -métal entoure le cryostat. En outre les parois de la jupe à ^3He entourant l'expérience sont recouvertes d'indium qui devient supraconducteur vers 3K, "gelant" ainsi le champ magnétique résiduel (de l'ordre du mGauss).

Nous allons maintenant décrire les différentes étapes de la circularisation qui sont schématisées à la figure 2.4 :

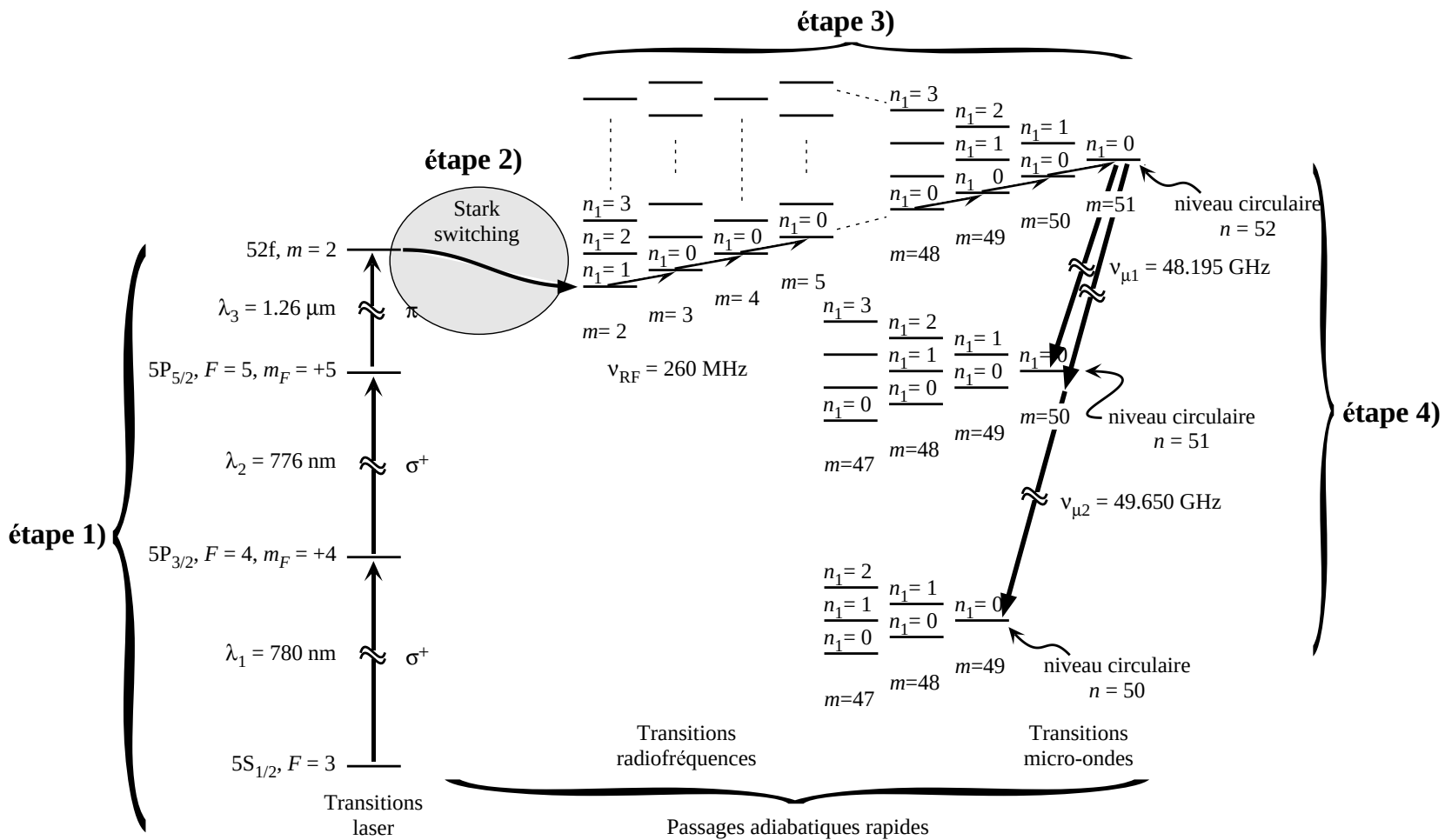


FIG. 2.4 – Circularisation des atomes de Rubidium vers les états circulaires $n = 51$ ou $n = 50$.

1) l'excitation vers les niveaux de Rydberg utilise trois transitions optiques, successivement à 780, 776 et 1258nm, en champ électrique nul. Elles transfèrent les atomes du niveau hyperfin $|F = 3, m_F = 3\rangle$ de l'état fondamental vers l'état $|52f, m = 2\rangle$, via les niveaux $|5P_{3/2}, F = 4, m_F = 4\rangle$ et $|5D_{5/2}, F = 5, m_F = 5\rangle$. Les longueurs d'onde requises sont facilement accessibles avec des diodes laser à cavité étendue. Tous les lasers sont résonants. La puissance de chaque laser est de quelques mW, ce qui sature les deux premières transitions. Le paramètre de saturation de la dernière transition est mal connu, mais sans doute entre 0.01 et 0.1. Les polarisations des lasers sont respectivement σ^+, σ^+, π . Les lasers sont focalisés au centre de la boîte à circulariser, là où le champ magnétique est le plus homogène, sur une zone d'environ $500\mu m$ de rayon.

2) On augmente ensuite le champ électrique $\mathbf{E}$ de manière à lever la dégénérescence entre l'état circulaire et les autres sous-niveaux de la multiplicité $n=52$. Ceci a pour effet de transférer plus ou moins adiabatiquement l'atome de l'état $|n = 52, l = 3, m = 2\rangle$ vers l'état $|n = 52, n_1 = 1, m = 2\rangle$ (procédure de "Stark switching").

3) Grâce à la structure hydrogénoïde pour $m \geq 3$ de la multiplicité Stark du niveau $n=52$, on peut passer de l'état $|n = 52, n_1 = 1, m = 2\rangle$ à l'état circulaire $n=52$ en absorbant une série de 49 photons radiofréquence polarisés σ^+ de fréquence 255MHz.

- un premier problème à résoudre concerne la polarisation. L'onde radiofréquence à 255 MHz arrive dans l'expérience avec une polarisation linéaire; or on souhaite que l'atome n'absorbe que des photons σ^+ , sinon il va se perdre quelque part dans la multiplicité Stark. Il faut donc lever la dégénérescence entre les transitions σ^+ et σ_- . C'est le rôle du champ magnétique de 18 Gauss créé par les bobines supraconductrices : la levée de dégénérescence est d'environ 50 MHz, suffisante pour avoir une sélectivité parfaite.

- ensuite il faut s'assurer que l'absorption des 49 photons a lieu efficacement. Au lieu de les absorber de manière résonante, ce qui aurait pour résultat de peupler à égalité les cinquante niveaux et de perdre beaucoup d'atomes, il est de beaucoup préférable de faire un passage adiabatique du niveau $|n = 52, n_1 = 1, m = 2\rangle$ au circulaire $n=52$. On est aidé en cela par le fait que toutes les transitions ont lieu à même fréquence. Si l'on allume la RF en faisant varier lentement la fréquence de transition atomique de part et d'autre de 255 MHz par effet Stark, on obtient un anticroisement géant des 50 niveaux atomiques "habillés" par les photons RF. Si la variation de $\mathbf{E}$ est suffisamment lente, on peut en théorie transférer tous les atomes dans l'état circulaire avec une efficacité proche de 1. Expérimentalement, l'efficacité de cette procédure de circularisation est d'environ 80%.

4) On a alors préparé des atomes du jet dans l'état circulaire $n=52$; or nous souhaitons préparer des atomes dans les états circulaires $n=51$ (état $|e\rangle$) et $n=50$ ($|g\rangle$). Si nous avons préparé les atomes directement dans l'état $n=51$ par exemple (au lieu d'aller vers $n=52$), nous aurions eu 20% d'atomes elliptiques dans l'état $n=51$ qui auraient pu interagir avec la cavité puisque certaines transitions elliptiques sont résonantes avec la transition circulaire-circulaire, contribuant ainsi de manière parasite au signal. Au contraire, nous pouvons maintenant augmenter la pureté de nos atomes circulaires

en appliquant une impulsion micro-onde de $n = 52$ vers $n = 51$ ou $n = 50$ *en champ électrique fort*. Le champ électrique lève la dégénérescence entre la transition circulaire-circulaire et les transitions entre niveaux elliptiques, et permet de ne transférer dans le niveau inférieur que les atomes circulaires. La pureté atteinte après cette étape dite de “purification” est d’environ 98%. Notons que l’on peut choisir d’effectuer une transition vers $|e\rangle$ (fréquence $\nu_e = 48.195GHz$) ou directement vers $|g\rangle$ à deux photons (fréquence $\nu_g = 49.65GHz$).

L’excitation des atomes est pulsée car le premier laser provient de l’ordre 1 de diffraction d’un Modulateur Acousto-Optique (MAO). Allumer ou éteindre l’onde acoustique revient à allumer et éteindre le laser d’excitation. Dans l’expérience, nous réalisons cela avec une résolution temporelle de 100 ns. Le choix de la durée de l’impulsion d’excitation permet de régler la probabilité d’excitation vers l’état de Rydberg circulaire et donc le flux atomique.

2.2.3 La sélection de vitesse

Pour connaître la position des atomes dans le dispositif au millimètre près, il nous faut connaître leur vitesse en plus de l’instant de leur préparation. Nous utilisons un dispositif de sélection de vitesse dont la description détaillée pourra être trouvée dans la thèse de X. Maître [63]. La méthode repose sur du pompage optique sélectif en vitesse par effet Doppler, combiné à une sélection en temps de vol.

Pompage optique sélectif en vitesse

Le principe du dispositif de sélection de vitesse est le suivant. Le processus de “circularisation” ne concerne que les atomes qui sont initialement dans le niveau hyperfin $F=3$ de l’état fondamental de l’atome de Rubidium $5S_{1/2}$ ². Lorsque les atomes sortent du four, ils sont répartis à égalité entre les deux niveaux hyperfins $F=2$ et $F=3$. Par pompage optique, on ne garde dans l’état $F=3$ que les atomes qui ont la vitesse voulue ; les autres ne seront pas circularisés.

L’interaction entre les atomes et les lasers de sélection de vitesse a lieu environ 15 cm après la sortie du four (voir figure 2.5 a)). Un premier laser L_1 commence par dépeupler complètement le niveau $F=3$. Il est accordé sur la transition $|5S_{1/2}, F = 3\rangle \rightarrow |5P_{3/2}, F' = 3\rangle$. Comme il est perpendiculaire au jet, il n’y a pas d’effet Doppler et les atomes sont tous résonants avec L_1 . Après quelques cycles de fluorescence, presque tous les atomes finissent pompés dans l’état $|5S_{1/2}, F = 2\rangle$. Nous avons vérifié expérimentalement par des mesures de flux d’atomes de Rydberg circulaires qu’après interaction avec le dépompeur, le rapport des populations des états $F=3$ et $F=2$ était inférieur à $1/2000$.

²La levée de dégénérescence entre les deux niveaux hyperfins $F = 2$ et $F = 3$ de l’état fondamental du Rubidium est de $3GHz$, qui est très largement résolu par le premier laser d’excitation dont la largeur est d’ $1MHz$.

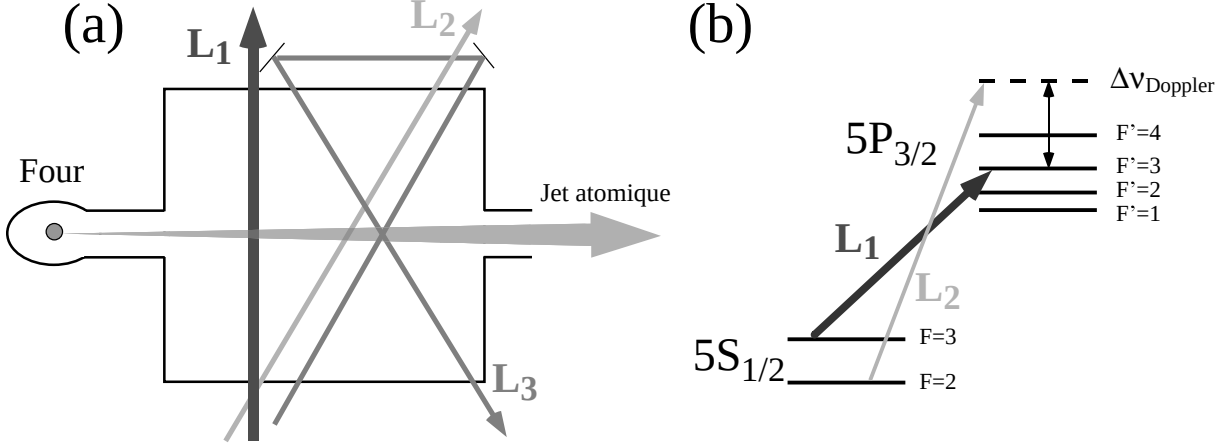


FIG. 2.5 – (a) Disposition des lasers de sélection de vitesse à la sortie du four atomique. (b) Structure hyperfine des niveaux $5S_{1/2}$ et $5P_{3/2}$ et fréquences des lasers L_1 et L_2 .

Un deuxième laser L_2 , appelé “repompeur”, est désaccordé d’une quantité réglable $\Delta\nu$ par rapport à la fréquence de la transition $F = 2 \rightarrow F' = 3$. Il fait un angle $\theta = 54^\circ$ avec le jet atomique. A cause de l’effet Doppler, la fréquence du laser “vue” par les atomes dépend alors de leur vitesse et vaut

$$\nu_{at} = \nu_{L_2} - \frac{1}{2\pi} \mathbf{k} \cdot \mathbf{v} = \nu_{F=2 \rightarrow F'=3} + \Delta\nu - \frac{v \cos \theta}{\lambda} \quad (2.3)$$

Seuls les atomes ayant une vitesse

$$v_{sel} = \frac{\lambda \Delta\nu}{\cos \theta} \quad (2.4)$$

verront un laser accordé sur la transition $F = 2 \rightarrow F' = 3$. Ces atomes seront repompés dans le niveau $F=3$ après quelques cycles de fluorescence, puis éventuellement excités dans les niveaux de Rydberg circulaires. On peut choisir $\Delta\nu$ en changeant la fréquence de l’onde acoustique envoyée dans un MAO, et ainsi sélectionner n’importe quelle classe de vitesse entre 100 et 700 m/s, ce qui couvre toute la distribution de vitesse de notre jet. Notons que, par cette procédure, la classe de vitesses $v = v_{sel} + 100 \text{ m/s}$ est aussi repompée par le laser L_2 . En effet, pour des atomes ayant cette vitesse, le laser est accordé sur la transition $F = 2 \rightarrow F' = 2$ qui est aussi repompante vers $F=3$.

On voit sur la formule 2.4 que la largeur de la classe de vitesse ainsi sélectionnée est directement proportionnelle à la largeur de la transition $F = 2 \rightarrow F' = 3$. Cette largeur dépend de l’intensité du laser repompeur. Lorsque l’intensité est faible, c’est la largeur naturelle de la transition $F = 2 \rightarrow F' = 3$ et la population repompée est proportionnelle à l’intensité du laser. On atteint alors une largeur d’environ 10 m/s sur la classe de vitesse sélectionnée. Lorsqu’au contraire on sature la transition, on élargit la raie ; la dispersion en vitesse peut alors atteindre environ 20 m/s avec une efficacité de transfert proche de 100% (voir la figure 2.6).

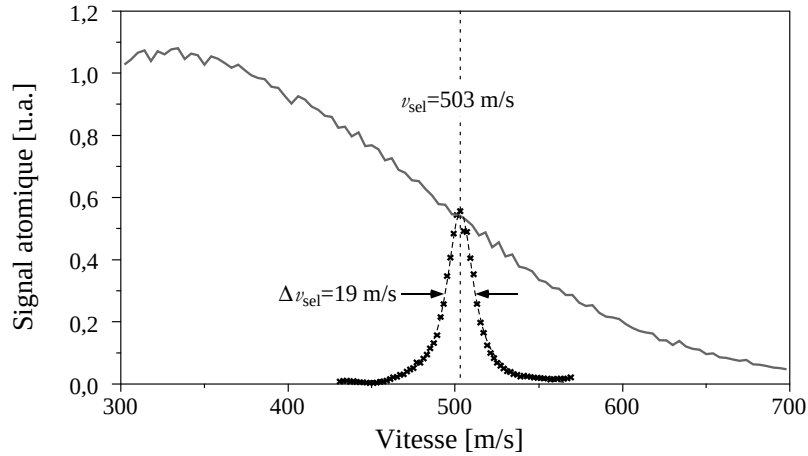


FIG. 2.6 – Mesure de la distribution de vitesse des atomes circulaires avant (trait continu) et après pompage optique. On constate que l'efficacité de transfert est bien de 100% au centre de la maxwellienne ; sa largeur est de 19m/s. La vitesse des atomes est déterminée par le temps qu'ils mettent à aller de la zone de circularisation au détecteur, dont nous connaissons la distance exacte à la zone de circularisation.

Sélection de vitesse en temps de vol

La dispersion de vitesse obtenue par le système de pompage optique étudié précédemment est trop importante pour nos expériences (elle est intrinsèquement limitée par la largeur naturelle de la transition atomique et ne peut être inférieure à 10m/s). En effet nous devons effectuer des opérations cohérentes sur les atomes et toute dispersion spatiale du jet est une cause supplémentaire d'inhomogénéités. Pour améliorer encore la sélectivité de notre dispositif, nous avons recours à un système de sélection de vitesse par temps de vol, schématisé à la figure 2.7 a).

Après pompage optique vers le niveau $F=3$ d'une certaine classe de vitesse, les atomes doivent parcourir une distance de 37.5cm qui sépare la zone de sélection de vitesse de la zone de circularisation. Pendant ce temps de vol, les atomes de différentes vitesses se dispersent, et arrivent dans la zone de circularisation à des instants différents. On peut alors, en n'allumant le laser d'excitation que lorsqu'une certaine classe de vitesse est présente, améliorer la sélectivité en vitesse du dispositif. C'est ce que montre la figure 2.8 où l'on a pu par cette méthode réduire d'un facteur 5 la largeur initiale de la distribution de vitesse.

Amélioration de la sélection de vitesse

Le dispositif de sélection de vitesse combinant pompage optique et sélection par temps de vol est celui que nous avons utilisé pour réaliser toutes les expériences de ce mémoire. Cependant il n'est pas parfait. L'utilisation du temps de vol a un inconvénient, qui devient apparent lorsqu'on considère la circularisation de plus d'un atome. On voit

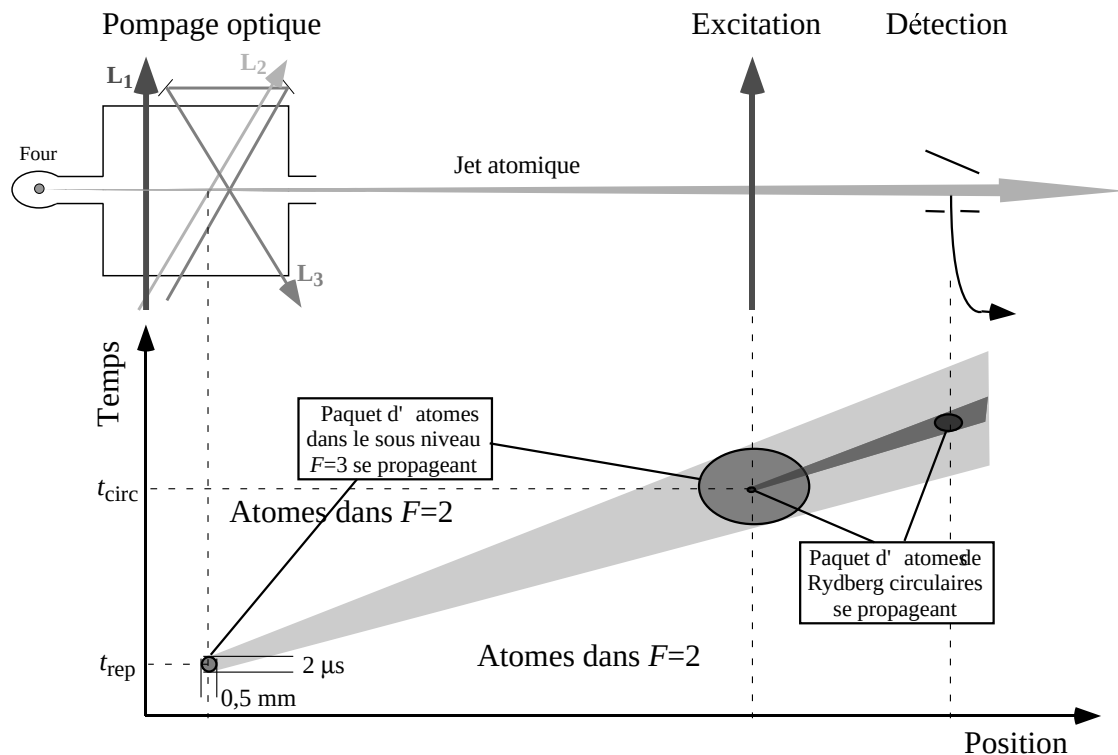


FIG. 2.7 – Principe de la sélection de vitesse par temps de vol. Un paquet d'atomes préparés dans $F=3$ après pompage optique s'élargit au cours du temps. Seule une fraction est excitée dans le niveau circulaire, ce qui résulte en une très faible dispersion de vitesse (environ 1%).

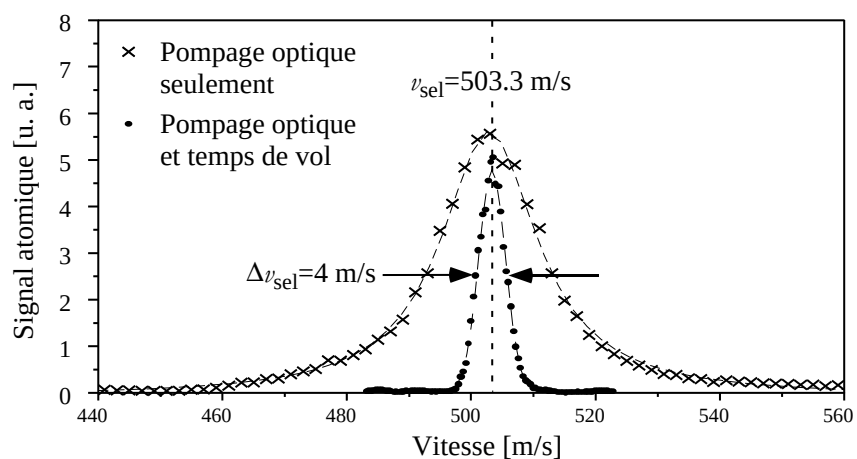


FIG. 2.8 – Résultat de la sélection de vitesse par temps de vol

sur la figure 2.9 que si l'on a deux paquets atomiques à préparer, il y a un risque de préparer des atomes ayant des classes de vitesse parasites : en effet, la circularisation de l'atome 2 préparera certes des atomes ayant une vitesse sélectionnée par la deuxième impulsion de repompage (ceux-là ont la bonne vitesse et sont indiqués en traits pleins sur la figure), mais risque également de circulariser des atomes qui ont été repompés par la première impulsion. Ces atomes (représentés en pointillés sur la figure) ont une vitesse qui n'est pas celle que l'on souhaite (ils sont plus lents), et sont susceptibles de créer des effets parasites, par exemple lors de leur interaction avec la cavité. On les appelle des "échos". Notons qu'un autre écho, plus rapide celui-là, est bien sûr généré par le repompage de l'atome 2 sur la circularisation de l'atome 1.

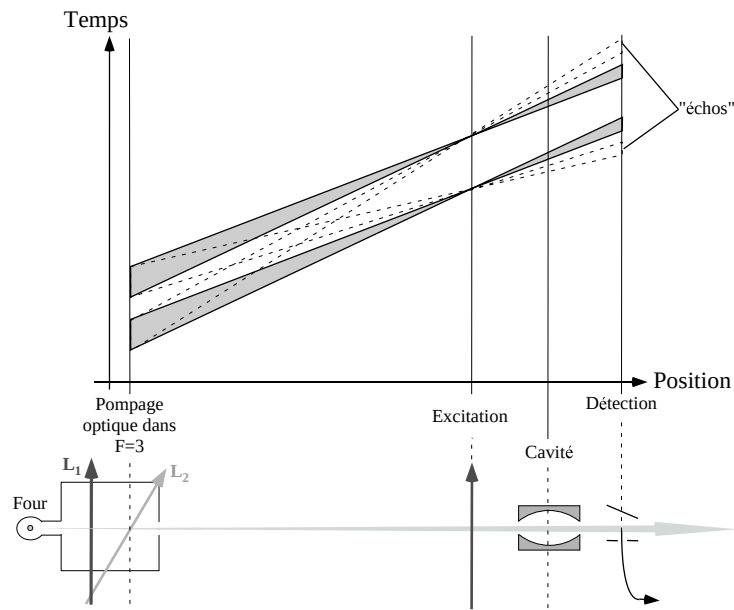


FIG. 2.9 – Des classes de vitesse parasites appelées "échos" peuvent être circularisées à cause par la sélection de vitesse par temps de vol.

Pour se débarrasser des échos, une solution est d'utiliser un laser dépompeur à angle accordé sur les vitesses des échos dont on souhaite se débarrasser. Nous avons utilisé cette solution, mais elle n'est pas pleinement satisfaisante :

- d'une part, on risque de dépomper aussi par effet de bord les atomes qui nous intéressent, et qui ont souvent des vitesses proches de celles des échos.
- d'autre part, lorsqu'on a beaucoup de paquets atomiques différents, les échos se multiplient et ont tous des vitesses différentes. Il ne peut donc pas être question de les dépomper à l'aide d'un laser sélectif en vitesse.

J'ai proposé au cours de ma thèse d'améliorer la sélectivité intrinsèque du pompage optique en faisant l'opération de repompage par une transition Raman entre les niveaux hyperfins $F=2$ et $F=3$. L'avantage est qu'une transition Raman a une largeur naturelle nulle puisque les niveaux hyperfins ont une durée de vie infinie. Expérimentalement,

la largeur de la transition observée dépend du bruit de phase relatif des lasers (qu'on réduit fortement en asservissant en phase les deux faisceaux Raman), de la puissance des lasers, et du temps d'interaction atome-champ. On pourrait ainsi avoir une sélectivité en vitesse suffisante, rien qu'en effectuant du pompage optique, et on se débarrasserait des phénomènes d'échos.

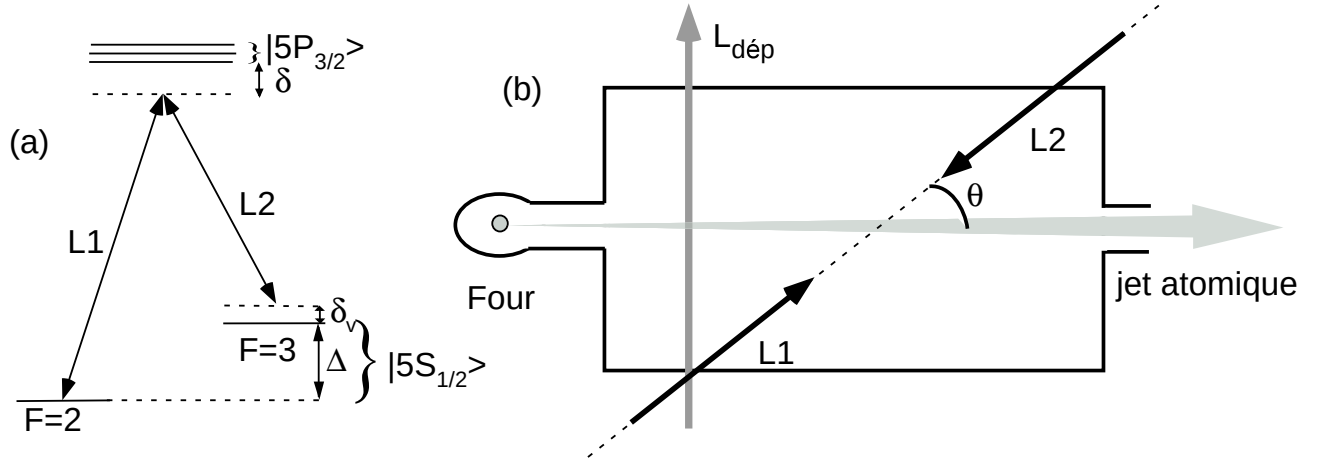


FIG. 2.10 – (a) Transition Raman de l'état $|F = 2\rangle$ vers l'état $|F = 3\rangle$ du fondamental de l'atome de Rubidium. (b) Pour avoir la meilleure sélectivité en vitesse, il faudra que les deux lasers $L1$ et $L2$ soient contrapropageants.

Le schéma de principe d'un repompage basé sur l'effet Raman est esquissé à la figure 2.10. La transition est effectuée par l'interaction de l'atome avec deux lasers $L1$ et $L2$, chacun étant désaccordé par rapport à la transition $5S - 5P$ mais dont la différence de fréquence est proche de l'écart entre les niveaux hyperfins du fondamental du Rubidium $\Delta = 3.036 GHz$. Plus précisément, cette différence de fréquence doit valoir $\Delta + \delta_v$ où δ_v est le désaccord qui permettra de choisir la fréquence. La fréquence vue par l'atome de chacun des lasers est : $\nu_{at}^{(L1)} = \nu_{L1} - \frac{1}{2\pi} \mathbf{k}_1 \cdot \mathbf{v}$ et $\nu_{at}^{(L2)} = \nu_{L2} - \frac{1}{2\pi} \mathbf{k}_2 \cdot \mathbf{v}$. Donc finalement

$$\nu_{at}^{(L1)} - \nu_{at}^{(L2)} = \nu_{L1} - \nu_{L2} + (\mathbf{k}_1 - \mathbf{k}_2) \cdot \mathbf{v} = \Delta + \delta_v + (\mathbf{k}_1 - \mathbf{k}_2) \cdot \mathbf{v} \quad (2.5)$$

Seuls les atomes dont la vitesse est telle que $\nu_{at}^{(L1)} - \nu_{at}^{(L2)} = \Delta$ seront repompés. On voit sur cette équation qu'on a intérêt à ce que $L1$ et $L2$ soient contrapropageants pour que la sélectivité en vitesse soit meilleure (on a alors $\mathbf{k}_1 = -\mathbf{k}_2$). Dans ces conditions la vitesse sélectionnée vaut :

$$v = \frac{\delta_v}{2k \cos \theta} \quad (2.6)$$

Le montage optique nécessaire a été réalisé par P. Maioli et T. Meunier au cours de son stage de DEA. Il sera décrit dans leur mémoire de thèse. Il reste à tester ce nouveau dispositif sur les atomes.

Contrôle des atomes dans l'expérience

Stabilité de l'excitation atomique Il est extrêmement important que toutes les séquences expérimentales soient le plus semblables possible. Or les acquisitions de données se font sur des périodes de plusieurs heures, et il y a beaucoup de paramètres qui peuvent fluctuer. Notamment, notre expérience comporte cinq diodes laser (trois pour l'excitation vers le niveau circulaire) et deux pour la sélection de vitesse. Ces diodes sont bien entendu asservies en fréquence, mais il peut arriver qu'une perturbation extérieure déstabilise transitoirement l'une de ces diodes. Pour stabiliser l'expérience, j'ai développé un système de surveillance automatique des lasers qui suspend l'acquisition de données lorsqu'une diode est perturbée.

Il est en outre crucial que le flux atomique moyen reste constant au cours d'une expérience. Or le flux peut fluctuer à cause de l'asservissement de notre dernier échelon d'excitation. Contrairement aux deux lasers à 780 et 776nm qui sont référencés sur une transition atomique, le laser à 1.258 μm est asservi sur une cavité Fabry-Pérot. Cette dernière a une dérive lente qui fait que le laser peut se désaccorder notablement de la raie Rydberg à l'échelle d'une heure, qui est précisément le temps caractéristique d'une acquisition de données. Le flux atomique diminue alors progressivement, ce qui nuit à la reproductibilité des séquences expérimentales. Pour éviter cela, j'ai mis en place un asservissement de la cavité Fabry-Pérot sur le flux atomique total. Ainsi le laser à 1.258 μm reste en permanence au sommet de la raie atomique.

Contrôle du nombre d'atomes par paquet atomique Grâce à la sélection de vitesse et à l'excitation impulsionnelle dans les niveaux circulaires, nous pouvons savoir à tout moment où se trouve le paquet atomique considéré dans l'expérience. Nous pouvons en outre choisir la probabilité d'excitation de manière à avoir entre 0.05 et 10 atomes par paquet en jouant sur plusieurs paramètres : puissance des lasers de repompage et d'excitation, durée des impulsions laser,... Le nombre exact d'atomes excités par séquence dépend bien sûr du nombre aléatoire d'atomes qui se trouvent dans la zone d'excitation. Nos paquets atomiques ont une statistique poissonnienne. Lorsqu'on souhaite étudier l'interaction d'un seul atome avec la cavité, on choisit les paramètres d'excitation de manière à ce que le nombre moyen d'atomes circulaires détectés soit très inférieur à 1 par séquence expérimentale. En général, on utilise des probabilités d'excitation d'environ 0.1.

2.2.4 La détection

Les détecteurs actuels

Le système de détection des atomes de Rydberg circulaires repose sur le fait qu'on peut les ioniser en leur appliquant des champs électriques statiques relativement faibles (une centaine de V/cm suffit).

La zone de détection est schématisée à la figure 2.11. Il s'agit d'un condensateur en forme de toit dont l'électrode horizontale (anode) est percée d'un trou de 2mm de diamètre. Grâce à la forme de toit, le champ électrique vu par les atomes augmente progressivement au fur et à mesure qu'ils avancent, jusqu'à atteindre la valeur limite d'ionisation pour laquelle l'électron de valence leur sera arraché. Ce seuil dépend de l'état circulaire considéré. Il vaut 136, 148 et 160 V/cm respectivement pour les états $|e\rangle$, $|g\rangle$ et $|i\rangle$.

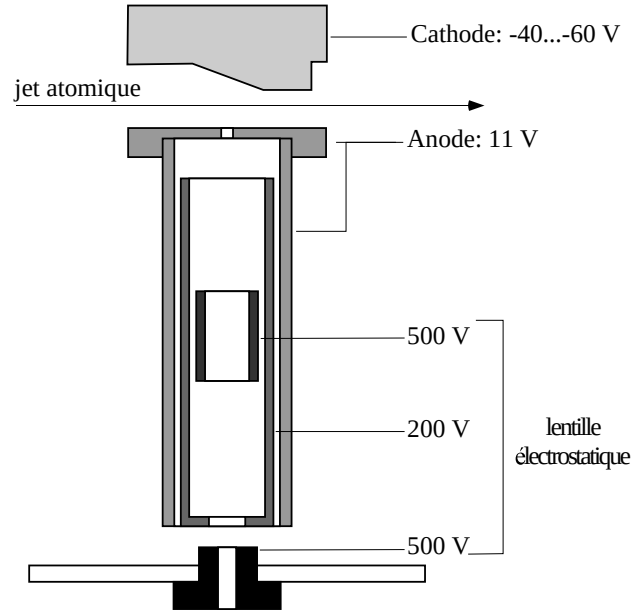


FIG. 2.11 – Zone de détection des atomes

On choisit la tension sur la cathode de manière à n'ioniser qu'un seul état circulaire en face du trou. Les autres états circulaires sont ionisés soit après, soit avant. L'électron émis est ensuite accéléré vers l'anode, puis focalisé par un montage de lentille électrostatique sur un multiplicateur d'électrons (qui n'est pas représenté sur la figure 2.11). L'impulsion résultante est amplifiée et mise en forme pour donner une impulsion TTL enregistrable. Lorsqu'on balaye la tension de la cathode, on obtient une série de pics d'ionisation qui correspondent aux ionisations successives des niveaux circulaires $n = 52$, $n = 51$, $n = 50$, $n = 49$ (voir la figure 2.12).

Lors des expériences présentées dans ce mémoire, nous ne disposons que d'un seul détecteur. Or nous nous intéressons en général à la probabilité de détecter l'atome dans l'un des deux niveaux $|e\rangle$ ou $|g\rangle$, ce qui signifie que l'on doit détecter deux niveaux circulaires différents. Pour ce faire, chaque séquence expérimentale est jouée deux fois, la tension sur la cathode étant réglée la première fois pour détecter $|e\rangle$ et la deuxième fois pour détecter $|g\rangle$. On répète ces deux séquences un grand nombre de fois pour reconstruire $P(e) = N_e / (N_e + N_g)$.

Nous avons mesuré l'efficacité quantique de notre chaîne de détection à 40%, due

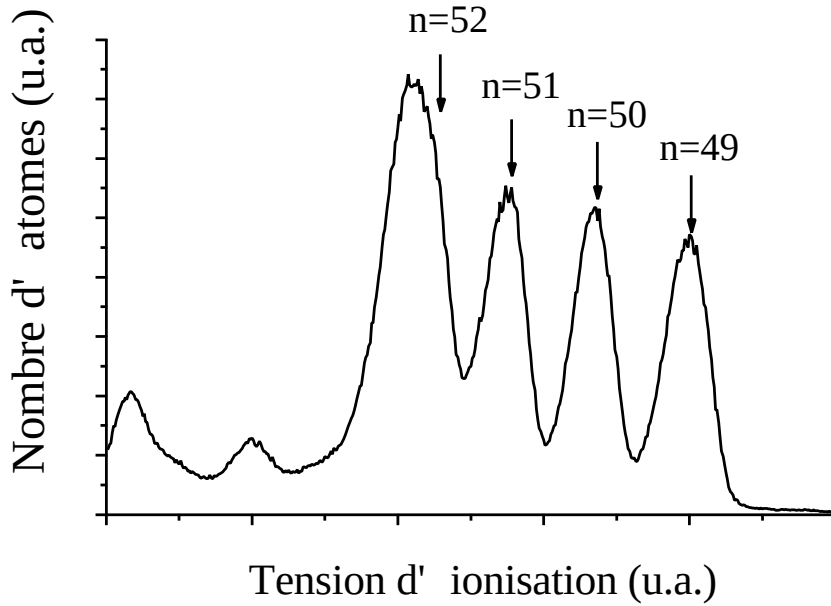


FIG. 2.12 – Signal d’ionisation des états de Rydberg circulaires. Lorsqu’on augmente le potentiel de la cathode, on observe des pics successifs correspondant à l’ionisation des différents états que nous pouvons préparer : d’abord $n=52$, puis $n=51$ (état $|e\rangle$), puis $n=50$ ($|g\rangle$), puis $n=49$ (état $|i\rangle$)

probablement à une focalisation imparfaite réalisée par les lentilles électrostatiques.

2.3 La cavité supraconductrice

La cavité supraconductrice est le coeur de notre dispositif expérimental. Pour assurer un couplage important aux atomes à l’échelle du photon, le volume du mode de la cavité doit être le plus faible possible. Une autre contrainte concerne la dissipation : nous voulons être dans le régime où le temps de vie du champ dans la cavité est beaucoup plus grand qu’une période de Rabi ($20\mu s$). Or la construction d’une cavité supraconductrice ouverte à $51.1GHz$ de très haut facteur de qualité représente un défi technologique très difficile à surmonter. La cavité avec laquelle nous avons fait les expériences que je présenterai aux chapitres suivants a été réalisée par Gilles Nogues avant mon arrivée. De nombreux détails techniques concernant ce processus pourront être trouvés dans sa thèse [68]. Dans cette partie je me contenterai de brièvement décrire les caractéristiques de notre résonateur.

2.3.1 Propriétés générales

Notre cavité est constituée de deux miroirs sphériques en Niobium se faisant face, en configuration Fabry-Pérot. La géométrie du résonateur est schématisée sur la figure

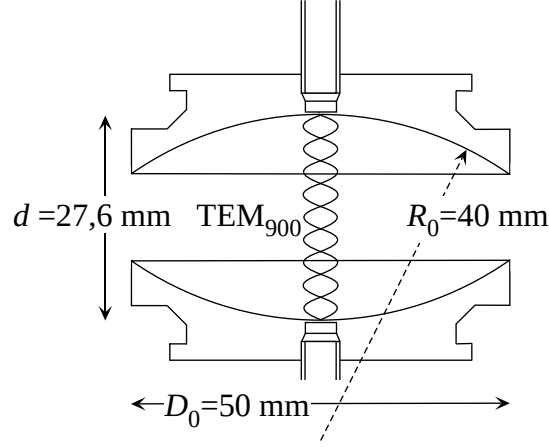


FIG. 2.13 – Géométrie de la cavité

2.13. Les miroirs ont un rayon de courbure $R_0 = 40\text{mm}$ et un diamètre $D_0 = 50\text{mm}$. Ils sont percés en leur centre par un trou de couplage qui permet d'injecter de la micro-onde et aussi d'observer la transmission de la cavité. On peut ainsi mesurer la fréquence des modes grâce à un banc de mesure micro-onde. Les trous de couplage sont bien sûr très petits devant la longueur d'onde (0.2mm de diamètre et de longueur pour une longueur d'onde de 6mm) pour que les pertes des miroirs par transmission à-travers ces trous soient négligeables.

Les miroirs sont séparés par une distance $d = 27.57\text{mm}$. Les deux modes $TEM_{9,0,0}$ de polarisation orthogonale sont alors quasi-résonants avec la transition e-g à 51.1GHz . La symétrie cylindrique de la cavité n'est cependant pas parfaite, et l'on observe par conséquent que les modes correspondant à deux polarisations linéaires orthogonales ont deux fréquences différentes. La différence de fréquence était de 128kHz pour le résonateur que nous avons utilisé, ce qui est largement supérieur au couplage $g = 25\text{kHz}$ entre les atomes et le champ. On peut donc considérer qu'un atome interagit exclusivement avec un seul mode lorsqu'il est à résonance avec lui.

La fréquence des modes est donnée par [51] :

$$\nu = \frac{\omega}{2\pi} = \frac{c}{2d} \left(q + \frac{2}{\pi} \arccos \sqrt{1 - \frac{d}{2R_0}} \right) \simeq 51,1 \text{ GHz}, \quad (2.7)$$

où c est la vitesse de la lumière et $q=9$ est le nombre de ventres de l'onde stationnaire entre les deux miroirs. Dans cette approximation paraxiale, la structure spatiale du champ électrique est gaussienne et est donnée en coordonnées cylindriques³ par la fonction $f(\mathbf{r})=f(r, z)$:

³L'axe de symétrie est porté par l'axe vertical (O, z) et l'origine O est prise au centre de la cavité.

$$f(r, z) = \frac{w_0}{w(z)} \cos \left(2\pi \frac{z}{\lambda} - \arctan \frac{\lambda z}{\pi w_0^2} + \frac{\pi}{r^2} \lambda R(z) + (q-1) \frac{\pi}{2} \right) \exp \left[- \left(\frac{r^2}{w^2(z)} \right) \right], \quad (2.8)$$

où $\lambda=c/\nu$ est la longueur d'onde associé au mode et w_0 est le col du mode. Il vaut ici :

$$w_0 = \left[\frac{\lambda d}{2\pi} \sqrt{\frac{2R_0}{d} - 1} \right]^{\frac{1}{2}} \simeq 5,96 \text{ mm}. \quad (2.9)$$

Les fonctions $w(z)$ et $r(z)$ décrivent respectivement la variation de l'extension du mode autour de l'axe (O, z) et celle du rayon de courbure des fronts d'onde. Elles sont données par les relations [51, 55] :

$$\begin{aligned} w(z) &= w_0 \sqrt{1 + \left(\frac{\lambda z}{\pi w_0^2} \right)^2} \\ R(z) &= z \left[1 + \left(\frac{\pi w_0^2}{\lambda z} \right)^2 \right] \end{aligned} \quad (2.10)$$

De l'expression de $f(\mathbf{r})$, nous déduisons le volume du mode :

$$V_{\text{mode}} = \iiint_{\text{cavité}} |\mathbf{f}(\mathbf{r})|^2 d^3\mathbf{r} = \frac{\pi w_0^2 d}{4} \simeq 769 \text{ mm}^3. \quad (2.11)$$

C'est ce volume de mode qui permet de définir la valeur du champ électrique par photon E_0 . Un petit volume, comme c'est le cas dans notre géométrie, assure le confinement du champ dans la cavité et une forte amplitude de E_0 :

$$\begin{aligned} E_0 &= \sqrt{\frac{\hbar\omega}{2\varepsilon_0 V_{\text{mode}}}} \\ E_0 &\simeq 1,58 \cdot 10^{-3} \text{ V} \cdot \text{cm}^{-1}. \end{aligned} \quad (2.12)$$

Une telle valeur du champ électrique est nécessaire pour assurer un bon couplage avec le dipôle atomique même en la présence d'un seul photon. En effet, on peut calculer la constante de couplage d'un atome de Rydberg dans l'état $|e\rangle$ avec la cavité vide

$$g_0 = \frac{\Omega_0}{2} = \frac{|\mathbf{d}|E_0}{\hbar} = 25.4 \text{ kHz} \quad (2.13)$$

Nous mesurons expérimentalement une fréquence de Rabi $\Omega_0 = 2g_0 = 49 \text{ kHz}$. La différence avec la valeur précédente est due à l'étendue transverse du jet (on mesure l'angle de Rabi moyen d'une collection d'atomes qui ne passent pas tous exactement par le centre du mode), et aussi à l'approximation paraxiale utilisée pour calculer la structure spatiale du mode. Pour atteindre le régime de couplage fort, il faut que ce couplage de 25 kHz soit beaucoup plus grand que le taux de relaxation d'un photon dans la cavité. Il est donc crucial d'avoir une cavité dont le temps de vie soit le plus élevé possible.

2.3.2 Le temps de vie de la cavité

Comment le mesurer ? On mesure le facteur de qualité de la cavité en observant la décroissance de l'intensité du champ intra-cavité par transmission à-travers les miroirs.

Les mesures de transmission de la cavité nécessitent un dispositif élaboré. En effet la cavité est très peu couplée à l'extérieur, puisque l'on souhaite minimiser les pertes par transmission. Le trou de couplage de chacun des miroirs atténue le signal transmis de $80dB$. Le pic de transmission associé au passage de la source par un mode de la cavité peut donc être de faible amplitude, sur un fond qui peut être important. Nous utilisons un analyseur vectoriel développé pour la société ABmm par Philippe Goy et Michel Gross [39], dont la dynamique est de $120dB$ (dans la suite, il y sera fait référence en parlant du "banc micro-onde"). Une description détaillée de cet appareil peut être trouvée dans [27]. Il contient deux sources YIG, un pour la source et l'autre pour la détection, asservis en phase l'un sur l'autre et décalés d'une fréquence de $9.01MHz$. Le YIG source est quadruplé en fréquence par une diode Schottky ; on mesure l'amplitude et la phase du signal transmis à-travers la cavité par hétérodynage avec le YIG de détection lui aussi quadruplé en fréquence.

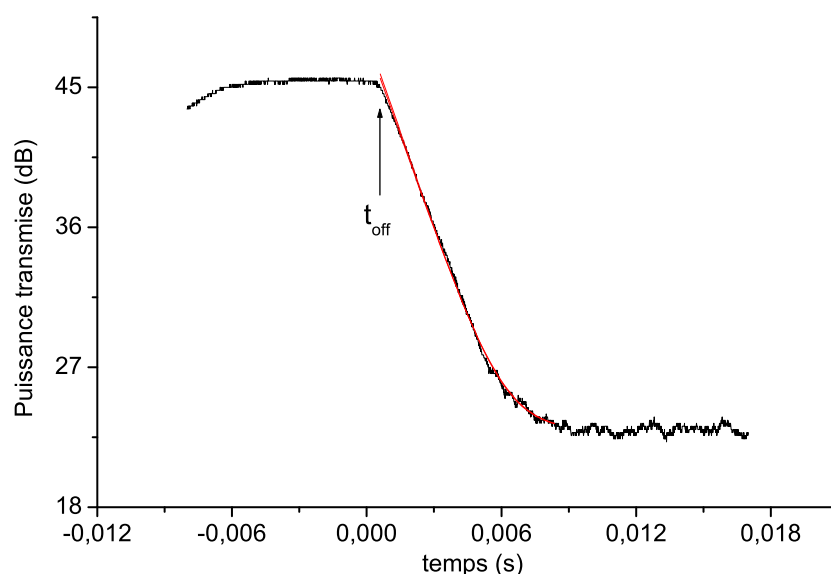


FIG. 2.14 – Mesure du temps de vie du mode HF à $51.1GHz$. La source est allumée jusqu'à l'instant t_{off} , puis l'on mesure la décroissance du signal transmis. Cette décroissance est exponentielle avec un temps caractéristique de $1ms$. Cela correspond à un facteur de qualité $Q = 3.2 \cdot 10^8$.

On montre à la figure 2.14 une mesure du temps de vie du mode HF effectuée

à l'aide du banc micro-onde. Le YIG source est asservi à la fréquence du mode HF. Il est allumé jusqu'à l'instant t_{off} . La décroissance du signal à 9.01MHz délivré par le banc reflète la décroissance du champ dans la cavité. Elle s'effectue en un temps caractéristique $T_{cav} = 1\text{ms}$. Le mode HF a donc un facteur de qualité $Q_{HF} = 3.2 \cdot 10^8$.

Les difficultés rencontrées De nombreux facteurs limitent le facteur de qualité de notre cavité. Parmi ceux-ci, on peut citer les pertes dues aux trous de couplage, la taille finie des miroirs, la résistance résiduelle du Niobium à température non-nulle pour un champ électromagnétique variable dans le temps, les défauts et impuretés de la phase cristalline du Niobium. On trouvera une analyse quantitative et détaillée de tous ces effets dans la thèse de Gilles Nogues [68]. Cette analyse montre que le facteur de qualité théorique auquel on peut prétendre malgré ces imperfections reste plus d'un ordre de grandeur au-dessus du facteur de qualité de notre résonateur ($Q_{th} \simeq 5 \cdot 10^9$ alors que nous n'avons expérimentalement "que" $Q = 10^7$). Il y a donc un autre facteur limitant.

