



Tomographie par trajectoires d'états délocalisés du champ micro-onde de deux cavités

Valentin Métillon

► To cite this version:

Valentin Métillon. Tomographie par trajectoires d'états délocalisés du champ micro-onde de deux cavités. Physique Quantique [quant-ph]. Université Paris sciences et lettres, 2019. Français. NNT : 2019PSLEE051 . tel-02467898v2

HAL Id: tel-02467898

<https://hal.science/tel-02467898v2>

Submitted on 27 Nov 2020

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.



THÈSE DE DOCTORAT
DE L'UNIVERSITÉ PSL

Préparée à l'École normale supérieure,
au sein du laboratoire Kastler-Brossel

Tomographie par trajectoires d'états délocalisés du champ micro-onde de deux cavités

Soutenue par

Valentin MÉTILLON

Le 12 novembre 2019

Ecole doctorale n° 564

Physique en Île-de-France

Spécialité

Physique quantique

Composition du jury :

Agnès, MAÎTRE Professeure, Sorbonne Université	<i>Présidente</i>
Olivier, BUISSON Directeur de recherche, Institut Néel	<i>Rapporteur</i>
Christian, ROOS Senior scientist, IQOQI Innsbruck	<i>Rapporteur</i>
Klaus, MØLMER Professeur, Université d'Aarhus	<i>Examineur</i>
Pierre, ROUCHON Professeur, Mines Paristech	<i>Membre invité</i>
Igor, DOTSENKO Maître de conférences, Collège de France	<i>Membre Invité</i>
Michel, BRUNE Directeur de recherche, École normale supérieure	<i>Directeur de thèse</i>

Table des matières

I	Atomes et photons en cavité	3
I.1	Description quantique d'un mode du champ électromagnétique	3
I.1.1	Quantification du champ	3
I.1.2	Couplage du mode à un courant classique	4
I.1.3	La fonction de Wigner	5
I.1.4	État du champ à deux modes fonction de Wigner	6
I.1.5	Décohérence du champ	9
I.2	L'atome à deux niveaux	11
I.2.1	Le pseudo-spin atomique	11
I.2.2	Évolution temporelle	11
I.2.3	Interaction de l'atome à deux niveaux avec un champ classique . .	12
I.3	Interaction entre atomes et lumière	15
I.3.1	Modèle de Jaynes et Cummings et théorie de l'atome habillé . . .	15
I.3.2	Interaction résonnante	18
I.3.3	Interaction dispersive	19
I.3.4	Le passage adiabatique	20
I.3.5	Pertinence du modèle d'un atome couplé à un mode du champ . .	21
I.4	Le dispositif expérimental	23
I.4.1	Des cavités de grande finesse	24
I.4.2	Atomes de Rydberg	30
	a) Propriétés des atomes de Rydberg circulaires	30
	b) Préparation sélective d'échantillons atomiques	33
	c) Détection des atomes	36
I.5	Techniques d'électrodynamique quantique	39
I.5.1	Calibration du champ électrique et contrôle de la fréquence atomique	39
I.5.2	Initialisation du mode dans l'état vide	39
I.5.3	Oscillations de Rabi et préparation d'états de Fock à un photon . .	40
	a) Mesure de la fréquence de Rabi du vide	41
	b) Oscillations de Rabi à proximité de la résonance : effet de l'habillage initial de l'atome	45
I.5.4	Interférométrie de Ramsey et mesure de la fonction de Wigner . .	50
II	Préparation et mesure d'états intriqués en électrodynamique quantique à deux cavités	63
II.1	Préparation d'un état délocalisé d'un photon	63
II.1.1	La séquence d'impulsions	64
II.1.2	Les états <i>noon</i>	64
II.1.3	Propriétés de l'état $ 10\rangle + 01\rangle$	65
II.1.4	Préparation d'un état à deux photons	71
II.2	La mesure résonnante	72
II.2.1	Séquence de mesure	73

II.2.2	Optimalité de la mesure résonnante : états propres et information de Fisher	74
II.3	Résultats	77
II.3.1	Un interféromètre de 6000 km	77
II.3.2	Contraste de l'interféromètre	78
a)	Erreurs de détection	78
b)	Erreurs de préparation	78
c)	Contraste de Rabi	79
d)	Échantillons à deux atomes	80
e)	Bilan	80
II.3.3	Mesure de coupes de la fonction de Wigner de l'état $ 10\rangle + 01\rangle$. .	81
III	Principes de la tomographie d'états quantiques	89
III.1	Quelques rappels de théorie de l'estimation	90
III.1.1	Cadre de l'estimation d'un paramètre réel	90
III.1.2	Efficacité de l'estimateur	91
a)	Information de Fisher et borne de Cramer-Rao	91
b)	Estimation par maximum de vraisemblance	92
III.2	Principe de la tomographie quantique et des méthodes de reconstruction .	94
III.2.1	Rappels sur l'opérateur densité	94
a)	Définitions et propriétés	94
b)	Ensemble $\mathcal{D}(\mathcal{H})$ des matrices densité	95
c)	Paramétrisation de $\mathcal{D}(\mathcal{H})$ pour un qubit	96
III.2.2	Méthodes usuelles de reconstruction	98
III.2.3	Limitations des méthodes usuelles	103
III.3	Tomographie par trajectoires	104
III.3.1	Trajectoires quantiques, opérateurs de Kraus et matrice d'effet . .	104
a)	Mesure généralisée	104
b)	Description de mesures imparfaites	106
c)	Cartes quantiques et somme de Kraus	106
d)	Interaction d'un système avec un environnement	108
e)	Équation pilote de Lindblad	109
f)	Matrices d'effet	111
g)	Conclusion de la section	114
III.3.2	Tomographie par trajectoires : principe et mise en application . .	114
a)	Principe de la méthode et construction des matrices d'effet	114
b)	Le choix des mesures et leur nombre	116
c)	Calcul du gradient et détermination de l'estimateur . . .	116
d)	Algorithme de montée de gradient	118
e)	Interprétation de la méthode : cas du qubit	120
III.3.3	Prise en compte des imperfections et pertinence des résultats . .	121
a)	Les résultats sont-ils physiques ?	121
b)	Le cas des états purs ou presque	122
III.4	Incertitudes de la tomographie	123
III.4.1	Surfaces d'isovraisemblance et volumes de confiance	123
III.4.2	Information de Fisher et calcul des barres d'erreur	128

III.4.3 Incertitudes sur la reconstruction d'un qubit	131
IV Reconstruction d'états intriqués	135
IV.1 Matrices d'effet théoriques pour nos mesures	135
IV.1.1 La mesure résonnante	136
a) Mesure résonnante à un atome	136
b) Mesure résonnante à plusieurs atomes	138
IV.1.2 La mesure dispersive	141
a) Mesure de la parité	141
b) Mesure de plusieurs atomes successifs	145
c) Mesure de la fonction de Wigner	147
d) Effet de la relaxation sur les matrices d'effet	150
IV.1.3 Prise en compte des imperfections de la détection	152
IV.2 Reconstruction de l'état $ 10\rangle + 01\rangle$	153
IV.2.1 Mesure résonnante seule	153
IV.2.2 Mesures locales de la fonction de Wigner	162
IV.2.3 Mesure complète	165
IV.2.4 Reconstruction avec plusieurs atomes résonnants	169
IV.2.5 Reconstruction de l'état $ 20\rangle + 02\rangle$	174
IV.3 Reconstruction partielle	177
IV.3.1 Motivation	177
IV.3.2 Préparation d'états de phases différentes et résultats obtenus.	177
a) Protocole et calibration de la phase	177
b) Incertitudes sur la phase	179
A Quelques ingrédients théoriques sur l'information de Fisher	186
B Construction de l'intervalle de confiance de Clopper-Pearson	189
C Purification d'états intriqués	191
D Interaction de deux atomes résonnants avec un mode du champ	194
E Calcul de fonctions de Wigner à deux modes	198
Bibliographie	201

À Louis Métillon, Edmée Billeron et Jean-Claude Hansen

Remerciements

Les travaux présentés dans ce manuscrit sont le fruit de presque cinq années passées au sein du groupe d'électrodynamique quantique en cavité. Les expériences de l'équipe sont d'une grande complexité et reposent sur un long travail d'équipe pour aller de la conception des dispositifs à leur réalisation, puis leur exploitation afin d'étudier les aspects les plus fondamentaux de la physique quantique. De nombreuses personnes sont impliquées à chaque étape et aucun des résultats présentés ici n'aurait pu être atteint sans elles. Je souhaite remercier l'ensemble des physiciens et physiciennes de l'équipe qui ont travaillé sur cette expérience et permis de mener plus de dix années de recherche intense couronnée par une multitude de résultats fascinants. J'ai pu bénéficier de l'expérience de celles et ceux qui m'ont précédé grâce aux nombreuses thèses de doctorat écrites avant mon arrivée, détaillant tout le savoir-faire expérimental nécessaire pour piéger des photons micro-onde, suivre leur évolution et les contrôler avec des atomes de Rydberg circulaires.

Comme une thèse n'est jamais qu'un (bref) résumé du travail entrepris et du dispositif mis en place, il faut voir l'expérience en vrai pour se rendre compte de sa complexité. La salle E.17 du Collège de France regorge de matériel du sol au plafond. Au centre se trouve le cryostat, qui, lorsqu'on a de la chance et que tout est en place, contient la boîte magique dans laquelle toute la physique étudiée a lieu. Tout autour s'entassent des dizaines d'appareils contrôlant chaque aspect du dispositif et permettant d'effectuer des séquences expérimentales complexes. Afin de s'y retrouver dans les centaines de câbles coaxiaux, de guides d'ondes et d'éléments optiques, ainsi que dans les multiples ordinateurs et les milliers de lignes de code nécessaires pour effectuer la moindre mesure sur l'expérience, il faut apprendre à connaître intimement la manip. Je n'aurais jamais pu en arriver à ce stade sans l'aide d'Igor. Sa connaissance des différents rouages du dispositif et de l'histoire des travaux qu'il a permis est essentielle au fonctionnement de notre équipe. Sa curiosité pour tous les aspects de la physique expérimentale et son incroyable capacité à expliquer ce qu'on a besoin de comprendre à l'instant t ont joué un rôle essentiel dans ma formation de physicien. Je lui suis profondément reconnaissant de m'avoir transmis une manière de pratiquer la physique ouverte autant sur les aspects pratiques et les moyens les plus efficaces d'atteindre les objectifs qu'on s'est fixés, que sur la prise de recul sur les résultats obtenus et la réflexion conceptuelle.

Je suis aussi grandement redevable aux autres personnes qui m'ont accueilli en E.17. J'ai pris beaucoup de plaisir à travailler avec Stefan, avec qui nous avons passé des heures affublés de charlottes, à monter et démonter les entrailles du cryostat que nous connaissons sur le bout des doigts, sans avoir besoin d'échanger un mot tant nous nous comprenions et partagions les mêmes exigences de rigueur. J'ai aussi énormément appris de Mariane, qui a toujours pris à cœur de transmettre sa connaissance de tous les éléments du dispositif qu'elle a fabriqués elle-même, avec une grande pédagogie et une clarté irréprochable. Ne pas parvenir à mettre les cavités à résonance tant que vous étiez encore là a été un véritable déchirement au début de ma thèse. Si je n'ai pas regretté les multiples refroidissements que nous avons effectué tous ensemble, je me souviens des pizzas partagées le soir lorsque nous mettions plus de quinze heures à atteindre 4 degrés au-dessus du zéro absolu (on s'est un peu amélioré après votre départ et on n'a plus mangé de pizzas...)

Le dispositif à deux cavités que nous avons installé est l'ultime aboutissement d'une recherche à très long cours, initiée dès les travaux de thèse de Serge et poursuivie avec

Jean-Michel et Michel pendant plusieurs décennies. Je les remercie sincèrement d'avoir mené à bien des expériences aussi riches et de m'avoir fait confiance pour y participer. La façon collégiale de travailler dans l'équipe m'a permis de prendre confiance en moi et d'oser partir dans des directions bien éloignées de celles envisagées initialement pour ma thèse. Écrire des articles et des projets et travailler en compagnie de physiciens aussi expérimentés a aussi été extrêmement formateur pour moi.

L'objectif initial de ma thèse n'avait rien à voir avec la tomographie d'états quantiques. Cette question est arrivée au cours du temps par une collaboration très fructueuse avec Pierre Rouchon qui nous a aidés à appliquer une méthode originale pour reconstruire l'état quantique que nous avions préparé. Grâce à son aide, nous avons pu soulever de nombreuses questions conceptuelles et j'ai beaucoup apprécié de travailler à comprendre les subtilités mathématiques cachées derrière ces problèmes. Merci pour les nombreuses discussions sur des sujets assez éloignés des compétences habituelles de l'équipe!

Je suis aussi redevable aux membres du jury d'avoir accepté de lire dans le détail un manuscrit de 200 pages en français, truffé de mathématiques et d'avoir fait preuve d'autant de curiosité concernant tous les aspects de mon travail lors de la soutenance de thèse.

Je souhaite aussi remercier tous les stagiaires qui sont passés en E.17. J'ai pris beaucoup de plaisir à essayer de vous transmettre un peu de cette expérience sur laquelle je serai vraisemblablement le dernier thésard. Merci à Guillaume pour son aide pour monter les miroirs! Merci à Rémi pour son travail sur la reconstruction d'états quantique et pour ses méthodes de calcul astucieuses! Merci à Gautier pour tout le code qu'il a écrit et pour son enthousiasme inébranlable! Merci à Luis pour sa bonne humeur, son travail acharné et sa présence incongrue mais réconfortante à des horaires tardifs pendant que je rédigeais ma thèse! Je vous souhaite beaucoup de réussite dans les projets de recherche dans lesquels vous vous êtes lancés.