Ce facteur limitant a été identifié comme étant l'état de surface des miroirs de la cavité. A petite échelle (quelques dixièmes de millimètre) la rugosité de la surface est de l'ordre de la centaine de nanomètres, ce qui est suffisant. Par contre, à une échelle de l'ordre de la longueur d'onde (quelques millimètres), on a observé des écarts importants par rapport à la surface idéale. Nous interprétons donc le facteur de qualité que nous mesurons comme étant dû à la diffraction de la micro-onde par les défauts de surface à grande échelle de nos miroirs. Un autre fait confirme cette interprétation : nous avons effectué des mesures de facteur de qualité des miroirs "en configuration fermée", c'est-à-dire que la cavité était complètement close, et les facteurs de qualité trouvés étaient largement supérieurs à ceux que l'on obtient lorsque la cavité est ouverte.

L'anneau Pour surmonter ces problèmes de diffraction, nous avons essayé de nous approcher le plus possible géométriquement d'une cavité fermée, tout en conservant la possibilité indispensable d'appliquer un champ directeur à l'intérieur de la cavité. Nous entourons donc la cavité d'un anneau métallique, qui limite la diffraction du mode de la cavité vers le continuum, isolé électriquement des miroirs. Cet anneau a été conçu et réalisé par Gilles Nogues et Arno Rauschenbeutel. Lorsqu'il est monté autour de la cavité, le facteur de qualité de celle-ci passe de 10^7 à $3.2 \cdot 10^8$, ce qui donne un temps de vie d' 1ms pour un photon stocké dans un mode (voir la figure 2.14).

On peut voir une photographie de l'anneau sur la figure 2.15. Il est constitué de deux demi-anneaux en Aluminium anodisé pour assurer l'isolation électrique. Les faces vues par les atomes ont été réusinées pour faire apparaître le métal et ainsi éviter toute charge électrique parasite à proximité des atomes. Deux trous circulaires de 3mm de diamètre ont été percés pour laisser passer les atomes.

L'inconvénient majeur du montage avec anneau est la taille de ces trous de couplage. En effet, le champ électrique directeur au niveau de l'anneau est très inhomogène comme on peut le constater sur la figure 2.16. Le résultat est que toute cohérence atomique est déphasée par ces champs inhomogènes.



FIG. 2.15 – Cet anneau est monté autour de la cavité et permet d'obtenir un facteur de qualité de $Q = 3.2 \cdot 10^8$. Sur la photo, on voit les trous laissant passer le jet atomique, et aussi le trou sur le côté qui permet d'effectuer les impulsions de Ramsey à l'intérieur de la structure "cavité+anneau".

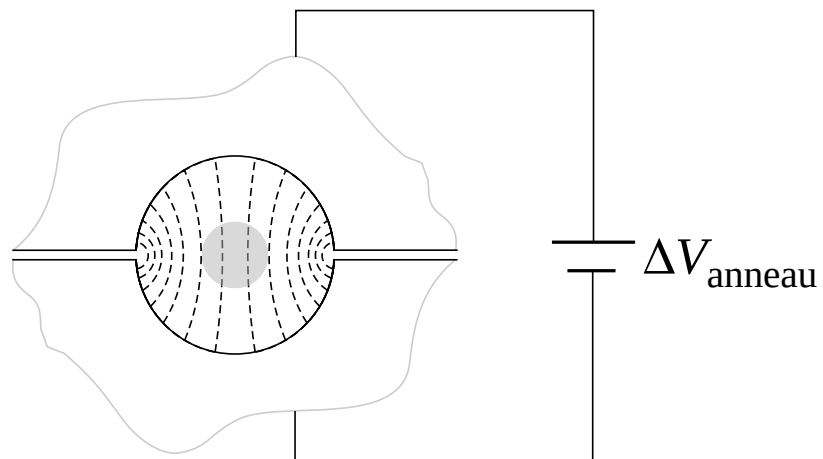


FIG. 2.16 – Champ électrique créé par l'anneau au niveau du trou.

Du coup, nous sommes contraints d'effectuer toutes les opérations cohérentes, y compris les impulsions de Ramsey, à l'intérieur de la cavité. Pour coupler la micro-onde des impulsions de Ramsey dans la cavité, un trou a été percé sur le côté de l'anneau (voir la photo 2.15). Un guide d'onde débouche au niveau de ce trou et nous permet d'effectuer les impulsions de Ramsey. On trouvera plus de détails sur le fonctionnement et le réglage de notre interféromètre de Ramsey au chapitre 3.

Une cavité sans trou ? En parallèle à l'utilisation de la cavité avec anneau pour les expériences dont nous parlons par la suite, une nouvelle technique de fabrication des miroirs était testée par Stefano Osnaghi. Elle est décrite en détail dans son mémoire de thèse [73] et je vais juste en parler très brièvement. Le Niobium est un matériau très difficile à usiner et à polir car il est très mou. Au contraire, il est beaucoup plus facile d'obtenir une surface sphérique de qualité optique avec du cuivre. On peut finalement vaporiser une fine couche de Niobium sur ces miroirs en cuivre de manière très propre (les cavités supraconductrices des accélérateurs de particules sont construites de cette manière).

Les tests effectués montrent que la qualité de la couche supraconductrice en Niobium est suffisante pour obtenir un facteur de qualité de 10^9 de manière reproductible en cavité ouverte. Par contre, le perçage des trous de couplage semble détériorer systématiquement la qualité de la surface. Nous envisageons donc d'utiliser ces miroirs sans percer de trou de couplage. Ainsi, nous pourrions atteindre un régime permettant de nouvelles expériences.

Chapitre 3

Techniques d'électrodynamique quantique en cavité

Nous venons de décrire en détail le dispositif expérimental d'électrodynamique quantique en cavité qui nous a permis de réaliser les expériences présentées aux chapitres 4 et 5. Avant d'aborder la description de ces expériences et des résultats obtenus, il nous a semblé opportun de présenter quelques techniques expérimentales que nous avons développées de manière à résoudre un certain nombre de problèmes qui se posent à chaque expérience.

Un outil essentiel de toutes les expériences d'électrodynamique quantique en cavité réalisées dans ce groupe est l'oscillation de Rabi quantique. Un atome dans $|e\rangle$ résonant avec la cavité vide échange de l'énergie périodiquement avec elle, à une fréquence qui caractérise la force de ce couplage : la fréquence de Rabi. On trouvera un traitement théorique de ce problème à l'annexe A ; dans ce chapitre nous nous intéressons uniquement à la réalisation expérimentale de l'oscillation de Rabi dans le vide. Ce sera l'occasion de montrer de quelle manière nous contrôlons l'interaction atome-cavité, et aussi comment nous préparons la cavité dans l'état vide $|0\rangle$ partant du champ thermique résiduel.

Dans la deuxième partie nous présenterons un autre outil essentiel de toutes nos expériences : l'interféromètre de Ramsey. Nous commençons par expliquer le principe de l'interférométrie de Ramsey. Il s'agit d'appliquer deux impulsions classiques R_1 et R_2 à un atome à deux niveaux $|a\rangle$ et $|b\rangle$. L'on observe alors une modulation sinusoïdale de la probabilité de détecter l'atome dans le niveau $|b\rangle$, révélant l'interférence entre deux chemins quantiques C1 et C2 indiscernables. Nous détaillons par la suite la réalisation expérimentale de cet interféromètre.

Enfin dans la troisième partie nous montrons comment cet interféromètre permet de mesurer avec une grande précision le nombre moyen de photons présents dans la cavité. Lorsque nous aurons besoin d'injecter un champ $|\alpha\rangle$ connu dans la cavité, cette propriété nous fournira une procédure de calibration du champ.

3.1 Observer les oscillations de Rabi quantiques

Rappelons les résultats de l'annexe A. Lorsqu'on envoie un atome dans l'état $|e\rangle$ à résonance avec la cavité initialement vide, la probabilité de détecter l'atome dans $|g\rangle$ est une fonction oscillante du temps d'interaction entre l'atome et la cavité : $P(g) = (1 - \cos(\Omega_0 t_{eff}))/2$. Le temps effectif t_{eff} prend en compte la forme gaussienne du couplage avec la cavité. Il est calculé à l'annexe A et vaut

$$t_{eff} = \int_{-\infty}^t \exp\left(-\frac{(vu)^2}{w^2}\right) du \quad (3.1)$$

pour un atome de vitesse v interagissant avec le mode jusqu'au temps t . Lorsque $\Omega_0 t_{eff} = \pi$ l'on parle d'"impulsion π "; lorsque $\Omega_0 t_{eff} = \frac{\pi}{2}$ d'"impulsion $\frac{\pi}{2}$ ". Pour observer les oscillations de Rabi, nous devons d'abord mettre la cavité à résonance avec la fréquence atomique, puis trouver un moyen de régler le temps d'interaction atome-cavité.

3.1.1 Contrôler l'accord atome-cavité

3.1.1.a Accorder la cavité sur la transition atomique

Les atomes de Rydberg dans le niveau $|e\rangle$ ont une durée de vie de $30ms$, donc la largeur naturelle de la transition atomique est d'environ $30Hz$. On pourrait penser qu'il est par conséquent nécessaire de contrôler la fréquence de la cavité à beaucoup mieux que $30Hz$ pour l'accorder sur les atomes. Heureusement il n'en est rien, car la raie atomique dans la cavité est élargie à cause de la durée finie de l'interaction entre l'atome et la cavité ¹. Typiquement, pour des atomes ayant une vitesse de $500m/s$, le temps effectif d'interaction est de $20\mu s$, ce qui donne une largeur de $50kHz$ comme nous allons le constater par la suite. Il est donc largement suffisant dans notre expérience de contrôler la fréquence de la cavité à l'échelle du kiloHertz.

Nous pouvons accorder la cavité sur une plage de $400kHz$ en appliquant une tension V_{PZT} aux cales piézoélectriques situées sous l'un des deux miroirs. Ainsi, on change légèrement la distance entre les deux miroirs, et donc également la fréquence des deux modes qui nous intéressent selon la formule 2.7. Nous obtenons une résolution de l'ordre d'une dizaine de Hertz. La cavité a en outre des dérives lentes qui font que, sur une durée de quelques heures, sa fréquence n'est maîtrisée qu'à l'échelle du kiloHertz dont nous venons toutefois de voir qu'elle était suffisante pour bien accorder la cavité sur les atomes. (Notons tout de suite que ce mode d'accord est "lent" car à cause de l'inertie due à la masse importante du miroir, la cale piézoélectrique ne répond pas à une modulation plus rapide que le kiloHertz. Or la résolution temporelle nécessaire, pour accorder la

¹La raie atomique est aussi élargie à cause du couplage fort avec la cavité. Cet effet est cependant moins important que la durée finie de l'interaction. A cause de la forme gaussienne du mode, les atomes se couplent et se découplent adiabatiquement au mode lorsqu'ils ne sont pas parfaitement résonants avec celui-ci.

cavité au cours d'une séquence expérimentale, serait de 100ns. Donc nous ne modifions V_{PZT} qu'entre deux séquences expérimentales.)

Pour accorder la cavité sur les atomes, nous effectuons une expérience similaire au micromaser [65, 20]. Nous envoyons des atomes dans l'état excité $|e\rangle$ à-travers la cavité initialement à l'équilibre thermique. Si l'un des deux modes de la cavité est résonant avec les atomes, ils vont échanger de l'énergie par des oscillations de Rabi, et l'atome sortira dans $|g\rangle$ avec une probabilité importante, alors qu'il restera dans $|e\rangle$ si aucun mode n'est résonant. La figure 3.1 montre un exemple d'un tel accord de la cavité sur les atomes, qui est une opération quotidienne lors du réglage d'une expérience.

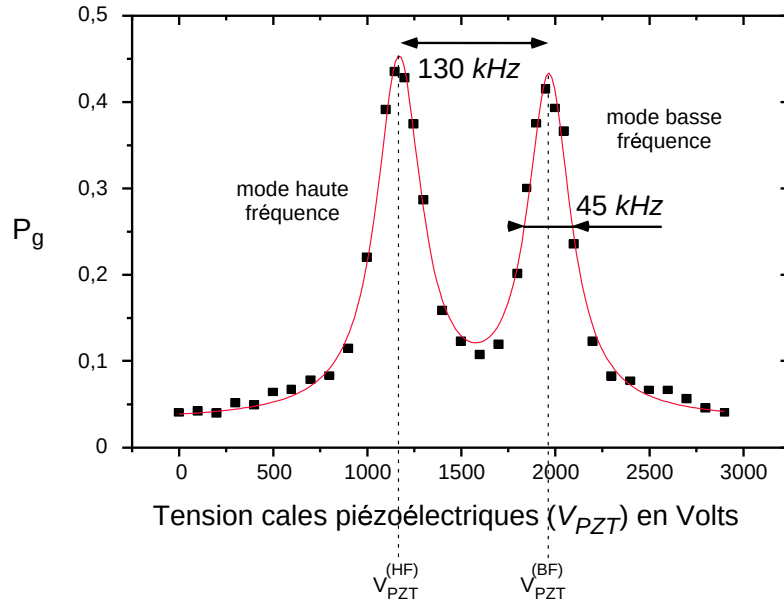


FIG. 3.1 – Probabilité pour détecter l'atome dans $|g\rangle$ en fonction de la tension V_{PZT} appliquée aux cales piézoélectriques (la fréquence atomique est fixée par le choix du champ Stark $E_{Stark}^{(1)} = 0.25V/cm$, $\nu_{eg} = \nu(E_{Stark}^{(1)})$). Lorsqu'on augmente V_{PZT} , les miroirs se rapprochent, et la fréquence de résonance des deux modes augmente. Les deux pics observés correspondent aux valeurs de V_{PZT} pour lesquelles le mode HF, puis le mode BF, passent à résonance avec l'atome. Le trait continu est un ajustement numérique réalisé avec deux lorentziennes.

Sur cette figure on voit nettement qu'il y a deux valeurs de V_{PZT} pour lesquelles l'atome est résonant soit avec le mode Haute Fréquence (HF), soit avec le mode Basse Fréquence (BF). On constate aussi que l'écart entre les deux modes est suffisant pour négliger un mode lorsque l'atome est résonant avec l'autre. Cette méthode permet d'accorder la cavité à la fréquence atomique avec une précision de l'ordre du kHz. Selon que l'on souhaite effectuer l'oscillation de Rabi sur le mode Haute Fréquence (HF) ou Basse Fréquence (BF) on fixera dans la suite la tension V_{PZT} à l'une des deux valeurs visibles sur la figure 3.1 $V_{PZT}^{(BF)}$ ou $V_{PZT}^{(HF)}$.

Remarquons enfin qu'il est bien sûr important de connaître la relation entre tension piézo et fréquence des modes. Nous effectuons cette calibration en mesurant la fréquence des modes de la cavité grâce à notre banc micro-onde (voir chapitre 2) pour diverses tensions V_{PZT} . Le résultat, montré à la figure 3.2, est une relation linéaire $\Delta\nu_{cav}(kHz) = 0.159\Delta V_{PZT}(V)$.

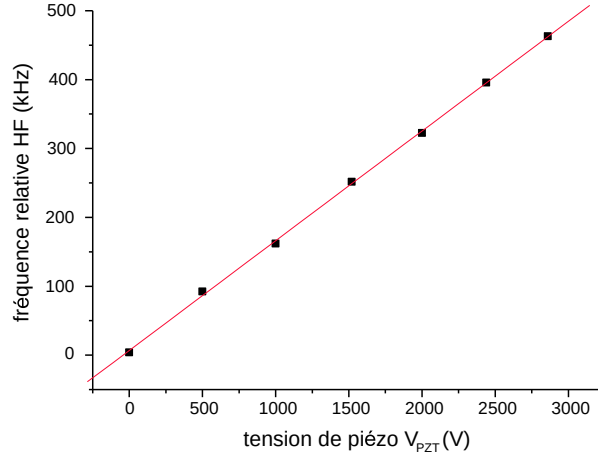


FIG. 3.2 – Fréquence mesurée du mode Haute Fréquence HF en fonction de la tension de V_{PZT} appliquée aux cales piézoélectriques sous un des miroirs de la cavité. L'ajustement linéaire donne $\Delta\nu_{cav}(kHz) = 0.159\Delta V_{PZT}(V)$

3.1.1.b Contrôler le temps d'interaction

L'idée la plus simple pour choisir un temps d'interaction effectif donné t_{eff} consiste à choisir des atomes dont la vitesse est telle qu'ils n'interagissent avec la cavité que pendant t_{eff} . Malheureusement cette idée n'est pas praticable dans notre dispositif, car pour avoir des temps d'interaction courts ($10\mu s$ par exemple, ce qui correspondrait à une impulsion π), nous aurions besoin d'avoir des atomes allant à $1000m/s$; or la température de notre four ne nous fournit que des atomes ayant des vitesses entre $100m/s$ et $600m/s$.

Nous utilisons donc le champ électrique directeur, chargé de maintenir l'orbite circulaire pendant le passage des atomes dans la cavité, pour modifier la fréquence atomique par effet Stark au cours de son trajet. De cette manière nous pouvons accorder la fréquence atomique en temps réel, au cours de chaque séquence expérimentale. Notamment, pour observer une oscillation de Rabi avec des temps d'interaction courts, il suffit de laisser l'atome en interaction avec la cavité le temps voulu, puis de le désaccorder brusquement en augmentant le champ électrique à un instant t_{Stark} , afin de "geler" l'évolution du système couplé atome-champ. De manière équivalente, on pourrait choisir de laisser l'atome hors résonance à son entrée dans le mode, et de ne l'accorder qu'à l'instant t_{Stark} .

Reste à choisir pour quelle valeur du champ E_{Stark} on souhaite que l'atome soit résonant avec la cavité. A priori, n'importe quelle valeur conviendrait. Dans la pratique pourtant il est important que la valeur du champ $E_{Stark}^{(1)}$ pour laquelle l'atome est à résonance avec la cavité soit la plus faible possible (mais elle doit être malgré cela suffisante pour assurer le maintien de l'orbite circulaire), car l'atome est beaucoup plus sensible aux inhomogénéités de champ électrique lorsque le champ moyen est important. Supposons en effet que le champ ait une valeur moyenne $\bar{\mathbf{E}}$ et des fluctuations $\delta\mathbf{E}$. La fréquence atomique varie comme

$$\Delta\nu = -\alpha\mathbf{E}^2, \quad (3.2)$$

α valant $255kHz/(V/cm)^2$ (voir chapitre 2). Dans le champ total $\mathbf{E} = \bar{\mathbf{E}} + \delta\mathbf{E}$, la fréquence de l'atome varie de

$$\Delta\nu = -\alpha(\bar{\mathbf{E}} + \delta\mathbf{E})^2 \simeq -\alpha\bar{\mathbf{E}}^2 - 2\alpha\delta\mathbf{E} \cdot \bar{\mathbf{E}} = \overline{\Delta\nu} + |\bar{\mathbf{E}}|\delta\nu(\delta\mathbf{E}) \quad (3.3)$$

D'après l'équation précédente, doubler la tension Stark sur les miroirs revient à doubler la sensibilité des atomes aux gradients de champ. Lorsque l'atome est dans un état stationnaire, cela n'a aucune conséquence sur son évolution ; par contre lorsque l'atome est dans une superposition d'états, au cours d'une oscillation de Rabi ou dans un interféromètre de Ramsey, il faut impérativement éviter les facteurs de déphasage, sous peine d'une réduction importante du contraste du signal.

On applique donc généralement la plus faible valeur du champ électrique Stark compatible avec un maintien de l'orbite circulaire lorsque l'atome interagit avec la cavité, $E_{Stark}^{(1)} = 0.25V/cm$.

3.1.2 Refroidir la cavité avant une expérience

Nous avons mesuré, par une méthode décrite dans [68], le champ thermique à l'équilibre comme étant d'environ 1 photon dans chaque mode. Cette valeur importante est due à deux causes distinctes :

- les miroirs étant refroidis à environ $1.3K$ se comportent comme un corps noir à cette température et rayonnent à la fréquence de la cavité ($51.1GHz$) 0.2 photons thermiques par mode.
- d'autre part, l'isolation du guide d'ondes "Ramsey" qui arrive sur le côté de l'anneau est imparfaite et occasionne des fuites de champ thermique de l'extérieur vers le mode de la cavité.

Il est important de se débarrasser de ce champ avant chaque expérience ; en effet, nous souhaitons contrôler l'état quantique de la cavité le mieux possible, or en présence des photons thermiques la cavité est dans un mélange statistique d'états. Pour cela, nous envoyons plusieurs paquets contenant chacun en moyenne 3 atomes dans l'état fondamental et on les fait interagir avec la cavité (ces paquets ont été rapidement baptisés "serpillères" dans notre équipe). Nous ne cherchons pas à détecter ces atomes ;

tout ce que nous pouvons en dire est qu'ils constituent, pour la cavité, un thermostat à température nulle. Lors de l'interaction avec la cavité, il y a contact thermique entre les deux milieux, donc la température moyenne de la cavité baisse un peu à chaque passage de serpillères. Après passage des serpillères, la cavité relaxe vers le champ thermique à l'équilibre en un temps de l'ordre du temps de vie de la cavité T_{cav} . Toutes nos expériences se déroulent dans un temps qui doit être plus court que le retour à l'équilibre de la cavité. (cette contrainte existe aussi pour les ions piégés, où les expériences doivent être précédées d'une séquence de refroidissement par bandes latérales, et avoir lieu avant que les ions ne soient réchauffés par des causes extérieures [64]).

Nous avons deux possibilités différentes pour faire interagir les serpillères avec la cavité :

- soit nous les laissons à résonance avec le mode que nous souhaitons refroidir pendant toute la durée de leur interaction. On pourra voir sur la figure 3.3 le déroulement typique d'une séquence de refroidissement. Six paquets atomiques, contenant entre 6 et 9 atomes chacun, sont envoyés successivement à résonance avec le mode qui est refroidi. Un atome sonde l'état du champ $100\mu s$ plus tard. On notera sur cette figure l'existence des "échos" déjà mentionnés au chapitre précédent à-propos de la sélection de vitesse. On trouvera plus de détails sur l'efficacité de cette procédure de refroidissement dans la thèse de Gilles Nogues [68]. Contentons-nous de dire que l'on a mesuré, après le passage des "serpillères résonantes", un champ thermique résiduel d'environ 0.05 photons.

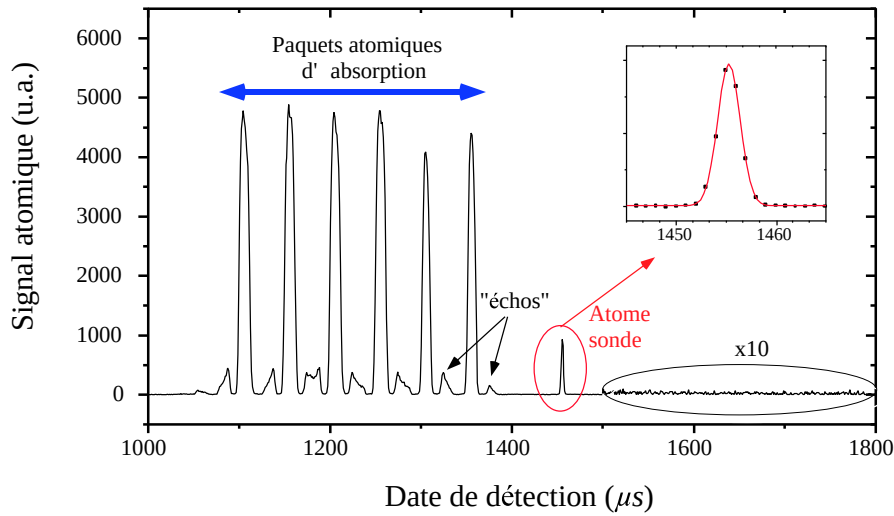


FIG. 3.3 – Signal atomique en fonction du temps d'arrivée sur le détecteur. On distingue clairement 6 paquets atomiques séparés de $50\mu s$ et contenant entre 6 et 9 atomes en moyenne. Ils servent à l'absorption du champ. Un septième paquet (l'insert le montre en détail) contient 0,3 atomes en moyenne et sert à sonder l'état du champ $100\mu s$ après le dernier paquet absorbant. On distingue entre chaque paquet absorbant des "échos", parasites environ 15 fois moins grands.

- Pourtant, cette procédure de refroidissement n'est pas extrêmement efficace. Lors-

qu'on envoie des atomes à résonance avec le mode, on sature la transition atomique. Lorsque la cavité contient des photons, une moitié seulement des atomes sort dans $|e\rangle$ et participe donc au refroidissement. En envoyant successivement plusieurs paquets atomiques, on arrive à refroidir la cavité. Mais il serait préférable de supprimer le champ thermique résiduel en une seule fois. C'est possible, en faisant effectuer au paquet atomique de refroidissement un passage adiabatique à-travers le mode. Le principe est exposé à la figure 3.4.

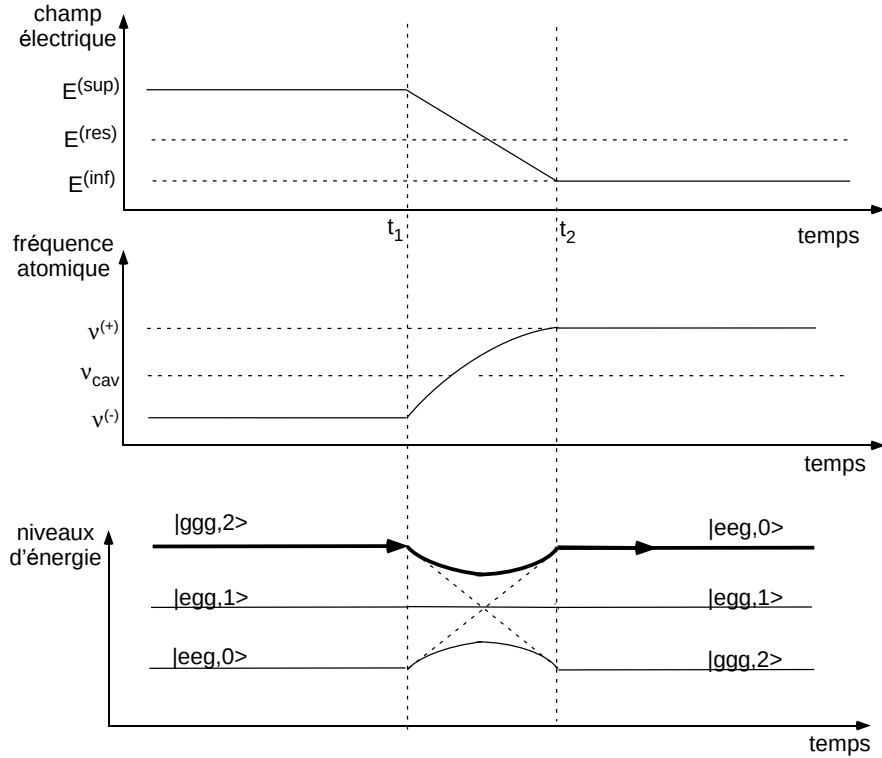


FIG. 3.4 – Principe des “serpillères adiabatiques”. On envoie un paquet atomique contenant des atomes à 500m/s initialement dans l'état fondamental. Le champ électrique Stark est initialement fixé à une valeur $E^{(sup)} = 0.78\text{V/cm}$ qui fixe la fréquence des atomes du paquet de serpillères à une valeur $\nu^{(-)}$ inférieure à la fréquence de la cavité de 60kHz . A l'instant $t_1 = -10\mu\text{s}$ (origine des temps prise au centre de la cavité), on abaisse lentement le champ électrique jusqu'à atteindre $E^{(inf)} = 0.25\text{V/cm}$ à l'instant $t_2 = +10\mu\text{s}$. La fréquence de la transition atomique est alors supérieure de 60kHz à la fréquence de la cavité. L'état $|ggg, 2\rangle$ se couple adiabatiquement à l'état $|eeg, 0\rangle$.

On envoie un seul paquet atomique à-travers la cavité. Il contient en moyenne 3 atomes dans l'état fondamental. Le champ électrique Stark est initialement fixé à une valeur $E^{(sup)} = 0.78\text{V/cm}$ qui fixe la fréquence de la transition atomique à une valeur $\nu^{(-)}$ inférieure à la fréquence de la cavité de 60kHz . A l'instant $t_1 = -10\mu\text{s}$ (l'origine des temps est prise au centre de la cavité), on abaisse lentement le champ

électrique jusqu'à atteindre $E_{(inf)} = 0.25V/cm$ à l'instant $t_2 = +10\mu s$. La fréquence de la transition atomique est alors supérieure de $60kHz$ à la fréquence de la cavité. Si le passage à résonance se fait suffisamment lentement, tous les photons devraient être absorbés avec une efficacité de 100%. En effet, supposons que la cavité contient N photons, et qu'il y ait un nombre n supérieur à N d'atomes dans le paquet. L'état $|ggg\dots g, N\rangle$ se couple adiabatiquement par un anticroisement de N niveaux à l'état dans lequel N atomes ont absorbé un photon $|eee\dots g, 0\rangle$ ². On a montré sur la figure 3.4 l'exemple de l'anticroisement du niveau $|ggg, 2\rangle$ avec les autres niveaux de la multiplicité qui le couple adiabatiquement au niveau $|eeg, 0\rangle$.

Nous avons testé expérimentalement cette idée en envoyant successivement un paquet de "serpillère adiabatique" que nous avons fait interagir comme dit précédemment avec la cavité, suivi $100\mu s$ plus tard d'une sonde atomique initialement dans $|g\rangle$ et résonante avec le mode (voir la figure 3.5). En mesurant son taux de transfert vers le niveau $|e\rangle$, on obtient une mesure du champ résiduel au moment du passage de la sonde dans la cavité.

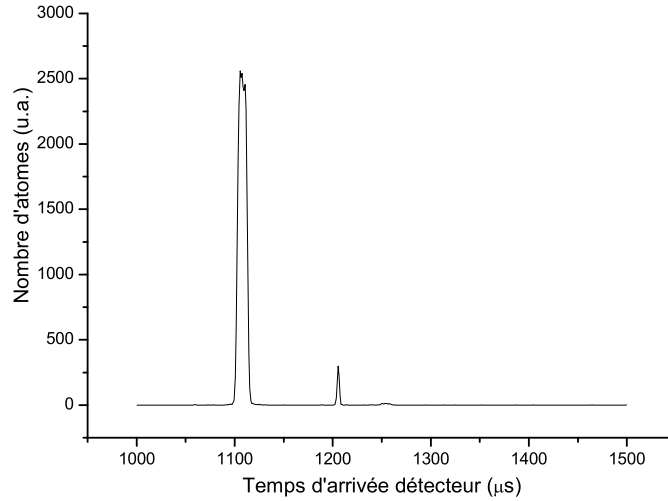


FIG. 3.5 – Un seul paquet contenant en moyenne 3 atomes est suffisant pour vider le mode de la cavité de ses photons thermiques. Le deuxième paquet atomique, beaucoup plus petit (en moyenne 0.15 atomes), sert à sonder l'état du champ $100\mu s$ après le passage de la serpillère.

Les résultats sont exposés à la figure 3.6. On y a reporté l'absorption de la sonde en fonction du nombre d'atomes présents dans le paquet de serpillères. On constate

²En fait, l'état ne serait pas celui qui est annoncé mais plutôt l'état symétrique où n quanta d'énergie sont distribués parmi les atomes. Par exemple, il faut lire $|eeg\rangle \equiv (|eeg\rangle + |ege\rangle + |gee\rangle)/\sqrt{3}$. Mais ce qui importe ici, c'est que la cavité soit vide à la fin du processus.

qu'au-dessus d'un certain flux, l'absorption de la sonde reste stable à 14%. Le champ dans la cavité est donc de 0.15 photons au passage de la sonde, et compte tenu de la relaxation du champ dans la cavité vers l'équilibre thermique, cette valeur indique qu'il y a environ 0.05 photons thermiques juste après le passage des serpillères (c'est la valeur la plus faible que nous obtenions avec les serpillères résonantes). Ainsi, nous avons montré qu'il était possible de refroidir efficacement la cavité grâce à ces serpillères adiabatiques. Cette méthode est particulièrement utile pour les expériences où nous devons travailler avec des atomes lents (par exemple 150m/s). En effet, nous sommes alors obligés d'utiliser des serpillères plus rapides (par exemple 450m/s), car cela nous permet de réduire autant que possible le délai entre la dernière serpillère et l'atome "de mesure". Dans ces conditions, il n'est pas possible d'envoyer suffisamment de paquets atomiques pour pouvoir utiliser des serpillères résonantes.

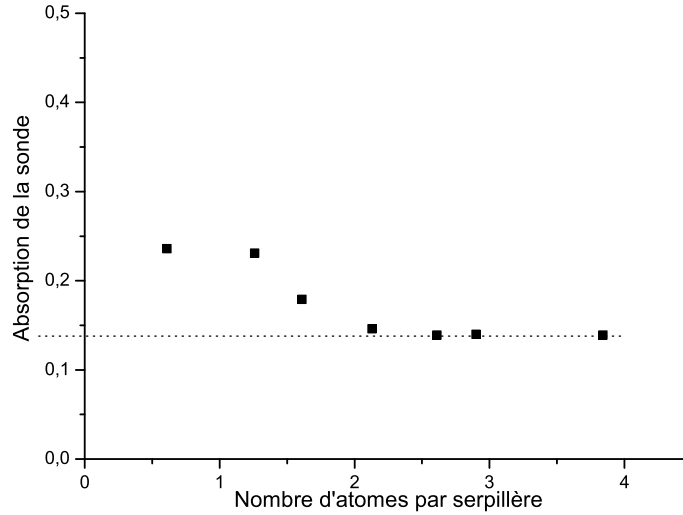


FIG. 3.6 – Absorption du paquet sonde $100\mu\text{s}$ après le passage de la serpillaire en fonction du nombre moyen d'atomes N_{serp} dans le paquet. On voit que lorsque $N_{\text{serp}} > 2.5$, l'absorption reste stable à 14%. Cette valeur est compatible avec le fait que le champ résiduel juste après le passage des serpillères est d'environ 0.05 photons.

3.1.3 Oscillations de Rabi

Nous pouvons maintenant décrire toutes les étapes qui nous permettent d'observer une oscillation de Rabi quantique. On commence par accorder le mode Basse Fréquence (BF) à la fréquence atomique dans $E_{\text{Stark}}^{(1)}$ comme expliqué dans le paragraphe 3.1.1.a. Puis on envoie un atome préparé dans $|e\rangle$ à 500m/s dans la cavité, précédé d'une séquence de refroidissement résonante du type de celle montrée à la figure 3.3. A l'instant

t_{Stark} on commute le champ électrique à la valeur $E_{Stark}^{(2)}$, et l'on étudie la probabilité de détecter l'atome dans $|g\rangle$ en fonction de t_{Stark} . La figure 3.7 montre les résultats obtenus.

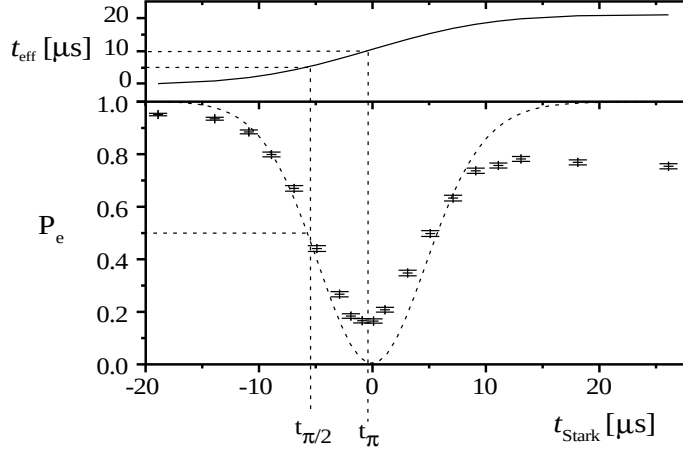


FIG. 3.7 – Oscillation de Rabi dans le vide. On envoie un atome dans $|e\rangle$, résonant avec la cavité jusqu'à l'instant t_{Stark} où l'on désaccorde l'atome brusquement par effet Stark. La courbe en pointillés représente la courbe théorique de l'oscillation de Rabi sans aucune dissipation. On a représenté dans le cadre supérieur le temps efficace d'interaction. Les temps d'interaction correspondant aux impulsions $\pi/2$ et π sont indiqués par des pointillés. Les temps effectifs sont alors respectivement de $5\mu s$ et $10\mu s$, ce qui correspond bien à une fréquence de Rabi au centre de la cavité de $50kHz$.

On observe un bon accord qualitatif entre les données et le résultat théorique. Sur la figure sont repérés les temps d'interaction t_{π} et $t_{\pi/2}$ correspondant respectivement aux impulsions π et $\pi/2$. Ils correspondent à des temps d'interaction effectifs respectivement de $5\mu s$ et de $10\mu s$.

A la suite d'une impulsion π l'atome devrait finir dans $|g\rangle$ dans 100% des cas. On constate sur la figure 3.7 que ce n'est pas tout-à-fait le cas. Dans 15% des cas environ l'atome ressort dans $|e\rangle$, à cause du champ thermique résiduel (lorsque l'atome effectue l'oscillation de Rabi dans la cavité, le champ thermique a déjà commencé à repousser puisque la dernière "serpillère" est passée $100\mu s$ plus tôt), et des dispersions géométriques de l'expérience.

3.2 L'interféromètre de Ramsey

L'oscillation de Rabi est la manifestation du couplage fort entre l'atome et la cavité. C'est ce couplage fort qui nous permettra de réaliser au chapitre 4 un interféromètre de Ramsey à la frontière classique/quantique. Pour l'instant, nous allons décrire le fonctionnement "normal" d'un interféromètre de Ramsey dans notre dispositif.

3.2.1 Principe de fonctionnement

Nous avons déjà évoqué le principe général de fonctionnement d'un interféromètre de Ramsey au chapitre 1. Une source externe S_{Ram} , accordée à la fréquence d'une transition atomique entre deux niveaux $|a\rangle$ et $|b\rangle$, est allumée brièvement à deux instants distincts t_1 et t_2 lorsque l'atome se trouve au niveau des zones de Ramsey. Notons que dans notre dispositif expérimental, les états $|a\rangle$ et $|b\rangle$ pourront être n'importe lequel des trois états $|e\rangle$, $|g\rangle$, ou $|i\rangle$. En effet, comme nous le verrons, nous sommes capables d'effectuer des franges de Ramsey sur les transitions g-e, mais aussi g-i et e-i. L'analyse ci-dessous est bien sûr la même pour chacune de ces transitions.

Faisons une étude quantitative de l'évolution de l'atome dans l'interféromètre. La source S_{Ram} émet une onde monochromatique $\mathbf{E}_i(t) = \mathbf{E}_i \sin(\omega t)$, avec $\omega = \omega_{ab}$. (l'indice i vaut 1 ou 2 selon que l'on considère la zone R1 ou R2). L'évolution atomique dans le champ $\mathbf{E}(t)$ est bien décrite par un modèle semi-classique où seul l'atome est quantifié, avec un hamiltonien $\hat{H}_{sc} = \hat{H}_{at} - \hat{\mathbf{d}} \cdot \mathbf{E}_i \cos(\omega t)$. En représentation d'interaction, le hamiltonien vaut

$$\hat{H}_{int} = \begin{pmatrix} 0 & \Omega_i \\ \Omega_i & 0 \end{pmatrix} \quad (3.4)$$

Le paramètre Ω_i (que l'on prendra réel et positif) est la pulsation de Rabi de l'atome dans le champ créé par S_{Ram} . Il caractérise la force du couplage entre l'atome et le champ, et vaut $\Omega_i = |\langle a | \hat{\mathbf{d}} \cdot \mathbf{E}_i | b \rangle| / \hbar$.

L'évolution de l'atome se déduit facilement de la relation 3.4. On a :

$$\begin{aligned} |a\rangle & \xrightarrow{S_{Ram}} \cos\left(\frac{\Omega_i t_i^{int}}{2}\right) |a\rangle - i \sin\left(\frac{\Omega_i t_i^{int}}{2}\right) |b\rangle \\ |b\rangle & \xrightarrow{S_{Ram}} \cos\left(\frac{\Omega_i t_i^{int}}{2}\right) |b\rangle - i \sin\left(\frac{\Omega_i t_i^{int}}{2}\right) |a\rangle \end{aligned} \quad (3.5)$$

Dans les zones de Ramsey R1 et R2 on règle les temps d'allumage de la source t_i^{int} et l'intensité du couplage Ω_i de manière à égaliser les populations atomiques dans les deux niveaux; on parle d'une "impulsion $\frac{\pi}{2}$ " (en effet le produit $\Omega_i t_i^{int}$ vaut alors $\frac{\pi}{2}$). Ecrivons l'évolution du système lors d'une impulsion $\frac{\pi}{2}$ en utilisant l'équation 3.5 :

$$\begin{aligned} |a\rangle & \xrightarrow{\frac{\pi}{2}} \frac{1}{\sqrt{2}}(|a\rangle - i|b\rangle) \\ |b\rangle & \xrightarrow{\frac{\pi}{2}} \frac{1}{\sqrt{2}}(|b\rangle - i|a\rangle) \end{aligned} \quad (3.6)$$

Etudions l'évolution de l'atome dans l'interféromètre. Il est initialement préparé dans $|a\rangle$. Puis, comme nous l'avons vu plus haut,

$$|a\rangle \xrightarrow{R_1} |\psi_1\rangle = \frac{1}{\sqrt{2}}(|a\rangle - i|b\rangle) \quad (3.7)$$

Cet état est stationnaire en représentation d'interaction. Si l'on introduit un déphasage supplémentaire $\Delta\Phi$ entre les deux composantes de la fonction d'onde, par une méthode

que nous présenterons par la suite, l'état de l'atome juste avant R2 est

$$|\psi_2\rangle = \frac{1}{\sqrt{2}} (e^{i\Delta\Phi}|a\rangle - i|b\rangle) \quad (3.8)$$

(notons qu'il serait bien sûr équivalent de changer la phase du champ d'une quantité $-\Delta\Phi$ entre R1 et R2 en laissant la phase du dipôle atomique inchangée). Calculons l'état final, après R2 :

$$|\psi_{fin}\rangle = \frac{1}{2} [(e^{i\Delta\Phi} - 1)|a\rangle - i(1 + e^{i\Delta\Phi})|b\rangle] \quad (3.9)$$

et la probabilité de détecter l'atome dans l'état $|b\rangle$ vaut

$$P(b) = |\langle b|\psi_{fin}\rangle|^2 = \frac{1}{2}(1 + \cos(\Delta\Phi)) \quad (3.10)$$

Nous allons maintenant présenter en détail la réalisation expérimentale de cet interféromètre.

3.2.2 Les franges de Ramsey

Description générale de l'interféromètre Rappelons que toute cohérence atomique est brouillée lors du passage des atomes par le trou de l'anneau à cause de champs électriques inhomogènes comme nous l'avons vu au chapitre 2. Par conséquent toutes les opérations cohérentes, y compris bien sûr les impulsions R1 et R2, doivent être faites à l'intérieur de la structure "anneau+cavité". Selon la transition choisie, on asservit la source externe S_{Ram} à la fréquence $\nu_{ge} = 51.099GHz$, ou bien $\nu_{gi} = 54.259GHz$, ou encore $\nu_{ei} = (\nu_{eg} + \nu_{gi})/2 = 52.794GHz$. S_{Ram} est couplée aux atomes via le guide d'onde Ramsey, comme schématisé sur la figure 3.8. Notons sur le schéma la présence d'un interrupteur à commande digitale, la "diode PIN", qui permet d'allumer ou éteindre la source avec un temps de réponse d'une dizaine de nanosecondes, largement inférieure à la résolution temporelle de nos expériences. Lorsque la source est allumée, elle excite des modes de la structure quasi-fermée "cavité+anneau", et crée donc une onde stationnaire avec des noeuds et des ventres. Comme ces modes ont un très faible facteur de qualité, le champ s'établit et disparaît instantanément lorsque l'on allume et éteint la source S_{Ram} . Reste à choisir précisément les zones de Ramsey.

Les sources micro-onde Les impulsions micro-onde de Ramsey créent ou analysent des cohérences entre les niveaux atomiques. Elles sont donc des outils précieux et doivent obéir à des exigences techniques rigoureuses. J'ai travaillé à l'amélioration du dispositif micro-onde, notamment en réalisant avec l'aide de Mathieu Perrin, étudiant de Polytechnique, un troisième dispositif d'asservissement qui nous a servi pour réaliser les expériences présentées au chapitre 5.

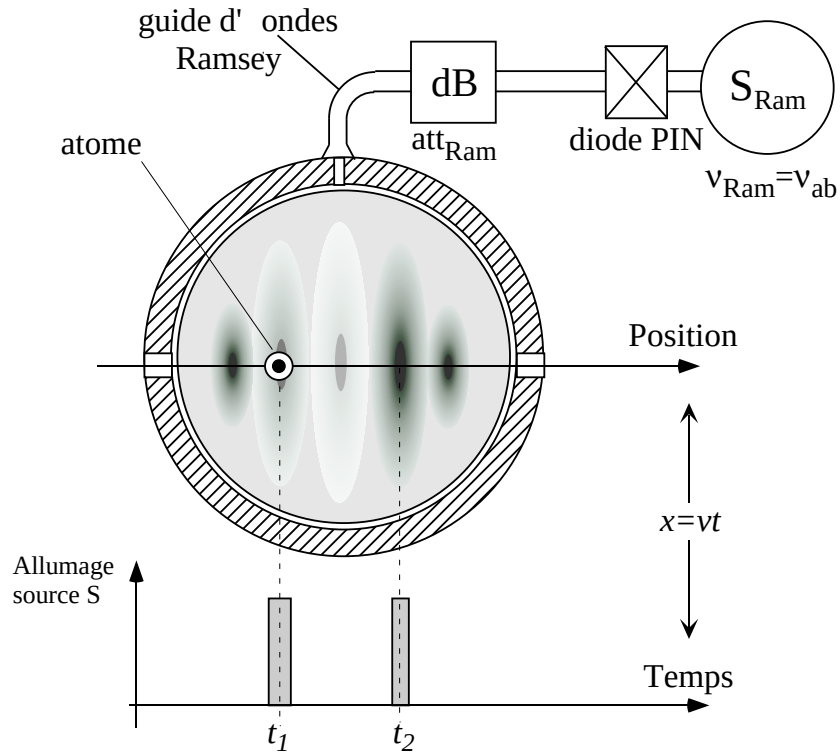


FIG. 3.8 – Schéma de l'interféromètre de Ramsey. La source S_{Ram} est couplée aux atomes via un guide d'onde latéral. La diode PIN est un interrupteur à commande digitale qui allume et éteint S_{Ram} . Lorsqu'elle est allumée, S_{Ram} crée une onde stationnaire dans la structure "cavité+anneau". La source est allumée aux instants t_1 et t_2 où l'atome est respectivement au niveau des première et deuxième zone de Ramsey. On peut régler l'intensité de l'impulsion grâce à l'atténuateur de précision att_{Ram} situé à l'extérieur du cryostat.

- Pour réaliser des franges de Ramsey, il est indispensable que la phase de la source S_{Ram} qui effectue les impulsions de Ramsey soit stable plus longtemps que la durée entre deux impulsions. En général, cette durée est de l'ordre de $100\mu s$. La source doit donc avoir une largeur spectrale inférieure au kiloHertz. En outre, il faut que l'amplitude des impulsions fluctue le moins possible d'une expérience à l'autre.

- Par ailleurs, nos atomes sont extrêmement sensibles aux champs micro-onde, donc nous devons les protéger le mieux possible contre les fuites micro-onde aux moments où ils ne doivent pas voir de radiation.

La solution adoptée dans notre équipe consiste à utiliser des sources à des fréquences typiques de $13GHz$ (soit quatre fois moins que les fréquences des transitions atomiques). Nous utilisons une source YIG (Yttrium-Grenat) commerciale qui est accordable entre 8 et $18GHz$ par simple application d'un champ magnétique externe et qui a une largeur d'environ $200kHz$ lorsqu'elle oscille librement. On l'asservit en phase sur un quartz à $100MHz$ dont le bruit de phase est très faible ($-130dBc$ à $100Hz$), lui-même asservi sur un quartz à $10MHz$ ultrastable (stabilité de 10^{-11} par jour) qui sert de référence de fréquence à toute l'expérience. Le schéma de l'asservissement que j'ai mis au point est montré à la figure 3.9. Il permet d'asservir le YIG à n'importe laquelle des fréquences atomiques qui nous sont utiles.

On commence par régler la fréquence du YIG à proximité de la fréquence souhaitée. Une tête d'échantillonnage de la société "Omega Technologies" génère des harmoniques d'ordre élevé du quartz à $100MHz$, ainsi que le battement de ce "peigne de fréquences" avec le YIG. Nous filtrons le signal de sortie de la tête d'échantillonnage de manière à ne conserver que le battement entre le YIG et l'harmonique du quartz $100MHz$ la plus proche. Par exemple, si l'on souhaite asservir le YIG à la fréquence de la transition $g-i$ (divisée par 4) soit $\nu_{g-i} = 13.564GHz$, on gardera le battement à $36MHz$ du YIG avec la 136^{me} harmonique du quartz à $100MHz$. Ce signal entre ensuite sur une carte PLL (Phase Lock Loop) fournie et conçue par Michel Gross pour la société ABmm. Il est comparé à un signal de référence issu d'un synthétiseur numérique référencé lui aussi sur le quartz $100MHz$. Le signal d'erreur généré par la carte est envoyé dans la bobine qui commande la fréquence du YIG. Comme la fréquence du synthétiseur est programmée par l'ordinateur de contrôle de l'expérience, on peut balayer la source YIG sur une plage d'environ $5MHz$ soit $20MHz$ pour le signal qui entre dans l'expérience.

Le signal délivré par le YIG asservi, à une fréquence ν_s , est ensuite multiplié par 4 à l'entrée du cryostat grâce à un mélangeur harmonique de la société "Pacific Millimeter". Avant le mélangeur, une diode PIN commande l'allumage de la source. Soulignons l'intérêt de n'avoir aucune source micro-onde directement à la fréquence d'une transition atomique ν_{at} , et d'utiliser un mélangeur harmonique dont la réponse est très non-linéaire. D'une part, lorsque la source est allumée, le mélangeur sature, et donc l'amplitude du signal émis à ν_{at} fluctue moins que le signal émis par le YIG. D'autre part, lorsque la source est éteinte, la non-linéarité du mélangeur améliore encore l'extinction assurée par la diode PIN. Si nous avions des sources émettant directement aux fréquences ν_{at} , nous serions incapables d'éteindre correctement ces sources, et cela

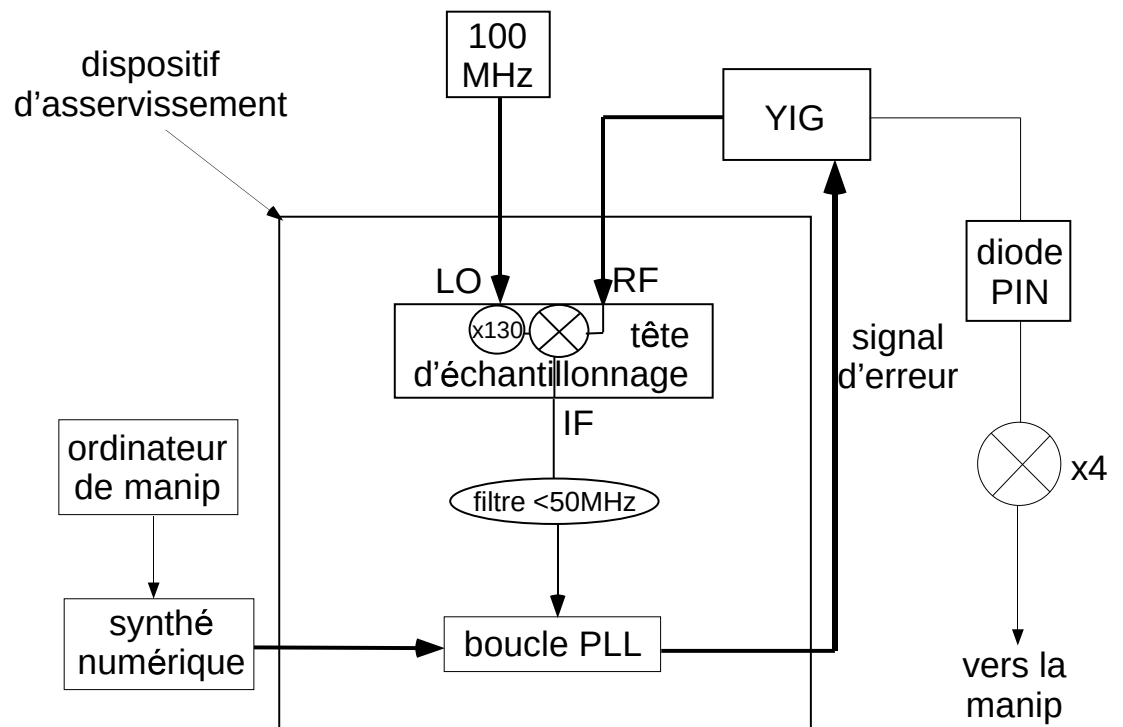


FIG. 3.9 – Schéma de principe de l'asservissement des sources micro-onde utilisées pour les impulsions de Ramsey. Une partie du signal de sortie du YIG est prélevée. On génère son battement avec une harmonique d'ordre élevé du quartz 100 MHz. Après filtrage, ce signal est comparé avec un signal de référence issu d'un synthétiseur numérique à une fréquence choisie par l'expérimentateur sur une carte PLL. Le signal d'erreur généré est réinjecté dans les bobines qui contrôlent la fréquence du YIG.s

conduirait inmanquablement à des fuites micro-onde que nos atomes détecteraient.

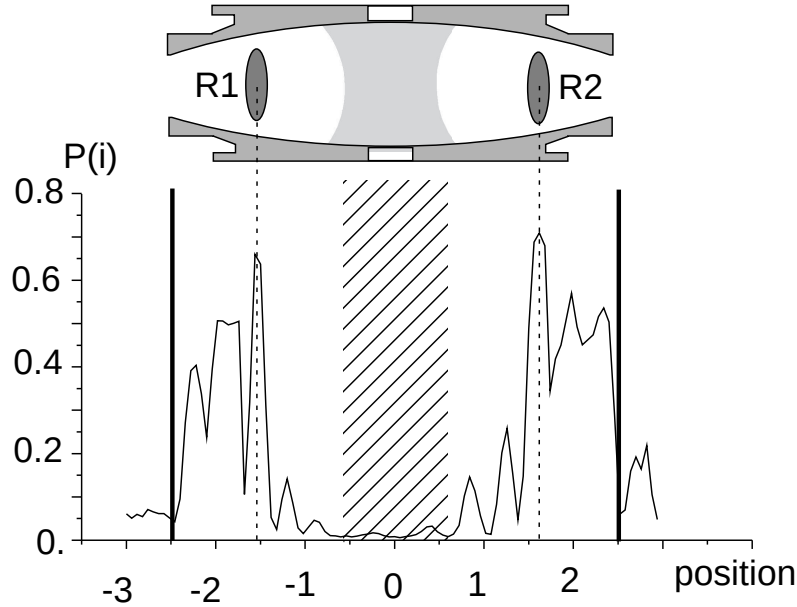


FIG. 3.10 – Carte du champ sur la transition $e-i$. On sonde l'intensité du champ micro-onde en mesurant la probabilité de détecter dans $|i\rangle$ des atomes initialement dans $|e\rangle$ en fonction de leur position dans la cavité au moment d'une impulsion très brève ($2\mu s$). L'origine est prise au centre de la cavité; l'extension du waist du mode est représentée par une zone hachurée. On a aussi montré en pointillés les ventres qui ont été effectivement choisis comme zones R_1 et R_2 .

Choix des zones de Ramsey Nous venons de décrire les outils à notre disposition pour effectuer les impulsions de Ramsey. Nous allons maintenant expliquer les opérations nécessaires au réglage de l'interféromètre. Nous commençons par le choix des zones de Ramsey. Le champ doit y être le plus homogène possible. En effet le paquet atomique a une certaine dispersion spatiale, donc deux atomes successifs ne verront pas le champ micro-onde tout-à-fait à la même position. Si ce champ est très inhomogène, cela causera des dispersions d'angle de Rabi qui réduiront le contraste des franges. C'est pourquoi nous souhaitons que les zones de Ramsey au niveau des ventres de l'onde stationnaire. Il nous faut donc être capables de mesurer la structure de l'onde stationnaire dans la cavité, d'établir une carte du champ micro-onde dans la cavité.

Pour ce faire, nous tirons profit du rapport entre les dimensions d'un "paquet" atomique et la longueur d'onde du champ. Rappelons en effet que nous connaissons la position d'un atome dans le dispositif à un millimètre près, alors que la taille typique des structures de l'onde stationnaire dans la structure "cavité+anneau" est $\frac{\lambda}{2} \simeq 3mm$. Nous pouvons donc utiliser les atomes comme des sondes pour déterminer les noeuds et les ventres de l'onde stationnaire. Nous établissons la carte du champ micro-onde dans

la cavité en envoyant des atomes de vitesses différentes dans l'état $|a\rangle$. Nous allumons brièvement la source à un instant t où tous les atomes sont dans la cavité, mais à des positions différentes. Puis nous relevons leur taux de transfert vers l'état $|b\rangle$ en fonction de leur position dans la cavité à l'instant t . Le résultat obtenu est montré sur la figure 3.10 pour la transition e-i.

Au vu de la carte du champ on choisit la position des zones de Ramsey. Sur la carte de champ 3.10 on voit clairement deux de ces ventres. Pour choisir précisément les zones de Ramsey il est ensuite nécessaire d'enregistrer des cartes de champ plus "locales" (la carte que je montre couvre en effet toute l'extension spatiale de la cavité), et également avec une source plus atténuée (en effet les transitions de la figure 3.10 semblent saturées).

Réglage des impulsions $\pi/2$ Une fois la position des zones de Ramsey déterminées, il reste à régler les deux impulsions $\pi/2$. Pour cela, nous appliquons une seule impulsion à l'atome, et nous ajustons son temps d'allumage et son atténuation de manière à ce que le transfert de population atomique du niveau $|a\rangle$ vers le niveau $|b\rangle$ soit de 50%. On procède de même avec la deuxième impulsion.

Balayage des franges Pour faire varier la phase $\Delta\Phi$ qui permet de balayer les franges d'interférence, nous avons deux possibilités équivalentes. Soit nous modifions la phase du dipôle atomique (analyse du paragraphe 3.2.1), soit nous modifions la phase de l'impulsion R2 par rapport à R1 en laissant le dipôle atomique inchangé.

Pour déphaser le dipôle atomique d'une quantité $\Delta\Phi$ réglable, nous avons développé une technique que nous avons baptisée "virgule" et qui est schématisée sur la figure 3.11. Elle consiste à appliquer une brève impulsion de durée $T_{virg} = 2\mu s$ de champ électrique directeur E_{Stark} . Appelons la valeur atteinte par le champ électrique E_{virg} . Au cours de cette impulsion, la fréquence de Bohr de la transition a-b a changé d'une quantité $\Delta\nu$ à cause de l'effet Stark différentiel entre les niveaux $|a\rangle$ et $|b\rangle$. Par conséquent, la différence de phase entre le dipôle atomique et la source S_{Ram} varie d'une quantité $\Delta\Phi = \Delta\nu T_{virg}$ (partie hachurée sur la figure 3.11). En balayant E_{virg} on parvient donc à balayer la phase $\Delta\Phi$.

On peut aussi bien laisser la phase du dipôle atomique inchangée et faire varier la phase de l'impulsion R2 par rapport à l'impulsion R1. Pour ce faire, nous désaccordons légèrement S_{Ram} de la transition atomique d'une quantité $\delta\nu = \nu_{ab} - \nu_{source}$:

- tant que $\delta\nu \ll 1/T_{on}$, la puissance vue par l'atome lors de chaque impulsion $\frac{\pi}{2}$ est pratiquement inchangée, donc elles transfèrent toujours 50% des atomes de $|a\rangle$ vers $|b\rangle$, malgré le désaccord $\delta\nu$.

- par contre la phase relative entre le champ et le dipôle atomique varie de $\Delta\Phi = \delta\nu(t_2 - t_1)$ que l'on peut faire varier facilement en changeant la fréquence de la source et donc $\delta\nu$.

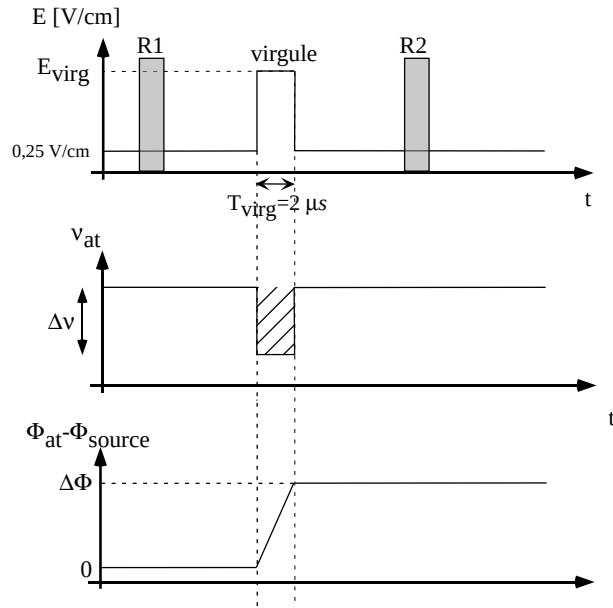


FIG. 3.11 – Principe de la "virgule", utilisée pour balayer la phase des franges de Ramsey. Une impulsion de $T_{virg} = 2\mu s$ de champ électrique Stark ralentit transitoirement la rotation du dipôle atomique par rapport au champ micro-onde en changeant sa fréquence de Bohr d'une quantité $\Delta\nu$. Le déphasage accumulé vaut $\Delta\Phi = T_{virg}\Delta\nu$. En changeant E_{virg} on balaye $\Delta\Phi$.

Signaux expérimentaux Finalement, on peut enregistrer le signal de franges de Ramsey. On mesure le transfert de population entre $|a\rangle$ et $|b\rangle$ en fonction du déphasage $\Delta\Phi$. Des signaux sont montrés sur la figure 3.12 des franges sur les 3 transitions considérées.

Les contrastes des franges de Ramsey sont entre 70% et 80%. Ce contraste est sans doute essentiellement limité par les inhomogénéités de champ micro-onde au moment des impulsions de Ramsey, ainsi que par les inhomogénéités de champs statiques électrique et magnétique.

Nous allons maintenant voir comment l'on peut utiliser l'interféromètre de Ramsey que nous venons de décrire pour calibrer un champ cohérent injecté dans la cavité, avec une grande précision, ce qui nous servira pour réaliser les expériences des chapitres 4 et 5.

3.3 Injecter un champ cohérent $|\alpha\rangle$

Lorsqu'on allume une source S_{cav} accordée à la fréquence de la cavité, initialement vide, pendant une durée T_{on} , un champ se construit dans la cavité. On peut montrer que le champ injecté est dans un état cohérent $|\alpha\rangle$, avec α proportionnel à T_{on} .

3.3.1 Propriétés d'un champ cohérent

Nous rappelons tout d'abord brièvement quelques propriétés des états cohérents qui nous seront utiles pour la suite de ce mémoire. Un état cohérent $|\alpha\rangle$ (α est un nombre complexe quelconque) peut être décomposé sur la base des états de Fock :

$$|\alpha\rangle = \sum_n C_n |n\rangle \quad (3.11)$$

Les coefficients C_n de cette décomposition valent $C_n = \exp^{-\frac{|\alpha|^2}{2}} \alpha^n / \sqrt{n!}$. La probabilité d'avoir n photons est donnée par le module carré de C_n :

$$P(n) = |C_n|^2 = \exp^{-|\alpha|^2} \frac{|\alpha|^{2n}}{n!} \quad (3.12)$$

Ainsi, la statistique des états cohérents est poissonnienne. Le nombre moyen de photons vaut $\bar{n} = |\alpha|^2$ et la phase moyenne est celle du nombre complexe α . La variance du nombre de photons vaut $\Delta n = \sqrt{\bar{n}} = |\alpha|$; la variance de la phase vaut $1/\sqrt{\bar{n}}$ (c'est la valeur minimale autorisée par la relation d'incertitude de Heisenberg $\Delta N \Delta \Phi \geq 1$ (voir paragraphe 1.1.3)). On peut montrer que la valeur moyenne du champ électrique dans un état cohérent $|\alpha\rangle$ est $\langle E \rangle = \alpha$, et qu'un état cohérent à un instant $t = 0$ reste cohérent au cours d'une évolution libre (seule sa phase augmente à la pulsation ω du mode considéré). Ainsi, $\langle E(t) \rangle = |\alpha| \cos(\omega t + \arg \alpha)$.

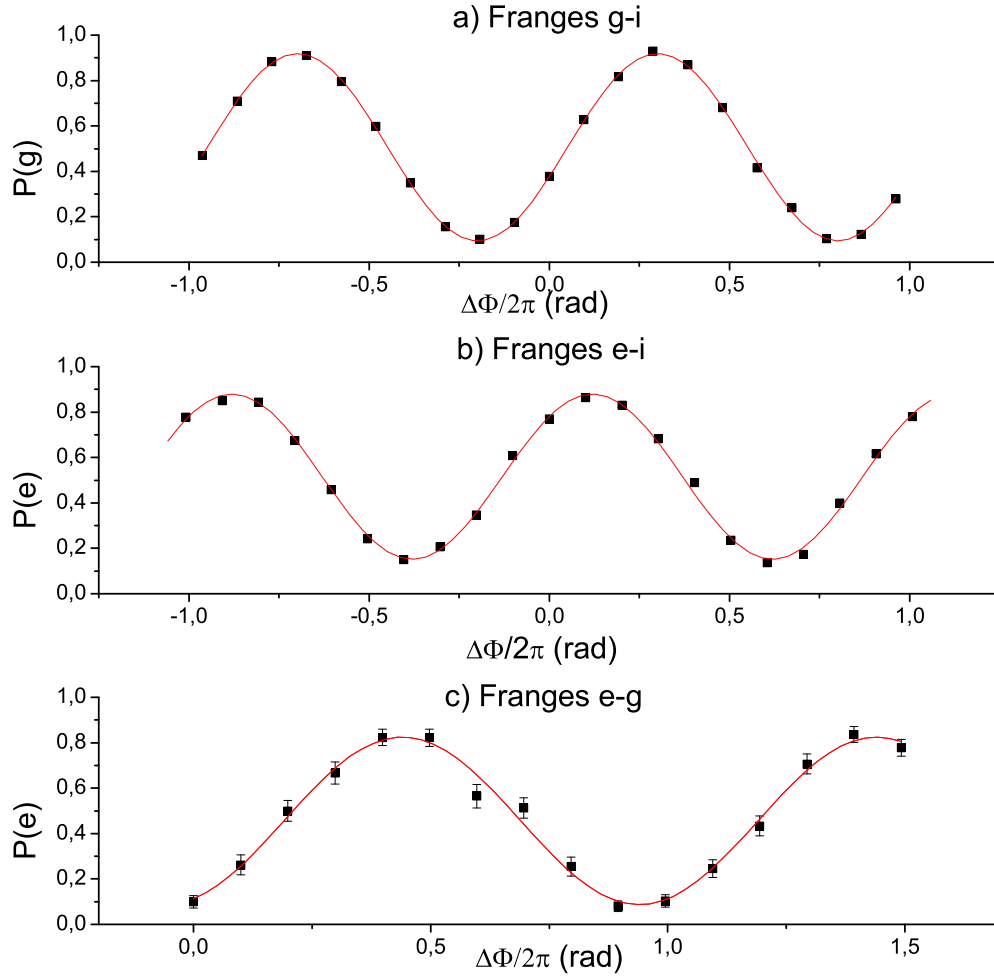


FIG. 3.12 – (a) Franges de Ramsey sur la transition $g-i$. Le contraste est d'environ 80%. (b) Franges de Ramsey sur la transition $e-i$ (franges à 2 photons). Le contraste est de 75%. (c) Franges de Ramsey sur la transition $e-g$, de contraste 72% (cette baisse de contraste est dû à l'utilisation d'atomes très lents pour ce signal). Dans toutes ces expériences la cavité était très désaccordée par rapport à la transition atomique considérée.

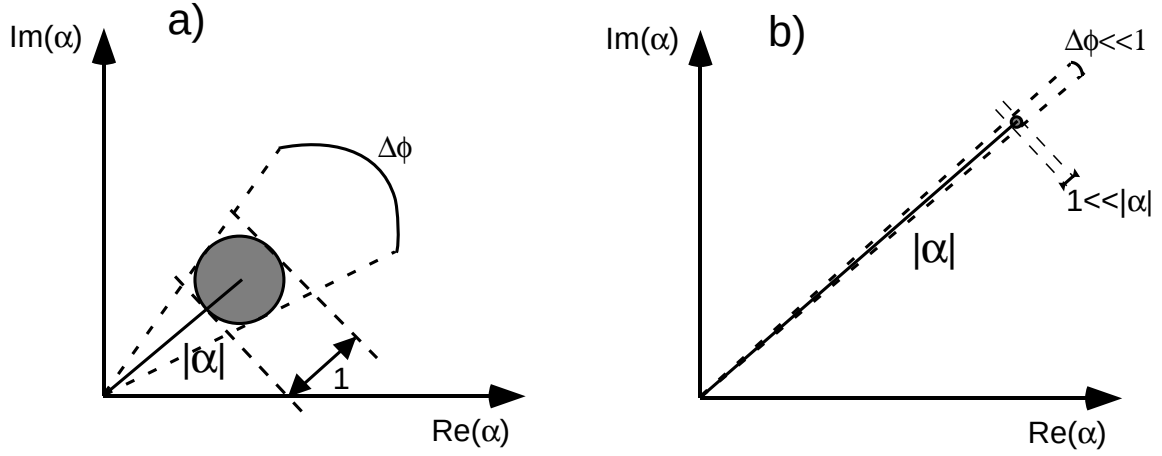


FIG. 3.13 – Représentation d'un état cohérent dans l'espace des phases. (a) Cas d'un petit champ cohérent. Les dispersions sur l'amplitude et la phase du champ sont importantes. (b) Cas d'un grand champ cohérent. Les dispersions sur la phase et l'amplitude sont négligeables : le champ est dans un état quasi-classique.

- Lorsque $\alpha \gg 1$, les incertitudes sur l'intensité et la phase sont négligeables devant les valeurs moyennes de ces deux quantités. On peut donc faire l'approximation semi-classique de considérer que le champ a une amplitude et une phase parfaitement définies. Comme en outre nous avons vu que la valeur moyenne du champ électrique était la même que celle d'un champ classique d'amplitude $|\alpha|$, on comprend que les états cohérents soient souvent appelés “états quasi-classiques” et qu'ils soient utilisés pour décrire l'état quantique d'une onde monochromatique (un laser par exemple).

- Lorsque $\alpha \simeq 1$, les fluctuations quantiques deviennent du même ordre de grandeur que les valeurs moyennes de l'amplitude et de la phase et ne sont plus négligeables. Nous en verrons des conséquences au chapitre 4.

Les deux cas ci-dessus sont représentés dans l'espace des phases à la figure 3.13

3.3.2 Comment calibrer un champ cohérent ?

Nous utilisons l'interféromètre de Ramsey que nous venons de décrire pour calibrer le champ cohérent $|\alpha\rangle$ injecté dans le mode de la cavité [19]. La source S_{cav} doit bien sûr être à la fréquence ν_{cav} du mode de la cavité avec lequel on souhaite travailler. Pour ce faire, on commence par mesurer ν_{cav} grâce au banc micro-onde, puis on asservit la source à cette même fréquence. La figure 3.14 montre le dispositif qui permet de coupler un champ à la cavité et de le calibrer. S_{cav} est couplée à la cavité par un guide d'onde (physiquement distinct du guide Ramsey) qui débouche au niveau du trou de couplage de l'un des miroirs. Un interrupteur (diode PIN cavité) contrôlé par un

signal digital permet de l'allumer pendant la durée T_{on} souhaitée.³ On peut modifier la puissance injectée en réglant l'atténuateur de précision att_{cav} à une valeur $A_{cav}^{(dB)}$ exprimée en décibels. Ainsi, on peut connaître la puissance relative d'un champ injecté à une certaine atténuation par rapport à un autre, injecté dans les mêmes conditions mais avec une atténuation différente. Mais nous voulons connaître l'intensité injectée dans la cavité dans l'absolu (nombre moyen de photons).

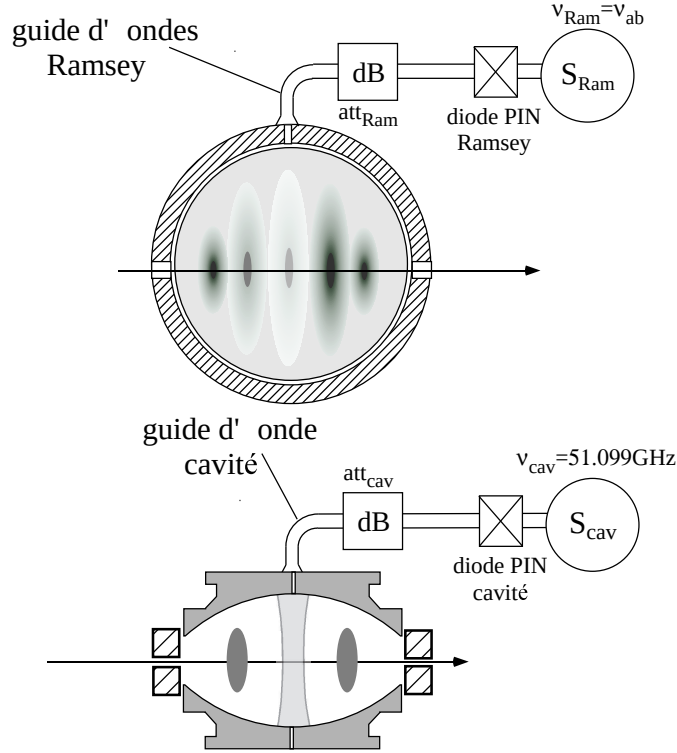


FIG. 3.14 – Dispositif utilisé pour la calibration d'un champ injecté dans la cavité. Deux sources sont nécessaires. La première, S_{cav} , est asservie à la fréquence de la cavité $\nu = 51.099\text{GHz}$, et couplée par un guide d'onde nommé "cavité" qui débouche derrière l'un des miroirs. Lorsqu'elle est allumée, un champ micro-onde se construit dans le mode qui est résonant avec elle. La deuxième source, S_{Ram} , sert à appliquer les impulsions de Ramsey à l'atome. Elle est couplée par un autre guide d'ondes, le guide "Ramsey", qui débouche sur le côté de l'anneau.

S'il est facile de mesurer la puissance émise par S_{cav} à $\nu = 51.099\text{GHz}$, il est en revanche beaucoup plus difficile de mesurer l'intensité intra-cavité due à l'allumage de S_{cav} pendant la durée T_{on} , car on ne connaît pas précisément l'atténuation du signal entre la source et la cavité, qui est principalement due aux pertes dans les guides d'onde,

³En réalité, comme nous l'avons vu au chapitre 3, la diode PIN est située entre la source et le mélangeur harmonique; le schéma 3.14 n'est donc pas tout-à-fait correct. Néanmoins, l'effet de l'interrupteur est bien d'éteindre la source S_{Ram} .

aux différents atténuateurs, et au trou de couplage de la cavité. On sait seulement que l'atténuation totale est supérieure à $80dB$ mais on ignore la valeur exacte.

Les seuls détecteurs qui soient suffisamment sensibles pour détecter des champs micro-onde de quelques photons, et qui aient accès à l'intérieur de la cavité sont les atomes de Rydberg eux-mêmes. Nous utilisons l'interaction *dispersive* entre les atomes et le champ stocké dans la cavité pour mesurer le nombre moyen de photons. Dans ce régime où l'atome est désaccordé d'une quantité δ par rapport à la cavité avec $\delta \gg \Omega_0 \sqrt{n+1}/2$ on calcule à l'annexe A l'évolution des états $|e, n\rangle$ et $|g, n\rangle$ et l'on montre que ces états sont déphasés d'une quantité proportionnelle au nombre de photons. L'état $|i, n\rangle$ d'autre part reste bien sûr stationnaire puisque le niveau $|i\rangle$ n'est pas couplé à la cavité.

$$\begin{aligned} |e, n\rangle &\xrightarrow{disp} \exp^{i\Phi_0(n+1)} |e, n\rangle \\ |g, n\rangle &\xrightarrow{disp} \exp^{-i\Phi_0 n} |g, n\rangle \\ |i, n\rangle &\xrightarrow{disp} |i, n\rangle \end{aligned} \quad (3.13)$$

où $\Phi_0 = \Omega_0^2 t_{eff}/4\sqrt{2}\delta$ est le déphasage par photon (t_{eff} est le temps effectif introduit à l'annexe A). Si l'on est capable de mesurer le déphasage de la fonction d'onde atomique, on peut en déduire le nombre de photons stockés dans la cavité. Un interféromètre de Ramsey, réglé sur une des trois transitions atomiques g-i, e-i ou e-g, peut parfaitement effectuer cette mesure.

Considérons en effet la situation représentée sur la figure 3.15 a), où l'interaction dispersive avec un état de Fock $|n\rangle$ est située dans l'un des bras d'un interféromètre de Ramsey sur la transition a-b (qui peut être n'importe laquelle de nos trois transitions habituelles). Il est clair, au vu des relations 3.13, que la seule modification du signal de franges de Ramsey dans le vide, sera un déphasage qui dépend à la fois de n et de la transition a-b considérée, et que l'on notera $\Phi_a(n) - \Phi_b(n)$ (les phases $\Phi_a(n)$ et $\Phi_b(n)$ sont définies par les relations 3.13). La probabilité finale de détecter l'atome dans l'état $|b\rangle$ sera $P_b(\Delta\Phi) = \frac{1}{2}(1 + \cos(\Phi_a(n) - \Phi_b(n) + \Delta\Phi))$. Le contraste des franges reste maximal. Nous pouvons aussi comprendre cette situation en termes de complémentarité. La cavité en effet peut être considérée comme un détecteur WP. Lorsqu'elle est initialement dans un état de Fock, l'information enregistrée par la cavité sur l'état de l'atome lors de son interaction avec le champ est nulle, puisque l'état de la cavité est inchangé, que l'atome soit dans l'état $|a\rangle$ ou $|b\rangle$. Il n'y a aucune intrication entre l'atome et la cavité dans l'état final. La visibilité des franges reste donc maximale.

La situation est différente lorsque la cavité contient un état cohérent $|\alpha\rangle$. En effet, nous allons voir que l'état final de la cavité est modifié par l'état atomique initial.

Prenons pour exemple le cas où l'atome est initialement dans l'état $|e\rangle$.

$$\begin{aligned}
|\psi_{in}\rangle = |e, \alpha\rangle &= |e\rangle \otimes (\sum_n C_n |n\rangle) \xrightarrow{disp} \sum_n C_n \exp^{i\Phi_0(n+1)} |e, n\rangle \\
&= \exp^{-\frac{|\alpha|^2}{2}} \sum_n \frac{\alpha^n}{\sqrt{n!}} \exp^{i\Phi_0(n+1)} |e, n\rangle \\
&= \exp^{i\Phi_0} |e\rangle \otimes (\exp^{-\frac{|\alpha|^2}{2}} \sum_n \frac{(\alpha \exp^{i\Phi_0})^n}{\sqrt{n!}} |n\rangle) \\
&= \exp^{i\Phi_0} |e\rangle \otimes |\alpha \exp^{i\Phi_0}\rangle = |e\rangle \otimes |\alpha_e\rangle
\end{aligned} \tag{3.14}$$

L'interaction dispersive d'un atome dans $|e\rangle$ fait ainsi évoluer l'état de la cavité vers l'état $|\alpha_e\rangle$ qui est aussi un état cohérent mais déphasé de $+\Phi_0$. On montrerait de même que, si l'atome était dans $|g\rangle$, l'état final de la cavité serait $|\alpha_g\rangle = |\alpha \exp^{-i\Phi_0}\rangle$, alors qu'un atome dans $|i\rangle$ laisserait l'état cohérent initial inchangé $|\alpha_i\rangle = |\alpha\rangle$. Par conséquent, si l'atome est initialement préparé dans une superposition cohérente des états $|a\rangle$ et $|b\rangle$, le système évolue comme suit

$$|\psi_{in}\rangle = |a, \alpha\rangle \xrightarrow{R1} \frac{1}{\sqrt{2}}(|a\rangle - i|b\rangle) \otimes |\alpha\rangle \xrightarrow{cav|\alpha\rangle} \frac{1}{\sqrt{2}}(|\alpha_a\rangle|a\rangle - i|\alpha_b\rangle|b\rangle) \tag{3.15}$$

vers un état final intriqué entre l'atome et la cavité. Nous savons, d'après l'analyse du chapitre 1, que le terme habituel d'interférence est alors multiplié par le produit scalaire des deux états du détecteur WP. Cela se traduit à la fois par un déphasage $\Phi_\alpha^{(a-b)}$ des franges par rapport au cas où la cavité est vide, donné par $\arg(\langle\alpha_a|\alpha_b\rangle)$, et par une réduction de contraste $C_\alpha^{(a-b)}$ qui vaut $|\langle\alpha_a|\alpha_b\rangle|$. On peut aisément calculer les valeurs de ces paramètres pour les trois transitions atomiques à l'aide de la formule qui donne le produit scalaire de deux états cohérents $|\alpha\rangle$ et $|\beta\rangle$:

$$\langle\alpha|\beta\rangle = \exp^{-\frac{1}{2}(|\alpha|^2 + |\beta|^2 - 2\beta\alpha^*)} \tag{3.16}$$

Les résultats sont résumés sur le tableau 3.1. Notons que ces réductions de contraste

transition a-b	$\Phi_\alpha^{(a-b)}$	$C_\alpha^{(a-b)}$
g-i	$- \alpha ^2 \sin \Phi_0$	$\exp^{-2 \alpha ^2 \sin^2(\frac{\Phi_0}{2})}$
e-i	$ \alpha ^2 \sin \Phi_0$	$\exp^{-2 \alpha ^2 \sin^2(\frac{\Phi_0}{2})}$
e-g	$ \alpha ^2 \sin 2\Phi_0$	$\exp^{-2 \alpha ^2 \sin^2 \Phi_0}$

TAB. 3.1 – Déphasage et réduction de contraste des franges de Ramsey sur les différentes transitions atomiques en fonction de α et du déphasage par photon Φ_0

peuvent aussi être interprétées de manière équivalente comme étant dûes à la dispersion du nombre de photons comme on peut le voir à la figure 3.15.

Lorsque les états $|\alpha_a\rangle$ et $|\alpha_b\rangle$ sont très distincts, le produit scalaire $\langle\alpha_a|\alpha_b\rangle$ est presque nul. L'état produit est alors semblable à un “chat de Schrödinger” mésoscopique [16] et les franges de Ramsey ont un contraste nul. Nous sommes plutôt intéressés ici

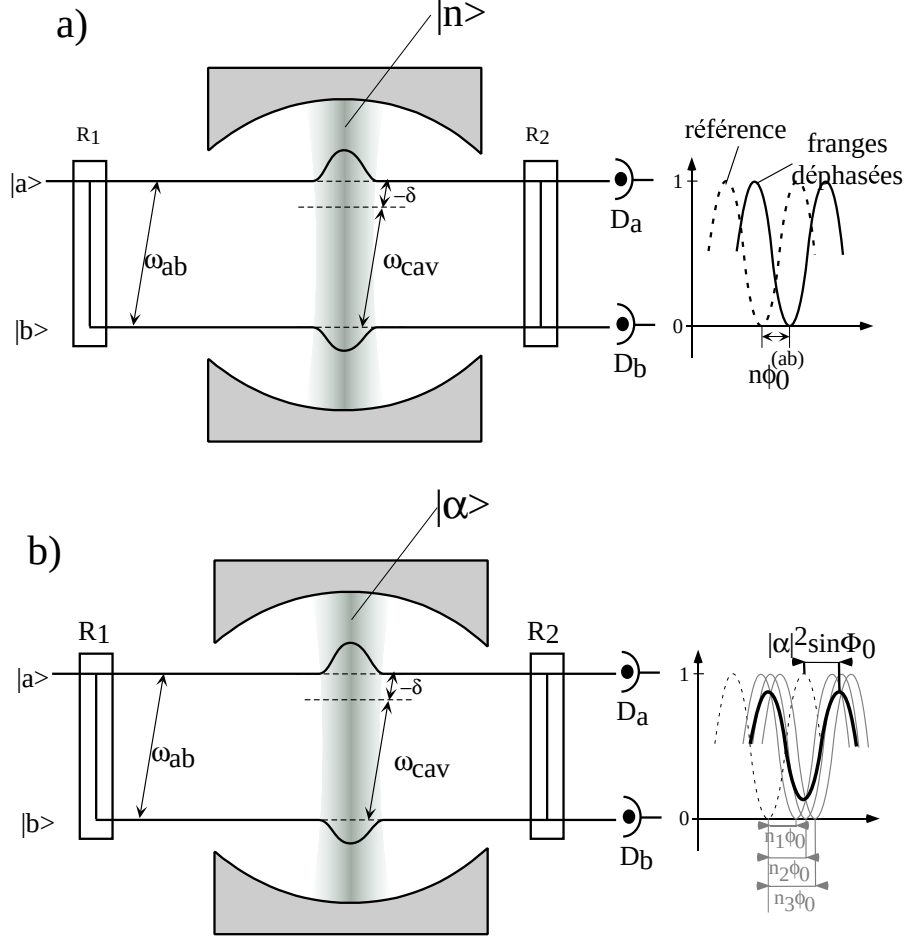


FIG. 3.15 – Le déphasage de la fonction d'onde atomique causé par l'interaction dispersive avec la cavité se manifeste par un déphasage des franges de Ramsey. Sur la figure (a) la cavité contient un état de Fock $|n\rangle$. Les franges de Ramsey sont déphasées de $\Phi_a(n) - \Phi_b(n)$ par rapport au cas où la cavité est vide; leur contraste reste égal à 1. Sur la figure (b) la cavité contient un champ cohérent $|\alpha\rangle$. Les franges sont alors déphasées de $|\alpha|^2 \sin \Phi_0$, et leur contraste est réduit à cause de la dispersion du nombre de photons dans un champ cohérent.

au régime où le contraste des franges reste important. En effet, nous tirons parti du déphasage des franges de Ramsey qu'il est facile de mesurer pour connaître le champ dans la cavité grâce aux formules du tableau 3.1.

Supposons par exemple que nous souhaitions mesurer un champ $|\alpha\rangle$ avec un interféromètre de Ramsey sur la transition e-i. Nous commençons par enregistrer des franges dites de “référence” avec une cavité vide. Puis nous prenons un deuxième signal de franges dans les mêmes conditions, mais cette fois-ci précédées de l'injection du champ $|\alpha\rangle$ dans la cavité. Ces franges auront un contraste réduit par rapport aux précédentes, mais elles seront surtout déphasées d'une quantité $\Phi_\alpha^{(e-i)}$ qui nous permettra de connaître $|\alpha|$.

Chapitre 4

Réalisation d'un interféromètre à la frontière quantique/classique

Dans ce chapitre nous décrivons la réalisation expérimentale de l'interféromètre de Ramsey à la frontière quantique/classique, tel qu'il a été brièvement présenté au chapitre 1, grâce aux outils expérimentaux décrits dans les deux chapitres précédents. Nous commençons par rappeler le principe de l'expérience dans la partie 4.1.

Une des difficultés expérimentales consiste à calibrer le champ présent dans la cavité avec la meilleure exactitude et précision. La procédure adoptée, suivant les idées développées au chapitre précédent et utilisant un interféromètre de Ramsey annexe, est expliquée au paragraphe 4.2.2.

Le paragraphe suivant expose le déroulement complet d'une séquence expérimentale. Les résultats obtenus sont présentés à la suite, et comparés aux prédictions théoriques.

Nous avons effectué une autre expérience, présentée au paragraphe 4.3, dans le prolongement de la précédente, prouvant qu'il est possible d'effacer de manière inconditionnelle l'information stockée par l'interféromètre sur le chemin suivi par l'atome. On retrouve alors des franges d'interférence avec un contraste maximal.

Enfin nous interprétons tous ces résultats du point de vue du champ dans la section 4.4. En particulier nous tentons de répondre précisément à la question : quand est-ce qu'un champ est quantique ou classique ?

4.1 Principe de l'expérience

4.1.1 La lame séparatrice quantique/classique

On a rappelé sur la figure 4.1 le principe de notre expérience. Un atome, initialement dans $|e\rangle$, effectue une impulsion de Ramsey R1 dans le mode de la cavité contenant un champ cohérent $|\alpha\rangle$, puis une deuxième impulsion R2 classique. Dans ce paragraphe, nous allons décrire quantitativement l'interaction entre un atome de Rydberg et le

champ cohérent $|\alpha\rangle$ stocké dans la cavité lors de la première impulsion de Ramsey, à l'aide du modèle de Jaynes-Cummings.

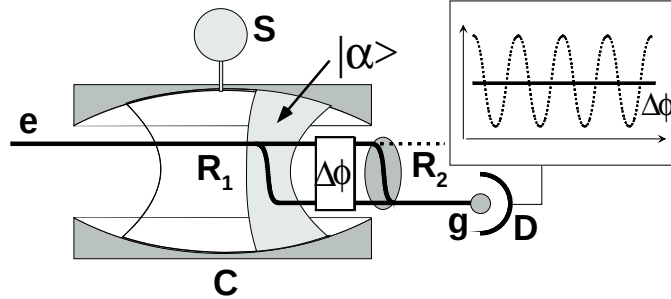


FIG. 4.1 – Principe de l'expérience : la première impulsion $\pi/2$ d'un interféromètre de Ramsey est réalisée dans le mode de très haut facteur de qualité de la cavité contenant un champ cohérent $|\alpha\rangle$ qui y a été préalablement injecté par une source S . La deuxième impulsion R_2 est classique. On balaye la phase des franges en faisant varier la phase du dipôle atomique d'une quantité $\Delta\Phi$ réglable.

Les propriétés d'un état cohérent ont été présentées au chapitre 3. Chaque composante à n photons de l'état cohérent induit une oscillation de Rabi entre $|e\rangle$ et $|g\rangle$ à la fréquence $\Omega_0\sqrt{n+1}$ (voir annexe A). L'évolution globale du système résulte de la superposition de ces oscillations de Rabi :

$$|\Psi\rangle = \sum_n C_n \left[\cos\left(\frac{\Omega_0}{2}\sqrt{n+1}t_{eff}\right) |e, n\rangle - i \sin\left(\frac{\Omega_0}{2}\sqrt{n+1}t_{eff}\right) |g, n+1\rangle \right] \quad (4.1)$$

Pour effectuer une impulsion $\pi/2$, on règle le temps d'interaction effectif t_α de manière à ce que la probabilité de détecter l'atome dans $|e\rangle$ ou $|g\rangle$ soit égale à $1/2$. Or $P_g = |\langle g|\Psi\rangle|^2$ donc la condition sur le temps d'interaction s'écrit

$$\sum_n |C_n|^2 \cos^2(\Omega_0\sqrt{n+1}t_\alpha/2) = \frac{1}{2} \quad (4.2)$$

Le système est alors dans l'état

$$|\Psi\rangle = \frac{1}{\sqrt{2}}(|e\rangle|\Psi_e\rangle - i|g\rangle|\Psi_g\rangle) \quad (4.3)$$

avec

$$\begin{aligned} |\Psi_e\rangle &= \sqrt{2} \left[\sum_n C_n \cos\left(\frac{\Omega_0}{2}\sqrt{n+1}t_\alpha\right) |n\rangle \right] \\ |\Psi_g\rangle &= \sqrt{2} \left[\sum_n C_n \sin\left(\frac{\Omega_0}{2}\sqrt{n+1}t_\alpha\right) |n+1\rangle \right] \end{aligned} \quad (4.4)$$

On voit donc qu'après l'impulsion $\frac{\pi}{2}$ dans la cavité, l'état de l'atome et celui du champ sont corrélés. La quantité d'information enregistrée par la cavité sur l'état de l'atome dépend du nombre moyen de photons $N = |\alpha|^2$. Etudions les deux cas limites $N = 0$ d'une part et $N \gg 1$ d'autre part.

Si $N = 0$,

$$|\Psi\rangle = \frac{1}{\sqrt{2}}(|e, 0\rangle - i|g, 1\rangle) \quad (4.5)$$

donc $|\Psi_e\rangle = |0\rangle$ et $|\Psi_g\rangle = |1\rangle$ sont des états orthogonaux. La cavité agit comme une lame séparatrice qui enregistre une information non-ambigüe sur l'état de l'atome. Le système atome-cavité est dans un état maximalement intriqué, du type EPR. Insistons sur le fait que, même lorsque la cavité est vide, la fréquence de Rabi est suffisamment grande pour que l'atome y effectue une impulsion $\pi/2$. C'est donc le couplage fort qui nous permet de réaliser une lame séparatrice qui enregistre de manière cohérente une information sur l'état de l'atome.

Si maintenant $N \gg 1$, l'incertitude sur l'amplitude du champ devient négligeable car $\Delta N = \sqrt{N} \ll N$. Par conséquent on peut récrire l'équation 4.2

$$\sum_n |C_n|^2 \cos^2\left(\frac{\Omega_0}{2}\sqrt{n+1}t_\alpha\right) \simeq \left(\sum_n |C_n|^2\right) \cos^2\left(\frac{\bar{\Omega}}{2}t_\alpha\right) = \frac{1}{2} \quad (4.6)$$

avec $\bar{\Omega} = \Omega_0\sqrt{N}$. Donc $\cos(\Omega_0\sqrt{n+1}t_\alpha/2) \simeq \cos(\bar{\Omega}t_\alpha/2) \simeq \sin(\bar{\Omega}t_\alpha/2) \simeq 1/\sqrt{2}$. D'après 4.4,

$$\begin{aligned} |\Psi_e\rangle &\simeq \sqrt{2} \left(\sum_n \frac{C_n}{\sqrt{2}} |n\rangle \right) = |\alpha\rangle \\ |\Psi_g\rangle &\simeq \sqrt{2} \left(\sum_n \frac{C_n}{\sqrt{2}} |n+1\rangle \right) = \sum_n \frac{\sqrt{n+1}}{\alpha} C_{n+1} |n+1\rangle \\ &\simeq \frac{\sqrt{N}}{\alpha} \sum_n C_{n+1} |n+1\rangle = \exp^{-i\phi} |\alpha\rangle \end{aligned} \quad (4.7)$$

si $\alpha = |\alpha| \exp^{i\phi}$. Par conséquent, l'état final atome-champ est factorisable :

$$|\Psi\rangle \simeq |\alpha\rangle \otimes \frac{1}{\sqrt{2}}(\exp^{i\phi} |e\rangle - i|g\rangle) \quad (4.8)$$

On retrouve l'action d'une impulsion de Ramsey classique qui imprime la phase du champ sur le dipôle atomique sans être elle-même modifiée par l'interaction avec l'atome. En résumé, la cavité se comporte bien vis-à-vis de l'atome qui y effectue son impulsion $\pi/2$, dans l'espace des énergies internes, comme le diaphragme $D1$ imaginé par Bohr dans l'espace réel.

4.1.2 Le contraste des franges

Voyons maintenant comment se manifeste l'intrication entre l'atome et la cavité sur le contraste des franges. Après une impulsion R1 dans la cavité l'état du système est

$$|\Psi_{at-chp}\rangle = \frac{1}{\sqrt{2}}(|e\rangle|\Psi_e\rangle - i|g\rangle|\Psi_g\rangle) \quad (4.9)$$

L'atome est intriqué avec la cavité donc sa fonction d'onde n'est pas bien définie. La matrice densité réduite de l'atome est obtenue à partir de 4.9 en traçant sur les variables de champ de la cavité.

$$\begin{aligned}\rho &= \text{tr}_{\text{champ}}(|\Psi\rangle\langle\Psi|) \\ &= \frac{1}{2} \begin{pmatrix} 1 & -i\langle\Psi_e|\Psi_g\rangle \\ i\langle\Psi_e|\Psi_g\rangle^* & 1 \end{pmatrix}\end{aligned}\quad (4.10)$$

Le déphasage de $\Delta\Phi$ effectué sur la composante en $|e\rangle$ de $|\Psi\rangle$, puis la rotation de $\frac{\pi}{2}$ en R2 transforment la matrice densité ρ en $\rho_{fin} = \hat{U}_{R2}\hat{U}_{\Delta\Phi}\rho\hat{U}_{\Delta\Phi}^\dagger\hat{U}_{R2}^\dagger$ avec

$$\hat{U}_{\Delta\Phi} = \begin{pmatrix} \exp^{i\Delta\Phi} & 0 \\ 0 & 1 \end{pmatrix}\quad (4.11)$$

et

$$\hat{U}_{R2} = \begin{pmatrix} \frac{1}{\sqrt{2}} & -\frac{i}{\sqrt{2}} \\ -\frac{i}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{pmatrix}\quad (4.12)$$

Au final,

$$\rho_{fin} = \frac{1}{2} \begin{pmatrix} 1 + \text{Re}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi}) & -\text{Im}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi}) \\ -\text{Im}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi}) & 1 - \text{Re}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi}) \end{pmatrix}\quad (4.13)$$

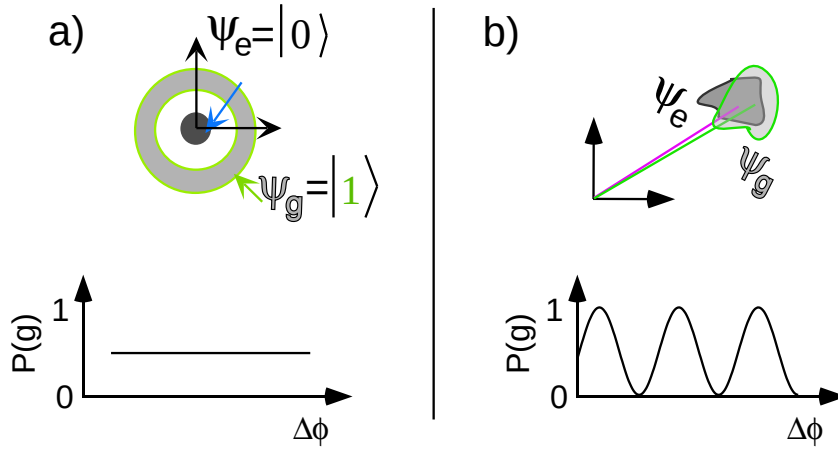


FIG. 4.2 – Lorsque les deux états $|\psi_e\rangle$ et $|\psi_g\rangle$ sont distincts (cas (a)) le contraste des franges est nul; au contraire quand le recouvrement de $|\psi_e\rangle$ avec $|\psi_g\rangle$ est important (cas (b)) les franges sont contrastées.

d'où l'on tire la probabilité de détecter l'atome dans l'état $|g\rangle$

$$P_g(\Delta\Phi) = \rho_{gg} = \frac{1}{2} [1 + \text{Re}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi})]\quad (4.14)$$

Elle présente des franges d'interférence dont le contraste est le module du produit scalaire des deux états $|\Psi_e\rangle$ et $|\Psi_g\rangle$. La figure 4.2 montre les deux cas limites $N = 0$ et $N \gg 1$. Notre dispositif permet donc de voir très directement les effets de l'intrication des atomes avec la lame séparatrice en énergie sur le contraste mesuré, et d'observer continûment la transition entre le fonctionnement quantique et le fonctionnement classique d'un même dispositif expérimental.

4.2 Réalisation de l'expérience de complémentarité

4.2.1 L'accord de la cavité

Nous devons tout d'abord choisir le mode de la cavité avec lequel nous allons travailler. Comme nous l'avons vu précédemment, il est important que l'atome soit à résonance avec le mode choisi dans un champ Stark faible ($E_{Stark}^{(1)} = 0.25V/cm$), sinon les inhomogénéités de champ électrique vont réduire le contraste des franges. Par ailleurs, nous devons utiliser la technique d'accord vue au chapitre précédent pour régler précisément l'impulsion $\pi/2$ dans la cavité (l'atome a une vitesse de $500m/s$ et ferait une impulsion de $\Omega_0 t_{eff} = \Omega_0 \sqrt{\pi} w/v \simeq 2\pi$ dans le vide, et plus encore dans un champ cohérent, si on le laissait à résonance durant tout son trajet dans le mode de la cavité). L'atome entrera donc dans la cavité dans un grand champ électrique ($E_{Stark}^{(1)} = 0.74V/cm$), suffisant pour le mettre hors-résonance, et sera accordé sur le mode à un instant $t_{Stark}^{(\alpha)}$ qu'il faudra régler pour chaque valeur de α . Cela nous oblige à faire l'expérience sur le mode Basse Fréquence BF. En effet, comme on peut le constater sur la figure 4.3, si l'on travaillait avec le mode Haute Fréquence, l'atome serait transitoirement résonant avec le mode BF au moment de son accord.