Les travaux que j'ai entrepris n'auraient pas pu être menés à bien sans les personnels techniques et administratifs du LKB et de l'institut de physique du Collège de France. Je remercie particulièrement Carmen qui a toujours traité mes demandes dans la journée. Je suis aussi profondément redevable à Pascal pour ses nombreux conseils lorsque nous tentions d'accorder les cavités et pour toutes les pièces qu'il a fabriquées pour nous. Merci d'avoir fait toutes ces cales en cuivre! Enfin, aucune expérience du groupe ne serait possible sans le travail d'Olivier et Florent grâce à qui nous avons toujours de l'hélium pour effectuer un transfert tous les deux jours, quelles que soient la période de l'année et les conditions. Depuis trois ans que l'expérience est à froid, nous n'avons jamais été en manque d'hélium et je les en remercie chaleureusement.

Au cours de mes années passées au Collège de France, j'ai eu la chance d'avoir des collègues exceptionnels qui ont enrichi mon expérience au laboratoire. Un grand merci en particulier à ceux qui ont partagé la plus grande partie de mon séjour. Merci à Clément et Sébastien d'avoir toujours répondu à mes questions et de m'avoir fait profiter de leurs connaissances sur mon expérience. Rodrigo, j'ai pris un immense plaisir à discuter de physique avec toi pendant toutes ces années, même si dans le fond je me fous bien du collapse de la fonction d'onde $\Rightarrow p$. Merci aussi pour ton incroyable énergie qui a tiré l'équipe dans les moments où aucune expérience ne tournait. Frédéric, heureusement que t'étais là pour que je me sente pas seul avec mes idées. Merci pour tous les mauvais coups dans lesquels on s'est retrouvés avec tes coloc! Arthur, t'as toujours été là pour discuter

ou faire une pause quand on se faisait des journées de dingue, surtout pendant la rédaction. Merci aussi pour tous les bons moments en dehors du labo. Bon courage pour finir, c'est bientôt ton tour ! Brice, merci pour ton arrivée au labo et tes histoires atypiques de krav-maga et de ruthénium qui nous ont sortis de la période de morosité pendant laquelle nous étions seulement quatre thésards au labo. Merci à tou.te.s les autres qui m'ont accueilli quand je suis arrivé ou qui ont égayé ma fin de thèse et bon courage aux suivant.e.s ! En plus de mon travail au laboratoire, j'ai eu la chance d'effectuer des missions doctorales qui m'ont permis de me changer les idées et qui m'ont énormément appris. Je remercie toute l'équipe du Palais de la Découverte de m'avoir accueilli dans ce lieu formidable. Je suis en particulier redevable à Romain de m'avoir laissé faire des exposés de physique quantique devant des enfants. Je remercie aussi toute l'équipe de la préparation à l'agrégation de Montrouge, où j'ai pris beaucoup de plaisir à enseigner. Je suis aussi très reconnaissant envers les membres de physique sans frontière et tout particulièrement François, pour leur engagement et pour m'avoir fait découvrir d'autres manières de pratiquer la physique. Lorsqu'on travaille sur une grosse expérience de physique, on se retrouve rapidement à avoir le moral intriqué avec l'état de fonctionnement du dispositif. Un appareil de cassé ou une rupture du vide sont des événements douloureux capable de gâcher l'existence pendant plusieurs mois. Je remercie les autres thésard.e.s en physique du laboratoire et d'ailleurs avec qui j'ai pu partager, en plus de discussions extrêmement pointues, le désarroi de ces moments où rien ne marche. Merci à mes collègues du LKB, notamment Hanna, les deux Adriens d'optique quantique avec qui j'ai beaucoup parlé d'intrication, Bertrand, Manel et les deux Raphaëls de l'équipe Dalibard et enfin les camarades du CGL, en particulier Tam, Léo et Thibault, pour avoir mis du rouge dans le noir des périodes où la manip faisait des caprices ! Merci aussi à mes camarades du master CFP avec qui nous avons suivi la même progression (à des rythmes différents certes, puisque je suis un des derniers à soutenir). Un grand merci en particulier à Mathieu d'avoir été mon colocataire au début. Notre visite à l'institut Néel et l'incroyable voyage en stop qui l'a accompagnée restent gravés dans ma mémoire.

Avant même de faire une thèse, ma scolarité a été semée d'embûches et je ne serais jamais parvenu aussi loin sans le rôle fondamental que mes parents et mes grands-parents ont joué dans mon éducation. Étudier a toujours été très important dans la famille et c'est grâce à la curiosité et au goût de l'apprentissage qui m'ont été transmis que j'ai pu garder la tête hors de l'eau dans les moments les plus difficiles. Merci à mes parents et ma grand-mère de m'avoir fait confiance lorsque je suis parti il y a plus de dix ans pour étudier des choses qui leur échappaient complètement ! Merci de m'avoir soutenu tout au long de cette thèse lorsque ma confiance en moi était au plus bas ! Merci à Annie et Georges de s'être toujours intéressés à ce que je faisais et de m'avoir continûment encouragé. Je sais la fierté qu'auraient aussi ressentie mes regrettés grands-parents en apprenant l'arrivée d'un docteur dans la famille. Je leur dédie ce travail.

Merci aussi à mes frères d'avoir été là pour tous les moments importants de ma vie (même si Schnikou est arrivé longtemps après =p) et d'avoir fait des super gâteaux pour mon pot de thèse ! J'ai quitté la maison un peu tôt pour aller chez les têtes d'ampoules. J'espère que vous comprenez pourquoi désormais. Maintenant que j'ai fini, il faudrait qu'on se refasse un voyage ensemble. Désolé Victor de t'avoir doublé sur les études, mais vu le temps que ça m'a pris déjà comme ça, heureusement que j'ai accéléré un peu au début. Grisse, t'es le meilleur, bon courage pour ta thèse ! Schnik, on t'aime bien quand

même, j'espère que tu trouveras bientôt quelque chose qui te fait vraiment vibrer. Merci à Bab' et Nick' d'être mes plus grands fans ! Merci à tout le reste de ma famille pour les tous les encouragements auxquels j'ai eu droit dans les moments importants !