La première opération à faire consiste donc à accorder le mode BF de la cavité sur la fréquence de la transition atomique e-g dans le champ $E_{Stark}^{(1)}$, en procédant comme nous l'avons vu au paragraphe 3.1.1.a. On fixe la tension de piézo V_{PZT} à la valeur $V_{PZT}^{(BF)}$. Nous aurons besoin au cours de l'expérience d'injecter un champ dans le mode BF. Pour cela, il nous faut connaître précisément la fréquence du mode ν_{BF} , afin d'asservir la source S_{cav} à cette fréquence. Nous utilisons le banc de mesure micro-onde dont nous avons déjà parlé au chapitre 2 pour mesurer ν_{BF} , et nous trouvons $\nu_{BF} = 51.099046GHz$. Il faut donc asservir la source S_{cav} qui va injecter un champ dans le mode BF au cours de l'expérience à cette même fréquence.

4.2.2 Calibration du champ

L'amplitude du champ injecté dans la cavité est réglée dans notre expérience en faisant varier l'atténuation $A_{cav}^{(dB)}$ de la source S_{cav} . Comme nous souhaitons étudier de façon quantitative le contraste en fonction du nombre de photons injectés, nous ne pouvons pas nous contenter de connaître le champ dans la cavité en relatif, il nous faut

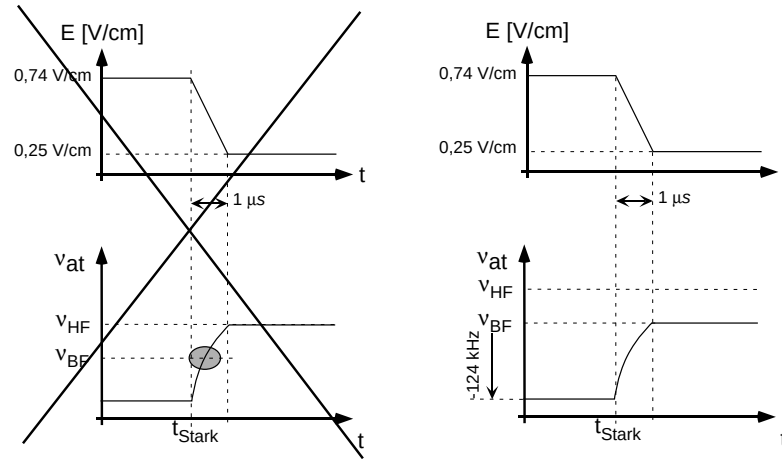


FIG. 4.3 – L'atome effectue son impulsion $\frac{\pi}{2}$ à la fin du mode de la cavité, donc on s'intéresse à une situation où E_{Stark} , initialement élevé (0.74V/cm), est amené à la valeur 0.25V/cm à un instant t_{Stark} . On constate sur cette figure que l'expérience ne peut avoir lieu que dans le mode Basse Fréquence.

calibrer le champ de manière absolue comme expliqué au chapitre 3, en utilisant un interféromètre de Ramsey annexe réglé sur la transition e-i.

4.2.2.a Description générale de la méthode

Nous aurions pu décider de calibrer le champ une seule fois, en enregistrant des franges de calibration à une certaine atténuation $A(\text{dB})$, puis considérer que l'atténuateur de précision att_{cav} est suffisamment linéaire pour que l'on puisse utiliser ses indications. Nous avons préféré enregistrer des franges de calibration pour chaque champ injecté, dans les mêmes conditions que l'expérience de franges quantiques que nous avons effectuée juste après. En effet nous avons jugé que cela nous permettait de minimiser les erreurs sur la calibration du champ. Notamment, la fréquence de la cavité dérive lentement dans le temps, à un rythme d'environ 1 kHz en quelques heures. En faisant l'expérience de franges quantiques dans la foulée de la calibration, il s'écoule moins d'une heure entre les deux opérations et l'on peut être certains que les dérives sont négligeables.

Pour calibrer le champ, nous avons besoin d'un interféromètre annexe réglé sur la transition e-i. Une source S_{Ram} est asservie à $\nu_{ei} = 52.794 \text{ GHz}$, et les franges d'interférence sont réglées comme nous l'avons expliqué dans le chapitre précédent.

4.2.2.b Les résultats expérimentaux

La calibration commence avec l'enregistrement d'un signal de franges e-i sans injecter de champ dans la cavité. Il s'agira de notre signal de référence par rapport auquel le déphasage induit par le champ dans la cavité sera mesuré. Puis, pour chaque valeur

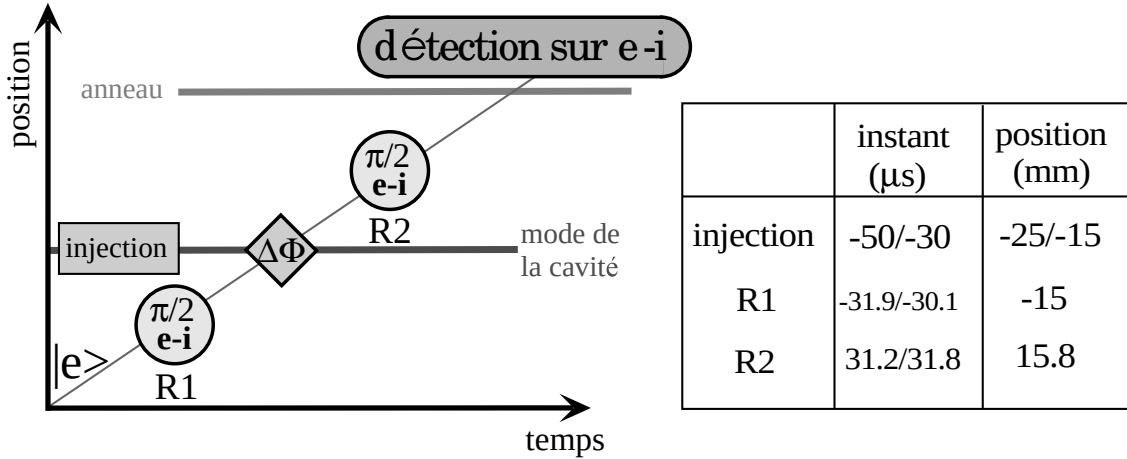


FIG. 4.4 – Déroulement de la séquence de calibration du nombre de photons injectés dans la cavité. La fréquence atomique est inférieure de 120kHz à celle du mode. La cavité, préalablement refroidie, est préparée dans l'état cohérent $|\alpha\rangle$ par l'allumage de S_{cav} (injection). L'atome subit deux impulsions $\pi/2$ de part et d'autre de la cavité. Les instants d'allumage des différentes sources, ainsi que la position de l'atome dans la cavité au moment des différents événements, sont indiqués dans le tableau joint. L'origine des temps et des espaces correspond au passage de l'atome au centre de la cavité.

de l'atténuation $A_{cav}^{(dB)}$, on enregistre un signal de franges. Le déroulement complet de la séquence expérimentale de calibration est exposé à la figure 4.4. Elle commence avec une séquence de refroidissement du mode BF (qui n'est pas représentée sur la figure 4.4) par des “serpillères résonantes” similaire à celle de la figure 3.3. Puis l'on injecte le champ dans la cavité par une impulsion de $20\mu s$, avant que l'atome n'entre dans le mode. On se place dans le régime dispersif en désaccordant l'atome par effet Stark dans un champ $E_{Stark}^{(2)} = 0.74\text{V/cm}$ qui met la transition e-g une centaine de kiloHertz sous la fréquence du mode BF. On applique les impulsions de Ramsey à l'atome de part et d'autre du mode, avec lequel il interagit dispersivement.

La figure 4.5 montre les franges e-i obtenues pour différents champs injectés (notons que pour l'instant nous ne connaissons que les valeurs $A_{cav}^{(dB)}$ pour lesquelles ont été enregistrées ces franges). On constate, comme prévu par les équations du tableau 3.1, que pour un champ important (atténuation faible), les franges sont déphasées et de contraste inférieur à celui des franges de référence.

4.2.2.c Exploitation des franges de calibration

Nous souhaitons maintenant exploiter les franges de calibration obtenues pour en tirer la relation entre l'atténuation $A_{cav}^{(dB)}$ et le nombre moyen de photons injectés N .

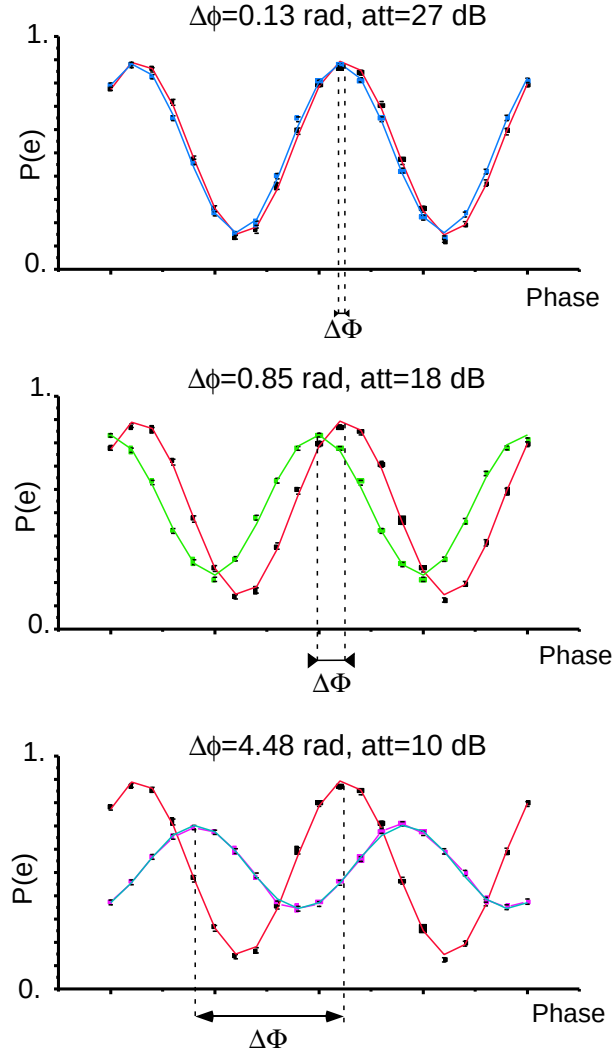


FIG. 4.5 – Franges e-i enregistrées pour trois champs différents dans la cavité, représentées avec les franges de référence. On constate que plus l'atténuation est faible, plus le champ dans la cavité est grand, plus les franges sont déphasées et plus la baisse de contraste est sensible.

Rappelons la formule établie au paragraphe 3.3 : $\Delta\Phi = N \sin \Phi_0$. Or

$$\Phi_0 = \frac{\Omega_0^2}{4\sqrt{2}\delta} t_{eff}, \quad (4.15)$$

avec $t_{eff} = \sqrt{\pi}w/v$. Du champ Stark dans lequel ont été effectuées les franges de calibration on peut déduire le désaccord atome-cavité par la formule $\delta = \nu(E_2) - \nu(E_1) = -255 * (E_2^2 - E_1^2) = 124kHz$. On obtient donc, d'après la relation 4.15, $\Phi_0 = 0.43rad$. Cette valeur de Φ_0 a été confirmée expérimentalement dans une expérience à deux atomes, où l'on mesurait directement et dans les mêmes conditions le déphasage des franges du deuxième atome conditionnées à l'émission d'un photon par le premier.

On pourrait avoir la tentation d'appliquer naïvement la formule

$$N = \Delta\Phi / \sin \Phi_0 \quad (4.16)$$

pour déterminer le nombre de photons injectés à chaque valeur de l'atténuation. Si l'on fait cela, on s'aperçoit que la relation entre l'atténuation imposée exprimée en linéaire $A^{(lin)} = 10^{-A^{(dB)}/10}$ et le nombre de photons ainsi déterminé n'est pas une droite. Il semble qu'en utilisant la relation 4.16, on sous-estime légèrement le nombre de photons injectés.

En effet, pour les champs injectés dont le nombre de photons est grand devant 1, le désaccord choisi n'est plus suffisant pour être vraiment dans le régime dispersif (la fréquence de Rabi dans un champ de n photons vaut $\Omega = \Omega_0\sqrt{n+1}$, donc la condition pour que l'interaction soit dispersive lorsque le nombre moyen de photons est N s'écrit $\delta \gg \Omega_0\sqrt{N}/2$) donc la formule 4.16 n'est plus correcte. On s'attend à une saturation du déphasage pour les champs importants.

Il est donc nécessaire d'effectuer un calcul numérique complet de la matrice densité atome-champ, pour différentes valeurs de $|\alpha|^2$, qui sera bien sûr non-perturbatif et tiendra compte de la saturation du déphasage. Nous avons utilisé un programme écrit dans ce but, et pu ainsi attribuer à chaque valeur de l'atténuation la valeur de $|\alpha|^2$ correspondante. Nous pouvons alors confronter les indications de l'atténuateur avec les valeurs mesurées par les atomes. Le résultat, montré sur la figure 4.6, confirme la linéarité de l'atténuation, et apporte une validation de la procédure de calibration. On constate sur cette figure que la procédure de calibration décrite ici semble précise et fiable pour des champs allant de 0.3 à 13 photons, soit une dynamique de 17dB.

Enfin nous montrons aussi la valeur du déphasage mesuré en fonction du nombre moyen de photons sur la figure 4.6. On voit clairement l'effet de saturation prédit précédemment. Cela nous permet d'établir la limite au-delà de laquelle le déphasage n'est plus proportionnel au nombre de photons, limite qui sera utile dans le chapitre suivant. On constate que, pour le désaccord de l'expérience $\delta = 120kHz$, l'interaction avec des champs de plus de 5 photons n'est plus dispersive puisque le déphasage n'est plus proportionnel au nombre moyen de photons. Pour des champs cohérents de moins de 5 photons en moyenne, la formule $\Delta\Phi = |\alpha|^2 \sin \Phi_0$ reste valable.

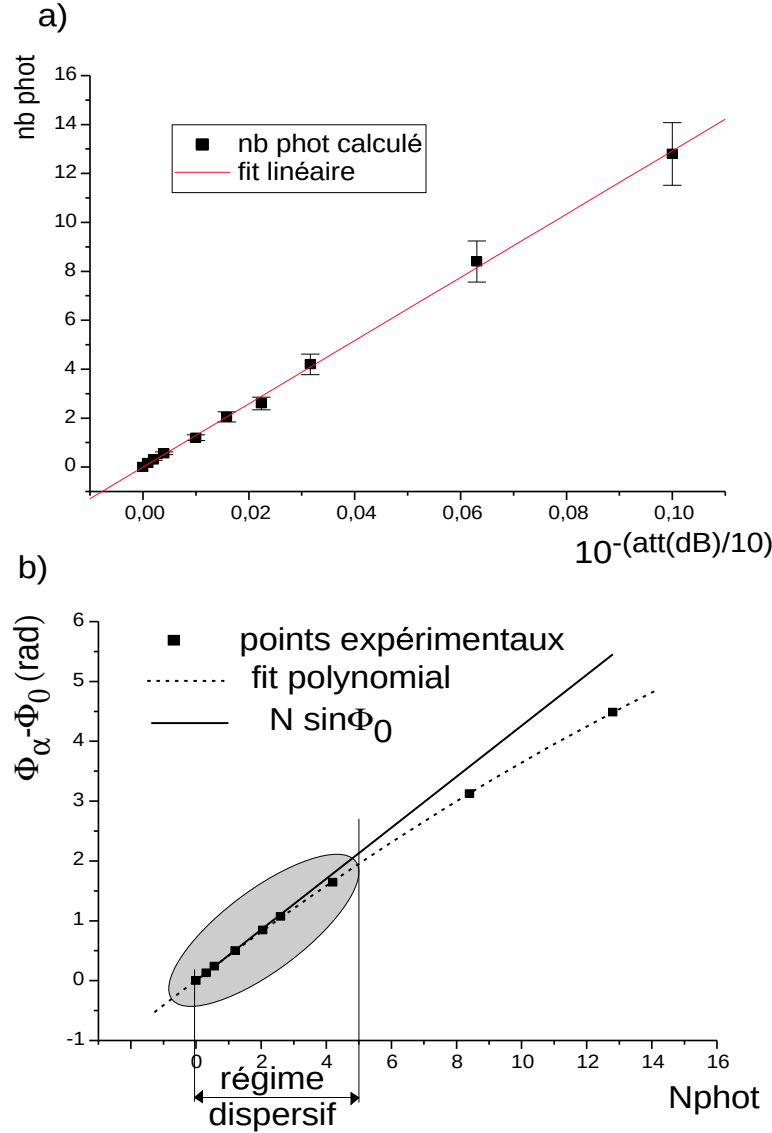


FIG. 4.6 – (a) Nombre de photons, calculé pour chaque point expérimental, à partir du déphasage mesuré, en fonction de l'atténuation relative $A^{(lin)} = 10^{-A^{(dB)}/10}$. On constate que tous les points sont parfaitement alignés. Cela valide notre procédure de calibration, et confirme expérimentalement la linéarité de la procédure d'injection du champ. (b) Déphasage mesuré en fonction du nombre de photons calculé. Les points marqués par des carrés sont les mesures. La courbe en pointillés est un ajustement numérique polynômial (degré 2) des valeurs expérimentales. En trait plein est tracée la courbe $\Phi = N \sin \Phi_0$, avec $\Phi_0 = 0.43 \text{ rad}$. On constate que cette loi n'est valable que pour les champs à petit nombre de photons, et on observe une saturation de la phase pour les grands nombres de photon. Tant que $N \leq 5$ l'écart entre les courbes en trait plein et en pointillés ne dépasse pas les 10%. On est alors dans le régime dispersif

4.2.3 L'interféromètre à lame séparatrice quantique

Après avoir détaillé la séquence de calibration, nous allons maintenant décrire le déroulement de la séquence expérimentale proprement dite.

- on commence par refroidir le mode Basse Fréquence en utilisant la séquence de “serpillères résonantes” montrée à la figure 3.3.

- puis nous envoyons un atome à 500m/s dans l'état excité $|e\rangle$ qui effectue une première impulsion de Ramsey R1 dans la cavité dans laquelle on a injecté précédemment un champ cohérent $|\alpha\rangle$, et une deuxième en-dehors du mode. On mesure le contraste des franges obtenues pour différentes valeurs de α .

Réglage de l'impulsion $\pi/2$ dans le mode On effectue l'impulsion R1 en accordant l'atome à résonance avec le mode BF comme indiqué à la figure 4.3. L'instant $t_{Stark}^{(\alpha)}$ où l'atome est mis en résonance avec le mode BF dépend bien sûr de α . Pour chaque champ injecté, nous devons donc régler $t_{Stark}^{(\alpha)}$.

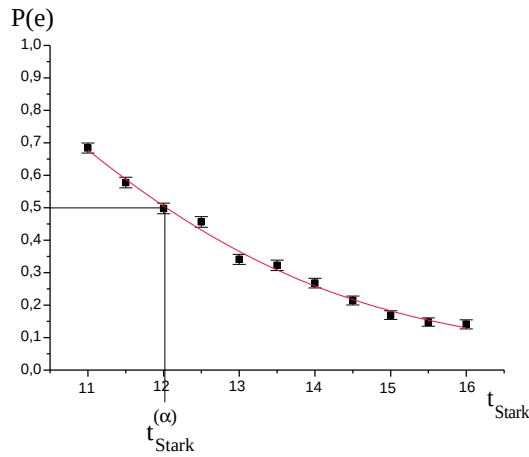


FIG. 4.7 – Réglage du paramètre $t_{Stark}^{(\alpha)}$ (ici l'atténuation est de 12dB ce qui correspond à $\alpha = 2.7$). On mesure la probabilité $P_e(t_{Stark})$ de détecter d'atome dans l'état $|e\rangle$ lorsqu'il est mis à résonance avec le mode à un instant t_{Stark} . On détermine $t_{Stark}^{(\alpha)}$ en disant que $P(t_{Stark}^{(\alpha)}) = 0.5$.

A la suite d'une impulsion $\frac{\pi}{2}$, les populations des niveaux $|e\rangle$ et $|g\rangle$ sont égalisées. Nous utilisons cette propriété pour régler l'instant de mise à résonance de l'atome $t_{Stark}^{(\alpha)}$ comme le montre la figure 4.7. On mesure la population atomique $P(e)$ dans le niveau $|e\rangle$ en fonction de l'instant t_{Stark} auquel l'atome est mis à résonance avec le mode BF, sans lui appliquer l'impulsion R2. Pour $t_{Stark} = t_{Stark}^{(\alpha)}$, $P(e) = 0.5$.

Réglage des franges Le déroulement précis de la séquence expérimentale est expliqué à la figure 4.8. L'atome entre dans la cavité $75\mu\text{s}$ après le passage de la dernière

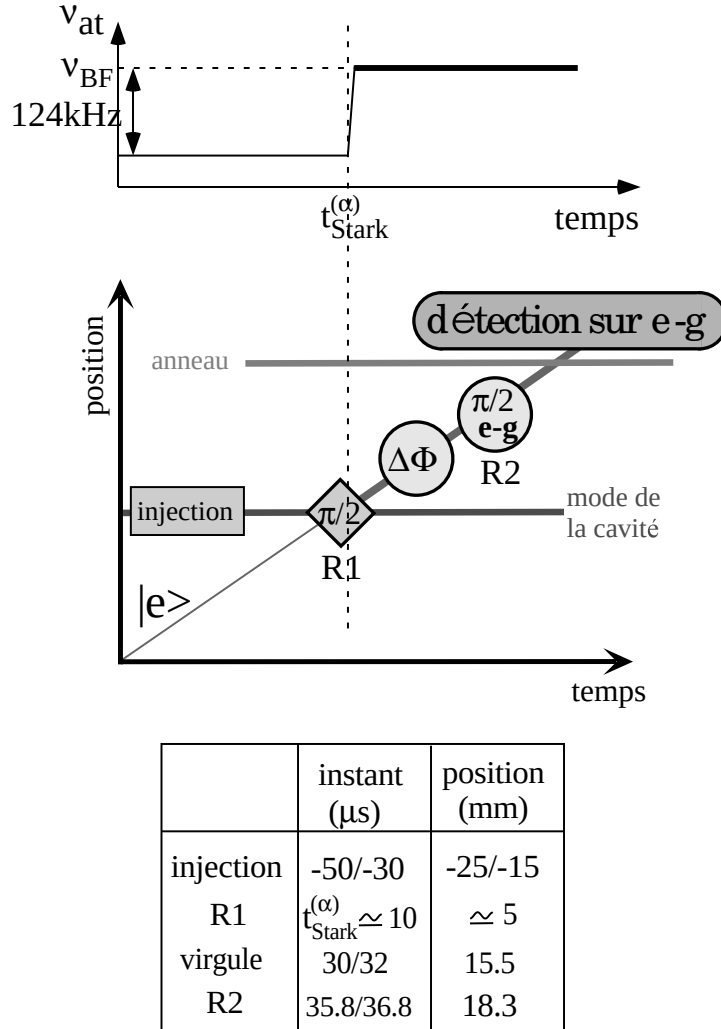


FIG. 4.8 – Séquence expérimentale des franges quantique/classique. L'atome, initialement hors résonance (trait fin), est mis à résonance avec le mode BF à un instant $t_{\text{Stark}}^{(\alpha)}$. Il y effectue une impulsion $\pi/2$ (zone R1). Puis il subit une deuxième impulsion dans la zone R2 classique. Entre les deux, une “virgule” permet de balayer la phase des franges. Les instants exacts de toutes ces opérations, ainsi que la position de l'atome dans la cavité à ce moment-là, sont donnés dans le tableau. (l'origine des temps et des positions correspond au passage de l'atome au centre de la cavité).

“serpillère atomique”. Il est initialement dans $|e\rangle$, désaccordé de $120kHz$ par rapport au mode BF à cause du champ Stark $E_{Stark} = E_{Stark}^{(2)}$. Avant qu’il ne pénètre dans le mode, la source S_{cav} est allumée pendant $20\mu s$. A l’instant $t_{Stark}^{(\alpha)}$, on commute le champ Stark à la valeur $E_{Stark}^{(1)}$ qui met l’atome à résonance avec le mode BF contenant le champ $|\alpha\rangle$. L’atome accomplit l’impulsion R1 pendant le reste de son interaction avec le mode.

Lorsqu’il est sorti du mode, on lui applique la “virgule” qui permet de balayer la phase des franges (voir paragraphe 3.2.2). Puis il est soumis à l’impulsion R2, produite par la source S_{cav} ¹. On a réglé préalablement la durée de l’impulsion R2 de manière à y effectuer une impulsion $\pi/2$. Enfin on détecte l’atome, une fois sur l’état $|e\rangle$, une fois sur $|g\rangle$, de manière à reconstituer la probabilité $P(g)$.

On obtient donc des franges pour chaque champ injecté. La figure 4.9 en montre trois, prises pour des champs différents. Comme prédit par le principe de complémentarité, les franges effectuées dans un champ $|\alpha\rangle$ faible sont peu contrastées, alors que dans un champ intense elles ont un contraste important (environ 75%).

4.2.4 Résultats expérimentaux

Nous avons étudié le contraste des franges quantiques pour des champs allant de 0.1 à 13 photons. La figure 4.10 montre les résultats que nous avons obtenus. Sur cette courbe, on constate que quand $N \gg 1$, le contraste des franges tend vers une valeur limite de 75%, et ne dépend plus du tout de l’état du champ dans la cavité. Ce contraste fini est du même ordre que celui des franges montrées à la figure 3.12, qui pourtant sont réalisées uniquement avec des impulsions classiques. Nous en concluons qu’il est parfaitement explicable par les diverses imperfections de notre dispositif : dispersion du paquet atomique, et champs inhomogènes. Lorsque le nombre de photons injectés est grand devant 1, rien ne distingue donc notre interféromètre d’un interféromètre de Ramsey classique.

En revanche, lorsque N_{phot} décroît et devient de l’ordre de 1, le contraste des franges diminue, ce qui prouve l’aspect quantique de l’interféromètre. Nous avons calculé au paragraphe 4.1.2 l’évolution du système à l’aide d’un modèle très simple, et nous avons trouvé que la probabilité de trouver l’atome dans $|g\rangle$ valait

$$P_g(\Delta\Phi) = \frac{1}{2} [1 + \text{Re}(\langle\Psi_e|\Psi_g\rangle \exp^{i\Delta\Phi})] \quad (4.17)$$

¹Notons que pour effectuer l’impulsion R2, la source S_{cav} n’est pas couplée par le même guide d’ondes que pour l’injection, mais par le guide d’ondes Ramsey. Nous utilisons un coupleur directionnel $3dB$ qui permet d’envoyer la micro-onde soit vers le guide “cavité”, soit vers le guide Ramsey. Sur chaque voie, une diode PIN contrôle la source si bien que nous pouvons coupler la source S_{cav} indépendamment par la cavité ou par l’anneau.

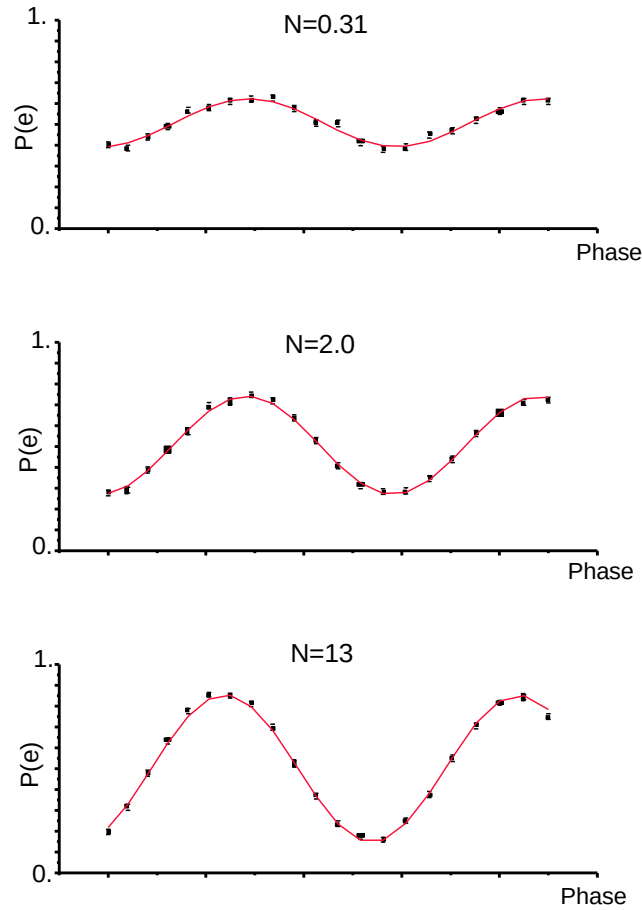


FIG. 4.9 – Franges quantiques obtenues pour trois champs différents. On constate que, pour les champs très atténués, le contraste est réduit, alors que dans un champ plus intense les franges sont très contrastées.

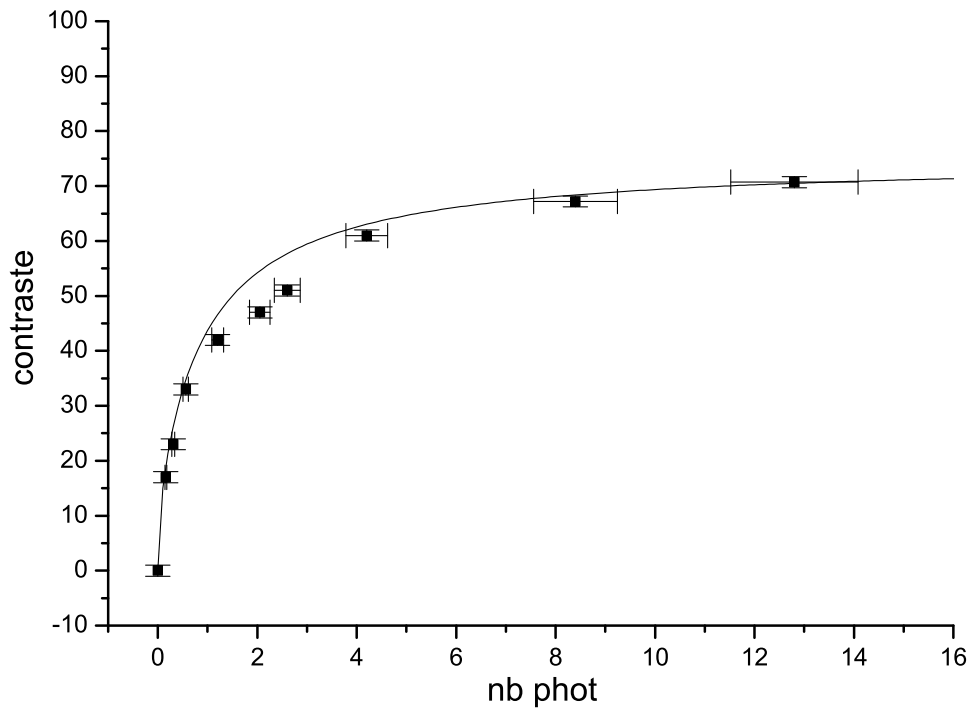


FIG. 4.10 – *Contraste des franges quantiques en fonction du nombre moyen de photons N dans la cavité. Quand $N \gg 1$, le contraste ne dépend plus de N , c'est le cas classique. Quand N au contraire diminue et tend vers 0, le contraste s'effondre et s'annule pour $N = 0$. La courbe en traits pleins est le résultat d'un calcul numérique sans aucun paramètre ajustable autre qu'un facteur multiplicatif global $f = 0.75$ ajusté sur le point dont le contraste est le plus élevé.*

les états $|\Psi_e\rangle$ et $|\Psi_g\rangle$ étant donnés par l'équation 4.4 que l'on rappelle ici :

$$\begin{aligned} |\Psi_e\rangle &= \sqrt{2} \left[\sum_n C_n \cos\left(\frac{\Omega_0}{2}\sqrt{n+1}t_\alpha\right) |n\rangle \right] \\ |\Psi_g\rangle &= \sqrt{2} \left[\sum_n C_n \sin\left(\frac{\Omega_0}{2}\sqrt{n+1}t_\alpha\right) |n+1\rangle \right] \end{aligned} \quad (4.18)$$

D'après la relation 4.17, le contraste des franges devrait être donné par $|\langle\Psi_e|\Psi_g\rangle|$. Ce produit scalaire peut être calculée numériquement. La courbe en traits pleins sur la figure 4.10 est le résultat de ce calcul sans aucun paramètre ajustable, hormis un facteur multiplicatif global qui rend compte du contraste limite de 75% dont nous avons parlé plus haut. L'accord excellent entre la courbe théorique et les points expérimentaux conforte notre interprétation de la réduction du contraste des franges en termes d'intrication avec la cavité [8]. Selon l'importance du champ dans la cavité, l'impulsion $\pi/2$ a pour effet d'intriquer l'atome avec la cavité, auquel cas la phase du dipôle atomique est mal définie et le contraste des franges est faible, ou bien, si le champ est grand, l'impulsion $\pi/2$ devient classique et l'état atomique est un état pur, auquel cas on retrouve un contraste important.

Notons pour finir qu'il n'a pas été possible de tester la relation de dualité énoncée au chapitre 1 entre visibilité des franges V et distinguabilité D des chemins dans l'interféromètre ($V^2 + D^2 \leq 1$). En effet, le paramètre D défini au chapitre 1 fait intervenir l'ensemble des observables qu'il est possible de mesurer avec le détecteur WP utilisé, or pour nous, le détecteur est le champ électromagnétique dans la cavité. Le seul moyen de mesurer D serait alors de caractériser complètement l'état quantique de la cavité corrélé à la détection de l'atome dans $|e\rangle$ ou dans $|g\rangle$, et de déterminer à partir de ces mesures quelle serait l'observable qui donnerait la meilleure information sur le chemin choisi par l'atome dans l'interféromètre. Au chapitre suivant, nous présenterons une méthode pour caractériser expérimentalement un état quantique par la mesure de sa fonction de Wigner, que l'on pourrait en théorie appliquer ici pour mesurer D . L'expérience serait en revanche trop complexe pour être réalisée avec le dispositif actuel.

4.2.5 Contraste des franges et fluctuations quantiques

Il est possible d'interpréter les résultats précédents en termes de la relation d'incertitude phase-nombre de photons. En effet, le contraste des franges mesure la cohérence de phase entre les deux impulsions, car $R1$ imprime sa phase sur la superposition d'états atomiques, qui est lue plus tard par l'impulsion $R2$. Le champ de l'impulsion $R2$ a une phase bien définie alors que celui de $R1$ présente des fluctuations quantiques $\Delta\Phi$, d'ordre $1/\Delta N$. Si l'on souhaite obtenir une information WP, il est nécessaire que $\Delta N < 1$, pour qu'il soit possible de détecter un changement de 1 photon. Par conséquent, on a alors $\Delta\Phi > 1$ et les franges disparaissent. De même, on peut rendre compte de la disparition des franges dans l'interféromètre de Mach-Zehnder quantique par la relation de Heisenberg $\Delta x \Delta p > \hbar$.

Ainsi, on peut fort bien rendre compte des contrastes observés dans les deux cas limite $N \rightarrow 0$ et $N \rightarrow \infty$. Pour comprendre quantitativement le contraste aux valeurs

intermédiaires de N , il est cependant nécessaire d'inclure l'intrication atome-cavité dans le modèle. En outre, l'utilisation de la relation d'incertitude semble dire que l'interaction entre l'atome et le champ dans la cavité conduit nécessairement à un déphasage inhomogène qui se traduit par le brouillage des interférences quantiques. Or il n'en est rien puisque, comme nous allons le voir, il est en réalité possible de retrouver des interférences si l'on efface de manière cohérente l'information enregistrée dans la zone de Ramsey.

4.3 La gomme quantique

4.3.1 Principe de l'expérience

Lien avec l'expérience EPR Plaçons-nous dans la situation où la cavité est initialement vide. Après l'impulsion $\pi/2$ effectuée dans la cavité, la présence ou non d'un photon dans la cavité témoigne du passage de l'atome par le chemin C_e ou le chemin C_g . Pour effacer cette information, il faut effectuer une mesure qui restaure l'ambiguïté sur l'état de l'atome après l'impulsion R1. Si l'on mesure simplement le nombre de photons contenus dans la cavité, il est clair que l'information n'est pas effacée, elle est au contraire révélée !

A l'inverse, si l'on mesure le champ dans la base $\{|\Phi_+\rangle = (|0\rangle + |1\rangle)/\sqrt{2}, |\Phi_-\rangle = (|0\rangle - |1\rangle)/\sqrt{2}\}$, on n'a plus aucune information sur l'état de l'atome après son interaction avec la cavité. L'état du système atome-champ après l'impulsion $\pi/2$ est

$$|\Psi\rangle = \frac{1}{\sqrt{2}}(|e, 0\rangle - i|g, 1\rangle) \quad (4.19)$$

que l'on peut récrire :

$$|\Psi\rangle = \frac{1}{\sqrt{2}}(|e, 0\rangle - i|g, 1\rangle) = \frac{1}{2\sqrt{2}}(|\Phi_+\rangle(|e\rangle - i|g\rangle) + |\Phi_-\rangle(|e\rangle + i|g\rangle)) \quad (4.20)$$

Sous cette forme, on constate que la détection du champ dans l'état $|\Phi_{\pm}\rangle$ "projette" l'atome sur l'état $(|e\rangle \mp i|g\rangle)/\sqrt{2}$. La cohérence atomique est donc restaurée conditionnellement à la détection du champ dans la base $|\Phi_-\rangle, |\Phi_+\rangle$.

Pour tester cela expérimentalement, on envoie un premier atome effectuer une première impulsion $\pi/2$ dans la cavité vide, puis on lui applique une impulsion $\pi/2$ classique. Nous avons déjà vu que dans ces conditions, on n'observe pas de franges d'interférence (il s'agit du point de la courbe 4.10 où $N = 0$). Le signal correspondant, montré à la figure 4.11 (courbe en losanges) confirme cela.

Puis on envoie un deuxième atome mesurer le champ dans la base $|\Phi_-\rangle, |\Phi_+\rangle$. Cet atome est initialement préparé dans $|g\rangle$. Il réalise une impulsion π dans le mode de la cavité. Puis on lui applique une impulsion micro-onde $\pi/2$. Après cette séquence, l'état de ce deuxième atome mesure le champ dans la base $|\Phi_-\rangle, |\Phi_+\rangle$. On peut alors

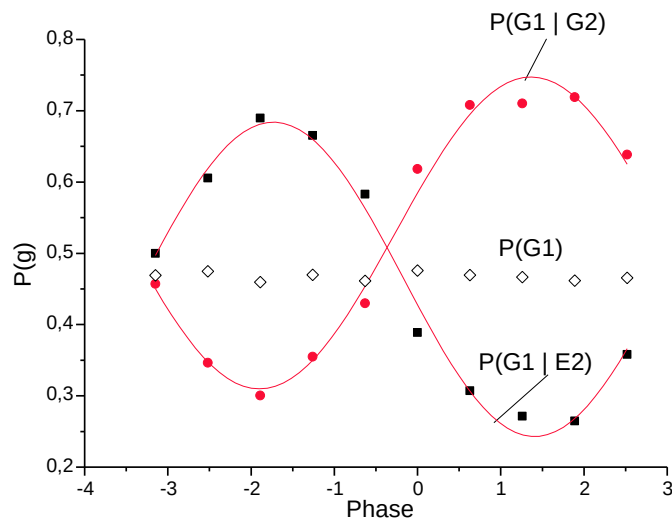


FIG. 4.11 – La courbe en losange montre la probabilité de détecter le premier atome dans $|g\rangle$, et ne présente pas de franges. Par contre, la probabilité de détecter le premier atome A1 dans $|g\rangle$ conditionnée à la détection d'un atome A2 dans $|g\rangle$ présente des franges d'interférence (cercles). La probabilité de détecter dans $|g\rangle$ conditionnée à la détection de A2 dans $|e\rangle$ présente aussi des franges d'interférence, mais en opposition de phase par rapport aux précédentes (carrés).

reconstruire la probabilité de détecter l'atome 1 dans $|g\rangle$ conditionnée à la détection du deuxième atome dans $|e\rangle$ ou dans $|g\rangle$ (ce qui revient, rappelons-le, à mesurer la cavité dans l'état $|\Phi_-\rangle$ ou dans l'état $|\Phi_+\rangle$) (respectivement représentée par des carrés ou des cercles à la figure 4.11). Ces deux courbes révèlent des franges d'interférence dont le signe est conditionné à l'état du deuxième atome. L'atome 2 se comporte bien comme une "gomme quantique" vis-à-vis de l'atome 1 : il restaure la cohérence initiale en effaçant de manière cohérente l'information contenue dans la cavité.

L'expérience dont nous présentons les résultats a en fait été réalisée dans notre groupe avant les résultats que nous avons présentés dans ce chapitre. Les signaux de la figure 4.11 peuvent en effet être interprétés comme la manifestation des corrélations transverses entre les deux atomes d'une paire EPR [40]. Ces deux interprétations d'une même expérience soulignent les liens très étroits entre le phénomène de l'intrication et la notion de complémentarité.

Gomme quantique non-conditionnelle Nous avons voulu aller plus loin que l'expérience présentée ci-dessus, où l'on retrouvait des franges pour l'atome 1 conditionnées à la détection d'un deuxième atome, et réaliser une gomme quantique non-conditionnelle. Le même atome qui a effectué l'impulsion $\pi/2$ va effacer de lui-même l'information laissée dans la cavité. Le principe de l'expérience est montré à la figure 4.12. L'atome entre dans l'état excité $|e\rangle$ dans la cavité, effectue une première impulsion $\pi/2$ R1 au début du mode, est désaccordé, puis remis à résonance à la fin du mode pour y effectuer une deuxième impulsion $\pi/2$ R2. Notons la similitude du principe de cette expérience avec un interféromètre de Mach-Zehnder dans lequel les deux lames séparatrices $L1$ et $L2$ font partie d'un même oscillateur harmonique quantique Ω . Au passage d'un photon P, l'oscillateur Ω reçoit une impulsion P_1 si le photon est réfléchi par $L1$ dans la voie b, et aucune impulsion s'il est transmis dans la voie a. Mais cette information est effacée au niveau de la lame $L2$ solidaire de $L1$. En effet, si on détecte un photon dans la voie indiquée sur la figure 4.12, il a communiqué une impulsion $P_2 = -P_1$ à Ω s'il provient de la voie a de l'interféromètre, et aucune impulsion s'il vient de la voie b. Quel que soit le chemin suivi dans l'interféromètre, l'état final de l'oscillateur quantique Ω reste inchangé. On retrouve un signal de franges non-conditionnelles sur le détecteur D.

De même, dans l'interféromètre de Ramsey, l'atome peut passer de $|e\rangle$ à $|g\rangle$ lors de la première impulsion de Ramsey ou lors de la deuxième, mais l'état final de la cavité ne donne aucune information. Nous allons voir cela en étudiant l'évolution de l'état du système atome-champ lors de cette séquence. Lors de la première impulsion R1,

$$|e, 0\rangle \xrightarrow{R1} \frac{1}{\sqrt{2}}(|e, 0\rangle - i|g, 1\rangle) \quad (4.21)$$

qui est l'état maximalement intriqué atome-cavité. Puis l'atome est désaccordé et déphasé d'une quantité $\Delta\Phi$ si bien qu'avant l'impulsion R2, le système est dans l'état

$$\frac{1}{\sqrt{2}}(\exp^{i\Delta\Phi} |e, 0\rangle - i|g, 1\rangle) \quad (4.22)$$

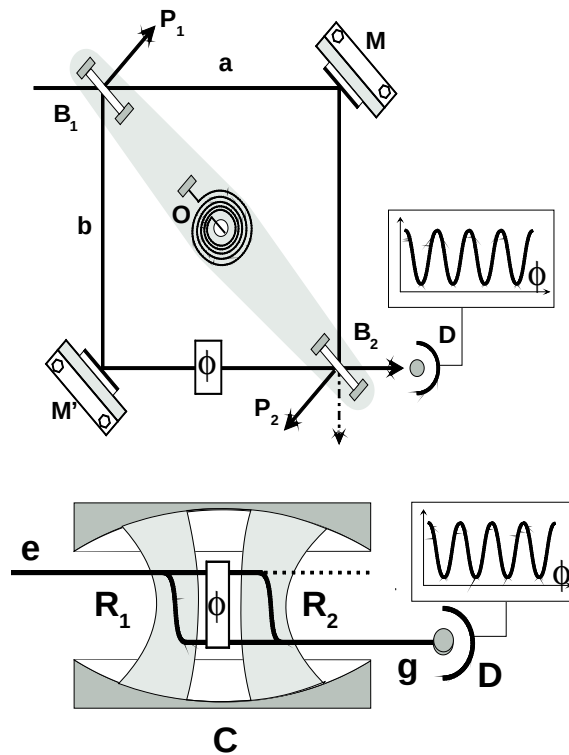


FIG. 4.12 – Principe de l'expérience de gomme quantique non-conditionnelle. L'atome, après avoir effectué une impulsion $\pi/2$ R_1 dans le mode de la cavité, est désaccordé puis réaccordé à la fin du mode pour réaliser une autre impulsion $\pi/2$ R_2 . Cette expérience est l'analogue d'un interféromètre de Mach-Zehnder avec deux lames séparatrices quantiques au lieu d'une seule.

Enfin l'impulsion R2 remélange les états e et g . L'état final est :

$$|\psi_{fin}\rangle = \frac{1}{2}((-1 + \exp^{i\Delta\Phi})|e, 0\rangle - i(1 + \exp^{i\Delta\Phi})|g, 1\rangle) \quad (4.23)$$

Ainsi, l'état final de la cavité ne conserve aucune trace de l'état atomique entre les deux impulsions $\pi/2$ qui ont lieu dans la cavité. Par contre, il est bien sûr corrélé à l'état final de l'atome, par la conservation de l'énergie. La matrice densité atomique réduite est

$$\begin{aligned} \rho &= \text{tr}_{\text{champ}}(|\Psi_{fin}\rangle\langle\Psi_{fin}|) \\ &= \frac{1}{4} \begin{pmatrix} |-1 + \exp^{i\Delta\Phi}|^2 & 0 \\ 0 & |1 + \exp^{i\Delta\Phi}|^2 \end{pmatrix} \end{aligned} \quad (4.24)$$

et par conséquent la probabilité de détecter l'atome dans $|g\rangle$ par exemple vaut

$$P(g) = \rho_{gg} = \frac{1}{4}|1 + \exp^{i\Delta\Phi}|^2 = \frac{1}{2}(1 + \cos \Delta\Phi) \quad (4.25)$$

et l'on s'attend bien à détecter des franges. Ceci ne contredit en rien le principe de complémentarité car après l'impulsion R2, nous n'avons plus aucun moyen de connaître l'état de l'atome entre R1 et R2.

4.3.2 Réalisation

Les réglages préliminaires, en vue de réaliser l'expérience de gomme quantique, sont les mêmes que pour l'expérience de complémentarité vue précédemment. Le mode Basse Fréquence est accordé à la fréquence des atomes dans le champ Stark $E_{Stark}^{(1)}$. La séquence de refroidissement avant l'expérience est également inchangée.

Le déroulement de la séquence expérimentale est montré à la figure 4.13. C'est le champ électrique E_{stark} qui permet de régler les deux impulsions $\pi/2$ dans le mode. L'atome entre dans la cavité dans le champ $E_{Stark}^{(1)}$ qui le met à résonance, puis il est désaccordé à un instant $t_{\pi/2}^{(1)}$ dans un champ $E_{Stark}^{\Delta\Phi}$. Pendant que l'atome est désaccordé par rapport à la cavité, l'état $|e, 0\rangle$ se déphase par rapport à l'état $|g, 1\rangle$ de la quantité $\Delta\Phi$. Le champ est remis à sa valeur initiale $E_{Stark}^{(1)}$, remettant ainsi l'atome à résonance à la fin du mode pour sa deuxième impulsion $\pi/2$, à un instant $t_{\pi/2}^{(2)}$. On balaye la phase des franges en faisant varier le champ $E_{Stark}^{\Delta\Phi}$ qui met l'atome hors résonance.

4.3.3 Résultats

Les résultats ainsi obtenus sont sur la figure 4.14. On retrouve effectivement des franges ayant un contraste de 60%. Ce contraste reste cependant légèrement inférieur au contraste maximal des franges de Ramsey obtenues avec un interféromètre classique (qui est plutôt de 80%). En effet, la présence dans le mode Haute Fréquence d'un champ thermique d'environ 1 photon provoque un déphasage inhomogène du dipôle atomique

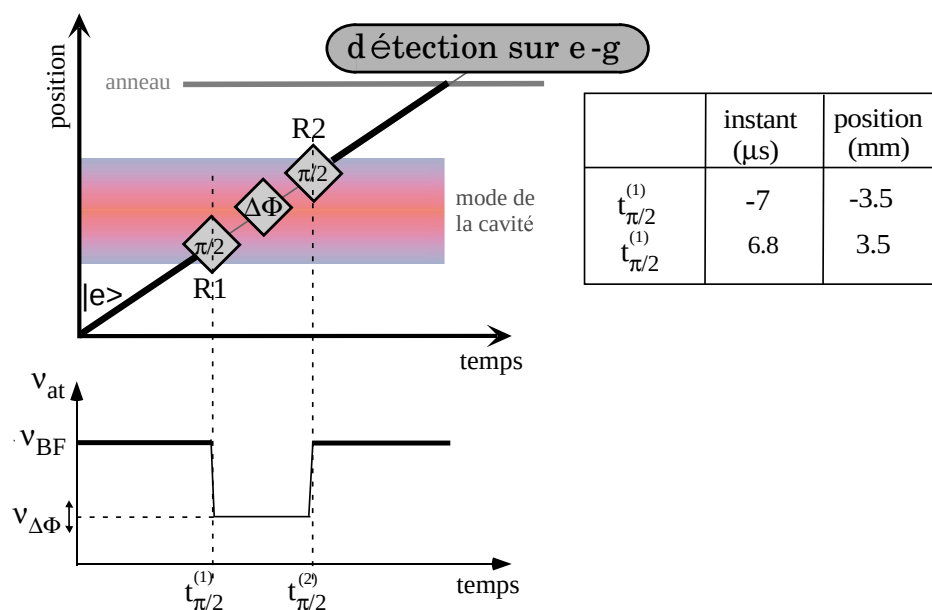


FIG. 4.13 – Réalisation expérimentale de l’expérience de “gomme quantique” non-conditionnelle. L’atome est initialement préparé dans $|e\rangle$. Il effectue une première impulsion $\pi/2$ dans le mode. La fréquence de la transition atomique est alors désaccordée brutalement par effet Stark, à une valeur $\nu_{\Delta\Phi}$. Il est ensuite réaccordé à la fin de son passage dans le mode, pour y effectuer une deuxième impulsion $\pi/2$. On fait varier la phase des franges en modifiant $\nu_{\Delta\Phi}$.

qui se traduit par une réduction de contraste ². Cependant les résultats précédents prouvent que la perte de cohérence qui se manifestait par la diminution du contraste des franges lorsque α devenait petit à la figure 4.10 n'était pas due à un brouillage irréversible de la phase du dipôle atomique, mais bien et uniquement à l'intrication de l'atome avec la cavité. Ils prouvent également que les fluctuations quantiques du champ électrique dans la cavité, même vide, ont un temps de corrélation long devant l'intervalle de temps qui sépare les deux impulsions quantiques de la figure 4.12 (ce temps de corrélation est en fait bien sûr T_{cav}).