Ces cinq dernières années ont été longues et ponctuées de phases douloureuses. Je ne serais sans doute pas parvenu jusqu'au bout si je n'avais pas bénéficié d'un soutien large et permanent de tou.te.s mes ami.e.s. Je vous remercie de m'avoir soutenu et d'avoir cru en moi, même quand je perdais espoir. Je suis particulièrement redevable à toute la coloc' d'Ivry, pour les moments que nous avons passés ensemble, les activités organisées, la stimulation intellectuelle et politique permanente et le réconfort que vous m'avez apporté. Marie, tu m'impressionnes par ta capacité à te poser toujours des questions et à ne rien accepter pour acquis. Continue à être fidèle à ta conscience et bon courage pour dynamiter l'ordre établi ! Antoine, tu es la personne la plus curieuse et la plus éclectique que je connaisse. Depuis le temps qu'on traine ensemble, je crois que c'est avec toi que j'ai partagé le plus de centres d'intérêts et de délires abscons, des hécatonchires aux théodolites ! Marine, je suis heureux de connaître une personne aussi attentive et enthousiaste que toi. Il faut vraiment qu'on se fasse un atelier couture (à Sepmes en janvier ?). Encore merci d'avoir relu cette thèse ! Amiel, depuis que je te connais tu n'as jamais cessé de me surprendre. J'espère qu'on aura plein d'autres occasions de faire les 400 coups ensemble. Romain, merci pour les discussions animées et pour la couleur ! Maguelone, merci d'avoir apporté de la diversité dans notre coloc' de scientifiques, c'était nécessaire ! Guillaume, merci pour ton enthousiasme et les discussions sur la façon de concilier la physique et ce qu'on veut faire de nos vies, continue d'être toi-même ! Pablo, merci pour les discussions, la politique, les gueuletons, les coups à boire et ta bonne humeur ; t'es gavé cool, change pas ! Barbara, ton entrain et ton rire ont suffi à me changer les idées dans les pires moments de doutes de la fin. Merci pour les retraites à Sepmes loin du cryostat, les spectacles de marionnettes, les fous rires inopportuns et merci de m'avoir montré qu'il y a une vie après la thèse ! Merci à Porval de me tenir chaud la nuit et à Calmar d'être notre étendard !

Merci aussi à tou.te.s les ami.e.s que je me suis faits à Paris et qui m'ont soutenu et aidé à me changer les idées quand j'étais trop absorbé par le travail. Cyrus, merci pour les discussions passionnées et pour les chantiers qui ont été mes seules vacances sur la fin de la thèse. On va la finir cette yourte et on trouvera des réponses aux questions qu'on se pose ! Mendes, merci pour ta gentillesse, ta motivation et tes questionnements qui font écho aux miens. Sam et Ange, je vous remercie tous les deux pour les séminaires HPS qui m'ont passionné, pour votre sympathie et l'ouverture d'esprit que vous m'avez apportée. Pierre, j'ai été très heureux de te retrouver au LKB quand j'y suis arrivé, après toutes les aventures qu'on avait déjà vécues ensemble, et de pouvoir discuter de physique avec toi à de multiples occasions. Merci aussi à Shirin ! Votre gentillesse et votre hospitalité me touchent à chaque fois que je passe vous voir à Rennes ! Merci à la coloc' de Bobigny pour les projections concurrentes et pour les balades matinales ! Merci à Juliette pour les discussions interminables et les consultations médicales gratuites (comme prévu, ma santé va mieux depuis la soutenance) ! Alex, ta sincérité et ta bienveillance me ravissent à chaque fois qu'on se voit. Je te souhaite de trouver un boulot qui te convienne ! Merci à Peni et Célian pour les cinés, des toits de la rue d'Ulm aux caves de la rue Champollion. Merci à Sévan pour toutes les blagues depuis qu'on se connaît (heureusement qu'elles se sont un peu améliorées depuis). Merci à Alexandre G. pour ses questions. Et enfin, un immense merci à toute la bande de gros.ses tas.ses pour tous les bons moments depuis

qu'on se connaît et le soutien par IRC pendant ces dernières années. Ce serait trop long de mentionner tout le monde et trop injuste de ne nommer qu'une partie, alors je vous fusionne en une seule entité =).

Merci aux potes de Rezé qui restent fidèles au poste après des années d'éloignement. Merci à Sylvain de se souvenir de tout ce qu'on a fait ensemble alors que j'ai tout oublié. Les vacances avec François et toi ont été un des meilleurs moments d'aventure de ces dernières années. Coralie, tu es une des personnes les plus courageuses que je connaisse et je chéris tout particulièrement les moments où je t'ai suivie sans réfléchir. Te rejoindre à Grande-Synthe pour aider des réfugiés a été l'expérience la plus formatrice de ma thèse. Maheva, depuis le temps qu'on se connaît, je continue de ne comprendre que la moitié de ce que tu racontes xd ! Merci de m'avoir recueilli quand je broyais du noir, ça m'a fait du bien que vous veniez vous installer aussi près ! Nico et Élo, merci pour votre écoute et pour les moments hors du temps qu'on a passés au ... RÊVE DE L'ABORIGÈNE!!! Baptou, merci de m'avoir redonné la pêche avec ton incroyable bonne humeur et tes rêves ! T'es trop cool, change rien ! Merci à tout le reste de la bande, Tim, Valou, Zouille, Marie, Lorine, François G. et j'en oublie. J'espère qu'on aura plus de temps pour se voir quand je reviendrai à Nantes !

Introduction

La physique quantique prédit l'existence d'états dont les propriétés vont à l'encontre de notre intuition de la vie de tous les jours. Une particule quantique peut se trouver dans une superposition de deux états et être, par exemple, à deux endroits en même temps. Si on effectue une mesure pour déterminer dans laquelle des deux positions elle se trouve, on obtiendra toujours l'une ou l'autre des réponses, avec des probabilités égales. Si on mesure en revanche l'observable conjuguée de la position, c'est-à-dire l'impulsion, on observe des interférences, dues aux phases relatives des deux paquets d'onde correspondant à chaque position. La présence d'oscillations est un indicateur de la cohérence quantique entre les deux membres de la superposition d'états. Cette cohérence est très sensible. Si un système quelconque, et en particulier l'environnement de la particule, enregistre une information sur la position de celle-ci, l'interférence est brouillée. Ce mécanisme, connu sous le nom de décohérence, fournit une interprétation de l'absence de phénomènes d'interférence quantique visibles à notre échelle. Pour mettre en évidence la cohérence quantique, il est donc nécessaire de contrôler très précisément l'interaction du système d'intérêt avec son environnement.

Parmi les états superposés, les états intriqués sont particulièrement intéressants [1]. Ils n'ont aucun équivalent classique. Un système constitué de deux sous-systèmes intriqués donne, lorsqu'on effectue des mesures locales sur chaque sous-système, des résultats aléatoires. Ce n'est qu'en observant les corrélations entre les mesures effectuées, dans plusieurs bases différentes, sur chaque partie, que l'on peut mettre en évidence l'intrication.

Le développement des techniques expérimentales de manipulation de systèmes quantiques isolés a ouvert la voie à l'utilisation d'états superposés et intriqués pour encoder et contrôler des états logiques. Ce champ disciplinaire est aujourd'hui appelé *information quantique* [2] et renouvelle la perspective sur les aspects les plus fondamentaux de la physique quantique. Il propose en effet d'utiliser les phénomènes d'interférences quantiques et d'intrication pour réaliser des opérations inaccessibles à des systèmes classiques de traitement de l'information, avec de nombreuses applications potentielles dans les domaines du calcul algorithmique, de la cryptographie, de la simulation de lois physiques ou de la métrologie.

Les différents protocoles d'information quantique reposent généralement sur la capacité à préparer avec une grande fidélité des états quantiques ressources, souvent des états intriqués, et à leur appliquer des portes quantiques bien contrôlées. Les premières expériences de traitement d'information quantique réalisées se contentaient généralement de systèmes simples, constitués d'un ou deux bits quantiques ou *qubits*, qui sont des systèmes à deux états généralisant la notion de *binary digit* des ordinateurs classiques. Les systèmes désormais utilisés et les protocoles envisagés reposent sur des états intriqués

de nombreux qubits, sur des systèmes dont l'espace de Hilbert est de grande dimension ou sur des séquences faisant intervenir de nombreuses portes quantiques, ou de longues séries de mesures. Les outils de caractérisation d'états, de portes et de mesures quantiques sont donc appelés à jouer un rôle important dans le développement de nouvelles technologies d'information quantique. Parmi eux, la tomographie d'état quantique, ou reconstruction d'état, qui consiste à déterminer l'état quantique complet du système, joue un rôle de premier plan. De multiples méthodes de reconstruction d'états quantiques ont été proposées [3–10] et appliquées [11–19].

De nombreuses plateformes expérimentales ont atteint le degré de contrôle requis pour réaliser des protocoles d'information quantique. Parmi celles-ci, on peut notamment citer les atomes froids [20], les ions piégés [21], les centres colorés du diamant [22], les circuits supraconducteurs [23], l'optique quantique [24], l'optomécanique [25], les boîtes quantiques [26] et l'électrodynamique quantique en cavité (CQED) [27]. Cette dernière vise à étudier l'interaction entre la matière et la lumière, en utilisant des résonateurs afin de renforcer leur couplage dans des proportions parfois gigantesques. Des prouesses technologiques ont permis, dans de nombreux systèmes, d'atteindre des régimes de couplage forts, dans lesquels l'interaction entre un atome et le champ électromagnétique d'une cavité dépassent de plusieurs ordres de grandeurs les phénomènes parasites d'interaction avec l'environnement, responsables de la décohérence.

Cette thèse est consacrée à une nouvelle méthode de reconstruction d'états quantiques, appelée *tomographie par trajectoires*. Cette méthode, a déjà été utilisée pour des systèmes simples constitués d'un seul bit quantique, ou d'un mode du champ électromagnétique dont on étudiait uniquement la distribution du nombre de photons [28]. Elle nécessite la préparation de nombreuses copies du système. Sur chaque copie, on effectue une longue série de mesures afin de suivre la trajectoire quantique du système. On utilise ensuite des outils mathématiques permettant d'extraire toute l'information sur l'état initial contenue dans l'ensemble de trajectoires quantiques mesurées y compris lorsque celles-ci sont dissipatives. Afin d'illustrer le potentiel de cette méthode pour étudier des systèmes de grande dimensionnalité, nous l'avons appliquée à un état intriqué de deux modes du champ électromagnétique, préparé grâce aux outils de l'électrodynamique quantique en cavité.

Le groupe de CQED du laboratoire Kastler-Brossel a mené depuis de nombreuses années des expériences pionnières dans l'utilisation conjointe d'atomes préparés dans un état très excité, dit de Rydberg, avec des cavités micro-onde de grande finesse, afin de manipuler des particules quantiques isolées. Une des briques élémentaires des travaux de l'équipe est l'interaction résonnante entre un atome et un mode du champ, qui donne lieu aux *oscillations de Rabi*. Celles-ci consistent en l'échange cohérent d'un quantum d'énergie entre un atome et une cavité micro-onde. Lorsqu'il est placé au centre d'une cavité de grande finesse, l'atome émet préférentiellement dans celle-ci. La présence des miroirs permet au photon de rester suffisamment longtemps au voisinage de l'atome pour être réabsorbé. L'atome peut ensuite émettre de nouveau le photon et on observe l'oscillation mentionnée précédemment. Celle-ci a une période suffisamment faible devant le temps de vie des atomes et celui du photon piégé dans la cavité pour que le système effectue de nombreuses oscillations avant que les phénomènes de décohérence ne se fassent sentir.

Le deuxième ingrédient central des expériences du groupe consiste en l'utilisation d'atomes proches de résonance mais suffisamment désaccordés pour ne pas pouvoir échanger d'énergie avec la cavité. Ces atomes subissent un déplacement lumineux causé par les

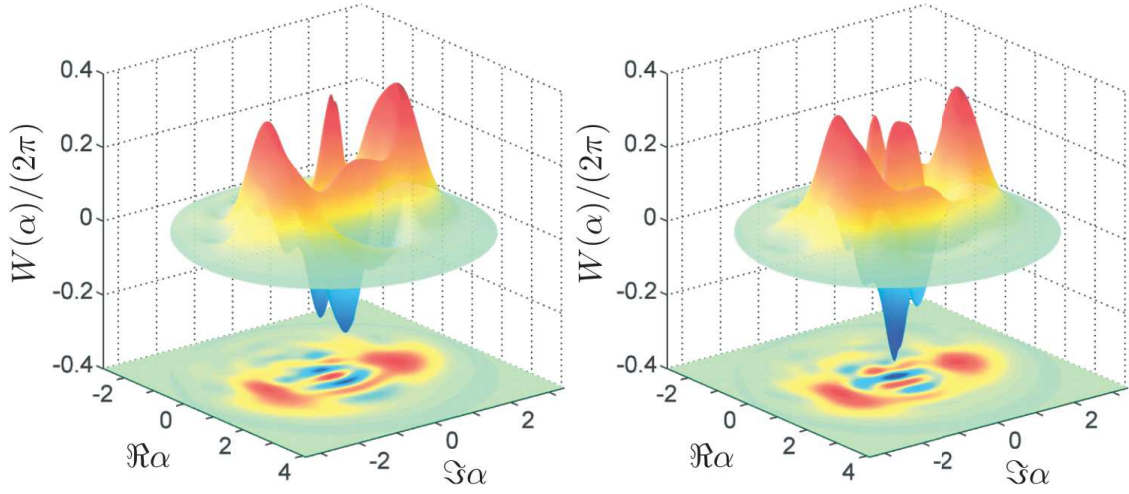


Fig. A Fonction de Wigner d'états chat de Schrödinger pair (à gauche) et impair (à droite) reconstruits par des mesures de la fonction de Wigner réalisées à l'aide de séries d'atome effectuant des mesures QND de la parité du champ. Figure extraite de [30]

photons piégés entre les miroirs, qui permet de compter ces derniers, sans modifier leur nombre. On peut donc réaliser des mesures quantiques non destructives (QND) du nombre de photons.