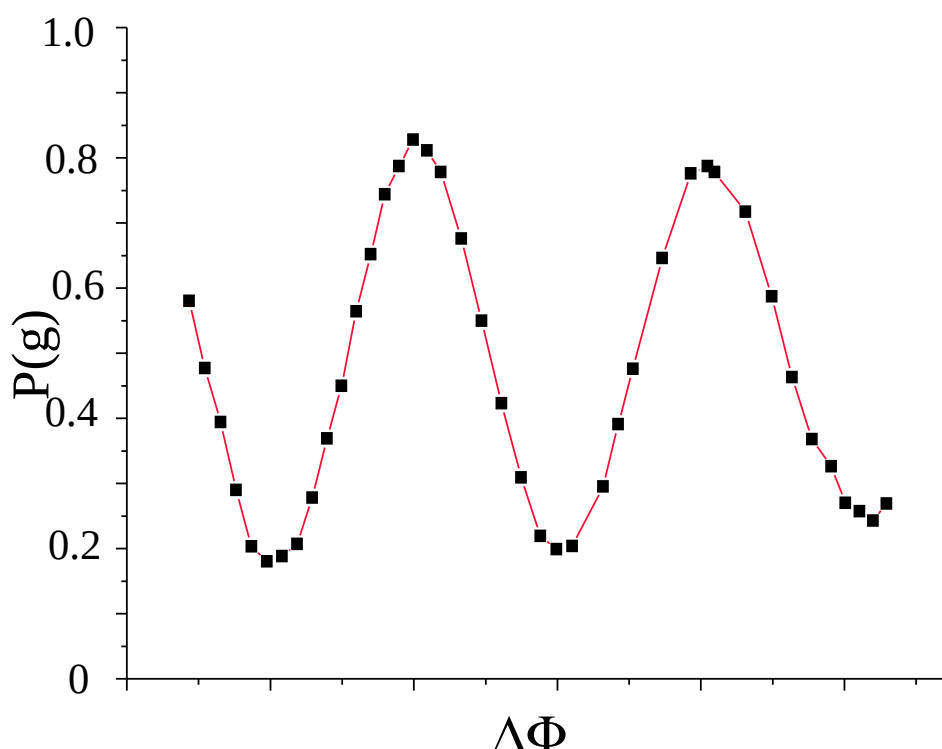


FIG. 4.14 – Résultats de l'expérience de gomme quantique.

4.4 Qu'est-ce qu'un champ classique ?

Dans ce paragraphe nous souhaiterions revenir sur la distinction entre un champ dit "classique" et un champ quantique, à la lumière des résultats précédents.

Dans le contexte de l'interférométrie de Ramsey, un champ classique doit agir sur l'atome comme une lame séparatrice classique dans l'espace des énergies internes. Plus

²Dans l'expérience précédente, cet effet ne jouait aucun rôle car l'atome n'était mis à résonance avec la cavité qu'à la fin de son interaction avec le mode ; pendant la plus grande partie de son interaction dispersive avec le mode HF il était dans $|e\rangle$, et donc insensible aux déphasages inhomogènes.

mathématiquement, la dynamique de l'atome doit être correctement décrite par un hamiltonien d'interaction semi-classique du type $\hat{H}_{sc} = \hat{H}_{at} - \hat{\mathbf{d}} \cdot \mathbf{E}_0 \cos(\omega t)$. Nous avons montré expérimentalement que tel était le cas pour un champ dans un état cohérent $|\alpha\rangle$ à la condition que $\alpha \gg 1$, car lorsque $\alpha \sim 1$, l'atome s'intrique à la cavité.

Cette constatation nous amène à poser une autre question. La description semi-classique est-elle réellement valable lorsque les impulsions de Ramsey sont appliquées, comme nous l'avons expliqué au chapitre 3, dans des modes de la structure "cavité+anneau" ?

En effet, on peut évaluer grossièrement le nombre moyen de photons dans une impulsion de Ramsey. Dans les expériences que je présente dans ce mémoire, les impulsions de Ramsey sont appliquées dans des modes de faible facteur de qualité de la structure cavité+anneau. Le volume de ces modes est difficile à connaître exactement, mais il est sans doute de quelques λ^3 . Au chapitre 5, nous verrons une expérience où la durée d'une impulsion $\pi/2$ était de $8\mu s$. Cela implique que le nombre moyen de photons dans le mode était de l'ordre de l'unité ! Cela signifierait-il, comme semble l'indiquer la courbe 4.10, que le contraste de ces franges doit être réduit ?

L'explication de ce paradoxe apparent peut être trouvée dans un article récent [49], et est simple à comprendre. Lorsque l'atome interagit avec la cavité supraconductrice, nous sommes dans le régime du couplage fort où la fréquence de Rabi est plus grande que le taux de dissipation : $\Omega_0 \gg \gamma_{cav}$, où $\gamma_{cav} = 1/T_{cav} \simeq 1kHz$ alors que $\Omega_0 = 50kHz$. Au contraire, les modes de la structure "cavité+anneau" ont un temps d'amortissement que nous n'avons pas mesuré explicitement, mais dont il semble raisonnable de dire qu'il n'excède pas $100ns$. Cela donne un taux d'amortissement γ_{Ram} d'au moins $10MHz$, et donc un rapport Ω/γ_{Ram} inférieur à $1/100$. Ainsi, au cours d'une impulsion $\pi/2$, il y a certes un photon en moyenne dans le mode, mais ce photon est "recyclé" une centaine de fois, donc pour l'atome, tout se passe comme si le champ était beaucoup plus grand. La dynamique du système est dominée par le couplage incohérent à un réservoir à mémoire courte. Les corrélations quantiques entre l'atome et la cavité Ramsey n'ont pas le temps d'apparaître, et l'on retrouve la dynamique semi-classique habituelle. C'est donc la dissipation qui permet de retrouver le comportement classique de l'atome dans une impulsion de Ramsey "classique".

Bien sûr, ce raisonnement est très qualitatif et insuffisant. On pourra trouver dans un article récent [49] un modèle détaillé qui part d'un traitement quantique de l'atome, du champ dans le mode, et du réservoir auquel il est fortement couplé par la dissipation, et qui montre comment l'on retrouve le hamiltonien semi-classique habituel dans la limite où $\Omega/\gamma \ll 1$, même lorsque le nombre moyen de photons est faible.

En résumé, le degré d'intrication entre l'atome et le champ est le bon critère pour déterminer si un champ est "classique" ou non. Lorsque le champ est très bien isolé, l'atome s'intrique avec lui si le nombre de photons est trop faible. Mais si le champ est couplé à un réservoir, il n'y aura pas d'intrication entre l'atome et le champ, même si le nombre moyen de photons est inférieur à 1.

Chapitre 5

Détermination directe de la fonction de Wigner d'un état de Fock à un photon

Dans ce chapitre j'aborde l'autre volet de ma thèse qui a porté sur la caractérisation complète de l'état quantique de la cavité par la mesure directe de sa fonction de Wigner. Les travaux des étudiants m'ayant précédé dans l'équipe ont montré que nous étions capables de produire des états non-classiques du champ : Jochen Dreyer a notamment présenté dans son mémoire la production d'un état dit “chat de Schrödinger”, superposition cohérente d'états mésoscopiquement distincts du champ [27]. Xavier Maître a réalisé une expérience démontrant la production d'un état $(|0\rangle + |1\rangle)/\sqrt{2}$ [63]. Enfin, nous avons montré comment un seul atome pouvait générer un état de Fock à deux photons en utilisant un couplage du troisième ordre entre l'atome et les deux modes de la cavité [7].

Jusque-là, nous sondions l'état produit dans la cavité en envoyant un atome de Rydberg circulaire qui mesurait généralement une toute petite partie de la matrice densité (voir par exemple [27]). Il serait intéressant de pouvoir *complètement* caractériser un état donné de la cavité. Cela permettrait notamment de reconstituer la dynamique de la relaxation d'un état “chat de Schrödinger”.

La fonction de Wigner du champ $W(\alpha)$, définie sur l'espace des phases d'un oscillateur, caractérise aussi complètement son état que la donnée de la matrice densité ρ . C'est l'analogue quantique de la distribution de probabilité dans l'espace des phases en mécanique statistique classique. Elle est très utile dans l'étude de la transition entre comportement quantique et comportement classique car elle est positive pour tout état classique tandis qu'elle prend des valeurs négatives pour une large classe d'états, ce qui est bien sûr interdit classiquement et constitue donc un critère fort de “non-classicité”.

Nous verrons dans la deuxième partie que la fonction de Wigner W est la valeur moyenne d'une observable en chaque point de l'espace des phases. Lutterbach et Davidovich [58] et Englert [36] ont suggéré indépendamment une méthode (dans la suite,

nous utiliserons l'abréviation LD) pour mesurer cette observable dans notre dispositif d'électrodynamique quantique en cavité que nous présentons en deuxième partie.

Le reste du chapitre est consacré à la présentation des résultats expérimentaux que nous avons obtenus en appliquant la méthode LD dans le cas du vide, et d'un état de Fock à un photon.

5.1 La fonction de Wigner

Dans cette partie, nous allons très brièvement présenter quelques propriétés fondamentales de la fonction de Wigner sans aucun souci d'exhaustivité. Le lecteur désireux d'approfondir cette question consultera les références [23] et [44].

5.1.1 Présentation générale

La fonction de Wigner appartient à la classe plus générale des quasi-distributions de probabilité, qui sont des fonctions définies sur l'espace des phases d'un oscillateur harmonique. Elle permet de calculer la valeur moyenne d'un opérateur $\hat{A}$, habituellement calculée par la formule

$$\langle \hat{A} \rangle = \text{Tr}(\hat{A}\hat{\rho}), \quad (5.1)$$

comme l'intégrale sur l'espace des phases d'une fonction $f_s^{(\hat{A})}(\alpha)$ pondérée par la fonction de Wigner $W(\alpha)$ (notons que dans toute la suite, nous considérerons uniquement le cas du champ électromagnétique, et nous utiliserons par conséquent la variable canonique complexe α plutôt que les variables p et q , ainsi que les opérateurs de création et d'annihilation associés a et $a^\dagger$) :

$$\langle \hat{A} \rangle = \int d^2\alpha W(\alpha, \alpha^*) f_s^{(\hat{A})}(\alpha, \alpha^*) \quad (5.2)$$

Dans cette formule, $W(\alpha)$ contient toute l'information sur l'état dans lequel se trouve le système, tandis que la fonction $f_s^{(\hat{A})}(\alpha, \alpha^*)$ est reliée à l'opérateur $\hat{A}$ par la règle proposée par Weyl en 1927 [95]. Elle consiste à développer $\hat{A}$ comme une série de termes contenant des produits d'opérateurs a et $a^\dagger$ ordonnés "symétriquement" :

$$\hat{A} = \sum_{n,m=0}^{\infty} f_{n,m}^{(\hat{A})} \cdot \{(\hat{a}^\dagger)^n \hat{a}^m\} \quad (5.3)$$

et de définir $f_s^{(\hat{A})}$ comme :

$$f_s^{(\hat{A})}(\alpha, \alpha^*) = \sum_{n,m=0}^{\infty} f_{n,m}^{(\hat{A})} \cdot (\alpha^*)^n \alpha^m \quad (5.4)$$

Par définition, l'opérateur $\{(\hat{a}^\dagger)^n \hat{a}^m\}$ est la moyenne de toutes les différentes manières d'ordonner les opérateurs $(\hat{a}^\dagger)^n$ et $(\hat{a})^m$. Prenons pour exemple l'opérateur $(\hat{a}^\dagger)^2 \hat{a}^2$. On a

$$\{(\hat{a}^\dagger)^2 \hat{a}^2\} = \frac{1}{6}((\hat{a}^\dagger)^2 \hat{a}^2 + \hat{a}^\dagger \hat{a} \hat{a}^\dagger \hat{a} + \hat{a}^\dagger \hat{a}^2 \hat{a}^\dagger + \hat{a} (\hat{a}^\dagger)^2 \hat{a} + \hat{a} \hat{a}^\dagger \hat{a} \hat{a}^\dagger + \hat{a}^2 (\hat{a}^\dagger)^2) \quad (5.5)$$

Supposons que l'on souhaite calculer la valeur moyenne de l'opérateur $\hat{N} = \hat{a}^\dagger \hat{a}$ à l'aide de la fonction de Wigner. On écrit

$$\hat{N} = \hat{a}^\dagger \hat{a} = \frac{1}{2}(\hat{a}^\dagger \hat{a} + \hat{a} \hat{a}^\dagger - 1) = \{\hat{a}^\dagger \hat{a}\} - \frac{1}{2} \quad (5.6)$$

Alors, d'après 5.2,

$$\langle \hat{N} \rangle = \int d^2 \alpha W(\alpha) \alpha^* \alpha - \frac{1}{2} \quad (5.7)$$

On peut montrer que les propriétés précédentes suffisent à déterminer W de manière unique pour un état donné. Le lien avec la matrice densité de l'état considéré est fourni par la formule suivante :

$$W(\alpha, \alpha^*) = \frac{1}{\pi} \int d^2 \lambda \exp^{\alpha \lambda^* - \alpha^* \lambda} Tr[\rho \exp^{-i(\lambda^* \hat{a} - \lambda \hat{a}^\dagger)}] \quad (5.8)$$

La forme de l'équation 5.2 est intéressante en ceci qu'elle ressemble tout-à-fait à la formule permettant de calculer la valeur moyenne d'une grandeur physique f pour une particule dont la densité de probabilité de présence dans l'espace des phases serait donnée par $W(\alpha)$; l'utilisation de la fonction de Wigner dans un problème est donc tout-à-fait indiquée pour en discuter la limite classique¹. On pourrait se demander ce qui distingue l'équation 5.2, qui est pourtant parfaitement exacte et quantique, de la formule analogue classique. La réponse doit être cherchée dans la forme des fonctions W autorisées par la mécanique quantique. Nous verrons en effet dans le paragraphe suivant que la fonction de Wigner d'un oscillateur harmonique quantique peut être notablement différente d'une distribution de probabilité classique.

5.1.2 Quelques propriétés de la fonction de Wigner

- La fonction de Wigner $W(\alpha)$ est normée. En effet, reprenons la formule 5.2 dans le cas où $\hat{A}$ est l'identité. On obtient

$$\int d^2 \alpha W(\alpha) = \langle \hat{I} \rangle = Tr(\rho) = 1 \quad (5.9)$$

¹elle fut d'ailleurs introduite en 1932 par Wigner pour calculer les corrections quantiques à la mécanique statistique classique [98].

- une propriété importante distingue tout particulièrement la fonction de Wigner parmi toutes les autres quasi-distributions de probabilité et lui donne un contenu physique particulièrement riche. La distribution marginale de probabilité de n'importe quelle quadrature du champ s'obtient en intégrant la fonction de Wigner selon la quadrature qui lui est orthogonale. Par exemple, si l'on écrit $\alpha = \alpha_1 + i\alpha_2$, on aura

$$p(\alpha_1) = \int W(\alpha) d\alpha_2 \quad (5.10)$$

Dans cette expression, $p(\alpha_1)$ est la densité de probabilité que la variable $Re(\alpha)$ prenne la valeur α_1 (elle est donc bien sûr toujours positive ou nulle). La figure 5.1 illustre cette propriété dans le cas de la fonction de Wigner d'un état de Fock à un photon. Cette propriété a été exploitée pour réaliser des mesures de W par une méthode dite "tomographique" que nous évoquerons plus loin.

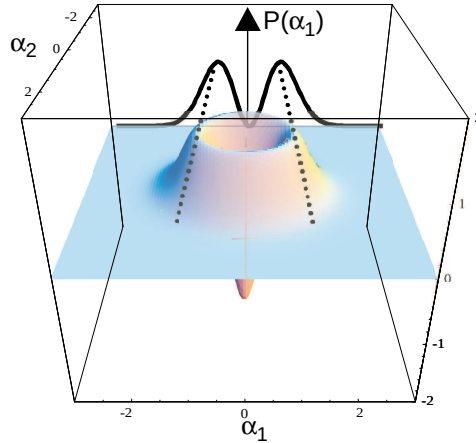


FIG. 5.1 – La distribution marginale de probabilité d'une quadrature du champ peut être obtenue en intégrant la fonction de Wigner selon l'axe de la quadrature complémentaire. Ici, on obtient la probabilité $P(\alpha_1)$ en intégrant $W(\alpha)$ selon α_2 .

- on peut montrer que $W(\alpha)$ doit être bornée :

$$-\frac{2}{\pi} \leq W(\alpha) \leq +\frac{2}{\pi} \quad (5.11)$$

- en outre, on peut montrer que la fonction de Wigner W doit être nécessairement une fonction continue de α . De ces trois propriétés (W doit être normalisée, bornée, et continue) on déduit que la fonction de Wigner d'un état quelconque du champ doit présenter une certaine régularité, ce qui la distingue déjà d'une densité de probabilité classique. Par exemple, l'état de plus basse énergie serait représenté en mécanique classique par une fonction delta de Dirac à l'origine. Une telle distribution n'a pas les propriétés de régularité évoquées plus haut et n'est donc pas une fonction de Wigner

acceptable. La fonction de Wigner de l'état fondamental, pour être normée et bornée, doit nécessairement être non-nulle sur une certaine étendue de l'espace des phases, ce qui traduit l'existence de fluctuations quantiques des deux quadratures du champ même dans l'état fondamental.

D'autre part, soit $|\phi\rangle$ et $|\psi\rangle$ deux états purs distincts. On peut exprimer le produit scalaire de ces deux états en termes des fonctions de Wigner $W_\phi(\alpha)$ et $W_\psi(\alpha)$ associées à chaque état :

$$|\langle\phi|\psi\rangle|^2 = \langle\phi|\psi\rangle\langle\psi|\phi\rangle = \langle|\psi\rangle\langle\psi|\rangle_\phi = \langle\rho_\psi\rangle_\phi \quad (5.12)$$

Dans cette relation, nous considérons l'opérateur densité ρ_ψ représentant l'état $|\psi\rangle$ comme une observable quelconque, dont on peut calculer la valeur moyenne dans l'état $|\phi\rangle$. Cet opérateur ρ_ψ peut être exprimé comme une somme de produits d'opérateurs $\hat{a}$ et $\hat{a}^\dagger$ écrits dans l'ordre symétrique (règle de Weyl) :

$$\rho_\psi = f_\psi(\hat{a}, \hat{a}^\dagger) \quad (5.13)$$

On peut alors calculer sa valeur moyenne en utilisant la propriété 5.2 :

$$\langle\rho_\psi\rangle_\phi = \int W_\phi(\alpha) f_\psi(\alpha) d^2\alpha \quad (5.14)$$

Or il se trouve que la fonction $f_\psi(\alpha)$ n'est rien d'autre que $W_\psi(\alpha)$ (cette propriété est spécifique de la fonction de Wigner W et la distingue des autres distributions dans l'espace des phases. Elle est démontrée dans l'article de Glauber [23]). Donc finalement,

$$|\langle\phi|\psi\rangle|^2 = \int W_\phi(\alpha) W_\psi(\alpha) d^2\alpha \quad (5.15)$$

La relation 5.15 est particulièrement intéressante lorsque les états $|\psi\rangle$ et $|\phi\rangle$ sont orthogonaux. Alors nécessairement, $\int W_\phi(\alpha) W_\psi(\alpha) d^2\alpha = 0$. Donc, supposons que la fonction de Wigner d'un état du champ soit strictement positive sur tout l'espace des phases (nous verrons que c'est le cas pour tous les champs "classiques"); on voit bien que la fonction de Wigner de tous les états orthogonaux (et dans l'espace de Hilbert de dimension infinie des états du champ il y en a beaucoup!) prendront nécessairement des valeurs négatives.

Ainsi, contrairement aux densités de probabilité classiques sur l'espace des phases, la fonction de Wigner de certains états peut prendre des valeurs négatives. On conçoit bien, à la lumière de tout ce qui vient de précéder, que le fait, pour un état du champ, d'avoir une fonction de Wigner qui prend des valeurs négatives peut être considéré comme un critère de non-classicalité important. Par exemple, un état cohérent de grande amplitude ne manifeste son caractère quantique que par ses fluctuations : sa fonction de Wigner pourrait être la densité de probabilité d'un oscillateur classique dont on ne connaît parfaitement ni la position ni le moment. En revanche, l'état de Fock $|n=1\rangle$ a une fonction de Wigner qui ne pourrait pas être la densité de probabilité dans l'espace des phases d'un oscillateur classique puisqu'il est orthogonal à l'état $|0\rangle$ dont la fonction de Wigner est partout positive.

5.1.3 Quelques fonctions de Wigner

Pour conclure cette première partie théorique, nous allons donner sans démonstration les expressions de fonctions de Wigner dont nous aurons besoin par la suite.

La fonction de Wigner du vide et plus généralement de n'importe quel état cohérent $|\beta\rangle$ est une gaussienne de largeur 1 centrée sur β comme on peut le constater sur la figure 5.2 :

$$W_\beta(\alpha) = \frac{2}{\pi} \exp^{-2|\alpha-\beta|^2} \quad (5.16)$$

Ainsi, pour un état cohérent du champ, le module du champ électrique vaut en moyenne $|\beta|$ mais présente des fluctuations quantiques d'amplitude moyenne 1. Dans certains cas, on peut négliger ces fluctuations et considérer que l'amplitude et la phase sont parfaitement bien définis; on traite alors le champ électromagnétique comme une variable classique. Ce traitement est bien sûr d'autant mieux justifié que l'amplitude moyenne du champ est importante.

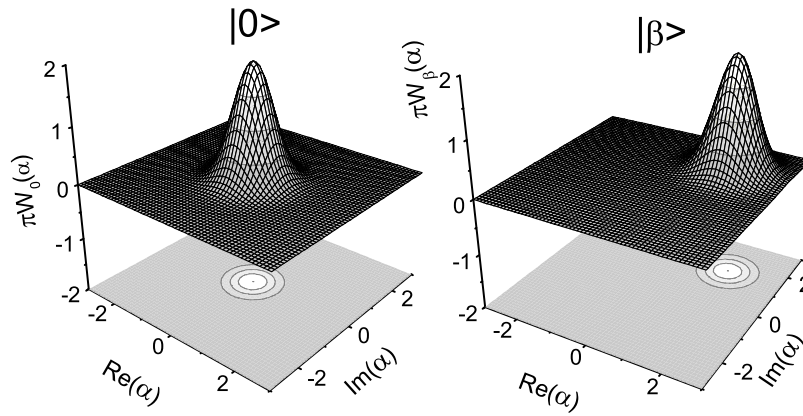


FIG. 5.2 – Fonctions de Wigner du vide (à gauche) et d'un état cohérent $|\beta\rangle$ avec $\beta = 1.5 + 1.5i$.

Donnons l'expression exacte des fonctions de Wigner des états de Fock $|n\rangle$:

$$\pi W_n(\alpha) = 2(-1)^n L_n(4|\alpha|^2) \exp(-2|\alpha|^2) \quad (5.17)$$

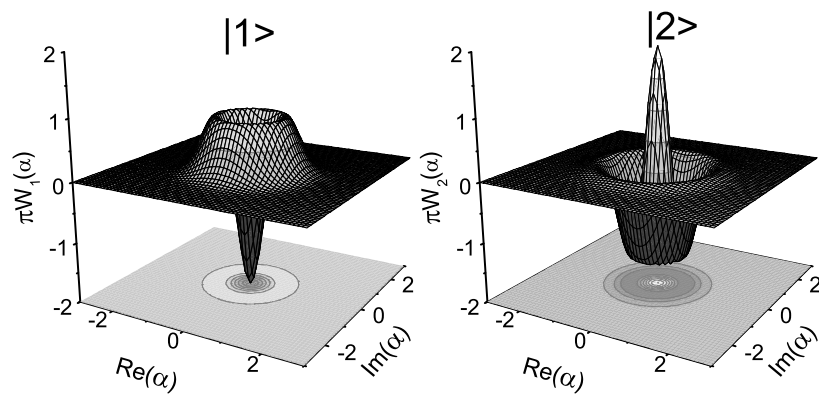


FIG. 5.3 – Fonctions de Wigner de l'état $|1\rangle$ (à gauche) et de $|2\rangle$

Dans cette expression, $L_n(x)$ est le n -ième polynôme de Laguerre (rappelons que $L_0(x) = 1$, $L_1(x) = 1 - x$, $L_2(x) = 1 - 2x + x^2$). On a représenté les fonctions de Wigner des états $|1\rangle$ et $|2\rangle$ à la figure 5.3. Puisque le polynôme de Laguerre d'ordre n a la propriété d'avoir n racines réelles, on voit que $W_n(\alpha)$ oscille d'autant plus que n est grand.

Il est également intéressant de comparer la fonction de Wigner des états de Fock, que nous venons de voir, avec celle d'un champ thermique au nombre moyen de photons N_{th} . La fonction de Wigner d'un mélange statistique d'états de Fock de poids p_n est bien sûr la somme des fonctions de Wigner W_n pondérées par les poids p_n . Vu le comportement très oscillant des fonctions W_n , on pourrait s'attendre à ce que la fonction de Wigner d'un champ thermique W_{th} oscille aussi. Il n'en est rien ; le résultat de la pondération est que W_{th} est une simple gaussienne de largeur $\sqrt{2N_{th} + 1}$.

Après ces préliminaires théoriques, nous allons rapidement passer en revue les expériences qui ont déjà été réalisées sur la mesure de la fonction de Wigner d'un système quantique dans un état donné.

5.1.4 Mesures de la fonction de Wigner déjà effectuées

5.1.4.a Méthode tomographique

Le principe Un grand nombre de mesures de la fonction de Wigner ont été réalisées dans le domaine optique, sur des champs propagatifs, par une méthode dite “tomographique”. Nous avons déjà mentionné au paragraphe 5.1.2 que la distribution de Wigner permettait de déterminer la distribution marginale de probabilité de n'importe quelle quadrature du champ. L'inverse est également vrai : si l'on connaît les distributions marginales $P(X_\theta)$, il est possible de reconstruire la fonction de Wigner [9].

L'idée des méthodes de reconstruction en optique quantique est de mesurer les distributions marginales $P(X_\theta)$ par détection homodyne balancée [92]. La figure 5.4 montre le principe d'un dispositif de détection homodyne balancée. Le champ que l'on souhaite mesurer (faisceau “signal”) est mélangé sur une lame séparatrice réfléchissante à 50% à un oscillateur local de même fréquence, ayant une différence de phase θ avec le signal. La différence des photocourants dans les deux voies est proportionnelle à la valeur de la quadrature du champ X_θ . En recommençant l'expérience un grand nombre de fois, et ce pour différentes valeurs de θ , on peut ainsi reconstituer les distributions marginales de probabilité des quadratures du champ. Il faut ensuite filtrer les données puisque l'on a échantillonné la phase θ puis appliquer une transformée de Radon inverse.

Cette méthode a été appliquée pour la première fois en 1993 par Smithey et al. [85] au vide squeezé. Schiller et al. [81] puis Breitenbach et al. [15] ont développé cette technique et fait des mesures des fonctions de Wigner de tous types d'états squeezés. Kurtsfeier et al. [52] ont mesuré la fonction de Wigner d'un ensemble d'atomes d'Helium métastables après passage à-travers deux fentes d'Young, par des méthodes tomographiques. Plus récemment, Lvovsky et al. [60] ont effectué une mesure tomogra-

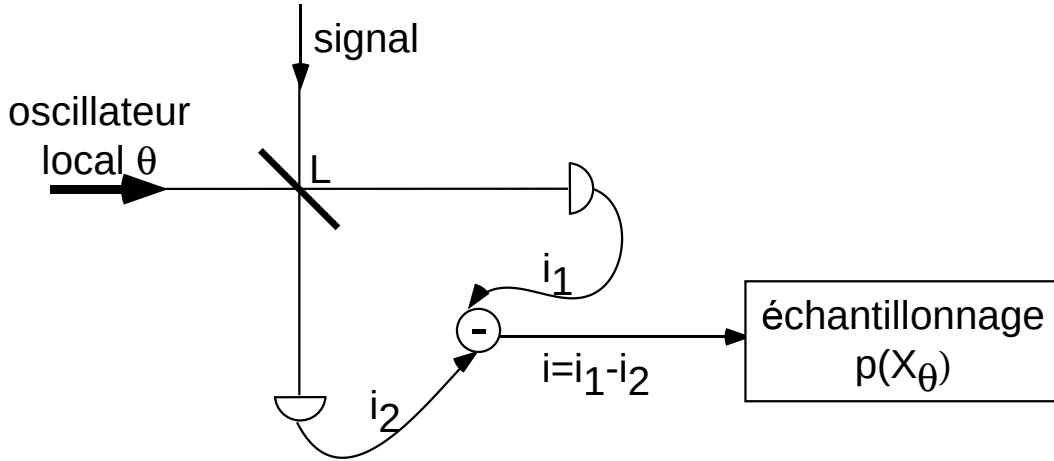


FIG. 5.4 – Schéma de principe d’une détection homodyne. Le faisceau “signal” est mélangé à un oscillateur local sur la lame séparatrice L réfléchissante à 50%. La différence des photocourants est proportionnelle à la valeur de la quadrature X_θ .

phique de la fonction de Wigner d’un état de Fock à un photon. Le même groupe a aussi reconstitué la fonction de Wigner d’une superposition cohérente $\alpha|0\rangle + \beta|1\rangle$ à poids α et β quelconques [59]. Nous allons décrire plus en détail l’expérience de caractérisation de l’état $|1\rangle$.

Mesure tomographique de $W_1(\alpha)$ Le dispositif expérimental utilisé est schématiquement représenté sur la figure 5.5 tirée de l’article [60]. La principale difficulté de l’expérience est de préparer un mode spatiotemporel bien défini dans l’état $|1\rangle$. Les physiciens allemands ont utilisé la conversion paramétrique d’un faisceau “pompe” intense en photons de moindre énergie dans un cristal non-linéaire, émis par paires dans des modes spatiaux différents. La détection d’un photon dans le mode “trigger” projette le mode signal dans un état à un photon pourvu que le filtrage spatial et spectral du mode “trigger” soit suffisant. Il faut ensuite homodyner ce signal par un oscillateur local, comme nous l’avons vu précédemment. Les alignements des différents faisceaux sont extrêmement critiques pour mener à bien l’expérience.

Nous montrons les résultats obtenus dans la partie droite de la figure 5.5. On voit nettement un creux négatif dans la fonction de Wigner ainsi reconstituée. La valeur à l’origine est $W(0) = -0.067 \pm 0.016$, une claire manifestation du caractère non-classique de l’état préparé. Pourtant, si l’état à un photon était parfaitement préparé, la fonction de Wigner devrait valoir -1 à l’origine (notons que la normalisation choisie par Lvovsky et al. pour W est différente de la nôtre). Mais toutes les imperfections d’alignement, ainsi que les comptes d’obscurité de la photodiode à avalanche, se traduisent par une préparation imparfaite de l’état $|1\rangle$, remplacé par un mélange statistique de 0 et 1 photon. La condition pour observer des valeurs négatives de W est que la probabilité d’avoir un photon soit supérieure à celle d’avoir 0 photon. L’histogramme des probabilités des

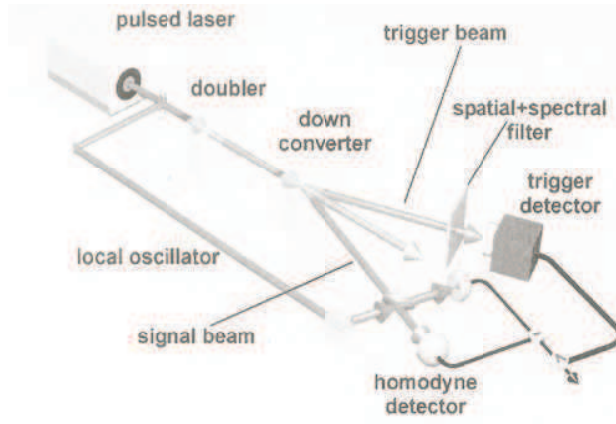


FIG. 2. Simplified scheme of the experimental setup.

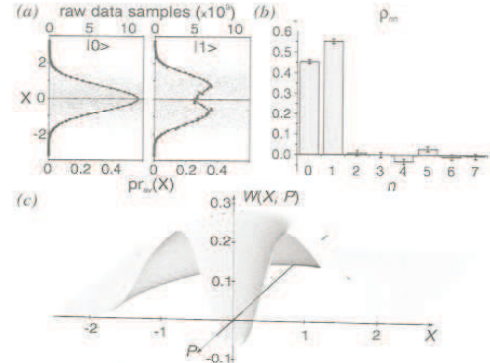


FIG. 4. Experimental results: (a) raw quantum noise data for the vacuum (left) and Fock (right) states along with their histograms corresponding to the phase-randomized marginal distributions; (b) diagonal elements of the density matrix of the state measured; (c) reconstructed WF which is negative near the origin point. The measurement efficiency is 55%.

FIG. 5.5 – (à gauche) Schéma du dispositif utilisé pour la mesure de $W_1(\alpha)$ par détection homodyne. (à droite) Résultats obtenus. (d'après [60])

différents nombres de photon reconstitué est : $p(0) = 0.45$ et $p(1) = 0.55$.

On voit également tout l'apport potentiel de l'électrodynamique quantique en cavité à ce domaine de l'optique quantique : certes, le dispositif est très lourd techniquement ; mais une fois l'investissement technique accompli, il est relativement aisé de préparer des états très non-classiques dans une cavité (états de Fock, mais aussi chats de Schrödinger). En outre, le mode de la cavité est parfaitement défini et il n'y a aucun problème de recouvrement spatial entre le signal et l'oscillateur local.

5.1.4.b Méthode directe

L'utilisation d'algorithmes sophistiqués pour reconstruire la fonction de Wigner à partir d'un grand nombre de mesures rend les résultats des mesures tomographiques particulièrement sensibles au bruit. A l'inverse, des méthodes permettant de mesurer *directement* la fonction de Wigner en chaque point de l'espace des phases ont été proposées [5, 94]. L'incertitude sur $W(\alpha)$ est alors tout simplement donnée par la précision d'une seule mesure.

Principe Toutes les méthodes directes qui ont déjà été employées pour mesurer la fonction de Wigner d'un champ électromagnétique reposent sur une expression simple de la fonction de Wigner démontrée dans un article de Cahill et Glauber [23] :

$$W(\alpha) = \text{Tr}(2\rho\hat{D}(\alpha)(-1)^{a^\dagger a}\hat{D}^{-1}(\alpha)/\pi) = (2/\pi)\text{Tr}(\hat{D}^{-1}(\alpha)\rho\hat{D}(\alpha)(-1)^{a^\dagger a}) \quad (5.18)$$

Posons

$$\rho(\alpha) = \hat{D}^{-1}(\alpha)\rho\hat{D}(\alpha) = \hat{D}(-\alpha)\rho\hat{D}^{-1}(-\alpha) \quad (5.19)$$

puisque $\hat{D}^{-1}(\alpha) = \hat{D}(-\alpha)$ (propriété bien connue de l'opérateur déplacement). D'après l'équation 5.19, l'opérateur $\rho(\alpha)$ est un opérateur densité, qui représente l'état initial ρ auquel on a appliqué un opérateur déplacement $\hat{D}(-\alpha)$. La formule 5.18 nous dit alors qu'à un facteur $2/\pi$ près, la fonction de Wigner n'est rien d'autre que la valeur moyenne de l'opérateur $(-1)^{\hat{a}^\dagger \hat{a}} = \hat{P}$ dans l'état $\rho(\alpha)$. Cet opérateur a de manière évidente pour valeurs propres -1 et $+1$, et pour vecteurs propres les états de Fock. Il mesure la parité du nombre de photons : lorsque celui-ci est pair il vaut $+1$, et -1 dans le cas inverse. Pour mesurer $W(\alpha)$, partant d'un état initial ρ quelconque, on peut donc appliquer l'opérateur déplacement $\hat{D}(-\alpha)$ au système, puis mesurer la valeur moyenne de l'opérateur parité. En répétant l'opération de nombreuses fois pour chaque valeur de α , on obtient

$$\frac{\pi W(\alpha)}{2} = Tr[\hat{\rho}(\alpha)\hat{P}] = \langle \hat{P} \rangle_\alpha = \sum_n (-1)^n \rho_{n,n}(\alpha) \quad (5.20)$$

Mesure de la fonction de Wigner d'un champ cohérent par comptage de photons Une application de cette méthode à la mesure de la fonction de Wigner d'un champ cohérent peut être trouvée dans [4]. Le champ que l'on étudie (faisceau "signal") est mélangé sur une lame séparatrice L très faiblement réfléchissante à un champ cohérent $|- \alpha\rangle$ (oscillateur local), réalisant ainsi l'opération de déplacement $\hat{D}(-\alpha)$. L'amplitude et la phase de l'oscillateur local déterminent directement le point de l'espace des phases en lequel on mesure W . Les physiciens polonais ont placé une photodiode à avalanche en régime de comptage de photons derrière L . Ils ont pu ainsi reconstituer la statistique des nombres de photons $\rho_{n,n}$ du champ déplacé. La formule 5.20 montre qu'il est alors très simple d'en déduire W .

Mesure de la fonction de Wigner d'un ion piégé dans l'état $|1\rangle$ de son mouvement Le groupe de D. Wineland a été le premier à reporter la mesure d'une valeur négative d'une fonction de Wigner [57]. Comme le système quantique ainsi caractérisé n'est pas un mode du champ électromagnétique mais un ion piégé dans un piège harmonique, nous ne décrivons que très brièvement cette très belle expérience. Ici l'espace des phases n'est plus bien entendu le plan de Fresnel mais l'espace des phases position-impulsion d'un oscillateur mécanique.

Les physiciens du NIST ont placé l'ion piégé dans un état de Fock $|n = 1\rangle$ du mouvement grâce à une impulsion laser résolvant les niveaux vibrationnels de l'ion, puis ils ont mesuré sa fonction de Wigner. Le principe de la méthode est similaire à celle qui a été exposée plus haut. Ils ont commencé par appliquer un opérateur déplacement $\hat{D}(\alpha)$ à l'ion. Puis ils ont mesuré les populations des différents états vibrationnels $\rho_{n,n}$ dans l'état déplacé, ce qui leur a donné accès à la valeur moyenne de l'opérateur

parité puisque $\langle \hat{P} \rangle = \sum (-1)^n \rho_{n,n}$. Pour mesurer les populations $\rho_{n,n}$, ils ont induit une oscillation de Rabi entre les deux niveaux hyperfins de l'ion. La probabilité de détecter l'atome dans le niveau excité donne par transformée de Fourier les coefficients $\rho_{n,n}$ [57].

5.2 Une méthode directe pour mesurer la fonction de Wigner d'un mode de la cavité

En électrodynamique quantique en cavité, il est très simple d'appliquer l'opérateur déplacement $\hat{D}(-\alpha)$ à un état ρ que l'on souhaite étudier. Il suffit d'injecter dans la cavité un champ cohérent calibré $|\alpha\rangle$, ce que notre dispositif expérimental permet de faire comme nous l'avons vu au chapitre 4. Pour mesurer W , il ne reste plus qu'à mesurer la valeur moyenne de l'opérateur parité dans cet état déplacé $\rho(\alpha)$. Lutterbach et Davidovich ont proposé une méthode très élégante pour mesurer $\langle (-1)^{a^\dagger a} \rangle$ dans notre dispositif [58]. Une alternative consiste à effectuer une oscillation de Rabi dans le champ déplacé, et à déduire les coefficients diagonaux de la matrice densité du champ $\rho_{n,n}$ par transformée de Fourier [50].

La mesure de $\langle (-1)^{a^\dagger a} \rangle$ proposée par LD repose sur l'interaction dispersive d'un atome de Rydberg avec la cavité *dans le régime où* $\Phi_0 = \pi/2$ (on rappelle que, dans un champ de n photons, une cohérence e-g est déphasée d'une quantité $(2n+1)\Phi_0$ par rapport au vide (voir l'appendice A). Alors une variation de l'intensité du champ dans la cavité d'un seul photon suffit à déphaser la cohérence e-g de $2\Phi_0 = \pi$.

Considérons en effet la situation décrite à la figure 5.6. On effectue des franges de Ramsey sur la transition e-g. Entre les deux impulsions, l'atome interagit dispersive-ment avec le champ $\rho(\alpha)$ dans le régime où $2\Phi_0 = \pi$. La probabilité p_g de détecter l'atome dans $|g\rangle$ présente une modulation lorsqu'on balaye la phase $\Delta\Phi$ de l'interféromètre. Choisissons l'origine des phases de manière à ce que, pour une cavité vide, on ait $p_g = (1 + \cos \Delta\Phi)/2$ (courbe en traits continus sur la figure 5.6). Le déphasage induit par un champ de n photons est $n\pi$ donc lorsque n est pair, les franges initiales sont déphasées d'un nombre entier de fois 2π . On retrouve donc la courbe en traits continus. Lorsque n est impair, le signal de franges est inversé (courbe en pointillés sur le schéma). Ainsi, p_g devient, en présence du champ déplacé de $-\alpha$ $\rho(\alpha)$,

$$\begin{aligned} p_g(\Delta\Phi, \alpha) &= [1 + \sum_n \rho_{n,n}(\alpha) \cos(\Delta\Phi + n\pi)]/2 \\ &= [1 + \sum_n (-1)^n \rho_{n,n}(\alpha) \cos \Delta\Phi]/2 = [1 + \langle \hat{P} \rangle \cos \Delta\Phi]/2 \end{aligned} \quad (5.21)$$

La valeur de $W(\alpha)$ est directement reliée au contraste “algébrique” des franges $c(\alpha)$ ($-1 \leq c(\alpha) \leq +1$) :

$$\pi W(\alpha) = 2\langle \hat{P} \rangle = 2c(\alpha) = 2[p_g(0, \alpha) - p_g(\pi, \alpha)] \quad (5.22)$$

$W(\alpha)$ est donc déterminé à partir de la moyenne d'une quantité mesurée directement sur l'atome, sans aucune procédure d'inversion sensible au bruit. Le reste du chapitre

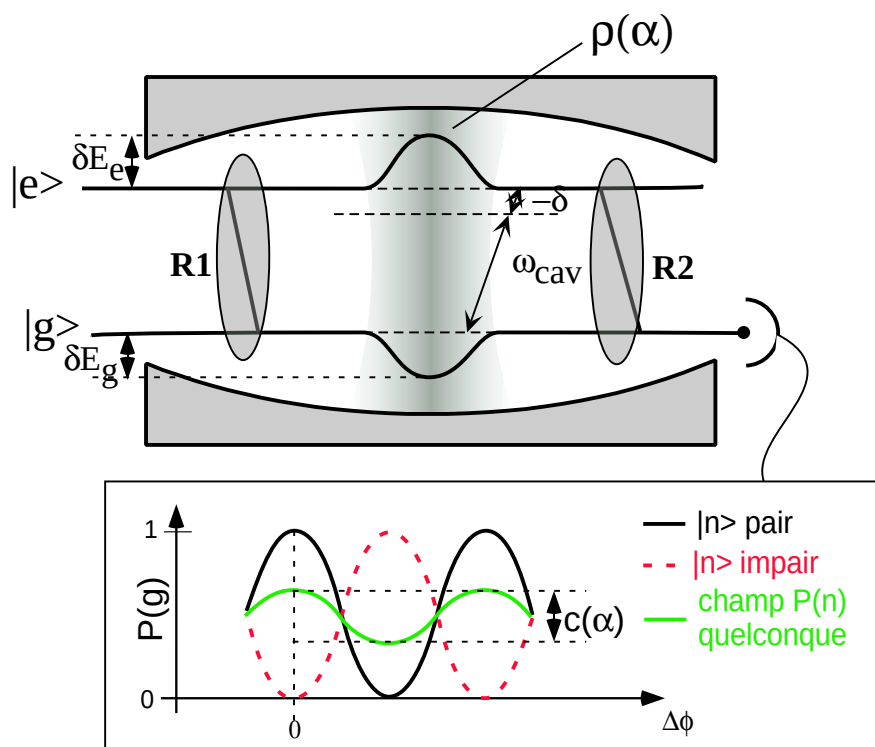


FIG. 5.6 – Principe de la mesure de l'opérateur parité. On réalise des franges de Ramsey sur la transition e - g . Entre les deux impulsions, l'atome interagit dispersivement avec le champ $\rho(\alpha)$. Lorsque la cavité est vide, on a un signal de franges représenté en traits pleins. Si $\Phi_0 = \pi/2$, les franges seront déphasées de $n\pi$ si la cavité contient n photons. En particulier, si le nombre de photons est impair, on obtient les franges en pointillés, alors que s'il est pair on retrouve les franges prises lorsque la cavité était vide (traits pleins). Si la cavité contient un mélange ou une superposition de plusieurs états de Fock de parité différente, le contraste des franges est une mesure de la valeur moyenne de l'opérateur parité.

montre comment nous avons utilisé cette méthode pour faire une mesure de la fonction de Wigner des états $|0\rangle$ et $|1\rangle$.

La méthode LD repose sur le fait que le déphasage dans le régime dispersif varie proportionnellement avec le nombre de photons. Néanmoins, une version de cette méthode restreinte au sous-espace $|0\rangle, |1\rangle$ a été appliquée dans notre groupe pour mesurer la valeur de la fonction de Wigner d'un état à 1 photon à l'origine de l'espace des phases [69]. On utilisait le fait qu'une impulsion de 2π résonante avec la cavité déphase la fonction d'onde atomique de π [70, 79]. Les expériences décrites dans la suite de ce mémoire constituent un prolongement naturel des résultats obtenus alors.

5.3 Expériences préliminaires

La principale difficulté à surmonter, pour réaliser cette méthode, est d'obtenir un déphasage de π par photon.

5.3.1 Le déphasage de π par photon

5.3.1.a Comment l'obtenir

Pour les paramètres de l'expérience de complémentarité, nous avons $\Phi_0 = 0.44rad$. Or nous souhaitons atteindre le régime où $\Phi_0 = \pi/2 = 1.57rad$. D'après l'expression établie au chapitre 3, $\Phi_0 = \Omega_0^2 t_{eff}/4\sqrt{2}\delta$, nous avons deux moyens d'augmenter le déphasage par photon du dipôle atomique :

- si l'on diminue le désaccord δ , Φ_0 va augmenter. Mais il y a deux obstacles qui nous empêchent de trop diminuer δ . D'une part nous avons vu au chapitre 4 que, pour un certain désaccord δ , il existait une limite sur le nombre de photons contenus dans le champ au-delà de laquelle le déphasage ne varie plus proportionnellement à N . Or la méthode que nous venons d'exposer repose de manière évidente sur cette propriété. Pour pouvoir explorer une partie importante de l'espace des phases, il est important que l'on reste dans ce régime même après que le champ dans la cavité ait été déplacé de $-\alpha$. Le désaccord δ doit donc rester suffisamment important. En outre, nous n'avons pas le droit de choisir un désaccord $\delta < 100kHz$ car nous avons constaté qu'un autre effet vient alors perturber l'interaction dispersive avec le mode que l'on souhaite caractériser. Il s'agit d'un processus du troisième ordre qui couple l'atome aux deux modes de la cavité, dont nous n'allons pas parler dans ce mémoire [7].

- comme l'on ne peut pas jouer sur δ pour atteindre le régime où $\Phi_0 = \pi/2$, il ne reste plus qu'à augmenter le temps d'interaction en sélectionnant uniquement des atomes très lents. Nous avons choisi d'utiliser des atomes ayant une vitesse de 150m/s, ce qui correspond à un temps effectif d'interaction $t_{eff} = \sqrt{\pi}w/v = 71\mu s$. Avec ces valeurs-là, on peut calculer qu'un désaccord $\delta = 118kHz$ produit un déphasage $\Phi_0 = \pi/2$. La fréquence atomique (51.09908GHz) est fixée par la valeur du champ électrique Stark que l'on souhaite appliquer pendant l'expérience : le plus faible compatible avec le

maintien de l'orbite circulaire pour les raisons exposées au chapitre 3. On accordera donc le mode HF à la fréquence $\nu_{HF} = 51.098962\text{GHz}$, de manière à avoir $\delta = 118\text{kHz}$.

Travailler avec des atomes “lents” (par rapport à ceux avec qui nous avons fait l'expérience de complémentarité qui avaient une vitesse de 500m/s) pose plusieurs problèmes. Tout d'abord il y a peu de signal. En effet, la distribution des vitesses à la sortie du four est centrée à 350m/s . Nous avons constaté que nous arrivions à préparer dix fois moins d'atomes de Rydberg circulaires à 150m/s qu'à 500m/s . En outre ces atomes mettent beaucoup plus de temps à traverser l'expérience et à atteindre les détecteurs; du coup, l'émission spontanée des atomes en-dehors de la cavité perturbe légèrement le signal. Malgré ces inconvénients techniques nous avons pu obtenir un signal convaincant.

5.3.1.b Réglage des franges e-g

Un problème technique Un autre problème que nous avons dû résoudre concerne les franges de Ramsey. Pour réaliser la méthode LD, il est nécessaire d'utiliser un interféromètre de Ramsey sur la transition e-g². Or lorsqu'on applique une impulsion de Ramsey sur la transition e-g en allumant une source S_{Ram} couplée aux atomes par le guide d'onde “Ramsey”, on risque aussi d'injecter, par effet de bord, quelques photons dans la cavité, car elle est quasi-résonante avec S_{Ram} (ce qui n'est pas le cas lorsque S_{Ram} était accordée sur les transitions g-i ou e-i). Pour l'expérience de complémentarité que nous avons présentée au chapitre précédent, cela n'avait pas d'importance car la seule impulsion de Ramsey sur la transition e-g était effectuée *après* l'interaction avec la cavité. On laissait suffisamment de temps entre deux séquences expérimentales pour que le champ injecté parasite se soit dissipé. Ici par contre, la première impulsion de Ramsey R1 doit bien sûr être réalisée avant que l'atome ne traverse la cavité, et ne doit surtout pas modifier le champ que l'on souhaite mesurer. En fait nous avons constaté que, sans les précautions que nous allons décrire, la “pollution” du mode de la cavité par l'impulsion R1 suffit à brouiller toute frange d'interférence³.