En utilisant les atomes résonnants et les mesures QND, le groupe a mené de nombreux travaux explorant les effets quantiques au niveau le plus fondamental : celui d'atomes isolés et de champs préparés dans des états de Fock ou des états mésoscopiques contenant quelques photons. Grâce à des cavités micro-ondes de grande finesse, mises au point dans l'équipe dans la deuxième moitié des années 2000, les temps de vie des photons ont atteint la centaine de millisecondes et il a été possible de réaliser des expériences dans lesquelles des centaines d'atomes étaient envoyés pour mesurer le champ, avant que la décohérence ne détériore totalement son état quantique. Ces conditions particulièrement favorables ont permis la réalisation d'états quantiques complexes, présentant des propriétés non classiques intéressantes. En particulier, il a été possible de préparer des états *chats de Schrödinger*, dans lesquels un champ mésoscopique constitué de quelques photons oscille avec deux phases opposées en même temps. Ces états ont ensuite été mesurés afin d'en faire la tomographie et de reconstruire leur fonction de Wigner, visibles sur la figure A [29].

Pour effectuer cette reconstruction, l'équipe a effectué des séries de mesures QND qui peuvent être répétées sur la même préparation de l'état puisqu'elles ne perturbent pas l'observable qui est mesurée. Cependant, le traitement appliqué à l'époque consistait à moyenner, pour chaque ensemble de mesures identiques, les résultats obtenus. Cette méthode amène à négliger une grande partie de l'information obtenue sur chaque trajectoire, contenue dans les corrélations entre les mesures d'atomes successifs.

Afin d'extraire plus d'information des trajectoires quantiques individuelles, d'autres expériences du groupe ont appliqué la méthode de *l'état passé* [31] à l'étude d'un champ

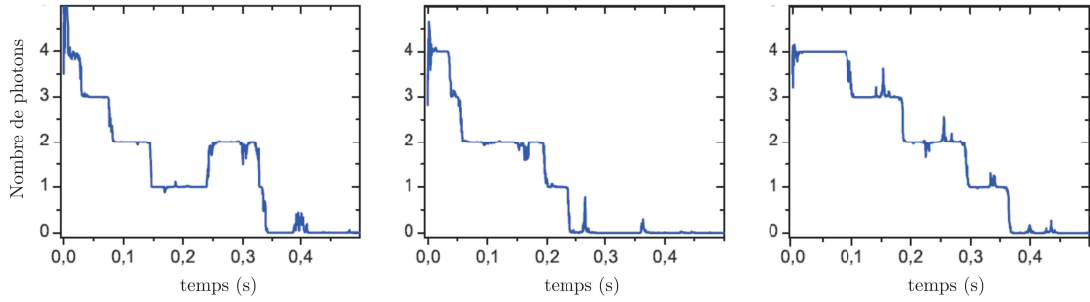


Fig. B Trajectoires quantiques individuelles du nombre de photons contenus dans une cavité micro-onde. Le nombre de photons à l'instant t est déterminé à l'aide de 110 mesures QND effectuées après t . Ces trajectoires font apparaître des sauts quantiques durant lesquels le nombre de photons varie d'une unité. Figure extraite de [32]

d'une cavité effectuant des sauts quantiques. Ces derniers correspondent à des événements de perte d'un photon ou d'injection dans le mode d'un photon de l'environnement, visibles sur la figure B, qui montre trois trajectoires quantiques du nombre de photons contenus dans une cavité. Le nombre de photons à l'instant t est déterminé *a posteriori* en utilisant un ensemble de mesures QND effectuées après t . Au lieu de se contenter de mesures après l'instant t , la méthode de l'état passé consiste à utiliser, dans une trajectoire quantique individuelle donnée, les mesures effectuées avant et après l'instant t . L'état passé combine la matrice densité du système avec un autre opérateur, appelé *matrice d'effet*, qui suit une rétro-évolution à partir de la dernière mesure effectuée et décrit l'information contenue dans les mesures effectuées après t .

Afin d'étudier les sauts quantiques, on envoyait un atome résonnant avec la cavité pour y injecter un photon, simulant un saut quantique de la lumière. La distribution de photons de la cavité était sondée avant et après ce saut par des atomes effectuant des mesures QND du nombre de photons. On utilisait l'information fournie par ces atomes pour déterminer l'instant de l'injection, sans rien inclure *a priori* sur cet instant dans la méthode de reconstruction. La méthode de l'état passé a permis de reconstruire les sauts quantiques en filtrant le bruit lié aux aléas de la mesure [33]. La figure C montre la résolution d'un saut quantique individuel à l'aide de la matrice densité, de la matrice d'effet et de l'état passé. La matrice densité détecte le saut quantique avec un retard. Il lui faut suffisamment de mesures après le saut pour certifier qu'il a bien eu lieu. La matrice d'effet le positionne quant à elle en avance. Les signaux obtenus grâce à ces deux opérateurs présentent de plus un bruit important à cause des fluctuations de mesure. Dans le cas de la méthode de l'état passé, en revanche, le saut est parfaitement résolu, bien qu'il s'agisse d'une trajectoire quantique unique. Le bruit est par ailleurs filtré. On voit tout l'intérêt d'extraire toute l'information contenue dans une réalisation unique de l'expérience.

Les travaux effectués dans ma thèse se situent à l'interface entre ces deux expériences. Ils se concentrent sur une méthode de tomographie par trajectoires, inspirée de la méthode de l'état passé, qui consiste à effectuer sur chaque réalisation de l'expérience de nombreuses mesures, afin de suivre la trajectoire du système et d'en tirer une matrice d'effet [28].

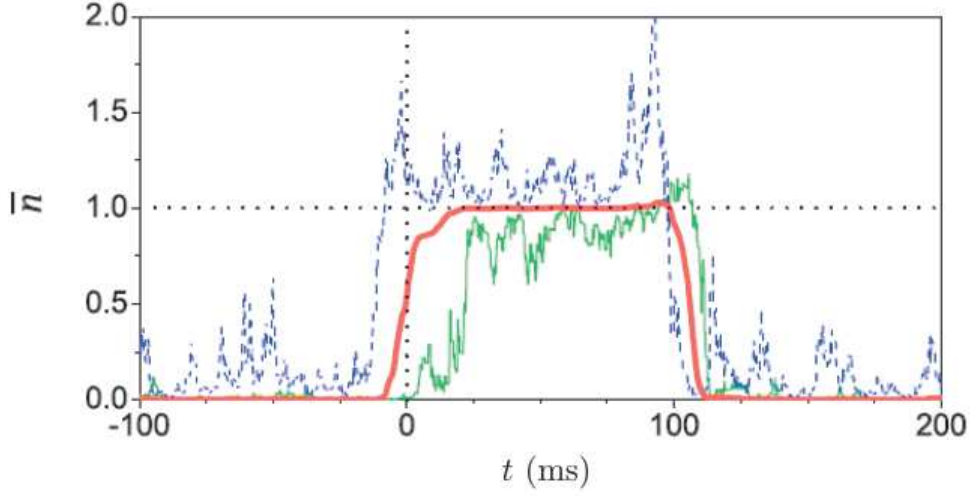


Fig. C Évolution du nombre moyen de photons d'un mode d'une cavité, en fonction du temps t , pendant un saut quantique. Les pointillés indiquent le moment auquel le saut a lieu. La courbe verte indique le nombre de photons obtenu en utilisant les mesures effectuées avant t , c'est-à-dire avec la matrice densité du système à l'instant t . La courbe bleue est obtenue avec les mesures effectuées après t , c'est à dire avec la matrice d'effet. La courbe rouge est obtenue avec l'état passé qui combine la matrice densité et la matrice d'effet. Figure extraite de [34]

L'ensemble des matrices d'effet obtenues permet d'effectuer une tomographie de la matrice densité initiale du système. Cette méthode permet de garder toute l'information contenue dans les corrélations entre les différentes mesures effectuées sur une même réalisation de l'état.

Afin de mettre en évidence l'intérêt de cette méthode pour l'étude de systèmes de grande dimension, nous l'avons appliquée à un état intriqué de deux modes du champ. Grâce à l'installation d'une deuxième cavité micro-onde dans le dispositif, un objectif planifié depuis longtemps dans l'équipe, nous avons pu préparer des états intriqués de deux modes du champ séparés d'une dizaine de centimètres. La préparation de ces états est inspirée d'une expérience effectuée au début des années 2000 avec deux modes d'une même cavité, de polarisations orthogonales et de fréquences voisines [35]. Dans ce travail, on faisait interagir un atome successivement avec les deux modes, de sorte qu'il émette un photon avec une probabilité égale dans chacun d'entre eux. Il était ainsi possible de préparer un état intriqué $|10\rangle + |01\rangle$ ¹ dans lequel un photon se trouve simultanément dans les deux cavités. Afin de mettre en évidence la cohérence de la superposition ainsi préparée, un deuxième atome était envoyé pour absorber le photon. Le battement des deux membres de la superposition, inscrit dans la phase quantique de l'état, donnait lieu à une oscillation de la probabilité d'absorption, visible sur la figure D. Le facteur de qualité de la cavité utilisée à l'époque ne permettait pas d'effectuer une reconstruction

1. Dans tout le manuscrit, on omettra les facteurs de normalisation lorsqu'on écrira des états superposés simples comme $(|10\rangle + |01\rangle)/\sqrt{2}$.

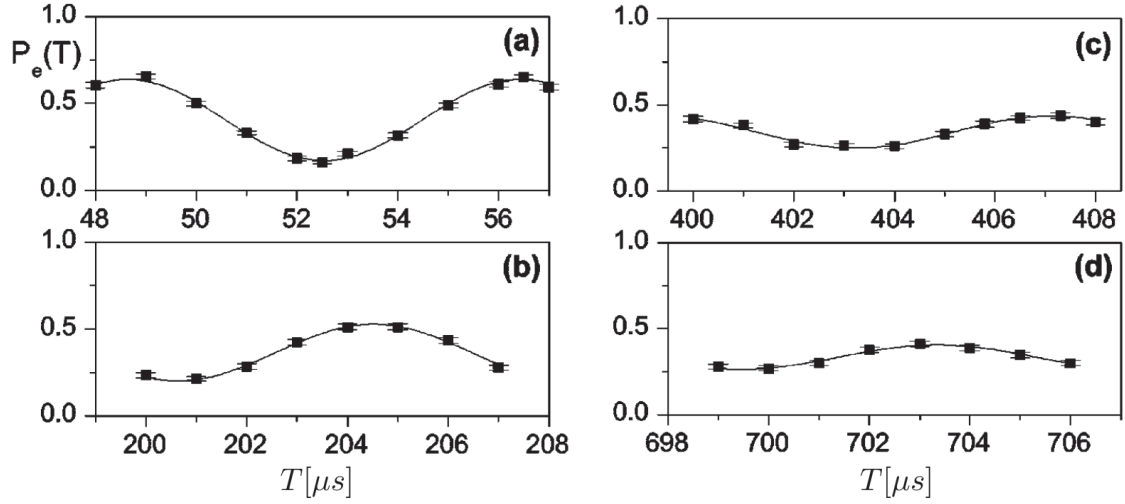


Fig. D Oscillation de la probabilité d'absorption, par un atome de Rydberg, d'un photon délocalisé sur deux modes d'une même cavité micro-onde. Le temps en abscisse est le temps entre la préparation de l'état et sa mesure. La probabilité d'absorber est une fonction oscillante de la phase quantique de l'état, qui traduit la cohérence de l'état superposé du photon. Figure extraite de [35]

complète de l'état.

Dans notre montage à deux cavités séparées, nous avons reproduit cette mesure. Elle permet de mettre en évidence les effets des cohérences non locales, mais pas de caractériser complètement l'état du système. Durant ma thèse, nous avons donc implémenté, en plus de ce protocole, d'autres mesures pour caractériser l'état quantique préparé. Dans un premier temps, nous avons reproduit l'expérience de délocalisation d'un photon sur deux modes et montré que la cohérence quantique subsiste après environ 20 millisecondes. Les modes étant désormais séparés, cette expérience peut être vue comme un interféromètre à un photon, dont les bras ont une longueur effective de 6000 kilomètres, qui correspondent à la distance parcourue par la lumière avant le brouillage des franges. Cette grande amélioration du temps de cohérence de l'état quantique par rapport à l'expérience originale nous a permis d'effectuer de nombreuses mesures successives sur chaque préparation de l'état afin d'appliquer la tomographie par trajectoires. Cette méthode a permis de reconstruire entièrement l'état du système. Nous avons en particulier pu en tirer des barres d'erreur sur les coefficients de la matrice densité reconstruite, ce qui constitue un avantage conséquent de la méthode [28]. Nous avons montré la pertinence de ces barres d'erreur en effectuant une étude statistique des résultats de reconstruction. Nous avons enfin montré que la tomographie par trajectoires permet d'effectuer des mesures successives, y compris quand celles-ci perturbent grandement l'état, afin d'accélérer la reconstruction de l'état quantique.