Ajustement du spectre de l'impulsion R1 La solution que nous avons adoptée consiste à choisir la durée de l'impulsion de manière à minimiser le couplage de la source à la cavité. Elle s'apparente aux techniques de “mise en forme d'impulsion” utilisées en Résonance Magnétique Nucléaire. L'enveloppe des impulsions de Ramsey telles qu'elles sont générées dans notre dispositif est à une bonne approximation un

²car, nous l'avons vu au chapitre 3, un photon déphase une cohérence e-g d'un angle $2\Phi_0$ alors qu'une cohérence e-i ou g-i ne serait déphasée “que” de Φ_0 puisque le niveau i n'est pas directement couplé à la cavité. Si nous avions voulu faire des franges sur la transition g-i par exemple, il nous aurait fallu des atomes encore deux fois plus lents (75m/s) pour atteindre le régime du π par photon. Or la vitesse de 150m/s est probablement la plus petite que l'on puisse obtenir avec le dispositif actuel.

³Nous avons présenté au chapitre 3 un signal de franges de Ramsey sur la transition e-g dont le contraste valait 70%, mais ces franges avaient été enregistrées lorsque la cavité était très désaccordée (environ 5MHz).

créneau aux flancs raides, grâce au faible temps de montée et de descente des diodes PIN, et à la multiplication de fréquence qui est faite après la diode. Supposons que l'impulsion ait une durée T et qu'elle soit appliquée entre les instants $-T/2$ et $T/2$.

$$\begin{aligned} E(t) &= E_0 \cos(\omega_{Ram}t) \text{ si } -T/2 \leq t \leq T/2 \\ E(t) &= 0 \text{ sinon} \end{aligned} \quad (5.23)$$

On obtient le spectre de puissance de ces impulsions par transformée de Fourier.

$$S(\omega) = \left| \frac{1}{2\pi} \int_{-T/2}^{+T/2} E_0 \cos(\omega_{Ram}t) \exp^{i\omega t} dt \right|^2 = \frac{T^2}{4\pi^2} E_0^2 \text{sinc}^2 \left(\frac{(\omega - \omega_{Ram})T}{2} \right) \quad (5.24)$$

La fonction $f(x) = \text{sinc}(x)$ s'annule lorsque x est un multiple non nul de π . Donc la puissance spectrale d'une impulsion de Ramsey s'annule aussi lorsque $\frac{(\omega - \omega_{Ram})T}{2} = k\pi$. Notamment, si l'on choisit la durée de l'impulsion

$$T = 2\pi / (\omega_{cav} - \omega_{Ram}) \quad (5.25)$$

la puissance spectrale de S_{Ram} à la fréquence de la cavité sera nulle. La réalité est rendue un peu plus complexe par l'existence des deux modes de la cavité séparés par $\Delta = 128kHz$, tous deux quasi-résonants avec la transition e-g et donc susceptibles d'être "pollués" par l'impulsion R1. Pour éviter cela, nous avons choisi d'asservir la source S_{Ram} 128kHz plus haut que le mode haute fréquence HF : $\nu_{Ram} = \nu_{HF} + \Delta = 51.09909GHz$. L'équation 5.25 impose alors la durée de l'impulsion $T = 7.8\mu s$. Avec ces choix, les fréquences des deux modes se situent dans des creux du spectre de puissance de S_{Ram} . La figure 5.8 résume les différentes contraintes et montre les choix adoptés. Puisque la fréquence de la source S_{Ram} est fixée une fois pour toutes, il sera impossible de balayer la phase des franges en changeant la fréquence de la source. Nous utiliserons donc le système de "virgule" décrit au chapitre 3 pour balayer les franges en changeant la phase du dipôle atomique.

Bien sûr, la puissance spectrale vue par la cavité n'est pas réellement nulle, puisque les flancs de nos impulsions ne sont pas infiniment raides (le temps de montée est d'environ $10ns$), et que la cavité a une largeur spectrale finie de $200Hz$. Nous avons, dans un premier temps, simulé la présence de la cavité par un analyseur de spectre, pour avoir une idée du rapport des puissances que l'on peut attendre. Nous avons donc enregistré le spectre d'une impulsion de Ramsey de $7.8\mu s$ à $51.1GHz$ grâce au banc micro-onde. Pour cette mesure, effectuée avec une bande passante de $300Hz$, on observe une différence de puissance entre la porteuse et le zéro du sinus cardinal à $128.5kHz$ de $24dB$ (voir la figure 5.7). Nous avons par ailleurs mesuré que cette différence était encore plus importante dans la cavité ($27dB$).⁴

⁴Pour cela, nous avons injecté un champ dans la cavité en asservissant S_{cav} à la fréquence du mode, et mesuré le taux d'absorption des atomes dans ce champ $p_0(e)$. Puis nous avons asservi S_{cav} à la fréquence du mode décalée de $128kHz$. Pour retrouver le même taux d'absorption $p_0(e)$, il a été nécessaire de désatténuer la source de $27dB$.

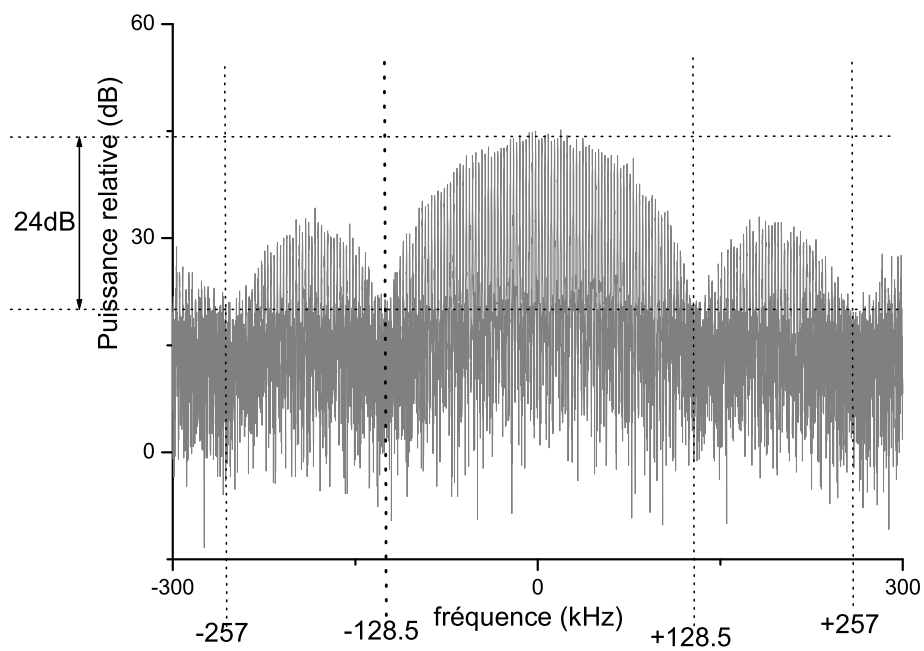


FIG. 5.7 – Spectre de l'impulsion de Ramsey d'une durée de $7.8\mu s$ enregistré à l'analyseur de spectre. On constate qu'à une fréquence de 128.5 kHz de la porteuse, la puissance injectée dans l'analyseur de spectre est 24 dB sous la porteuse. La bande passante de l'appareil était de 300 Hz

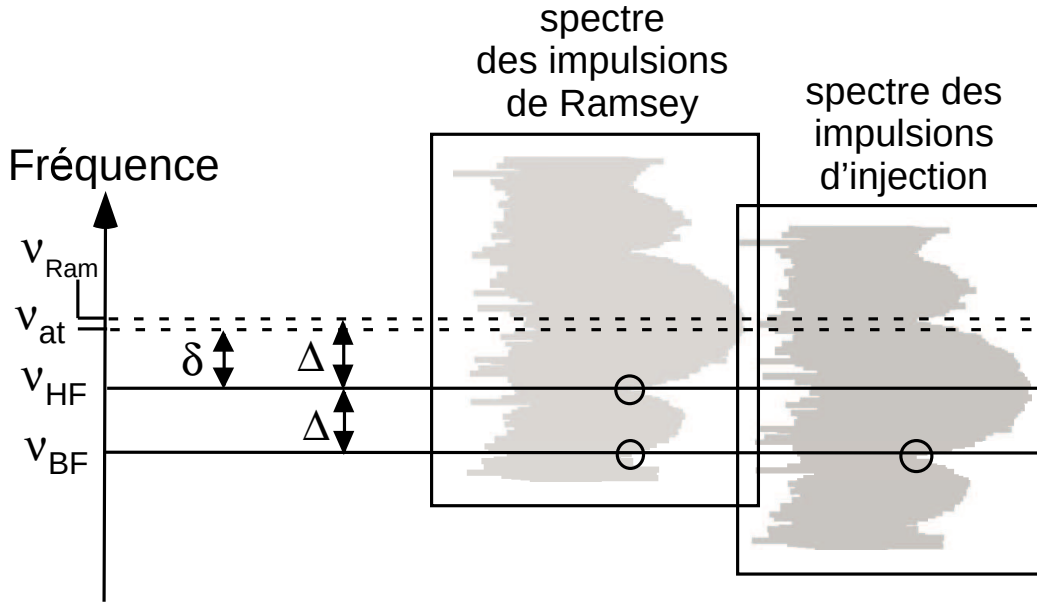


FIG. 5.8 – Résumé des différentes fréquences des sources. La source S_{Ram} est accordée à $\Delta = 128.3\text{kHz}$ au-dessus du mode HF, et donc à $2\Delta = 256.6\text{kHz}$ au-dessus du mode BF. Ainsi, le spectre de l'impulsion de Ramsey de durée $T = 1/\Delta$ a ses deux zéros exactement à la fréquence des modes HF et BF. La source qui permet d'injecter un champ dans le mode HF est accordée à la fréquence du mode HF. Les impulsions d'injection durent aussi la durée $T = 1/\Delta$. Elles ont donc un spectre qui a aussi des zéros à la fréquence du mode BF, et à une fréquence proche de la fréquence atomique. La fréquence atomique est imposée par la nécessité d'avoir un déphasage de π par photon du dipôle atomique, $\delta = 118\text{kHz}$ au-dessus du mode HF. Elle est légèrement différente de la fréquence de la source S_{Ram} , ce qui n'affecte en rien le signal puisque le spectre des impulsions de Ramsey est large d'une centaine de kiloHertz.

Mesure du champ injecté dans le mode de la cavité par R1 Nous avons ensuite effectué une mesure “in situ” du champ parasite injecté dans la cavité par une impulsion de Ramsey en envoyant un atome initialement préparé dans l’état $|g\rangle$ interagir avec la cavité juste après l’impulsion. On a mesuré la probabilité qu’il finisse dans $|e\rangle$ pour différents temps d’interaction, ce qui donne une mesure quantitative du nombre de photons contenus dans la cavité.

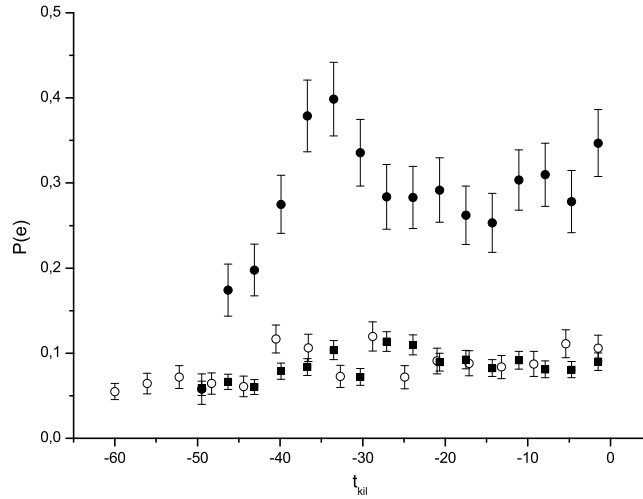


FIG. 5.9 – Un atome est préparé dans $|g\rangle$ et mis à résonance avec la cavité jusqu’à un instant t_{kil} . La probabilité qu’il soit transféré dans $|e\rangle$ nous permet de sonder le champ dans le mode HF dans trois cas de figure. (a) cercles vides : la cavité contient le champ thermique résiduel après passage des serpillères (b) cercles pleins : on a allumé la source S_{Ram} pendant $T = 7.8\mu s$, mais asservie à une fréquence différente de ν_{Ram} ($\nu = 51.09906GHz$). On constate que l’atome a une probabilité beaucoup plus importante de finir dans $|e\rangle$, ce qui montre que l’on injecte des photons dans la cavité. (c) carrés pleins : on a allumé S_{Ram} pendant $T = 7.8\mu s$ à la fréquence $\nu = \nu_{Ram} = 51.09909GHz$. On constate qu’on ne peut pas faire la différence entre la courbe (a) et la (c), ce qui montre que le champ dans la cavité n’est quasiment pas modifié par l’impulsion R1. Les données de cette figure montrent que le champ ajouté par l’impulsion R1 est inférieur à 0.05 photons.

Les résultats sont montrés à la figure 5.9. L’atome est laissé à résonance avec le mode HF jusqu’à l’instant t_{kil} où il est brusquement désaccordé. Auparavant, on a simulé une impulsion de Ramsey en allumant la source S_{Ram} pendant la durée T à la puissance nécessaire pour effectuer une impulsion $\pi/2$. L’expérience a été effectuée dans trois cas différents :

- On commence par effectuer une mesure du champ thermique résiduel, la source

S_{Ram} restant éteinte (cercles vides). L'atome sonde interagit avec la cavité $100\mu s$ après le dernier paquet atomique de refroidissement. La courbe expérimentale montre que le champ dans la cavité est alors d'environ 0.15 photons, comme nous l'avons déjà vu dans le chapitre 3.

- La courbe en cercles pleins montre l'absorption de l'atome sonde lorsqu'on a allumé S_{Ram} pendant la durée T à la fréquence $51.09906GHz$ (soit décalée de $30kHz$ par rapport à la fréquence à laquelle S_{Ram} sera asservie au cours de l'expérience). Comme prévu, la cavité est effectivement polluée par un champ injecté qui provient de S_{Ram} , et que l'on peut évaluer à environ 1 photon. Il serait impossible de voir des franges de Ramsey dans de telles conditions.

- La courbe représentée par des carrés montre l'absorption de la sonde pour la même puissance de la source, la même durée de l'impulsion, mais à la fréquence $\nu = \nu_{Ram} = 51.09909GHz$. Elle est exactement au même niveau que l'absorption du champ thermique résiduel. Par conséquent, le champ injecté par cette impulsion est largement inférieur à ce champ résiduel qui est lui-même de 0.15 photons. La courbe 5.9 nous permet donc d'affirmer que, grâce à l'ajustement spectral que nous avons réalisé, le champ parasite dû à l'impulsion R1 est inférieur à 0.05 photons.

Pour être certains que le champ injecté était minimal lorsque $\nu_{Ram} = 51.09909GHz$, nous avons balayé localement la fréquence de la source autour de cette valeur, en fixant le temps d'interaction de la sonde à une valeur $t_{kil} = -35\mu s$. On voit bien, sur la figure 5.10, un creux dans l'absorption atomique centré à la fréquence $\nu_{Ram} = 51.09909GHz$.

Grâce à ces précautions, nous avons été capables d'obtenir des franges e-g, dans les conditions de l'expérience de mesure de la fonction de Wigner, avec 40% de contraste, qui sont représentées à la figure 5.11. Nous tenterons à l'annexe B de comprendre quantitativement les effets qui réduisent le contraste des franges à 40%; pour l'instant nous allons poursuivre la description de l'expérience.

5.3.2 La calibration du champ

Le dispositif permettant d'injecter le champ $|\alpha\rangle$ pour déplacer l'état initial dans l'espace des phases est identique à celui qui a été décrit au chapitre 4. Une source S_{cav} est asservie à la fréquence du mode de la cavité avec lequel on souhaite travailler (en l'occurrence, le mode HF), et cette source est couplée à la cavité via le guide d'ondes "cavité". Le nombre de photons injectés est déterminé par une expérience préliminaire, similaire à celle effectuée lors de l'expérience de franges de Ramsey quantique. Nous avons choisi de faire des franges de Ramsey sur la transition g-i. Une troisième source S_{gi} est donc asservie à la fréquence de la transition g-i $\nu = 54.2594GHz$ pour effectuer les franges de Ramsey qui nous permettront de calibrer le champ injecté par S_{cav} . Comme l'expérience présentée au chapitre 4 a démontré que notre atténuateur était aussi précis en relatif que la mesure interférométrique, nous avons calibré le champ pour une seule valeur de $A_{dB}^{(cav)}$, en nous servant des valeurs de l'atténuation pour connaître les champs injectés au cours de l'expérience. Afin d'évaluer la précision de notre connaissance du

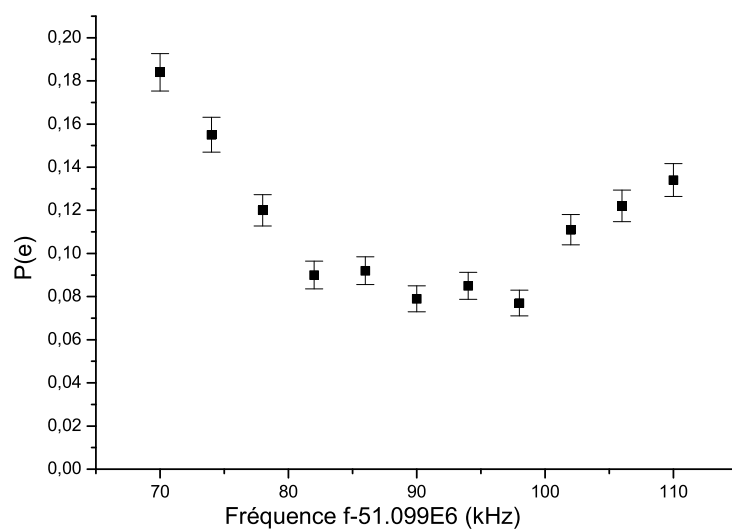


FIG. 5.10 – Probabilité d’absorption d’un atome sonde après une impulsion de Ramsey, en fonction de la fréquence de la source S_{Ram} . On voit nettement un minimum de l’absorption pour la fréquence $\nu = \nu_{Ram} = 51.09909GHz$. Le champ injecté dans la cavité par l’impulsion R1 est donc minimal à cette fréquence.

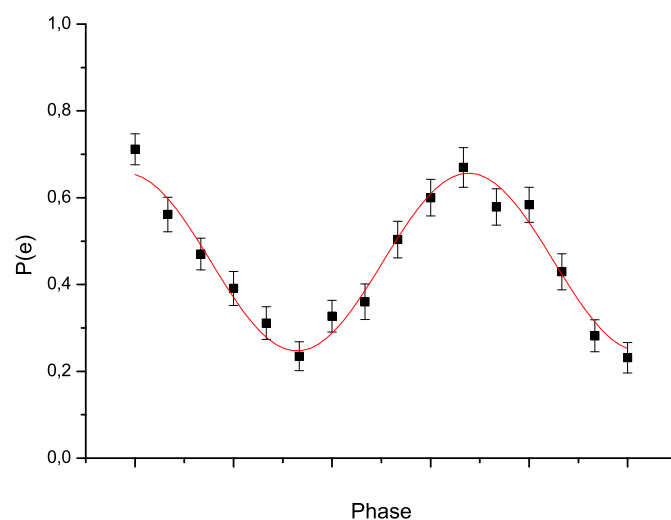


FIG. 5.11 – Franges de Ramsey sur la transition e-g, la cavité étant à 118 kHz de la fréquence atomique. Le contraste est de 40 %.

champ, nous avons enregistré des franges de calibration à deux reprises : avant et après l'ensemble des mesures. Nous avons constaté que nous obtenions le même déphasage à 5% près. Aussi y a-t-il une incertitude de 5% sur notre connaissance de l'intensité du champ injecté. Cela donne une incertitude de 2.5% sur $\alpha = \sqrt{n}$; donc dans la suite nous ne parlerons plus des erreurs statistiques faites sur la détermination de α qui sont négligeables.

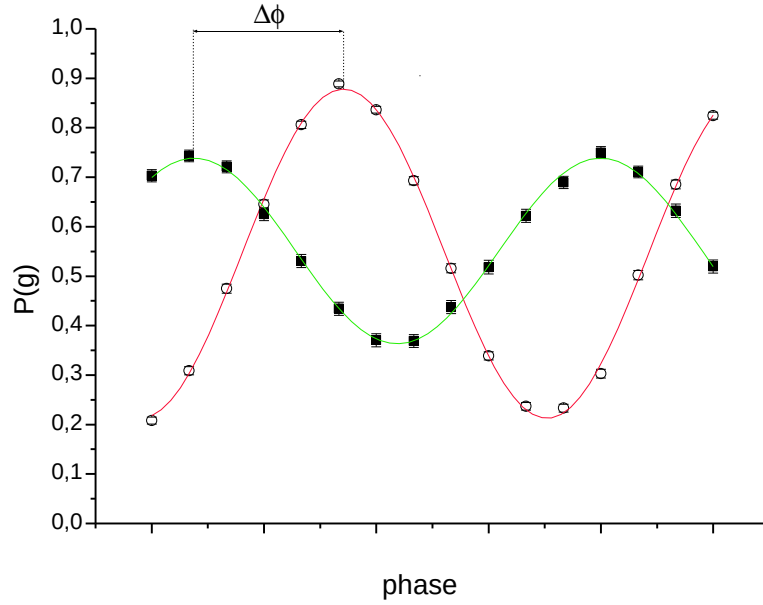


FIG. 5.12 – Franges de calibration sur la transition $g-i$. En cercles ouverts, les franges de référence. En carrés pleins, les franges obtenues lorsqu'on injecte un champ de 5.2 photons. On mesure un déphasage $\Delta\Phi = 2.31\text{rad}$

Donc, pour une atténuation $A_0 = 26\text{dB}$, on mesure un déphasage $\Delta\Phi = 2.31\text{rad}$ (voir la figure 5.12). D'autre part on connaît le désaccord $\delta = 118\text{kHz}$. On peut donc, comme nous l'avons expliqué au chapitre 4, calculer le nombre moyen de photons correspondant. On trouve $N_{ph} = 5.2 \pm 0.1$.

5.3.3 Le temps de vie de la cavité

A la suite d'une manoeuvre malheureuse qui a provoqué un dégazage des écrans thermiques qui entourent le coeur de l'expérience, la facteur de qualité de notre cavité a été dégradé. Ainsi, lors des mesures de fonctions de Wigner que je vais présenter, $T_{cavHF} = 800\mu\text{s}$ et $T_{cavBF} = 740\mu\text{s}$.

5.4 Fonction de Wigner du vide : démonstration de la méthode

Pour commencer nous avons voulu tester la méthode exposée précédemment avec l'état le plus simple possible à obtenir, le vide. Il s'agira en réalité d'un petit champ thermique, dû à la relaxation de la cavité, après refroidissement, vers le champ thermique à l'équilibre.

5.4.1 Réalisation expérimentale

Réglages préliminaires On cherche à reproduire la situation de la figure 5.8. On commence par accorder la cavité de manière à ce que le mode HF se trouve 118kHz sous la transition atomique e-g. Pour cela, on commence par déterminer la tension de piézo $V_{PZT}^{(HF)}$ qui accorde le mode HF sur la raie atomique comme nous l'avons déjà vu au chapitre précédent, puis on fixe la tension de piézo à la valeur $V_{PZT}^{(Wigner)} = V_{PZT}^{(HF)} - 0.1593 * 118kHz$. Le champ électrique E_{Stark} est fixé à la valeur minimale de 0.25V/cm.

L'expérience commence par une séquence de refroidissement de la cavité. Pour avoir suffisamment d'atomes dans les paquets atomiques de cette séquence, nous avons été contraints d'utiliser des atomes d'une autre classe de vitesse que ceux que nous employons pour faire la mesure de la fonction de Wigner et qui vont à 150 m/s. Nous avons choisi des atomes allant à 450 m/s. Nous envoyons trois paquets atomiques pour refroidir la cavité, séparés de $50\mu s$. Le premier fait un refroidissement adiabatique du mode HF, le deuxième du mode BF, puis un troisième vient refroidir encore une fois le mode HF (voir chapitre 3).

Mesure de W Une fois le champ refroidi, l'atome lent traverse la cavité et mesure la fonction de Wigner du mode HF. Il est initialement préparé dans l'état $|g\rangle$, et désaccordé de 118kHz par rapport à la cavité. L'injection du champ $|- \alpha\rangle$ qui fixe, par son amplitude et sa phase, le point de l'espace des phases en lequel nous allons mesurer W est effectuée bien avant que l'atome ne pénètre dans le mode. (les instants exacts d'allumage des différentes sources sont indiqués à la figure 5.13). Puis l'atome subit l'impulsion R1, dont on a préalablement réglé la puissance de manière à ce que l'atome effectue une rotation d'angle $\pi/2$. Il interagit ensuite dispersivement avec la cavité. Après qu'il soit sorti du mode, on lui applique la "virgule" qui permet de balayer la phase des franges, puis l'impulsion de Ramsey R2. On détecte finalement son état interne dans l'espace des niveaux $|e\rangle$ et $|g\rangle$.

On répète l'expérience un grand nombre de fois, ce qui nous permet de reconstituer $p_\phi(g)$ pour différentes valeurs de α . La figure 5.14 présente les résultats obtenus pour trois valeurs différentes de α .

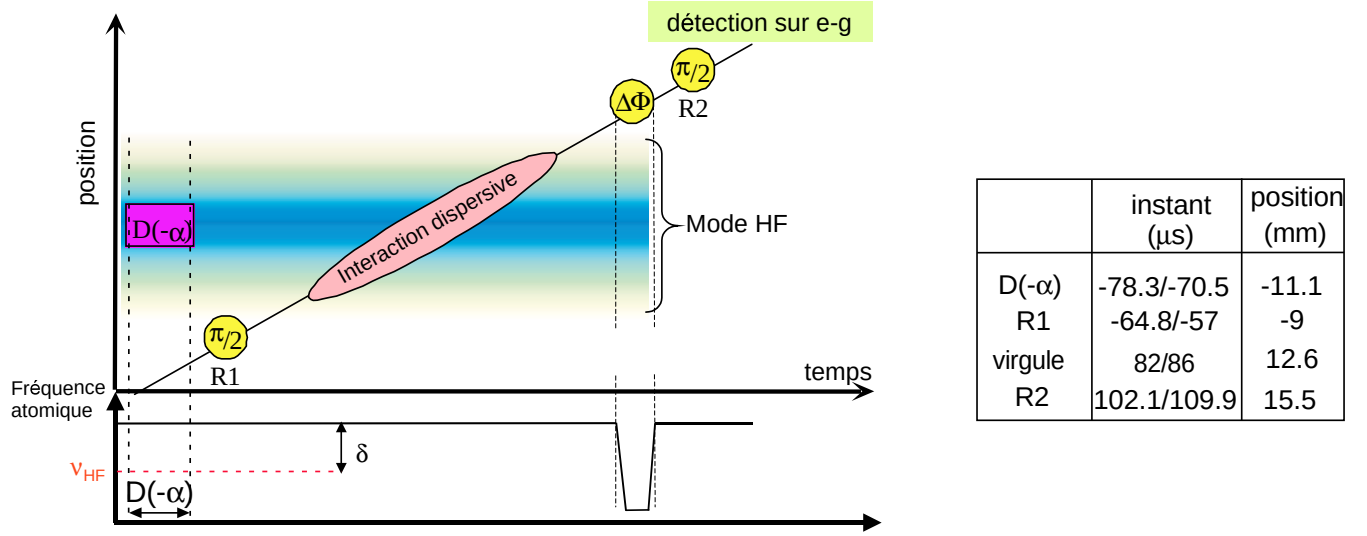


FIG. 5.13 – Déroulement temporel de l'expérience de mesure de la fonction de Wigner d'un petit champ thermique. L'atome est initialement dans $|g\rangle$, désaccordé de $\delta = 118\text{kHz}$ par rapport au mode HF. On commence par injecter le champ $|\alpha\rangle$, puis on applique la première impulsion de Ramsey R1 à l'atome. Il interagit dispersivement avec la cavité, puis subit la deuxième impulsion de Ramsey avant d'être détecté dans l'état $|g\rangle$ ou dans $|e\rangle$. La valeur de W sera déterminée par le contraste des franges.

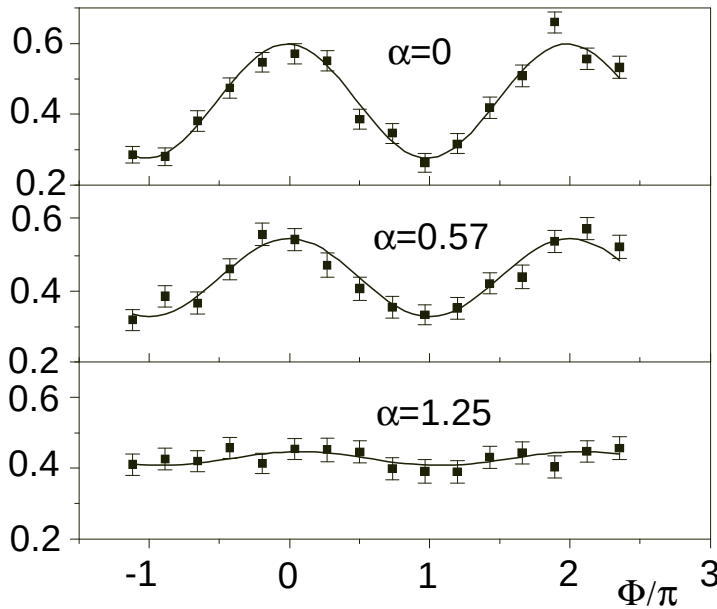


FIG. 5.14 – Résultats de l'expérience de mesure de la fonction de Wigner d'un petit champ thermique pour trois valeurs différentes de α .

5.4.2 Discussion des résultats

5.4.2.a La phase des franges

Un déphasage de π par photon On commence par réaliser un ajustement numérique des différents signaux expérimentaux de franges de Ramsey en laissant la phase globale Φ comme paramètre libre. La figure 5.15 montre la phase relative des franges de Ramsey par rapport aux franges de référence en fonction du nombre de photons injectés. (L'incertitude sur la mesure de la phase augmente avec le nombre de photons à cause de la réduction de contraste progressive, pour atteindre π lorsque le contraste des franges est nul.) On constate que la phase des franges reste constante, et ce quel que soit le champ injecté dans la cavité.

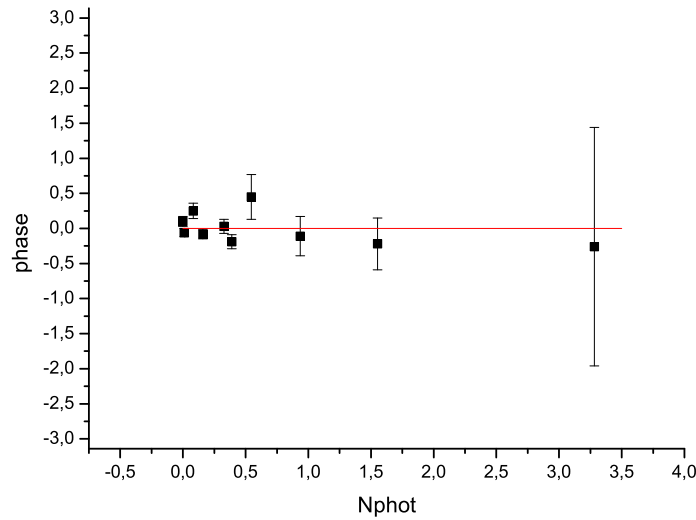


FIG. 5.15 – Phase des frange de Ramsey en fonction du nombre de photons injectés dans le mode HF. On constate que la phase reste quasiment constante (ligne pointillée). Un ajustement numérique linéaire sur les données donne une pente de $A = -0.17 \pm 0.15$, ce qui est compatible avec $A=0$. On est donc bien dans le régime où $2\Phi_0 = \pi$.

Nous allons exploiter ce résultat pour calculer le déphasage par photon Φ_0 . Remarquons tout d'abord qu'au cours de cette expérience, nous sommes dans le régime dispersif défini au chapitre précédent. En effet, nous avons constaté (voir la figure 4.6) que, lorsque $\delta = 120\text{kHz}$, ce qui est le cas aussi dans cette expérience, pourvu que le nombre de photons injectés reste inférieur à 5, on pouvait appliquer la relation $\Delta\Phi = N \sin 2\Phi_0$ ⁵. Nous avons donc procédé à un ajustement linéaire des données de

⁵Rappelons que dans cette expérience, les franges sont faites sur la transition e-g donc le déphasage dans un champ cohérent $|\alpha\rangle$ est donné par $\Delta\Phi = |\alpha|^2 \sin(2\Phi_0)$

la figure 5.15 pour déterminer Φ_0 . On trouve : $\sin 2\Phi_0 = -0.17 \pm 0.15$, ce qui donne $2\Phi_0 = 3.3rad \pm 0.15$, résultat compatible avec $2\Phi_0 = \pi$. Par conséquent, nous avons bien atteint le régime où un seul photon dans la cavité déphase le dipôle atomique de π , avec une précision d'environ 5% sur la valeur de Φ_0 .

Essayons d'évaluer dans quelle mesure une imprécision de 5% sur Φ_0 affecte les valeurs de la fonction de Wigner mesurées. Pour cela, nous rappelons la formule que nous avons établie au chapitre 3 concernant le contraste des franges de Ramsey sur la transition e-g déphasées par un champ cohérent $|\alpha\rangle$:

$$C = \exp^{-2|\alpha|^2 \sin^2(\Phi_0)} \quad (5.26)$$

Pour $\Phi_0 = (1 + \epsilon)\pi/2$, cela donne une variation de contraste relative

$$\frac{\Delta C}{C} = 2|\alpha|^2 \left(\epsilon \frac{\pi}{2}\right)^2 \quad (5.27)$$

On voit que l'incertitude sur le contraste (et donc sur W) est proportionnelle au nombre de photons moyen du champ cohérent, et quadratique en ϵ . Pour $\epsilon = 0.05$ (valeur expérimentale), on trouve une erreur de 5% pour un champ cohérent de 4 photons. Or nous verrons que W ne prend des valeurs significatives que lorsque $0 < \alpha < 1$. Dans cette zone, il est largement suffisant de réaliser $\Phi_0 = \pi/2$ avec une précision de 5%. En conclusion l'erreur sur W due à l'incertitude sur Φ_0 est largement inférieure aux barres d'erreur statistiques.

Une preuve directe de la quantification du champ dans le régime dispersif

Le fait que la phase des franges reste constante quelle que soit l'intensité du champ injecté dans la cavité provient bien sûr d'un effet de moiré entre le déphasage du dipôle atomique et l'intensité du champ dans la cavité, qui est quantifiée. Cet effet est néanmoins intéressant à souligner, car il permet d'interpréter la courbe 5.15 comme une preuve directe de la quantification du champ dans la cavité (l'équivalent de l'oscillation de Rabi quantique, mais dans le régime dispersif). Pour cela, nous pouvons comparer les résultats de la figure 5.15 aux données sur la phase des franges de calibration que nous avons montrées au chapitre précédent (figure 4.6). Les deux courbes représentent le déphasage des franges en fonction de l'intensité du champ injecté ; mais dans un cas $\Phi_0 \ll \pi$ alors que dans l'autre $\Phi_0 = \pi/2$.

Les résultats de la figure 4.6 sont parfaitement compatibles avec la théorie semi-classique d'un atome à deux niveaux en interaction avec un champ classique, qui prévoit un déphasage du dipôle atomique proportionnel à l'intensité lumineuse I dans le régime dispersif ($\Delta\Phi = kI/\delta$). Rien ne montre sur cette courbe que l'intensité I ne peut prendre que des valeurs discrètes, puisque lorsqu'on fait varier I il est clair que $\Delta\Phi$ varie continûment. La chute de contraste observé est explicable par une certaine dispersion de l'intensité I , mais ne donne aucun renseignement sur sa quantification.

Au vu des résultats de la figure 4.6 on pourrait penser qu'accroître le couplage entre l'atome et le champ en allongeant le temps d'interaction se traduirait simplement par

un accroissement de la pente de la droite, et ce serait effectivement le cas si l'intensité n'était pas quantifiée. Par exemple, si un champ de un photon provoque un déphasage de π , un champ d'un-demi photon devrait produire un déphasage de $\pi/2$.

Au contraire, la courbe 5.15 montre de manière explicite que l'intensité du champ qui provoque le light-shift est quantifiée. On constate que, quel que soit l'intensité *moyenne* du champ qui déphase le dipôle atomique, celui-ci est toujours déphasé d'un nombre entier de demi-tours. Sur la courbe 5.15 on ne voit pas les valeurs intermédiaires de la phase car les valeurs intermédiaires de l'énergie sont interdites.

5.4.2.b Résultats sur le contraste

Valeur de W Maintenant que nous nous sommes placés à $2\Phi_0 = \pi$, nous pouvons exploiter les signaux de franges de Ramsey pour la mesure de W. On commence par faire un ajustement numérique des franges en fixant la valeur du déphasage à 0. On mesure alors, pour chaque valeur de α , le contraste $c_0(\alpha)$. En théorie, on devrait alors avoir $W_0(\alpha) = 2c_0(\alpha)/\pi$. En fait, les diverses imperfections de l'expérience réduisent le contraste des franges pour des raisons tout autres que la présence d'un champ dans le mode étudié. On peut supposer que toutes ces imperfections sont indépendantes du champ injecté et de α (nous allons discuter ce point par la suite). Il suffit alors d'introduire un facteur de normalisation global F_0 pour en tenir compte. On aura donc $W_0(\alpha) = 2F_0c_0(\alpha)$. Il est très simple de mesurer F_0 puisque la fonction W doit être normalisée selon la relation 5.9. Par conséquent,

$$F_0 = \frac{1}{2 \int d^2\alpha c_0(\alpha)} \quad (5.28)$$

Pour notre expérience on trouve $F_0 = 2.44$. La figure 5.16 montre les résultats obtenus pour $W_0(\alpha)$ après multiplication de $c_0(\alpha)$ par F_0 . Soulignons que dans notre mesure de $W(\alpha)$, la phase de α n'a joué aucun rôle, c'est pourquoi les résultats sont présentés uniquement en fonction de $|\alpha|$, et non pas en trois dimensions dans le plan de Fresnel. Ce n'est pas à cause d'une quelconque limitation de notre dispositif (nous avons démontré à maintes reprises, et par exemple dans l'expérience du chapitre 4 que la phase du champ cohérent injecté était bien contrôlée). En raison de notre procédure de préparation du champ, qui se fait entièrement par interaction avec des atomes initialement dans un état propre de l'énergie, la phase de l'impulsion d'injection n'est tout simplement pas relevante. La distribution de Wigner d'un champ préparé de cette manière est nécessairement invariante par rotation. Ce dernier point sera aussi valable pour la mesure de $W_1(\alpha)$ que nous allons présenter dans la partie suivante.

Il s'agit maintenant d'exploiter la fonction obtenue pour en tirer des informations sur le champ. Notamment, nous avons voulu connaître la statistique du champ dont nous avons mesuré la fonction de Wigner. Pour cela, nous avons effectué un ajustement numérique de nos mesures de $W(\alpha)$ par la fonction de Wigner d'un mélange statistique

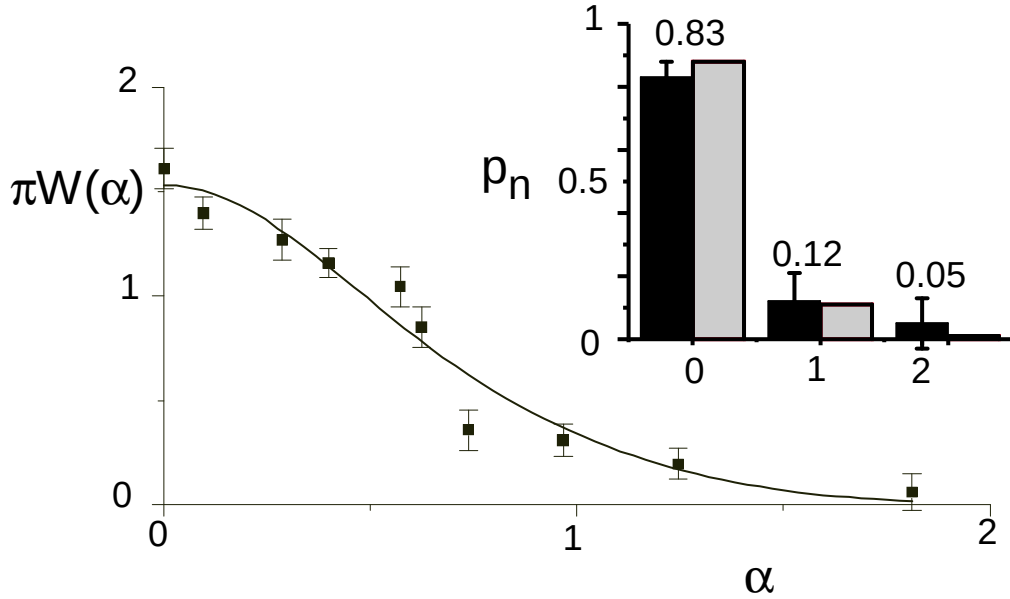


FIG. 5.16 – *Mesure directe de la fonction de Wigner d'un petit champ thermique. On a multiplié le contraste des franges par le facteur de normalisation $F_0 = 2.44$.*

de 0,1 et 2 photons dont la formule est

$$W^{(P_0, P_1, P_2)}(\alpha) = \sum_{i=0}^{i=2} P_i W_i(\alpha) = (2/\pi) \sum_{i=0}^{i=2} P_i (-1)^i L_i(4|\alpha|^2) \exp(-2|\alpha|^2) \quad (5.29)$$

La courbe en traits pleins de la figure 5.16 est le résultat de cet ajustement numérique. Les paramètres P_i sont représentés à gauche, avec leur barre d'erreur, sur l'histogramme qui accompagne la figure 5.16. Leurs valeurs sont données au-dessus. A droite des P_i , on a représenté le résultat d'un calcul de la distribution théorique du champ, étant donnée la méthode de préparation (un champ thermique résiduel de 0.05 photons relaxant vers un champ de 1 photon à l'équilibre pendant $75\mu s$). Les deux distributions sont très voisines, ce qui confirme la validité de nos résultats expérimentaux.

5.4.2.c Discussion du facteur de normalisation

Nous devons maintenant essayer de comprendre pourquoi nos données ont besoin d'être normalisées par un facteur qui est important (environ 2.5). Notamment il est important de pouvoir connaître quantitativement l'effet de chacune des imperfections du dispositif afin de pouvoir l'améliorer par la suite. Nous avons mis en annexe la discussion détaillée du facteur de normalisation de nos données expérimentales (voir l'annexe B), car elle est un petit peu longue et technique. Dans ce paragraphe, nous nous contenterons d'en rappeler les conclusions. Elles sont reprises dans le tableau 5.1.

Nous avons dégagé trois imperfections indépendantes qui contribuent à dégrader le signal. Soulignons qu'il s'agit uniquement d'imperfections de notre procédure de mesure. Nous ne considérons pas la qualité de la préparation de l'état, qui est une chose différente. Les deux imperfections les plus importantes (champ thermique résiduel dans le mode BF et contraste intrinsèque limité) sont évidemment indépendantes de α , ce qui est crucial pour l'interprétation des données. Nous n'avons pas de preuve directe que la troisième (les événements à deux atomes) le soit aussi mais nous avons de bonnes raisons théoriques de le croire (voir l'annexe B), et de plus, il s'agit d'un effet d'au plus quelques pour cent qui est en-deçà des barres d'erreur de la figure 5.14. Pour chacune d'elles, nous donnons la contribution au facteur de normalisation total. Nous précisons en outre si le chiffre avancé pour la contribution au contraste total provient d'une mesure indépendante ou non :

cause	contribution	mesuré ?
contraste intrinsèque fini de l'interféromètre de Ramsey	$f_1 = 1.43$	oui
événements à deux atomes	$f_2 = 1.09$	oui+théorie
champ thermique résiduel dans le mode BF	$f_3 = 1.5$	non, plausible

TAB. 5.1 – *Imperfections de la mesure de la fonction de Wigner du vide et leur contribution relative au facteur de normalisation total*

L'effet combiné de ces imperfections donne un facteur de normalisation $F = f_1 f_2 f_3 = 2.3$, ce qui est très proche du facteur de normalisation mesuré $F_0 = 2.44$. C'est pourquoi nous pensons comprendre de manière satisfaisante les contrastes observés, ce qui en retour confirme la validité de notre mesure de la fonction de Wigner, obtenue après normalisation des données par un facteur global mesuré.

5.5 Fonction de Wigner d'un état de Fock à un photon

5.5.1 Réalisation expérimentale

Utilisation d'un seul atome Pour effectuer la mesure de la fonction de Wigner W d'un état de Fock à un photon, on aurait pu penser à une simple transposition de l'expérience précédente, mais avec deux atomes, le premier préparant le mode de la cavité dans l'état $|1\rangle$ en effectuant une impulsion π de Rabi dans le mode considéré, et le deuxième reproduisant les opérations vues plus haut. Dans cet esprit, on conditionnerait les franges de Ramsey sur le deuxième atome à la détection d'un premier atome dans $|g\rangle$, qui assure que la cavité est dans l'état $|1\rangle$. Une telle expérience pose cependant deux problèmes :

- à cause du délai entre le passage des deux atomes, le champ dans la cavité a un peu de temps pour relaxer vers le champ thermique à l'équilibre, donc la partie négative

de la fonction de Wigner sera moins marquée.

- plus important encore, nous avons vu que le nombre moyen d'atomes par séquence expérimentale doit être très faible si l'on veut éviter les effets parasites à deux atomes (typiquement $f = 0.04$). Cela impose des temps d'acquisition importants : pour enregistrer un signal de franges de Ramsey avec ces atomes, il faut environ une heure. Si en plus il faut conditionner les franges de Ramsey à la détection d'un premier atome, qui doit aussi être envoyé avec un flux faible (0.15 at/c), sans quoi il risque d'injecter plus d'un photon dans la cavité, il faut compter 7 heures par point de la fonction de Wigner. Pour mesurer 10 points il faudrait donc 70 heures ! or la stabilité de l'expérience n'excède pas pour l'instant 10 heures.

Nous avons donc eu recours à une astuce qui nous a permis d'effectuer l'expérience dans de bonnes conditions. Nous avons utilisé le même atome *à la fois* pour préparer la cavité dans l'état $|1\rangle$ et pour en mesurer la fonction de Wigner. L'atome est initialement préparé dans $|e\rangle$, la cavité étant vide. Au début de son interaction avec le mode, l'atome est mis à résonance avec ce dernier afin d'y effectuer une impulsion π . Après cette impulsion, le système atome-cavité est dans l'état produit $|g, 1\rangle$. On désaccorde l'atome par effet Stark, et l'on effectue alors la mesure de la fonction de Wigner : injection du champ $-\alpha$, et franges de Ramsey sur la transition e-g. On minimise ainsi la relaxation du champ entre les deux atomes, et on effectue l'impulsion π avec un flux si faible que les événements où deux atomes présents dans le paquet atomique émettent deux photons sont, comme nous le verrons, négligeables.

Déroulement d'une séquence expérimentale A cause de la procédure adoptée, l'atome "perd" un peu de temps d'interaction dispersive avec la cavité puisqu'il commence par effectuer une impulsion π . Pour compenser, nous avons légèrement réduit le désaccord δ ⁶. On peut calculer qu'il suffit d'imposer $\delta' = 105\text{kHz}$ au lieu de 118kHz.

L'atome est préparé dans $|e\rangle$, et le champ électrique E_{Stark} est réglé à $E_{\text{Stark}} = E_{\text{Stark}}^{(HF)} = 0.74\text{V/cm}$ de manière à ce que l'atome soit mis à résonance avec le mode HF au début de son interaction avec la cavité. On commence par régler l'instant t_π auquel l'atome doit être désaccordé pour effectuer une impulsion π dans le mode. Pour cela, dans une expérience préliminaire, on mesure le transfert de population du niveau $|e\rangle$ vers le niveau $|g\rangle$ en fonction du temps t_{Stark} sans appliquer les impulsions micro-onde de l'expérience. La figure 5.17 montre le résultat de cette mesure préliminaire. L'instant t_π est celui pour lequel le transfert de population est maximal. Les sources micro-onde S_{cav} et S_{Ram} sont asservies à la même fréquence que dans l'expérience précédente (voir la figure 5.8).

Le déroulement d'une séquence expérimentale est décrit à la figure 5.18. Après que l'atome ait effectué son impulsion π dans la cavité, le champ Stark est com-

⁶Notons qu'il suffit de réduire légèrement δ : en effet, on joue sur le fait que le couplage résonant varie comme $\Omega(x) = \Omega_0 \exp(-x^2/w^2)$ alors que le couplage dispersif est en $\Omega(x)^2/\delta$. Donc la gaussienne qui représente le couplage dispersif est $\sqrt{2}$ fois plus étroite que l'autre, et lorsqu'on désaccorde l'atome à la fin de l'impulsion π dans le mode, le couplage dispersif est encore très faible.

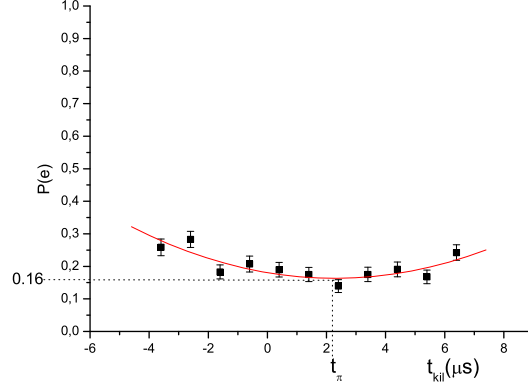


FIG. 5.17 – Réglage de l’impulsion π dans le mode HF de la cavité. L’atome entre dans la cavité dans l’état $|e\rangle$, il est laissé à résonance avec le mode HF jusqu’à l’instant t_{kil} où on le met hors résonance en commutant le champ Stark de la valeur $E_{Stark}^{(HF)}$ à la valeur $E_{Stark}^{(Wig)}$. On a mesuré la probabilité de détecter l’atome dans $|e\rangle$ en fonction de l’instant t_{kil} (l’origine des temps est arbitraire). On voit qu’à l’instant t_π $P(e) = 16\%$

muté à la valeur $E_{Stark} = E_{Stark}^{(Wig)} = 0.32V/cm$ qui met la transition atomique $\delta' = 105kHz$ au-dessus de la fréquence du mode HF. Le système est alors dans l’état $|g, 1\rangle$. Immédiatement après l’impulsion π dans le mode, on injecte le champ $-\alpha$ dans la cavité, en même temps que l’on applique la première impulsion de Ramsey R1 à l’atome (nous discuterons plus loin des effets parasites que cela peut provoquer). Puis l’atome interagit dispersivement avec le champ déplacé. Après qu’il soit sorti du mode, on lui applique une “virgule” pour balayer la phase des franges, et la deuxième impulsion de Ramsey R2. Après quoi on détecte son état interne. On répète la séquence expérimentale un grand nombre de fois afin de reconstituer les probabilités $P_e(\phi)$ pour chaque champ injecté α .

La figure 5.19 montre des franges obtenues pour 3 valeurs différentes de α . On constate que la courbure des franges s’inverse entre $\alpha = 0$ et $\alpha = .81$, en s’annulant pour une valeur intermédiaire de α . Ce comportement est une signature du fait que la fonction de Wigner de l’état $|1\rangle$ est négative au centre de l’espace des phases mais redevient positive lorsque $|\alpha|$ augmente. On voit aussi très clairement la valeur de α pour laquelle $W(\alpha)$ vaut zéro : elle correspond à des franges de contraste nul.

5.5.2 Discussion des résultats

5.5.2.a Résultats sur la phase

Nous commençons par faire un ajustement numérique sinusoïdal de toutes les franges obtenues pour les différentes valeurs de α (en imposant que le contraste soit positif).

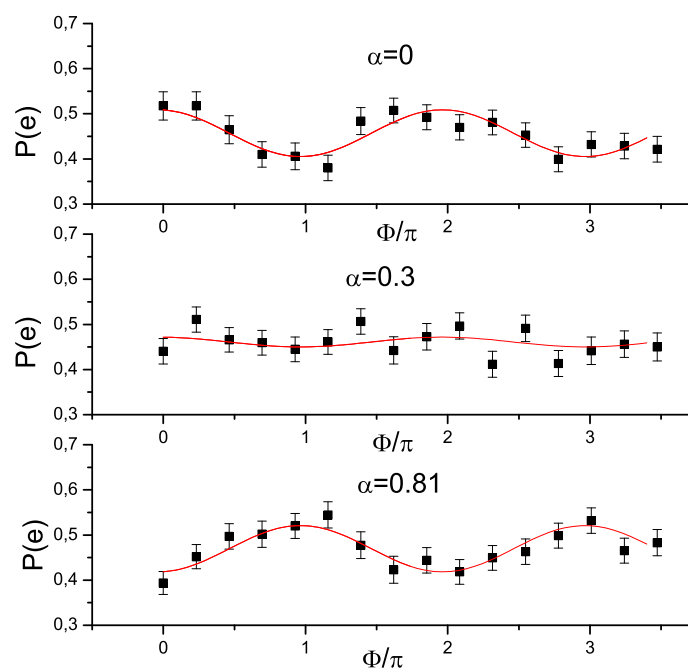


FIG. 5.19 – Signal de franges de Ramsey obtenues pour différentes valeurs de α . On constate que la phase des franges s'inverse entre $\alpha = 0$ et $\alpha = 0.8$ en s'annulant pour $\alpha = 0.3$.

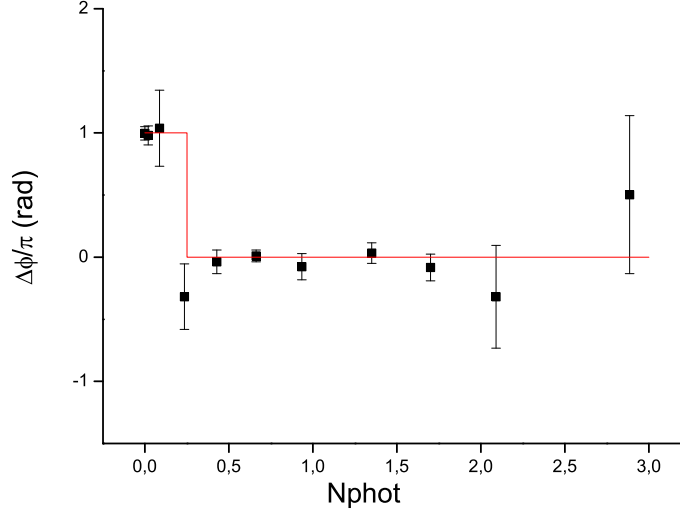


FIG. 5.20 – Phase des franges de Ramsey en fonction du nombre moyen de photons $N_{phot} = |\alpha|^2$

On a mesuré $W_1(0) = -0.8 \pm 0.2$.

Puis nous avons effectué un ajustement numérique de la fonction de Wigner obtenue par un mélange statistique des états 0, 1 et 2 avec un certain poids P_i , comme au paragraphe 5.4.2. On obtient ainsi la courbe en traits pleins. Les paramètres P_i trouvés sont montrés sur l'histogramme de la figure 5.21 avec leurs barres d'erreur (partie gauche de l'histogramme). Nous avons par ailleurs calculé la distribution attendue pour le champ compte tenu de la manière dont il a été préparé. Nous sommes partis d'un champ thermique résiduel de 0.05 photons, relaxant vers le champ thermique à l'équilibre. Nous avons calculé l'effet sur le champ de l'interaction avec un atome dans $|e\rangle$, dans les conditions exactes de l'expérience.

Nous avons en outre tenu compte d'une imperfection supplémentaire dans la préparation de l'état $|1\rangle$, et qui explique en partie l'importance de P_0 (nous avons mesuré $P_0 = 0.25$). Nous avons mesuré par ailleurs qu'un atome ayant une vitesse de 150m/s , initialement préparé dans $|e\rangle$, avait 12% de chances de finir dans $|g\rangle$, même lorsque la cavité est désaccordée. Sur ces 12%, 10% proviennent de l'émission spontanée de l'atome vers le continuum, légèrement stimulée par la présence d'un champ thermique ambiant de 1 photon en moyenne comme nous l'avons déjà dit précédemment, et 2% aux collisions avec les paquets de refroidissement initialement dans $|g\rangle$ qui dépassent l'atome dans $|e\rangle$ et peuvent, par interaction de Van der Wals, provoquer sa transition vers $|g\rangle$. Comme la distance entre la boîte à circulariser et la cavité est moindre qu'entre la boîte à circulariser et les détecteurs, on en déduit que 8% des atomes arrivent dans la cavité déjà dans $|g\rangle$. Cela cause une augmentation de 8% de P_0 dans la statistique

du champ mesuré.

Une fois cet effet pris en compte, l'accord entre la distribution mesurée et la théorie est très satisfaisant.

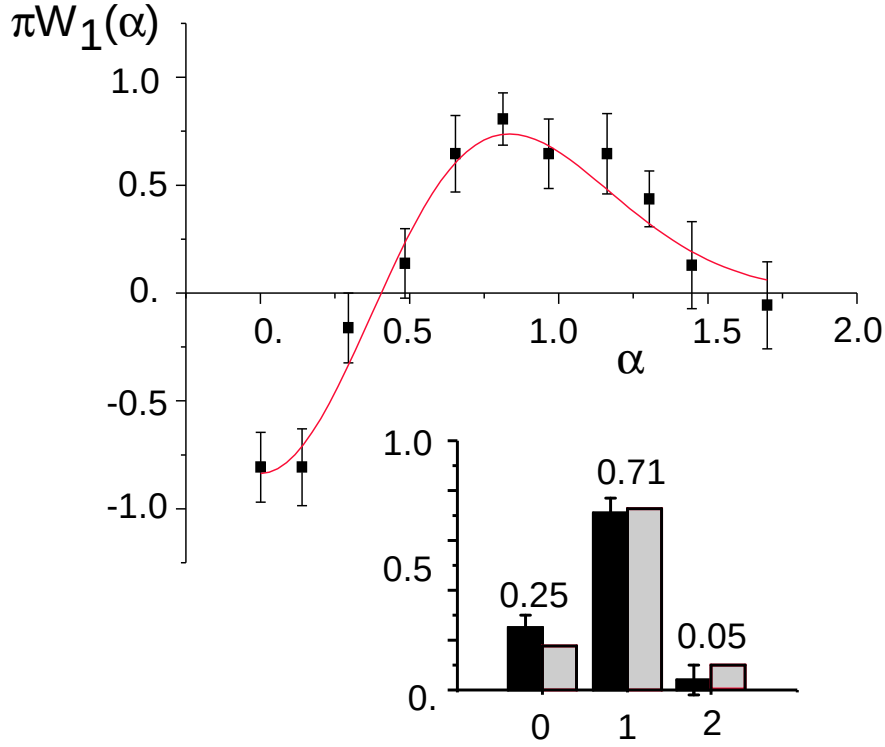


FIG. 5.21 – Fonction de Wigner mesurée après multiplication du contraste des franges par le facteur $F_1 = 4.0$. La courbe en traits continus est le résultat d'un ajustement numérique par la fonction de Wigner d'un mélange statistique de 0,1 et 2 photons avec les poids p_i qui sont donnés par l'histogramme.

5.5.2.c Discussion du facteur de normalisation

Nous venons de voir que les résultats obtenus étaient en bon accord avec les résultats théoriques. Il reste à comprendre la valeur du facteur de normalisation $F_1 = 4.0$, qui est largement supérieure au facteur de normalisation $F_0 = 2.44$ dans le cas du vide. On trouvera l'analyse détaillée de ce facteur de normalisation à l'annexe B. Parmi les imperfections, l'efficacité limitée de l'impulsion π est celle dont les effets sont le moins bien maîtrisés bien qu'elle ne concerne qu'un nombre limité d'atomes. Notamment, nous devons nous contenter de faire l'hypothèse que ces effets ne dépendent pas de α . Les justifications sont exposées à l'annexe B. Ses conclusions sont résumées sur le tableau 5.2

cause	contribution	mesuré ?
facteur normalisation Wigner vide	$F = 2.3$	(voir plus haut)
facteur supplémentaire (R1 inhomogène, flux supérieur)	$f_4 = 1.28$	oui
imperfections de l'impulsion π	$f_5 = 1.19$	oui

TAB. 5.2 – *Imperfections de la mesure de la fonction de Wigner de l'état $|1\rangle$ et leur contribution relative au facteur de normalisation total*

Au total, nous obtenons un facteur de normalisation $F' = F f_4 f_5 = 3.6$, ce qui est très proche de la valeur mesurée $F_1 = 4.0$.

5.6 Discussion et perspectives

Nous allons maintenant conclure sur l'apport des résultats précédents, et sur les perspectives qu'ils ouvrent pour la caractérisation d'états quantiques plus complexes que ceux présentés dans ce mémoire.

5.6.1 Intérêt des résultats précédents

Les résultats montrés précédemment ont le mérite de montrer que le dispositif d'électrodynamique quantique en cavité permet réellement, malgré les difficultés techniques que nous avons dû surmonter, d'employer la méthode proposée par LD pour caractériser un état quantique par sa fonction de Wigner. Les résultats montrés sur la fonction de Wigner du vide suffisent à démontrer la validité de cette méthode.

Le régime où le dipôle atomique est déphasé de π par photon dispersivement n'avait à notre connaissance jamais été atteint auparavant. Le verrouillage de la phase des franges de Ramsey à une valeur constante quel que soit le nombre moyen de photons est un effet spectaculaire de la quantification de l'énergie dans la cavité. Outre les mesures de fonction de Wigner, l'obtention de ce régime rend possible d'autres applications également intéressantes comme les mesures QND du nombre de photons (voir plus loin). Nous avons en outre éclairci certains mécanismes complexes qui limitent le contraste des franges à fort flux (effets à plusieurs atomes).

Enfin nous avons mesuré la fonction de Wigner d'un état proche d'un état de Fock $|1\rangle$ sur tout l'espace des phases. Le contraste est frappant avec les résultats obtenus pour un petit champ thermique : alors que dans ce dernier cas, le contraste des franges se réduisait lorsqu'on augmentait α , dans le cas de l'état de Fock on observe de façon évidente un "retournement" des franges lorsqu'on augmente le champ injecté α , avec un passage par 0 pour $\alpha = 0.40$. Ainsi l'on constate tout l'intérêt d'une méthode directe de mesure de W : le caractère non-classique du champ se voit de manière extrêmement manifeste sur un signal expérimental simple.

5.6.2 Perspectives

5.6.2.a Améliorations possibles

Malgré tout l'intérêt de ces résultats, ils souffrent d'un rapport signal sur bruit limité dû à de nombreuses imperfections techniques. Un certain nombre de directions peuvent être explorées pour grandement les améliorer.

La plus évidente est la suppression du champ thermique de fuite dans la cavité, qui est pour une bonne part dans le contraste réduit que l'on a observé. Un autre point crucial que nos résultats ont dégagé est l'efficacité de détection : on constate en effet que, tant que cette efficacité sera limitée à sa valeur actuelle de 40%, les effets parasites à deux atomes réduiront drastiquement le contraste. Grâce aux nouveaux détecteurs mis au point par Michel Brune et Alexia Auffèves, on peut espérer que cette efficacité sera accrue, ouvrant la voie à des expériences plus complexes puisqu'autorisant l'utilisation de flux plus importants, et donc d'expériences à plusieurs atomes.

Enfin il est possible d'améliorer considérablement la méthode utilisée pour la mesure de $W_1(\alpha)$ et qui consistait à utiliser le même atome à la fois pour préparer le champ dans l'état recherché et pour mesurer sa fonction de Wigner. Le plus grand problème, au cours de la discussion des résultats de cette expérience est venu de l'efficacité limitée de l'impulsion π dans la cavité qui laissait des atomes dans e , lesquels contribuaient de manière non-contrôlée au signal final. Une possibilité serait de transférer tous les atomes restés dans l'état $|e\rangle$ dans un autre niveau (de manière analogue au "shelving" utilisé pour les ions piégés) que l'on ne détecterait pas, avant la première impulsion de Ramsey. Par exemple, il suffirait d'appliquer aux atomes une impulsion π à deux photons sur la transition $e-i$ pour se débarrasser du problème. Nous n'avons pas exploité cette possibilité car nous ne disposons pas d'une source suffisamment puissante pour induire l'impulsion π $e-i$ à un endroit de la cavité où de plus la carte de champ n'est pas favorable. Avec de nouveaux quadrupleurs actifs qui fournissent une puissance de $15dBm$ (soit $30mW$) à $51GHz$ achetés auprès de la société Spacek Labs, cela sera peut-être possible.

De manière analogue, il serait possible d'améliorer la préparation de l'état $|1\rangle$ en se débarrassant des atomes qui sont passés dans $|g\rangle$ avant d'interagir avec la cavité par émission spontanée vers le continuum. Il suffirait d'appliquer une impulsion π sur la transition $|g\rangle \rightarrow |n=48\rangle$ (en effet le niveau $|i\rangle$ serait déjà utilisé) avant que l'atome n'interagisse avec le mode.

5.6.2.b Extensions possibles

Reconstruction de $\alpha|0\rangle + \beta|1\rangle$ Un certain nombre d'expériences peuvent être réalisées en employant des variantes de la méthode que nous avons utilisée. Par exemple, il serait possible de mesurer la fonction de Wigner d'une superposition cohérente $|\psi_{cav}\rangle = \alpha|0\rangle + \beta|1\rangle$ en modifiant très légèrement la séquence expérimentale de la figure 5.18. Il suffirait que l'atome entre dans la cavité dans l'état $|\psi_{at}\rangle = \alpha|g\rangle + \beta|1\rangle$,

qui peut être très simplement réalisé grâce à une impulsion micro-onde sur la transition e-g. L'impulsion π dans la cavité transfère alors la cohérence atomique en une cohérence entre les états $|0\rangle$ et $|1\rangle$ et prépare l'état $|\psi_{cav}\rangle$. La suite de la séquence peut être reprise directement de la figure 5.18.

La reconstruction de cet état démontrerait notre capacité à mesurer des fonctions de Wigner non-invariantes par rotation. En effet, dans nos expériences précédentes, la phase du champ cohérent injecté α n'était pas un paramètre relevant puisque les états de Fock préparés n'ont pas de phase. Pour mesurer la fonction de Wigner de l'état $|\psi_{cav}\rangle$ dans tout l'espace des phases, au contraire, la phase relative entre l'impulsion de préparation de la superposition des états atomiques et le champ cohérent injecté doit être balayée au même titre que l'amplitude du champ α . Il sera intéressant de comparer les résultats de cette reconstruction à ceux obtenus très récemment par Lvovsky et al. par une méthode tomographique [59].

Reconstruction de la fonction de Wigner de $|2\rangle$ Si l'efficacité de détection est améliorée et le champ thermique supprimé, il devient alors possible de mesurer la fonction de Wigner d'un état de Fock à deux photons - ce qui n'a encore jamais été réalisé. Il suffirait d'envoyer un premier atome qui émettrait un photon par une impulsion π dans la cavité initialement vide, et un deuxième atome par la suite viendrait ajouter un deuxième photon dans la première partie de son interaction avec le mode, puis mesurer $W_2(\alpha)$ par la suite.

5.6.3 Perspectives à plus long terme

Reconstruction de la fonction de Wigner d'un état "chat de Schrödinger" L'étude et l'observation directe du processus de décohérence ont été une des motivations importantes de la réalisation du dispositif expérimental d'électrodynamique quantique en cavité. La dynamique de la relaxation des états du type "chat de Schrödinger" ($|\psi_{chat}\rangle = (|\alpha\rangle + |-\alpha\rangle)/\sqrt{2}$) est le paradigme de beaucoup de modèles de décohérence. La fonction de Wigner théorique d'un tel état est représentée à la figure 5.22, ainsi que son évolution temporelle.

La fonction de Wigner montrée à la figure 5.22 présente deux traits distinctifs :

- d'une part, la présence de deux gaussiennes centrées respectivement en α et en $-\alpha$, qui sont bien les fonctions de Wigner de deux états cohérents $|\alpha\rangle$ et $|-\alpha\rangle$.
- d'autre part, des franges d'interférence dont le contraste est maximal à l'origine. Ce sont ces franges qui distinguent l'état $|\psi_{chat}\rangle$ d'un simple mélange statistique de deux états cohérents de phase opposée. On constate que les franges d'interférence disparaissent très rapidement : au bout d'un temps $T_{cav}/2$, elles ont complètement disparu, alors même que les deux gaussiennes sont à peine déplacées vers l'origine.

Ce comportement illustre la fragilité des cohérences quantiques entre états macroscopiquement distincts. La théorie [25] prévoit que la disparition des franges a lieu en un temps caractéristique t_ϕ que l'on nomme "temps de décohérence", et qui est d'autant

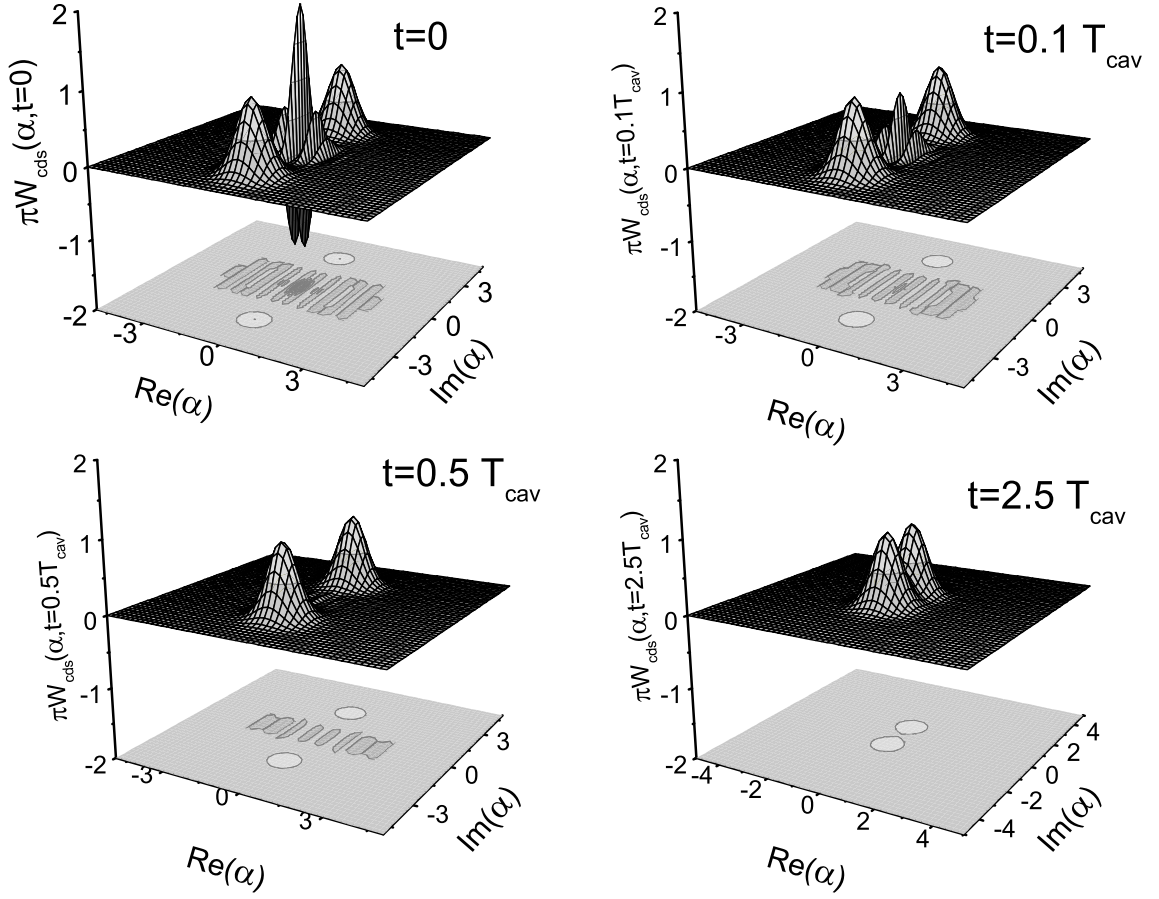


FIG. 5.22 – Fonctions de Wigner d'un état "chat de Schrödinger" $|\psi_{chat}\rangle = (|\alpha\rangle + |-\alpha\rangle)/\sqrt{2}$ avec $\alpha = 3$. De gauche à droite et de bas en haut on peut suivre la relaxation théorique de cet état. Au bout d'un temps $T = T_{cav}/2$, les franges d'interférence ont complètement disparu, alors même que les gaussiennes ont à peine relaxé vers le vide.

plus court que α est grand :

$$t_\phi = \frac{T_{cav}}{2|\alpha|^2} \quad (5.30)$$

Or nous savons préparer un tel état dans la cavité (en effet, la détection d'un atome dans l'état $|g\rangle$ ou $|e\rangle$ dans l'expérience de mesure de la fonction de Wigner du vide projetait précisément le champ dans la cavité dans un tel état). Il serait donc extrêmement intéressant de mesurer directement sa fonction de Wigner, et d'en étudier la dynamique.

Pour réaliser cette expérience, il suffirait d'envoyer successivement deux atomes dans la cavité, chacun effectuant exactement la séquence expérimentale décrite à la figure 5.13, et de faire varier la phase et l'amplitude du champ injecté la deuxième fois $|\beta\rangle$. En reconstruisant la probabilité de détecter le deuxième atome dans $|e\rangle$ conditionnée à la détection du premier dans $|g\rangle$ par exemple, pour chaque valeur de β , on mesurerait la fonction de Wigner $W_{chat}(-\beta)$ avec $|\psi_{chat}\rangle = (|\alpha\rangle + |-\alpha\rangle)/\sqrt{2}$ (α serait le champ injecté avant le passage du premier atome).

Il est clair qu'une telle expérience est très difficile et occasionnera des temps d'acquisition très longs : en effet, il faudra travailler avec des probabilités d'excitation très faibles pour chacun des atomes pour éviter les effets à deux atomes ; en outre il faudra faire varier l'amplitude *et* la phase du champ β . Mais si l'efficacité de détection est vraiment améliorée dans la nouvelle version du montage, il sera sans doute possible de réaliser cette expérience pour des chats allant jusqu'à $\alpha = 2$.

Remarquons d'ailleurs qu'un problème se posera lors de cette expérience. Nous avons constaté au chapitre 4 que le nombre moyen de photons devait rester inférieur à 5 pour réaliser la condition de couplage dispersif, nécessaire pour que le déphasage reste proportionnel au nombre moyen de photons. Lorsqu'on homodyne un état chat de Schrödinger $|\psi_{chat}\rangle = (|\alpha\rangle + |-\alpha\rangle)/\sqrt{2}$ par un champ β , lorsque $\beta = -\alpha$, on ramène l'une des composantes du chat vers l'origine, tandis que l'autre vaut désormais 2α , ce qui signifie que la matrice densité du champ a des éléments non nuls pour un nombre de photons $n \simeq 4|\alpha|^2$. Or, si $\alpha = 2$ par exemple, $4|\alpha|^2 = 16$, ce qui nous fait largement sortir du régime dispersif. Une solution serait de restreindre l'étude de la fonction de Wigner autour de l'origine, qui reste la région la plus intéressante puisque c'est là qu'on s'attend à observer les franges d'interférence manifestant la cohérence mésoscopique entre les deux composantes du chat. [16]

Dans notre cavité, la relaxation du champ est habituellement très bien décrite par un couplage linéaire à un bain d'oscillateurs à température nulle. Ce modèle est connu, on peut le résoudre analytiquement. Il est important de vérifier expérimentalement ses prédictions pour la relaxation d'un état "chat de Schrödinger" comme je l'ai évoqué plus haut. Mais le dispositif d'électrodynamique quantique en cavité permet aussi de réaliser des environnements "sur mesure", qui devraient produire des effets très différents de ceux dus à la dissipation habituelle du champ par diffraction sur la surface des miroirs. En effet, les atomes de Rydberg eux-mêmes peuvent être considérés comme un environnement pour le champ dans la cavité. Notre technique de refroidissement le

prouve bien, à un niveau très rudimentaire. Il est aussi possible d'avoir des couplages plus "exotiques" que l'interaction résonante d'un atome avec un mode. Par exemple, nous savons réaliser des couplages en $\hat{a}^2$ ou en $(\hat{a}^\dagger)^2$ en utilisant un couplage au troisième ordre entre l'atome et les deux modes de la cavité [7].

Mesure non-destructive du nombre de photons Lorsqu'on mesure la fonction de Wigner à l'origine de l'espace des phases (donc sans injecter de champ dans la cavité), on mesure en fait la parité du nombre de photons présents dans la cavité. Cette mesure est bien sûr non-destructive : si l'on envoyait un deuxième atome, il donnerait le même résultat que le premier pour peu que le champ n'ait pas évolué entre-temps. C'est tout l'intérêt d'une interaction dispersive. Avec un seul atome, cette information est en quelque sorte la plus complète que l'on puisse avoir sur le champ, puisque la détection de l'atome dans l'un ou l'autre niveau effectue une projection *orthogonale* de l'état de la cavité sur l'espace des états dont la parité du nombre de photons est bien définie⁷.

Si l'on se restreint au sous-espace des champs dont le nombre de photons est inférieur ou égal à 1, cette mesure de la parité est équivalente à une mesure non-destructive du nombre de photons. Nous avons déjà prouvé que nous étions capables de faire une telle mesure en utilisant une interaction résonante avec la cavité [70].

La mesure de la parité du champ ne constitue qu'un seul bit d'information sur le nombre de photons contenus dans la cavité. Il serait très intéressant de poursuivre la mesure du champ avec d'autres atomes, chacun apportant un bit d'information supplémentaire. Pour déterminer de manière non-ambigüe n'importe quel nombre de photons entre 0 et $N - 1$, il faut utiliser $\log_2 N$ atomes.

Le principe de la méthode est exposé à la figure 5.23 pour mesurer un champ de 4 photons. Un premier atome mesure la parité du champ comme nous l'avons déjà expliqué. Supposons qu'il trouve que le nombre de photons est pair. On envoie ensuite un deuxième atome qui est deux fois plus rapide que le premier. Le déphasage par photon qu'il subira dans la cavité sera deux fois moindre. On choisit la phase de l'interféromètre pour que, si le champ contient 0, 4 ou 8 photons, l'atome soit détecté dans $|e\rangle$, et dans $|g\rangle$ si le champ contient 2 ou 6 photons. On obtient ainsi une information supplémentaire sur le nombre de photons contenus dans le champ, qui complète celle obtenue sur la parité. On conçoit que cette procédure puisse être itérée, et que l'on puisse déterminer complètement le nombre de photons contenus dans le champ [17, 18].

Cette expérience semble réalisable dans l'état actuel du montage puisque la partie la plus difficile concerne la mesure effectuée par le premier atome qui demande un déphasage de π par photon dans le régime dispersif, et nous avons montré dans ce mémoire que nous pouvions l'effectuer. En réalité, le temps de vie de notre cavité est encore trop faible pour réaliser complètement la méthode présentée ci-dessus. En effet,

⁷Si l'on faisait la même opération mais que $\Phi_0 < \pi/2$, on effectuerait une "mesure imparfaite" sur le champ en détectant l'état de l'atome.

cette méthode requiert de pouvoir choisir la phase des franges d'un atome donné *en fonction* du renseignement apporté par les atomes précédents. Par exemple, dans le cas exposé sur la figure 5.23, si le premier atome avait indiqué que le nombre de photons était impair, la détection du deuxième atome n'aurait apporté *aucun* renseignement. Or il faut environ $500\mu s$ au premier atome pour atteindre le détecteur, et dans ce laps de temps le champ dans la cavité est susceptible d'avoir nettement changé si $T_{cav} = 1ms$. Avec une cavité 5 fois meilleure en revanche, il devient envisageable d'effectuer l'expérience.

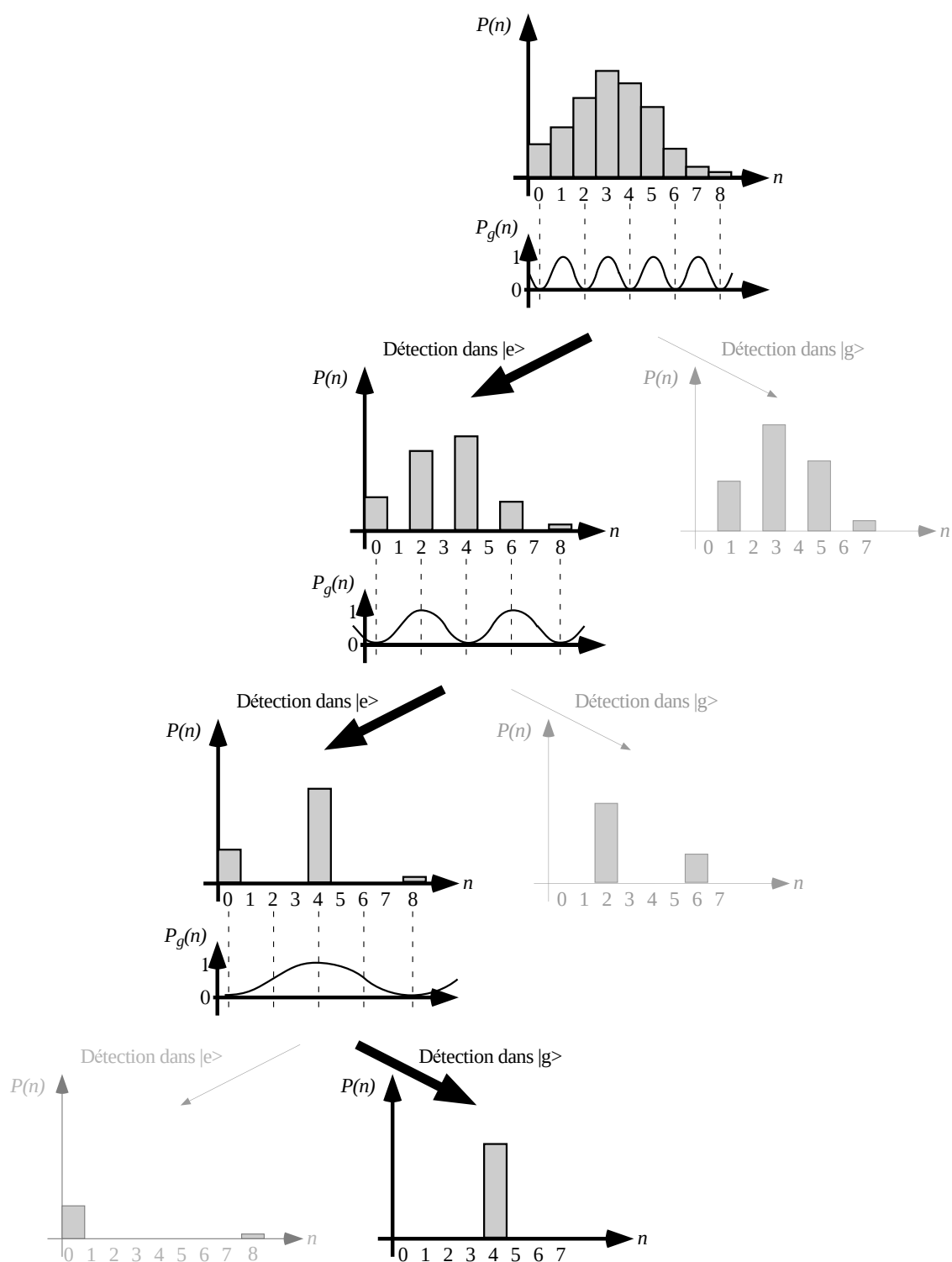


FIG. 5.23 – Mesure non-destructive d'un champ contenant un "grand" nombre de photons, exemple d'une séquence expérimentale conduisant à la mesure de $|4\rangle$.

Annexe A

Eléments de théorie

Le but de cet appendice est de retrouver les deux comportements caractéristiques de l'atome en interaction avec un mode de la cavité dans le modèle de Jaynes-Cummings : oscillations de Rabi quantiques dans le cas où la cavité est accordée à résonance avec la transition atomique e-g, et déphasage proportionnel au nombre de photons du dipôle atomique dans le cas où l'interaction est dispersive.

A.1 Le Hamiltonien de Jaynes-Cummings

Le modèle de Jaynes-Cummings s'applique particulièrement bien à notre dispositif expérimental, puisque c'est sa raison d'être ! Considérons un atome à deux niveaux $|g\rangle$ et $|e\rangle$ en interaction avec un mode du champ électromagnétique.

Le Hamiltonien atomique s'écrit

$$\hat{H}_{\text{at}} = \frac{\hbar\omega_{eg}}{2}(|e\rangle\langle e| - |g\rangle\langle g|) . \quad (\text{A.1})$$

L'origine des énergies a été prise entre les deux niveaux et ω_{eg} est la fréquence de la transition atomique $e \leftrightarrow g$. L'opérateur dipôle atomique, $\hat{d}$, quant à lui, peut être écrit sous la forme

$$\hat{d} = \mathbf{d}_{eg}|e\rangle\langle g| + \mathbf{d}_{eg}^*|g\rangle\langle e| , \quad (\text{A.2})$$

où $\mathbf{d}_{eg} = \langle e|q\hat{\mathbf{r}}|g\rangle$ est l'élément de matrice dipolaire (q représente la charge électrique de l'électron et $\hat{\mathbf{r}}$ est l'opérateur position).

On ne considère qu'un seul mode du champ dans la cavité. Son Hamiltonien s'écrit

$$\hat{H}_C = \hbar\omega_C(\hat{a}^\dagger\hat{a} + \frac{1}{2}) , \quad (\text{A.3})$$

où $\hat{a}^\dagger$ et $\hat{a}$ sont respectivement les opérateurs de création et d'annihilation et ω_C est la fréquence du mode. L'opérateur champ électrique $\hat{\mathbf{E}}(\mathbf{r})$ s'exprime en fonction de $\hat{a}$ et $\hat{a}^\dagger$ suivant l'équation

$$\hat{\mathbf{E}}(\mathbf{r}) = E_0 [\mathbf{f}(\mathbf{r}) \hat{a} + \mathbf{f}^*(\mathbf{r}) \hat{a}^\dagger] , \quad (\text{A.4})$$

où $\mathbf{f}(\mathbf{r}) = f(\mathbf{r}) \cdot \epsilon$ est une fonction vectorielle qui traduit la structure spatiale $f(\mathbf{r})$ et la polarisation ϵ du mode. Nous prenons $f(\mathbf{r})$ comme réelle et positive. Finalement, $E_0 = \sqrt{\hbar\omega_C/2\epsilon_0 V_{\text{mode}}}$ est le champ électrique associé à un photon, où $V_{\text{mode}} = \iiint_{\text{cavité}} |f(\mathbf{r})|^2 d^3\mathbf{r}$ est communément appelé le volume du mode. Si la longueur d'onde est grande devant l'extension spatiale de la fonction d'onde atomique et si l'élément de matrice dipolaire est non-nul, $|\mathbf{d}_{eg}| \neq 0$, le terme prédominant dans l'interaction entre l'atome et le champ est le couplage dipolaire électrique. Ce couplage est décrit par le Hamiltonien $\hat{H}_{\text{dip}} = -\hat{\mathbf{d}} \cdot \hat{\mathbf{E}}$. En utilisant les équations (A.2) et (A.4), on déduit

$$\hat{H}_{\text{dip}} = -E_0 f(\mathbf{r}) (\mathbf{d}_{eg} \cdot \epsilon |e\rangle\langle g| \hat{a} + \mathbf{d}_{eg} \cdot \epsilon^* |e\rangle\langle g| \hat{a}^\dagger + \text{c. c.}) \quad (\text{A.5})$$

L'approximation de l'onde tournante consiste maintenant à négliger les termes non-résonnants de ce Hamiltonien. On obtient le Hamiltonien de couplage atome-champ $\hat{H}_{\text{coupl}}$:

$$\hat{H}_{\text{coupl}} = \hbar G(\mathbf{r}) (|e\rangle\langle g| \hat{a} + |g\rangle\langle e| \hat{a}^\dagger) \quad (\text{A.6})$$

Le paramètre $G(\mathbf{r})$, que nous avons introduit ici, caractérise la force de l'interaction entre les deux sous-systèmes au point $\mathbf{r}$. Nous l'avons choisi réel et positif. Ce choix revient à dire que nous définissons les états atomiques $|e\rangle$ et $|g\rangle$ de telle façon que $-\mathbf{d}_{eg} \cdot \epsilon$ soit réel et positif.

La dépendance spatiale de $G(\mathbf{r})$ n'est fonction que de la structure du mode. On peut donc écrire

$$G(\mathbf{r}) = \frac{\Omega_0}{2} f(\mathbf{r}) , \quad (\text{A.7})$$

où $\Omega_0 = -2E_0 \mathbf{d}_{eg} \cdot \epsilon / \hbar$ est appelé la fréquence de Rabi du vide. Elle est également réelle et positive. $\Omega_0/2$ est le couplage maximal de l'atome au mode (en effet, nous choisissons une normalisation telle que $f(\mathbf{r}) \leq 1$).

Le Hamiltonien total du système est $\hat{H}_{\text{JC}} = \hat{H}_{\text{at}} + \hat{H}_C + \hat{H}_{\text{coupl}}$. Il est couramment appelé Hamiltonien de Jaynes-Cummings [47] :

$$\hat{H}_{\text{JC}} = \frac{\hbar\omega_{eg}}{2} (|e\rangle\langle e| - |g\rangle\langle g|) + \hbar\omega_C (\hat{a}^\dagger \hat{a} + \frac{1}{2}) + \hbar G(\mathbf{r}) (|e\rangle\langle g| \hat{a} + |g\rangle\langle e| \hat{a}^\dagger) . \quad (\text{A.8})$$

Les états stationnaires du système non couplé ($G(\mathbf{r}) = 0$) sont $|e, n\rangle = |e\rangle \otimes |n\rangle$ et $|g, n\rangle = |g\rangle \otimes |n\rangle$. Ils correspondent à l'atome dans l'état e (resp. g) en présence de n photons. Le Hamiltonien d'interaction H_{coupl} couple ces états deux par deux, $|e, n\rangle$ avec $|g, n+1\rangle$. Seul le niveau fondamental $|g, 0\rangle$ n'est couplé à aucun autre état.

A.2 L'évolution temporelle du système

A.2.1 Les états habillés

Nous allons déterminer les états propres qui résultent du couplage entre les états $|e, n\rangle$ et $|g, n+1\rangle$. Ecrivons la matrice du hamiltonien dans la base non-couplée :

$$\begin{aligned}
\hat{H}_n &= \hbar \begin{pmatrix} \omega_{eg}/2 + \omega_C(n+1/2) & G(\mathbf{r})\sqrt{n+1} \\ G(\mathbf{r})\sqrt{n+1} & -\omega_{eg}/2 + \omega_C(n+3/2) \end{pmatrix} \\
&= \hbar \begin{pmatrix} \omega_C(n+1) - \delta/2 & G(\mathbf{r})\sqrt{n+1} \\ G(\mathbf{r})\sqrt{n+1} & \omega_C(n+1) + \delta/2 \end{pmatrix},
\end{aligned} \tag{A.9}$$

où $\delta = \omega_C - \omega_{eg}$ est le désaccord entre l'atome et le champ.

En diagonalisant cette matrice, on trouve les nouveaux états propres du système couplé :

$$\begin{aligned}
|+, n\rangle &= \cos \theta_n |e, n\rangle - \sin \theta_n |g, n+1\rangle \\
|-, n\rangle &= \sin \theta_n |e, n\rangle + \cos \theta_n |g, n+1\rangle
\end{aligned} \tag{A.10}$$

L'angle θ_n est défini par la relation

$$\tan(2\theta_n) = 2 \frac{G(\mathbf{r})}{\delta} \sqrt{n+1} \text{ et } 0 \leq \theta_n \leq \frac{\pi}{2} \tag{A.11}$$

Les énergies propres de ces états valent

$$E_{\pm, n} = \hbar[\omega_C(n+1) \mp \frac{\sqrt{\delta^2 + 4G(\mathbf{r})^2(n+1)}}{2}] \tag{A.12}$$

Lorsque l'atome traverse la cavité, le couplage atome-champ varie à cause de la structure spatiale du mode qui est gaussienne avec un waist w (voir chapitre 2). Si l'on appelle l'axe du jet atomique l'axe des x , et qu'on en prend l'origine au centre du mode, on a donc

$$f(x) = \exp(-x^2/w^2) \tag{A.13}$$

(on néglige ici la variation du couplage en y et en z) Le mouvement des atomes à-travers le dispositif n'est pas affecté par leur interaction avec la cavité, il s'agit donc d'un mouvement rectiligne uniforme de vitesse v . On peut donc transformer la variation spatiale du couplage atome-champ en une variation temporelle

$$G(t) = \frac{\Omega_0}{2} f(vt) = \frac{\Omega_0}{2} \exp\left(-\frac{(vt)^2}{w^2}\right) \tag{A.14}$$

L'évolution temporelle du système est en général complexe et doit être intégrée numériquement dans le cas général. Nous pouvons cependant la déterminer de manière analytique dans deux cas limites : le cas résonant ($\delta = 0$) et le cas très désaccordé ($\delta \gg \Omega$).

A.2.2 L'interaction résonante

Nous nous intéressons dans ce paragraphe au cas $\delta = 0$ ($\omega_C = \omega_{eg}$). Les deux états propres que nous avons calculés précédemment sont alors :

$$\begin{aligned}
|+, n\rangle &= (|e, n\rangle - |g, n+1\rangle) \frac{1}{\sqrt{2}} \\
|-, n\rangle &= (|e, n\rangle + |g, n+1\rangle) \frac{1}{\sqrt{2}}
\end{aligned} \tag{A.15}$$

d'énergies propres $E_{\pm,n} = \hbar\omega_{eg}(n+1) \mp G(\mathbf{r})\sqrt{n+1}$. Il est donc particulièrement simple d'exprimer les états non couplés $|e, n\rangle$ et $|g, n+1\rangle$ en fonction des états habillés :

$$\begin{aligned} |e, n\rangle &= (|+, n\rangle + |-, n\rangle) \frac{1}{\sqrt{2}} \\ |g, n+1\rangle &= (|-, n\rangle - |+, n\rangle) \frac{1}{\sqrt{2}} \end{aligned} \quad (\text{A.16})$$

Soulignons que ces expressions ne dépendent pas de la position de l'atome dans la cavité. Supposons maintenant que l'on envoie un atome, initialement dans l'état $|e\rangle$, dans la cavité contenant n photons. L'état initial du système atome-champ est donc $|\psi(t=0)\rangle = |e, n\rangle$. On décompose l'état à l'instant t du système sur la base des états habillés :

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}}(a_{+,n}(t)|+, n\rangle + a_{-,n}(t)|-, n\rangle) \quad (\text{A.17})$$

On calcule les coefficients $a_{+,n}(t)$ et $a_{-,n}(t)$ grâce à l'équation de Schrödinger :

$$\frac{d|\psi\rangle}{dt} = \frac{i}{\hbar} \hat{H}|\psi\rangle \quad (\text{A.18})$$

On obtient

$$\begin{aligned} a_{+,n}(t) &= \exp(-i \frac{\Omega_0 \sqrt{n+1}}{2} t_{eff}) \\ a_{-,n}(t) &= \exp(i \frac{\Omega_0 \sqrt{n+1}}{2} t_{eff}) \end{aligned} \quad (\text{A.19})$$

où l'on a posé $t_{eff} = \int_{-\infty}^t \exp(-(vu)^2/w^2) du$. Il s'agit d'un temps d'interaction effectif, la durée de l'interaction dans le cas où le couplage atome-champ serait constant et égal à $\Omega_0/2$. Dans le cas où l'atome traverse le mode entièrement, ($t \rightarrow +\infty$), le couplage tend vers 0 et le système n'évolue plus. Le temps effectif d'interaction associé à ce transit est :

$$\begin{aligned} t_{\text{eff,transit}} &= \int_{-\infty}^{+\infty} \exp\left(-\frac{(vu)^2}{w_0^2}\right) du \\ &= \frac{\sqrt{\pi} w_0}{v} \end{aligned} \quad (\text{A.20})$$

Ainsi, à l'instant t correspondant à un temps effectif d'interaction t_{eff} , l'état du système sera :

$$|\psi(t)\rangle = \cos(\Omega_0 \sqrt{n+1} t_{eff}/2) |e, n\rangle - i \sin(\Omega_0 \sqrt{n+1} t_{eff}/2) |g, n+1\rangle \quad (\text{A.21})$$

De cette évolution on tire les probabilités de détecter l'atome dans $|e\rangle$ (resp. $|g\rangle$) en fonction du temps effectif d'interaction, qui valent

$$\begin{aligned} P_e(t_{\text{eff}}) &= \cos^2(\Omega_0 \sqrt{n+1} t_{\text{eff}}/2) = \frac{1}{2} [1 + \cos(\Omega_0 \sqrt{n+1} t_{\text{eff}})] \\ P_g(t_{\text{eff}}) &= \sin^2(\Omega_0 \sqrt{n+1} t_{\text{eff}}/2) = \frac{1}{2} [1 - \cos(\Omega_0 \sqrt{n+1} t_{\text{eff}})] \end{aligned} \quad (\text{A.22})$$

Ces probabilités oscillent en fonction de t_{eff} à la fréquence $\Omega_0 \sqrt{n+1}$. Pour n petit, on est dans le régime des oscillations de Rabi quantiques [21, 61].

On obtient le même genre d'évolution si l'état initial du système à l'instant t' est $|g, n+1\rangle$:

$$|g, n+1\rangle \rightarrow \cos(\Omega_0\sqrt{n+1}t_{\text{eff}}/2)|g, n+1\rangle - i\sin(\Omega_0\sqrt{n+1}t_{\text{eff}}/2)|e, n\rangle, \quad (\text{A.23})$$

avec les probabilités de détecter l'atome dans $|e\rangle$ (resp. $|g\rangle$)

$$\begin{aligned} P_e(t_{\text{eff}}) &= \sin^2(\Omega_0\sqrt{n+1}t_{\text{eff}}/2) = \frac{1}{2} [1 - \cos(\Omega_0\sqrt{n+1}t_{\text{eff}})] \\ P_g(t_{\text{eff}}) &= \cos^2(\Omega_0\sqrt{n+1}t_{\text{eff}}/2) = \frac{1}{2} [1 + \cos(\Omega_0\sqrt{n+1}t_{\text{eff}})] \end{aligned} \quad (\text{A.24})$$

A.2.3 L'interaction dispersive

Dans le cas où $\delta \gg \Omega_0\sqrt{n+1}/2$ (on suppose ici $\delta > 0$), l'angle de mélange θ_n est petit devant 1. Par conséquent, on peut faire l'approximation $|+, n\rangle \simeq |e, n\rangle$ et $|-, n\rangle \simeq |g, n+1\rangle$. Par ailleurs, on peut développer les énergies propres $E_{\pm, n}$ selon Ω/δ :

$$E_{\pm, n} = \hbar[\omega_C(n+1) \mp \frac{\sqrt{\delta^2 + 4G(\mathbf{r})^2(n+1)}}{2}] \simeq \hbar[\omega_C(n+1) \mp \frac{\delta}{2} \mp \frac{G(\mathbf{r})^2(n+1)}{\delta}] \quad (\text{A.25})$$

Ainsi, tout se passe comme si l'état $|e, n\rangle$ était juste déplacé en énergie d'une quantité $\delta E_e = -G(\mathbf{r})^2(n+1)/\delta$ à cause du couplage dispersif de l'atome à la cavité. De même, l'état $|g, n+1\rangle$ est déplacé en énergie d'une quantité opposée, $\delta E_g = +G(\mathbf{r})^2(n+1)/\delta$. Ces déplacements en énergie induisent des déphasages de chacun des états atomiques proportionnels au nombre de photons et au temps d'interaction effectif :

$$\begin{aligned} \delta\Phi_e &= \int_{-\infty}^t \delta E_e(t') dt' \\ \delta\Phi_g &= \int_{-\infty}^t \delta E_g(t') dt' \end{aligned} \quad (\text{A.26})$$

Lorsque l'atome interagit avec le mode de $-\infty$ à $+\infty$, on peut calculer ces déphasages. On obtient

$$\begin{aligned} \delta\Phi_e &= -\frac{\Omega_0^2(n+1)}{4\delta} \frac{t_{eff}}{\sqrt{2}} \\ \delta\Phi_g &= \frac{\Omega_0^2(n+1)}{4\delta} \frac{t_{eff}}{\sqrt{2}} \end{aligned} \quad (\text{A.27})$$

Posons $\Phi_0 = \Omega_0^2 t_{eff} / 4\sqrt{2}\delta$. L'effet du couplage atome–champ dans le cas très désaccordé est donc un déphasage des états par rapport au cas non-couplé. Si l'atome traverse le mode entièrement, ce déphasage vaut

$$\begin{aligned} |e, n\rangle &\xrightarrow{\text{transit}} \exp(-i\Phi_0(n+1)) |e, n\rangle \\ |g, n+1\rangle &\xrightarrow{\text{transit}} \exp(i\Phi_0(n+1)) |g, n+1\rangle \end{aligned} \quad (\text{A.28})$$

Annexe B

Discussion des facteurs de normalisation des mesures de fonctions de Wigner

Dans cet appendice, nous menons une analyse détaillée des imperfections des expériences sur la mesure des fonctions de Wigner, en vue de comprendre quantitativement les facteurs de normalisation mesurés. Les conclusions de cette analyse se trouvent dans le chapitre 5.

B.1 Analyse détaillée des imperfections de l'expérience de mesure de la fonction de Wigner du vide

Dans cette expérience, nous avons mesuré un facteur de normalisation $F_0 = 2.44$. Nous essayons de le retrouver quantitativement dans cette partie avec trois effets dont nous calculons les contributions f_i respectives (finalement, $F_0 = \Pi_i f_i$).

B.1.1 Contraste intrinsèque de l'interféromètre

Tout d'abord, l'interféromètre de Ramsey a un contraste qui ne vaut pas 100% même en l'absence de la cavité. Nous avons mesuré ce contraste intrinsèque en désaccordant beaucoup la cavité par rapport à la résonance atomique, de sorte que les déplacements lumineux inhomogènes dûs au couplage dispersif soient négligeables. Pour un désaccord de 5MHz, nous avons enregistré des franges contrastées à 70%. Plusieurs facteurs expliquent ce contraste, largement inférieur à celui mesuré par exemple avec des atomes à 500 m/s (pour lesquels nous avons déjà obtenu des franges dont le contraste valait 87%). Notons tout d'abord que la qualité des sources qui produisent les impulsions de Ramsey n'est pas en cause. En effet, la largeur spectrale de ces sources est largement

inférieure au Hertz, et la porteuse est plus de 40dB au-dessus du bruit à 100Hz. Par contre, la dispersion géométrique des atomes au moment des impulsions produit nécessairement des dispersions d'angle de Rabi qui réduisent le contraste. Enfin les atomes qui vont à 150m/s "vivent" beaucoup plus longtemps que des atomes à 500m/s ; notamment, l'émission spontanée dans les autres modes que ceux de la cavité commence à jouer un rôle significatif. De plus, les champs inhomogènes ont beaucoup plus d'effet sur des atomes lents que sur les atomes rapides, car ils voient ces inhomogénéités pendant une durée plus importante, et perdent donc plus que les autres leur cohérence.

Cet effet contribue donc pour un facteur $f_1 = 1/0.7 = 1.43$.

B.1.2 Événements à deux atomes

Un autre effet, plus complexe, réduit aussi le contraste des franges. Il existe même lorsque la cavité est vide (donc même si le champ thermique résiduel était supprimé). Il concerne les événements à deux atomes. Nous avons vu au chapitre 2 que nous ne savions pas préparer exactement 1 atome de Rydberg, et qu'au contraire notre processus de préparation avait une statistique poissonnienne. Par conséquent, même lorsque le nombre moyen d'atomes préparés est faible ($\ll 1$), le risque d'avoir deux atomes dans le même paquet est non nul. Or notre détecteur, qui a une efficacité de détection limitée à 40%, peut ne détecter qu'un seul de ces deux atomes. A première vue, on aurait tendance à penser que si le paquet atomique qui mesure la fonction de Wigner contenait deux atomes ou plus, le signal n'en serait pas perturbé. Supposons en effet que deux atomes interagissent dispersivement avec un mode du champ dans l'état $|0\rangle$. On pourrait penser que le système évolue comme suit

$$\frac{1}{\sqrt{2}}(|e_1\rangle + |g_1\rangle) \otimes \frac{1}{\sqrt{2}}(|e_2\rangle + |g_2\rangle) \otimes |0\rangle \xrightarrow{cav|0\rangle} \frac{1}{\sqrt{2}}(\exp^{i\Phi_0} |e_1\rangle + |g_1\rangle) \otimes \frac{1}{\sqrt{2}}(\exp^{i\Phi_0} |e_2\rangle + |g_2\rangle) \otimes |0\rangle \quad (\text{B.1})$$

auquel cas les franges sont inchangées quel que soit le nombre d'atomes puisque les fonctions d'onde atomiques se factorisent.

En fait B.1 est faux. Il existe en effet un processus du second ordre qui intrique les deux atomes via la cavité. Nous allons évoquer très brièvement ce processus dont on pourra trouver une étude plus détaillée dans [73], et qui nous a permis de préparer deux atomes dans un état maximalement intriqué dans une autre expérience [74]. Considérons l'état $|e_1, g_2, 0\rangle$ où un premier atome est dans e et l'autre dans g (la cavité étant vide). Deux types de processus virtuels peuvent exister (qui sont schématisés sur la figure B.1) :

- soit l'atome 1 émet et réabsorbe un photon. Ce processus (1) est responsable du déplacement lumineux subi par le niveau $|e\rangle$ en interaction dispersive avec le vide.
- soit l'atome 1 émet virtuellement un photon et l'atome 2 le réabsorbe (processus (2)). On voit que ce processus conserve l'énergie, et qu'il couple l'état $|e_1, g_2\rangle$ à $|g_1, e_2\rangle$. D'autre part, il est évident d'après notre discussion que les couplages des processus (1) et (2) sont strictement égaux (la seule différence réside dans le label de l'atome qui

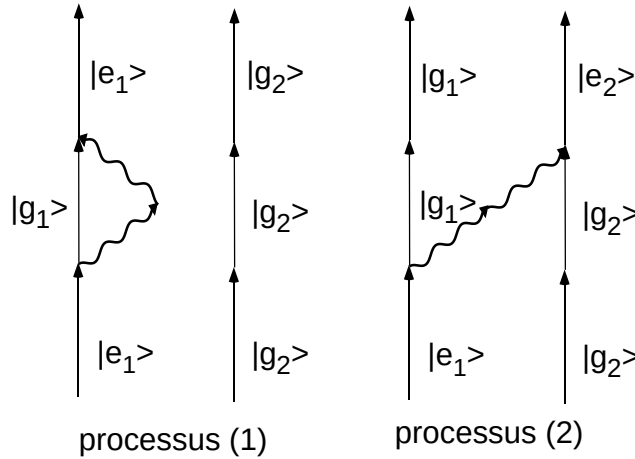


FIG. B.1 – Lorsque deux atomes sont couplés dispersivement à la cavité vide, deux processus virtuels distincts entrent en jeu : dans le processus (1), l'atome dans l'état $|e\rangle$ peut émettre et réabsorber virtuellement un photon. Ce processus est responsable d'un déplacement lumineux de Lamb lorsque la cavité est vide, et lorsqu'elle contient un champ, des déplacements lumineux qui déphasent le dipôle atomique de π par photon, permettant la mesure de la fonction de Wigner. Le processus (2) fait intervenir un échange de photons entre les deux atomes. La constante de couplage des deux processus est clairement la même.

réabsorbe le photon, les deux atomes étant bien évidemment distinguables par leurs variables externes). On ne peut donc certainement pas négliger cet effet par rapport à l'interaction dispersive avec la cavité. Le résultat est qu'une oscillation de Rabi a lieu entre les états $|e_1, g_2\rangle$ et $|g_1, e_2\rangle$. S'il n'y avait qu'un seul mode, l'angle de Rabi correspondant serait exactement $2\Phi_0 = \pi$ (l'oscillation de Rabi va deux fois plus vite que le couplage) pour les paramètres de l'expérience. En réalité, la présence du deuxième mode accélère l'oscillation et fait que l'angle de Rabi sera plutôt $2\Phi_1 = 3\pi/2$. Ecrivons l'évolution des états $|e_1, g_2, 0\rangle$ et $|g_1, e_2, 0\rangle$ sous l'effet de ce couplage :

$$\begin{aligned} |e_1, g_2, 0\rangle &\xrightarrow{\Phi_1} (\cos \Phi_1 |e_1, g_2, 0\rangle - i \sin \Phi_1 |g_1, e_2, 0\rangle) \exp^{i\Phi_1} \\ |g_1, e_2, 0\rangle &\xrightarrow{\Phi_1} (\cos \Phi_1 |g_1, e_2, 0\rangle - i \sin \Phi_1 |e_1, g_2, 0\rangle) \exp^{i\Phi_1} \end{aligned} \quad (\text{B.2})$$

Reprenons maintenant l'évolution de l'équation B.1 en tenant compte de ce processus. Les états $|e_1, e_2, 0\rangle$ et $|g_1, g_2, 0\rangle$ ne sont pas affectés par le processus (2).

$$\begin{aligned} |\Psi_{in}\rangle &= \frac{1}{\sqrt{2}}(|e_1\rangle + |g_1\rangle) \otimes \frac{1}{\sqrt{2}}(|e_2\rangle + |g_2\rangle) \otimes |0\rangle \xrightarrow{cav|0\rangle} \\ |\Psi_{fin}\rangle &= \frac{1}{2}[\exp^{i2\Phi_1} |e_1, e_2, 0\rangle + |g_1, g_2, 0\rangle + \\ &\quad \exp^{i\Phi_1} (\cos \Phi_1 - i \sin \Phi_1) |e_1, g_2, 0\rangle + \exp^{i\Phi_1} (\cos \Phi_1 - i \sin \Phi_1) |g_1, e_2, 0\rangle] \\ &= \frac{1}{2}[\exp^{i2\Phi_1} |e_1, e_2, 0\rangle + |g_1, g_2, 0\rangle + |e_1, g_2, 0\rangle + |g_1, e_2, 0\rangle] \end{aligned} \quad (\text{B.3})$$

On peut calculer la matrice densité réduite d'un atome pour voir dans quelle mesure il garde une cohérence qui se traduirait par des franges de Ramsey. On trouve

$$\begin{aligned} \rho_{at1} &= Tr_{at2}(|\Psi_{fin}\rangle\langle\Psi_{fin}|) \\ &= \frac{1}{4} \begin{pmatrix} 2 & 2 \cos \Phi_1 \exp^{i\Phi_1} \\ 2 \cos \Phi_1 \exp^{-i\Phi_1} & 2 \end{pmatrix} \end{aligned} \quad (B.4)$$

L'équation B.4 montre que le terme non-diagonal de cohérence entre les états $|e\rangle$ et $|g\rangle$ n'est pas nul, ce qui signifie qu'on observerait des franges d'interférence après une deuxième impulsion R2, de contraste réduit $|\cos \Phi_1| = 0.7$. Par contre, puisque $\cos \Phi_1 = -0.7$, la phase de ces franges est *opposée* à celles que l'on enregistre avec des paquets qui ne contiennent qu'un seul atome! (qui elles ont une phase $+\Phi_1$)

Essayons de déterminer quantitativement l'effet de ce processus sur les résultats expérimentaux. Nous allons commencer par faire une hypothèse que nous n'avons pas directement testée expérimentalement. Il s'agit de dire que le résultat B.4 reste valable quel que soit le champ dans la cavité, c'est-à-dire que les événements à deux atomes donnent des franges en opposition de phase et contrastées à 70% par rapport aux événements "normaux" à un seul atome quel que soit l'état de la cavité. Cette hypothèse a un fondement physique, qui est l'insensibilité du couplage entre $|g_1, e_2\rangle$ et $|e_1, g_2\rangle$ au champ dans la cavité dans le régime dispersif, dû à une interférence destructrice entre un processus virtuel du type (2) et un processus virtuel du même type, mais où l'atome commencerait par absorber un photon avant que le deuxième n'en émette un. Nous avons observé indirectement cette insensibilité dans l'expérience décrite dans la référence [74].

En partant de cette hypothèse, nous allons déterminer l'effet des événements à deux atomes sur le contraste des franges de Ramsey en fonction du flux atomique détecté f (qui est le paramètre auquel on a directement accès) et de l'efficacité de détection ϵ (dans l'expérience ϵ vaut 40%). Le flux atomique réel vaut $f_r = f/\epsilon$. La proportion d'événements à deux atomes par rapport aux événements à un atome est donné par le rapport $r = p(2)/p(1) = f_r/2$ pour une distribution poissonnienne. Or, sur $N = 100$ événements, il y a un nombre $N(1 - r)\epsilon$ d'événements à un atome qui ont été correctement détectés, et un nombre $2rN\epsilon(1 - \epsilon)$ d'événements à deux atomes dont un seul aura été détecté correctement. Donc la proportion de ces événements à deux atomes vaut $p = 2r\epsilon(1 - \epsilon)/((1 - r)\epsilon + 2r\epsilon(1 - \epsilon))$. La réduction de contraste des franges est directement proportionnelle à p si l'on admet que les événements à deux atomes donnent des franges en opposition de phase contrastées à 70% :

$$C = 1.4p = 1.4 \frac{f(1 - \epsilon)}{\epsilon(1 - \frac{f}{2\epsilon} + f\frac{1-\epsilon}{\epsilon})} \simeq 1.4f \frac{(1 - \epsilon)}{\epsilon} \quad (B.5)$$

On a représenté la valeur théorique de C en fonction du flux mesuré f pour une efficacité de détection de $\epsilon = 0.4$ sur la figure B.2. On constate qu'aux flux où l'on travaille d'habitude (entre 0.1 et 0.2 atomes par courbe), la réduction de contraste est très importante, allant de 18 à 30%. C'est pourquoi nous avons travaillé à très faible

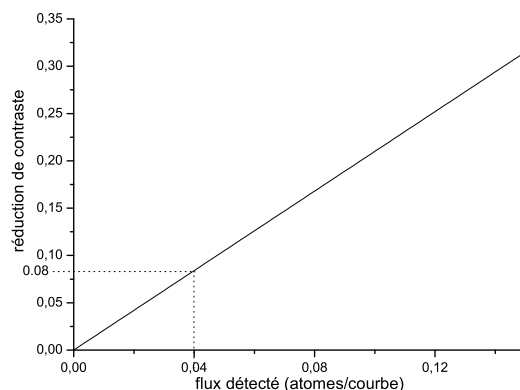


FIG. B.2 – Réduction de contraste en fonction du flux détecté pour une efficacité de détection de $\epsilon = 0.4$. On constate que pour un flux $f = 0.04at/c$, cette réduction est de 8%.

flux ($f=0.04at/c$) lors de la mesure de la fonction de Wigner montrée précédemment ; la réduction de contraste est alors néanmoins de 8%. Nous avons expérimentalement constaté l'effet drastique de réduction du contraste prédit plus haut : pour un flux de $f = 0.04at/c$ le contraste des franges est de 35% alors qu'il tombe à 11% lorsque $f = 0.12at/c$.

Ce dernier effet contribue pour $f_3 = 1/(1 - 0.08) = 1.09$.

B.1.3 Champ thermique dans le mode BF

Enfin un dernier effet implique le mode Basse Fréquence de la cavité. Bien que celui-ci soit lui aussi refroidi au début de chaque séquence expérimentale par le deuxième paquet des serpillères adiabatiques, nous avons de bonnes raisons de croire qu'il est un peu plus "chaud" que le mode HF au cours de l'expérience. En effet, il ne voit qu'un seul paquet atomique de refroidissement contrairement au mode HF qui en voit deux. En outre, le délai est plus long entre le passage de ce paquet atomique et de l'atome de mesure, donc le champ thermique a plus de temps pour repousser. De plus, nous contrôlons moins bien le passage adiabatique sur le mode BF que sur le mode HF, car pour mettre l'atome à résonance avec celui-ci il faut des champs Stark importants, et nous avons déjà expliqué que les atomes sont plus sensibles aux inhomogénéités dans ce cas de figure. Enfin n'oublions pas que le temps de vie de ce mode est plus court que l'autre. Or si la présence d'un seul photon dans le mode HF suffit à déphaser le dipôle atomique de π , un photon dans le mode BF déphasera le dipôle atomique de $\pi/2$ (en effet, $\Phi_0^{(BF)} = \Phi_0^{(HF)}\delta/(\delta + \Delta)$; comme $\delta = 118kHz \simeq \Delta$, finalement $\Phi_0^{(BF)} = \Phi_0^{(HF)}/2$).

Il est difficile d'évaluer précisément le champ thermique dans le mode BF au moment de l'expérience, et nous n'avons pas de données expérimentales sur ce point.

Néanmoins, nous avons pu calculer le champ thermique résiduel nécessaire pour expliquer la réduction de contraste observée. Si l'on suppose que le champ résiduel est de 0.35 photons après passage du paquet atomique de refroidissement, on retrouve bien le contraste observé. La contribution de cet effet est $f_2 = 1.5$.

Le produit $F = f_1 f_2 f_3$ vaut 2.3, ce qui est proche de la valeur mesurée $F_0 = 2.44$. Nous pensons donc comprendre quantitativement les réductions de contraste observées dans l'expérience.

B.2 Analyse détaillée du facteur de normalisation de la mesure de la fonction de Wigner de $|1\rangle$

Dans cette deuxième expérience, le facteur de normalisation était de $F_1 = 4.0$, supérieur à celui de l'expérience précédente. Nous analysons dans cette partie les imperfections qui s'ajoutent à celles qui ont été évoquées dans la partie B.1.

B.2.1 La première impulsion de Ramsey

Remarquons tout d'abord que la première impulsion de Ramsey R1 est appliquée en même temps que l'on injecte le champ dans la cavité, et à un moment où l'atome est déjà couplé au mode. Il nous faut justifier que cela ne cause pas d'effet systématique à notre mesure.

Lorsqu'on injecte un champ dans la cavité, on allume la source S_{cav} pendant la durée T_{on} . Le guide d'onde "cavité" couple la source à la cavité. Lorsqu'on allume S_{cav} , une partie de l'onde transmise par le guide est couplée dans le mode de la cavité, mais une partie importante est aussi diffractée par le trou de couplage hors du mode considéré jusqu'à présent. Ainsi, même si l'atome n'est pas dans le mode de la cavité, il voit une partie de l'impulsion micro-onde d'injection. Cela ne se traduit pas par une modification des populations des niveaux $|e\rangle$ ou $|g\rangle$ dans la mesure où l'atome est en général désaccordé par rapport à S_{cav} . Par contre, cette impulsion cause nécessairement des déplacements lumineux incontrôlés pendant l'injection. Lorsque l'atome est dans un état propre de l'énergie, ces déplacements lumineux n'ont aucune incidence sur la fonction d'onde atomique. Mais au cours de l'expérience de mesure de la fonction de Wigner de $|1\rangle$, l'injection avait lieu précisément au moment où l'impulsion de Ramsey générerait une superposition cohérente des niveaux $|e\rangle$ et $|g\rangle$. Nous avons donc dû vérifier que cela ne réduisait pas artificiellement le contraste des franges.

Nous avons mis au point un test pour vérifier que la cohérence de la superposition $(|e\rangle + |g\rangle)/\sqrt{2}$ n'était pas affectée par l'injection simultanée de champ dans le mode de la cavité. Le meilleur moyen de vérifier une cohérence entre deux niveaux atomiques est bien sûr de mesurer le contraste de franges de Ramsey. Nous avons réalisé des franges de Ramsey en effectuant l'impulsion R1 bien avant le mode, et l'impulsion R2 à l'instant exact où a lieu l'injection dans le mode dans l'expérience. Puis nous avons comparé le

contraste des franges obtenues avec et sans l'impulsion d'injection (pour le champ le plus important que nous injectons dans l'expérience). La figure B.3 montre les deux signaux obtenus. On constate que la phase des franges n'est pas affectée par l'injection. Le contraste par contre subit une réduction de 8% pour $\alpha = \alpha_{max}$. On peut faire l'hypothèse que le contraste est réduit d'un facteur dépendant linéairement de l'intensité du champ injecté (puisque cette réduction provient de déplacements lumineux qui sont eux-mêmes proportionnels au nombre moyen de photons injectés). On constate alors que cette erreur systématique n'affecte pas significativement les résultats obtenus sur la fonction de Wigner (figure 5.21).

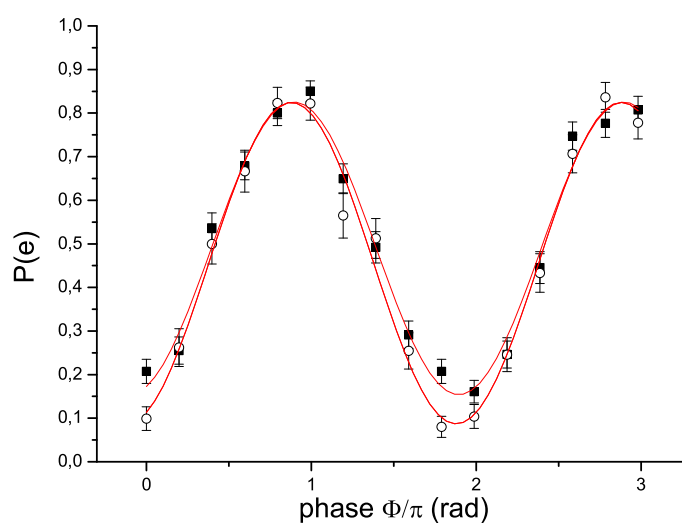


FIG. B.3 – Influence de l'injection simultanée d'un champ $\alpha = 1.8$ sur le contraste des franges. On a enregistré des franges de Ramsey en faisant l'impulsion R1 avant le mode et l'impulsion R2 au moment où l'on faisait R1 dans l'expérience. La courbe en cercles ouverts représentent les franges obtenues sans injection du champ α ; la courbe en carrés pleins montre les franges obtenues lorsqu'on injecte le champ le plus intense en même temps que l'on applique l'impulsion R2. La réduction de contraste est de 0.91. Cette erreur systématique sur le contraste n'affecte pas significativement les résultats de la figure 5.21

Une autre cause d'erreur provient de la structure des modes du système "anneau+cavité". On constatera en se reportant à la figure 3.10 que les ventres de l'onde stationnaire créée se situent surtout de part et d'autre du mode. Or nous avons eu besoin pour cette expérience d'appliquer une impulsion de Ramsey à l'atome lorsque ce dernier se trouvait déjà dans le mode. Donc nous avons eu besoin d'allumer la source S_{Ram} à une puissance plus importante que précédemment pour que l'atome effectue une impulsion $\pi/2$ pendant la même durée. Il nous a fallu vérifier à nouveau que cette

impulsion de Ramsey ne pollueait pas les modes de la cavité à cette puissance plus importante. Nous avons vérifié ce point en envoyant une sonde atomique absorber le champ résiduel après l'impulsion R1. On constate sur la figure B.4 que, même dans ces conditions, le champ injecté est complètement négligeable (inférieur à 0.05 photon).

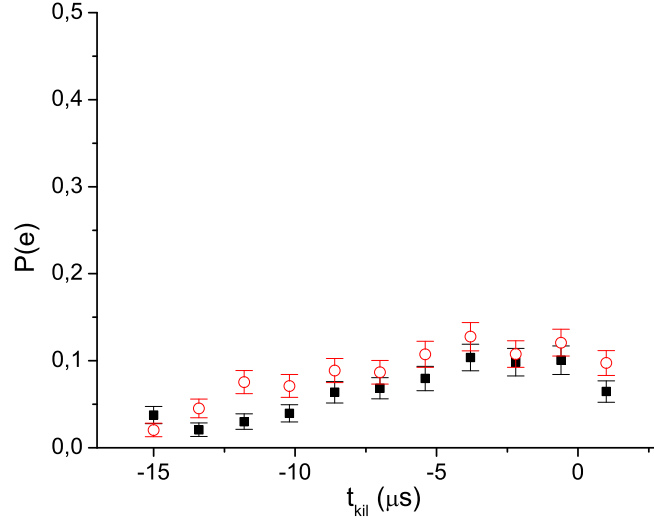


FIG. B.4 – Mesure du champ injecté dans le mode HF par l'impulsion R1 de l'expérience de mesure de $W_1(\alpha)$. Courbe en carrés pleins : absorption du champ thermique résiduel (environ 0.15 photons thermiques) sans impulsion R1 par une sonde atomique. Courbe en cercles vides : absorption du champ par la sonde après une impulsion R1. On constate qu'il y a un peu de champ injecté, mais cela reste un champ dont le nombre moyen de photons est inférieur à 0.2, ce qui ne perturbe pas les données.

Bien que le champ injecté par R1 dans les modes de la cavité soit négligeable, nous avons constaté une réduction du contraste des franges de référence, prises pour un atome préparé initialement dans l'état $|g\rangle$, tout le temps désaccordé : $c = 25\%$ au lieu des 32% que nous avons en effectuant l'impulsion R1 en-dehors du mode. Nous attribuons cette réduction du contraste à deux causes :

- d'une part ; l'expérience a été réalisée à un flux légèrement supérieur à l'expérience précédente $f = 0.06at/c$, ce qui augmente la probabilité d'événements à deux atomes.

- d'autre part, nous n'avons pas eu le choix de la zone où réaliser l'impulsion R1 puisqu'elle devait impérativement être effectuée juste après que l'atome ait effectué son impulsion π dans la cavité. Ainsi, le champ micro-onde dans lequel a été réalisée l'impulsion R1 était-il beaucoup plus inhomogène que les champs des zones R1 et R2 précédentes.

Cet effet contribue pour un facteur $f_4 = 1.28$.

B.2.2 Les imperfections de l'impulsion π dans la cavité

Une cause plus importante d'erreur provient du fait que l'impulsion π qui permet de préparer l'état $|1\rangle$ dans la cavité n'est pas efficace à 100% mais seulement à 84% comme nous l'avons déjà vu. Pour voir l'effet de cette imperfection, comparons deux cas envisageables : dans l'un, l'état du système à la fin de la procédure de préparation est $|\psi_1\rangle = |g, 1\rangle$; dans l'autre cas le système se trouve dans l'état $|\psi_2\rangle = |e, 1\rangle$. Alors il est clair que, quel que soit le champ injecté par la suite $|\alpha\rangle$, les franges données par l'atome dans le deuxième cas seront en *opposition de phase* par rapport à celles données par le premier atome. Cela résulte tout simplement de l'unitarité de l'évolution de la fonction d'onde atomique jusqu'à la deuxième impulsion R2. Donc, pour un même état du champ ρ , si l'atome commence la mesure de la fonction de Wigner dans l'état $|g\rangle$, il mesure $W_\rho(\alpha)$, mais s'il est dans $|e\rangle$ il mesurera $-W_\rho(\alpha)$! Or nous ne sommes pas en mesure de discriminer les événements où l'atome était initialement dans $|g\rangle$ de ceux où l'atome était dans $|e\rangle$. Ainsi l'imperfection de l'impulsion π n'affecte pas seulement la préparation du champ dans le mode HF (ce qui ne serait pas un réel problème), comme il semblerait à première vue, mais aussi la procédure de mesure en elle-même.

Essayons de quantifier cet effet. L'imperfection de l'impulsion π dans la cavité a deux causes distinctes qui ont des effets nettement différents :

(i) d'une part, les inhomogénéités diverses des champs électrique, magnétique, et la dispersion géométrique des atomes dans le jet, font que tous les atomes ne font pas exactement le même angle de Rabi. Par conséquent, certains atomes, au lieu d'avoir l'évolution $|e, 0\rangle \longrightarrow |g, 1\rangle$, restent dans l'état $|e\rangle$: $|e, 0\rangle \longrightarrow |e, 0\rangle$. Ceux-là donneront une contribution $-W_0(\alpha)$ au résultat final.

(ii) d'autre part, le champ thermique résiduel dans la cavité fait que le champ dans le mode HF n'est pas exactement $|0\rangle$ mais plutôt un mélange statistique de $|0\rangle$ et de $|1\rangle$ avec des poids d'environ $p(0) = 0.9$ et $p(1) = 0.1$. Or lorsque le champ dans la cavité est de un photon, pour le temps d'interaction t_π , les atomes n'effectuent pas une impulsion π mais $\sqrt{2}\pi$ puisque la fréquence de Rabi varie en $\Omega\sqrt{n}$ si n est le nombre de photons dans la cavité. Ils finiront dans environ la moitié des cas dans l'état $|e, 1\rangle$, et sinon dans l'état $|g, 2\rangle$. Les atomes qui ont fini dans $|e\rangle$ donneront donc une contribution en $-W_1(\alpha)$ au signal total.

Résumons-nous. Sur 100 atomes, 84 finissent dans l'état $|g\rangle$ après l'impulsion π , ils nous donneront un signal de franges que nous interprétons comme la fonction de Wigner du champ dans la cavité. 16 sont restés dans $|e\rangle$. Sur ces 16 atomes, nous ignorons quelle proportion est due aux imperfections de type (1) ou (2). Il semble raisonnable de supposer que cette proportion est de 50%. D'après l'analyse que nous venons de conduire, ces 16 atomes donneront une contribution au signal du type $-0.5W_0(\alpha) - 0.5W_1(\alpha)$. Or ce signal est essentiellement plat (en 0 notamment il vaut 0) puisque les fonctions de Wigner des états $|0\rangle$ et $|1\rangle$ prennent des valeurs opposées dans les parties de l'espace des phases où leur module est important.

Cette analyse permet d'étayer l'hypothèse selon laquelle la seule conséquence sur

le signal de l'imperfection de l'impulsion π de préparation du champ est de réduire le contraste des franges par un facteur global qui est précisément le pourcentage d'imperfections 16%. Cela nous donne une contribution au facteur de normalisation $f_5 = 1/0.84 = 1.19$.

Au total, en tenant compte de cette longue analyse des imperfections de notre expérience, on obtient un facteur de normalisation $F = F_0 f_4 f_5 = 3.6$; or nous avons mesuré $F_1 = 4.0$. Il semble donc que nous comprenions de manière globalement satisfaisante les contrastes observés. Cette analyse quantitative des (nombreuses) imperfections de l'expérience rend légitime la procédure de normalisation globale adoptée pour obtenir la fonction de Wigner du champ à partir de la mesure de la courbe $C(\alpha)$. Le fait que les valeurs expérimentales coïncident avec la fonction de Wigner théorique d'un champ qui a la statistique prévue les conforte encore plus.

Bibliographie

- [1] ALLEN, L. et EBERLY, J. H., *Optical resonance and two level atoms* (Wiley, New York, 1975).
- [2] ASPECT, A., GRANGIER, P. et ROGER, G., “Experimental realization of Einstein-Podolsky-Rosen-Bohm gedankenexperiment : A new violation of Bell’s inequalities”, *Phys. Rev. Lett.*, 49, p. 91 (1982).
- [3] BADUREK, G., RAUSCH, H. et TUPPINGER, D., “Neutron interferometric double resonance experiments”, *Phys. Rev. A*, 34, p. 2600 (1985).
- [4] BANASZEK, K., RADZEWICZ, C., WÓDKIEWICZ, K. et KRASISKI, J. S., “Direct measurement of the wigner function by photon counting”, *Phys. Rev. A*, 60, p. 674 (1999).
- [5] BANASZEK, K. et WÓDKIEWICZ, K., “Direct probing of quantum phase space by photon counting”, *Phys. Rev. Lett.*, 76, p. 4344 (1996).
- [6] BERNARDOT, F., *Électrodynamique en cavité : expériences résonnantes en régime de couplage fort*, Thèse de doctorat, Université Paris 6 (1994).
- [7] BERTET, P., OSNAGHI, S., MILMAN, P., AUFFEVES, A., MAIOLI, P., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., “Generating and probing a two-photon Fock state with a single atom in a cavity”, *Phys. Rev. Lett.*, 88, p. 143601 (2002).
- [8] BERTET, P., OSNAGHI, S., RAUSCHENBEUTEL, A., NOGUES, G., AUFFEVES, A., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., “A complementarity experiment at the quantum-classical boundary”, *Nature*, 400, p. 239 (1999).
- [9] BERTRAND, J. et BERTRAND, P., “A tomographic approach to wigner’s function”, *Found. Physics*, 17, p. 397 (1987).
- [10] BETHE, H. A. et SALPETER, E. E., *Quantum mechanics of one- and two-electron atoms* (Plenum, New York, 1977).
- [11] BOHR, N., “On the constitution of atoms and molecules”, .
- [12] BOHR, N., “The quantum postulate and the recent development of atomic theory”, *Nature*, 121, p. 580 (1928).
- [13] BOHR, N., “Can quantum mechanical description of physical reality be considered complete?”, *Phys. Rev.*, 48, p. 696 (1935).

- [14] BOHR, N., "Discussion with Einstein on epistemological problems in atomic physics", dans P. A. Schilpp, (rédacteur) *Albert Einstein : Philosophers-Scientist*, p. 200 (The Library of Living Philosophers, 1949).
- [15] BREITENBACH, G., MÜLLER, T., PEREIRA, S. F., POIZAT, J.-P., SCHILLER, S. et MLYNEK, J., *J. Opt. Soc. B*, 12, p. 2304 (1995).
- [16] BRUNE, M., HAGLEY, E., DREYER, J., MAÎTRE, X., MAALI, A., WUNDERLICH, C., RAIMOND, J. M. et HAROCHE, S., "Observing the progressive decoherence of the meter in a quantum measurement", *Phys. Rev. Lett.*, 77, p. 4887 (1996).
- [17] BRUNE, M., HAROCHE, S., LEFÈVRE, V., RAIMOND, J.-M. et ZAGURY, N., "Quantum nondemolition measurement of small photon numbers by rydberg-atom phase-sensitive detection", *Phys. Rev. Lett.*, 65, p. 976 (1990).
- [18] BRUNE, M., HAROCHE, S., RAIMOND, J.-M., DAVIDOVICH, L. et ZAGURY, N., "Manipulation of photons in a cavity by dispersive atom-field coupling : Quantum-nondemolition measurements and generation of "Schrödinger cat" states", *Phys. Rev. A*, 45(7), p. 5193 (1992).
- [19] BRUNE, M., NUSSENZVEIG, P., SCHMIDT-KALLER, F., BERNARDOT, F., MAALI, A., RAIMOND, J.-M. et HAROCHE, S., "From Lamb shift to light shift : vacuum and subphoton cavity field mesured by atomic phase sensitive detection", *Phys. Rev. Lett.*, 72, p. 3339 (1994).
- [20] BRUNE, M., RAIMOND, J.-M., GOY, P., DAVIDOVICH, L. et HAROCHE, S., "Realization of a two-photon maser oscillator", *Phys. Rev. Lett.*, 59, p. 1899 (1987).
- [21] BRUNE, M., SCHMIDT-KALER, F., MAALI, A., DREYER, J., HAGLEY, E., RAIMOND, J. M. et HAROCHE, S., "Quantum Rabi oscillation : a direct test of field quantization in a cavity", *Phys. Rev. Lett.*, 76.
- [22] BUKS, E., SCHUSTER, R., HEIBLUM, M., MAHALU, D. et UMANSKY, V., "De-phasing in electron interference by a "which-path detector", *Nature*, 391, p. 871 (1998).
- [23] CAHILL, K. E. et GLAUBER, R. J., *Phys. Rev.*, 177, p. 1882 (1969).
- [24] CHAPMAN, M., HAMMOND, T., LENEF, A., SCHMIEDMAYER, J., RUBENSTEIN, R., SMITH, E. et PRITCHARD, D., "Photon scattering from atoms in an atom interferometer : Coherence lost and regained", *Phys. Rev. Lett.*, 75, p. 3783 (1995).
- [25] DAVIDOVICH, L., ZAGURY, N., BRUNE, M., RAIMOND, J. et HAROCHE, S., "Teleportation of an atomic state between two cavities using non-local microwave fields", .
- [26] DIVINCENZO, D. P., *Science*, 270, p. 255 (1995).
- [27] DREYER, J., *Atomes de Rydberg et cavités : observation de la décohérence dans une mesure quantique*, Thèse de doctorat, Université Paris 6 (1997).

- [28] DURR, S., NONN, T. et REMPE, G., "Fringe visibility and which-way information in an atom interferometer", *Phys. Rev. Lett.*, 81, p. 5705 (1998).
- [29] DURR, S., NONN, T. et REMPE, G., "Origin of quantum-mechanical complementarity probed by a "which-way" experiment in an atom interferometer", *Nature*, 395, p. 33 (1998).
- [30] EICHMANN, U., BERGQUIST, J., BOLLINGER, J., GILLIGAN, J., ITANO, W. et WINELAND, D., "Young's interference experiment with light scattered from two atoms", *Phys. Rev. Lett.*, 70, p. 2359 (1993).
- [31] EINSTEIN, A., "On a heuristic point of view about the creation and conversion of light", *Ann. Phys.*, 17, p. 132 (1905).
- [32] EINSTEIN, A., PODOLSKY, B. et ROSEN, N., "Can quantum mechanical description of physical reality be considered complete?", *Phys. Rev.*, 47, p. 777 (1935).
- [33] EKERT, A. et JOZSA, R., "Quantum computation and Shor's factoring algorithm", *Rev. Mod. Phys.*, 68, p. 733 (1996).
- [34] ENGLERT, B.-G., "Fringe visibility and which-way information : An inequality", *Phys. Rev. Lett.*, 77, p. 2154 (1996).
- [35] ENGLERT, B.-G., SCULLY, M. O. et WALTHER, H., "Complementarity and uncertainty", *Nature*, 375, p. 367 (1995).
- [36] ENGLERT, B.-G., STERPI, N. et WALTHER, H., "Parity states in the one-atom maser", *Opt. Commun.*, 100, p. 526 (1993).
- [37] FEYNMAN, R., LEIGHTON, R. et SANDS, M., *The Feynman Lectures on Physics*, tome 3 (Addison-Wesley, 1965).
- [38] GORDON, W., *Ann. Phys.*, 2, p. 1031 (1929).
- [39] GOY, P., *Int. J. Infrared Millimeter Waves*, 3, p. 221 (1982).
- [40] HAGLEY, E., MAÎTRE, X., NOGUES, G., WUNDERLICH, C., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., "Generation of Einstein-Podolsky-Rosen pairs of atoms", *Phys. Rev. Lett.*, 79, p. 1 (1997).
- [41] HEISENBERG, W., "Quantumtheoretical reinterpretation of kinematic and mechanical relations", .
- [42] HEISENBERG, W., "Über den anschaulichen Inhalt der quantentheoretischen Kinematik und Mechanik", *Z. Phys.*, 43, p. 172 (1927).
- [43] HERZOG, T., KWIAT, P., WEINFURTER, H. et ZEILINGER, A., "Complementarity and the quantum eraser", *Phys. Rev. Lett.*, 75, p. 3034 (1995).
- [44] HILLERY, M., O'CONNELL, R. F., SCULLY, M. O. et WIGNER, E. P., *Phys. Rep.*, 108, p. 121 (1984).
- [45] HULET, R. G. et KLEPPNER, D., "Rydberg atoms in circular states", *Phys. Rev. Lett.*, 51, p. 1430 (1983).

- [46] J. M. RAIMOND, M. B. . S. H., "Manipulating quantum entanglement with atoms and photons in a cavity", *Rev. Mod. Phys.*, 73, p. 565 (2001).
- [47] JAYNES, E. T. et CUMMINGS, F. W., "Comparison of quantum and semiclassical radiation theories with application to the beam maser", *Proc. IEEE*, p. 89 (1963).
- [48] KAMADA, H., GOTOH, H., TEMMYO, J., TAKAGAHARA, T. et ANDO, H., "Exciton rabi oscillation in a single quantum dot", *Phys. Rev. Lett.*, 87, p. 246401 (2001).
- [49] KIM, J. I., FONSECA ROMERO, K. M., HORIGUTI, A. M., DAVIDOVICH, L., NEMES, M. C. et DE TOLEDO PIZA, A. F. R., "Classical behaviours with small quantum numbers : The physics of ramsey interferometry of rydberg atoms", *Phys. Rev. Lett.*, 82, p. 4737 (1999).
- [50] KIM, M. S., ANTESBERGER, G., BODENDORF, C. T. et WALTHER, H., "Scheme for direct observation of the wigner characteristic function in cavity qed", *Phys. Rev. A*, 58, p. R65 (1998).
- [51] KOGELNIK, H. et LI, T., "Laser beams and resonators", *Proc. IEEE*, 54, p. 1312 (1966).
- [52] KURTSFEIER, C., PFAU, T. et MLYNEK, J., "Measurement of the wigner function of an ensemble of helium atoms", *Nature*, 386, p. 150 (1997).
- [53] KWIAT, P. G., MATTLE, K., WEINFURTER, H. et ZEILINGER, A., "New high-intensity source of polarization-entangled photon pairs", *Phys. Rev. Lett.*, 75, p. 4337 (1995).
- [54] KWIAT, P. G., STEINBERG, A. M. et CHIAO, R. Y., "Observation of a "quantum eraser" : A revival of coherence in a two-photon interference experiment", *Phys. Rev. A*, 45, p. 7729 (1992).
- [55] LALANNE, J. R., DUCASSE, A. et KIELICH, S., *Interaction laser molecule* (Polytechnica, 1994).
- [56] LEGGETT, A. J., CHAKRAVARTY, S., DORSEY, A. T., FISCHER, M. P. A., GIARG, A. et ZWERGER, W., "The dissipative two-state system", *Rev. Mod. Phys.*, 59, p. 1 (1987).
- [57] LEIBFRIED, D., MEEKHOF, D., KING, B., MONROE, C., ITANO, W. et WINELAND, D., "Experimental determination of the motional quantum state of a trapped atom", *Phys. Rev. Lett.*, 77(21), p. 4281 (1996).
- [58] LUTTERBACH, L. G. et DAVIDOVICH, L., "Method for direct measurement of the Wigner function in cavity QED and ion traps", *Phys. Rev. Lett.*, 78.
- [59] LVOVSKY, A. et MLYNEK, J., "Quantum optical catalysis : Generating non-classical states of light by means of linear optics", *prl*, 88, p. 250401 (2002).
- [60] LVOVSKY, A. I., HANSEN, H., AICHELE, T., BENSON, O., MLYNEK, J. et SCHILLER, S., "Quantum state reconstruction of the single-photon fock state", *Phys. Rev. Lett.*, 87, p. 050402 (2001).

- [61] MAALI, A., *Oscillations de Rabi quantiques : test direct de la quantification du champ*, Thèse de doctorat, Université Paris 6 (1996).
- [62] MARTINIS, J. M., NAM, S., AUMENTADO, J. et URBINA, C., "Rabi oscillations in a large josephson-junction qubit", *Phys. Rev. Lett.*, 89, p. 117901 (2002).
- [63] MAÎTRE, X., *Une paire d'atomes intriqués : Expériences d'électrodynamique quantique en cavité*, Thèse de doctorat, Université Paris 6 (1998).
- [64] MEEKHOF, D. M., MONROE, C., KING, B. E., ITANO, W. M. et WINELAND, D. J., "Generation of nonclassical motional states of a trapped atom", *Phys. Rev. Lett.*, 76, p. 1796 (1996).
- [65] MESCHEDÉ, D., WALTHER, H. et MULLER, G., "One-atom maser", *Phys. Rev. Lett.*, 54.
- [66] MONROE, C., MEEKHOF, D. M., KING, B. E. et WINELAND, D. J., "A "Schrödinger cat" superposition state of an atom", *Science*, 272, p. 1131 (1996).
- [67] NAKAMURA, Y., PASHKIN, Y. A. et TAI, J. S., "Coherent control of macroscopic quantum states in a single-cooper-pair box", *Nature*, 398, p. 786 (1999).
- [68] NOGUES, G., *Détection sans destruction d'un seul photon. Une expérience d'électrodynamique quantique en cavité*, Thèse de doctorat, Université Paris 6 (1999).
- [69] NOGUES, G., RAUSCHENBEUTEL, A., OSNAGHI, S., BERTET, P., BRUNE, M., RAIMOND, J. M., HAROCHE, S., LUTTERBACH, L. G. et DAVIDOVICH, L., "Measurement of a negative value for the Wigner function of radiation", *Phys. Rev. A*, 62, p. 054101 (2000).
- [70] NOGUES, G., RAUSCHENBEUTEL, A., OSNAGHI, S., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., "Seeing a single photon without destroying it", *Nature*, 400, p. 239 (1999).
- [71] NUSSENZVEIG, P., *Mesures de champs au niveau du photon par interférométrie atomique*, Thèse de doctorat, Université Paris 6 (1994).
- [72] NUSSENZVEIG, P., BERNARDOT, F., BRUNE, M., HARE, J., RAIMOND, J. M., HAROCHE, S. et GAWLIK, W., "Preparation of high principal quantum numbers circular states of rubidium", *Phys. Rev. A*, 48, p. 3991 (1993).
- [73] OSNAGHI, S., *Réalisation d'états intriqués dans une collision atomique assistée par une cavité*, Thèse de doctorat, Université Paris 6 (2001).
- [74] OSNAGHI, S., BERTET, P., AUFFEVE, A., BERTET, P., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., "Coherent control of an atomic collision in a cavity", *Phys. Rev. Lett.*, 87, p. 037902 (2001).
- [75] PFAU, T., SPÄLTER, S., KURTSIEFER, C., EKSTROM, C. R. et MLYNEK, J., "Loss of spatial coherence by a single spontaneous emission", *Phys. Rev. Lett.*, 73, p. 1223 (1994).
- [76] PLANCK, M., *Verh. Deutsch. Phys. Ges.*, 2, p. 202 (1900).

- [77] RAMSEY, N. F., *Molecular beams*, International series of monographs on physics (Oxford (1985), 1956).
- [78] RAUSCHENBEUTEL, A., *Atomes et cavité : préparation et manipulation d'états intriqués complexes*, Thèse de doctorat, Université Paris 6 (2001).
- [79] RAUSCHENBEUTEL, A., NOGUES, G., OSNAGHI, S., BERTET, P., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., "Coherent operation of a tunable quantum phase gate in cavity QED", *Phys. Rev. Lett.*, 83, p. 5166 (1999).
- [80] RAUSCHENBEUTEL, A., NOGUES, G., OSNAGHI, S., BERTET, P., BRUNE, M., RAIMOND, J. M. et HAROCHE, S., "Step-by-step engineered multiparticle entanglement", *Science*, 288, p. 2024 (2000).
- [81] SCHILLER et MLYNEK, "Quantum statistics of the squeezed vacuum by measurement of the density matrix in the number state representation", *Phys. Rev. Lett.*, 77, p. 2933 (1996).
- [82] SCHRÖDINGER, E., "Die gegenwärtige situation in der quantenmechanik", *Naturwissenschaften*, 23, pp. 807, 823, 844 (1935), traduction anglaise dans [97].
- [83] SCULLY, M.-O., ENGLERT, B.-G. et WALTHER, H., "Quantum optical tests of complementarity", *Nature*, 351, p. 111 (mai 1991).
- [84] SCULLY, M. O. et ZUBAIRY, M. S., *Quantum Optics* (Cambridge University Press, 1997).
- [85] SMITHEY, D. T., BECK, M., RAYMER, M. G. et FARIDANI, A., *Phys. Rev. Lett.*, 70, p. 1244 (1993).
- [86] STIEVATER, T. H., LI, X., STEEL, D. G., GAMMON, D., KATZER, D. S., PARK, D., PIERMAROCCHI, C. et SHAM, L. J., "Rabi oscillations of excitons in single quantum dots", *Phys. Rev. Lett.*, 87, p. 133603 (2001).
- [87] STOREY, P., TAN, S., COLLET, M. et WALLS, D., "Complementarity and uncertainty", *Nature*, 375, p. 368 (1995).
- [88] THOMPSON, R. J., REMPE, G. et KIMBLE, H. J., "Observation of normal-mode splitting for an atom in an optical cavity", *Phys. Rev. Lett.*, 68, p. 1132 (1992).
- [89] TURCHETTE, Q. A., HOOD, C. J., LANGE, W., MABUCHI, H. et KIMBLE, H. J., "Measurement of conditional phase shifts for quantum logic", *Phys. Rev. Lett.*, 75, p. 4710 (1995).
- [90] UDEM, T., REICHERT, J., HOLZWARTH, R. et HÄNSCH, T. W., "Absolute optical frequency measurement of the cesium d1 line with a mode-locked laser", *Phys. Rev. Lett.*, 82, p. 3568 (1999).
- [91] VION, D., AASSIME, A., COTTET, A., JOYEZ, P., POTHIER, H., URBINA, C., ESTEVE, D. et DEVORET, M. H., "Manipulating the quantum state of an electrical circuit", 296, p. 886 (2002).

- [92] VOGEL, K. et RISKEN, H., "Determination of quasi-probability distributions in terms of probability distributions for the rotated quadrature phase", *Phys. Rev. A*, 40, p. 2847 (1989).
- [93] VAN DER WAL, C. H., TER HAAR, A. C. J., WILHELM, F. K., SCHOUTEN, R. N., HARMANS, C. J., ORLANDO, T. P., LLOYD, S. et MOOIJ, J. E., "Quantum superposition of macroscopic persistent-current states", 290, p. 773 (2000).
- [94] WALLENTOWITZ, S. et VOGEL, W., "Unbalanced homodyning for quantum state measurements", *Phys. Rev. A*, 53, p. 4528 (1996).
- [95] WEYL, H., *Z. Phys*, 46, p. 1 (1927).
- [96] WHEELER, J. A., *Mathematical foundations of quantum theory* (A. R. Marlow, Academic, 1978).
- [97] WHEELER, J. A. et ZUREK, W. H., *Theory of measurement* (Princeton University Press, 1983).
- [98] WIGNER, E. P., *Phys. Rev.*, 40, p. 749 (1932).
- [99] ZIMMERMAN, M., LITTMAN, M., KASH, M. et KLEPPNER, D., "Stark structure of the Rydberg states of alkali-metal atoms", *Phys. Rev. A*, 20, p. 2251 (1979).
- [100] ZUREK, W. H., "Pointer basis of quantum apparatus : Into what mixture does the wave packet collapse?", *Phys. Rev. D*, 24, p. 1516 (1981).
- [101] ZUREK, W. H., "Environment-induced superselection rules", *Phys. Rev. D*, 26, p. 1862 (1982).