Le plan de cette thèse est le suivant. Le chapitre I traite de l'interaction entre un atome et un mode du champ électromagnétique. Après avoir rappelé le cadre théorique de cette interaction et les différents régimes auxquels elle donne lieu, on présente le dispositif expérimental permettant en pratique de se placer dans les conditions où un atome

interagit successivement de manière individuelle avec deux modes du champ contenus dans des cavités supraconductrices de grande finesse. On explique ensuite les techniques d'électrodynamique quantique en cavité qui permettent de tirer profit de cette interaction pour réaliser des protocoles d'information quantique, dont en particulier les nouvelles méthodes mises au point durant cette thèse pour les expériences à deux cavités.

Le chapitre II décrit la méthode de préparation de l'état à un photon délocalisé et ses principales propriétés. On explique aussi la mesure des cohérences de l'état par des atomes résonnants. On donne une interprétation interférométrique de cette mesure et on explique, en la comparant à des mesures de la fonction de Wigner de l'état, en quoi elle est plus adaptée pour extraire de l'information sur les cohérences.

Le chapitre III est une introduction aux principes de la reconstruction d'états quantiques. Après avoir rappelé le cadre de l'estimation d'un paramètre réel et le concept d'information de Fisher qui permet de caractériser l'efficacité d'une estimation, on décrit le principe de la tomographie quantique et on présente les méthodes usuelles de tomographie. On explique ensuite en détails le fonctionnement de la tomographie par trajectoires. Enfin, on montre comment on peut extraire de cette dernière des barres d'erreur sur des quantités dépendant linéairement de la matrice densité.

Enfin, le chapitre IV est consacré à l'application de la méthode de tomographie par trajectoires à des états intriqués des deux cavités. On commence par présenter les calculs de matrices d'effet pour les mesures effectuées. On montre ensuite les résultats obtenus dans les différents cas. On montre en particulier qu'on parvient à reconstruire l'état complet du système et à obtenir des barres d'erreur sur les paramètres estimés. Enfin, on propose une méthode de tomographie partielle qui permet de reconstruire un paramètre de la matrice densité, en l'occurrence la phase quantique de la superposition $|10\rangle + e^{i\phi} |01\rangle$, et on utilise cette tomographie pour montrer la pertinence des barres d'erreur extraites de la reconstruction.



Chapitre I

Atomes et photons en cavité

Toutes nos expériences reposent sur l'interaction entre deux éléments essentiels : des cavités supraconductrices de grande finesse et des atomes de Rydberg circulaires. L'interaction entre ces deux systèmes permet d'utiliser les atomes pour modifier et mesurer l'état de la lumière piégée dans deux cavités. Dans ce chapitre, nous commençons par rappeler les principaux éléments théoriques de description d'un mode du champ électromagnétique quantifié, d'un atome à deux niveaux et de leur interaction. Nous présentons ensuite le montage permettant de préparer les atomes de Rydberg circulaires et de les faire interagir avec les modes de deux cavités. Puis nous montrons les techniques usuelles de manipulation de l'état d'une cavité utilisées dans l'équipe, ainsi que les nouvelles techniques développées pour le système à deux cavités micro-onde.

I.1 Description quantique d'un mode du champ électromagnétique

I.1.1 Quantification du champ

En physique quantique, le champ électromagnétique est quantifié : il est constitué de photons, dont le caractère corpusculaire peut être mis en évidence dans des expériences de mesure de fonction de corrélation du champ [36]. On peut rendre compte de ce résultat théoriquement en choisissant une base des modes du champ électromagnétique, qu'on quantifie [37]. Chaque mode est alors décrit par un oscillateur harmonique quantique. Les états quantiques accessibles au mode j appartiennent à un espace de Hilbert $\mathcal{H}_j$, de dimension infinie, dont une base est donnée par les états de Fock $\{|n\rangle\}_{n \in \mathbb{N}}$. Deux opérateurs $\hat{a}_j$ et $\hat{a}_j^\dagger$ décrivent respectivement l'annihilation et la création d'un photon dans le mode j . Ils sont définis par leur action sur la base de Fock¹ :

$$\hat{a}_j |n\rangle = \sqrt{n} |n-1\rangle, \quad (\text{I.1.1a})$$

$$\hat{a}_j^\dagger |n\rangle = \sqrt{n+1} |n+1\rangle, \quad (\text{I.1.1b})$$

et leur commutateur vaut

$$[\hat{a}_j, \hat{a}_j^\dagger] = 1. \quad (\text{I.1.2})$$

1. Pour simplifier, on n'indiquera pas la base de Fock lorsque le mode auquel elle se rapporte est évident.

Le hamiltonien du mode j , de pulsation ω_{cj} , s'écrit [27] :

$$\hat{H}_{cj} = \hbar\omega_{cj} \left(\hat{a}_j^\dagger \hat{a}_j + \frac{1}{2} \right) = \hbar\omega_{cj} \left(\hat{N}_j + \frac{1}{2} \right), \quad (\text{I.1.3})$$

où $\hat{N}_j \equiv \hat{a}_j^\dagger \hat{a}_j = \sum_n n |n\rangle \langle n|$ est l'opérateur nombre. Les états propres de ce hamiltonien sont les états de Fock $|n\rangle$, associés à l'énergie $n\hbar\omega_{cj}$, correspondant à n quanta d'énergie. L'effet de ce hamiltonien sur un champ dans l'état initial $|\psi_0\rangle = \sum_n c_n |n\rangle$ est donc d'opérer une rotation des phases quantiques associées à chaque n , d'un angle proportionnel à n :

$$|\psi\rangle(t) = e^{-\frac{i\hat{H}_{cj}t}{\hbar}} |\psi_0\rangle = \sum_n c_n e^{-in\omega_{cj}t} |n\rangle. \quad (\text{I.1.4})$$

Considérons une base de modes, indicés par j , vérifiant les équations de Maxwell et les conditions aux limites. Le champ produit en $\mathbf{r}$ ² s'écrit :

$$\hat{\mathbf{E}}_{c,j}(\mathbf{r}) = i\mathcal{E}_{0,j} \left[\mathbf{f}_j(\mathbf{r}) \hat{a}_j - \mathbf{f}_j^*(\mathbf{r}) \hat{a}_j^\dagger \right], \quad (\text{I.1.5})$$

où $\mathbf{f}_j(\mathbf{r}) \equiv \boldsymbol{\epsilon}_{cj} f_j(\mathbf{r})$ est la fonction, normalisée à 1, décrivant la structure spatiale du mode j . $\mathcal{E}_{0,j}$ a la dimension d'un champ électrique et est donné par ([27], chapitre 3.1) :

$$\mathcal{E}_{0,j} = \sqrt{\frac{\hbar\omega_{cj}}{2\epsilon_0\mathcal{V}_j}}, \quad (\text{I.1.6})$$

où $\mathcal{V}_j$ est le volume du mode j , donné par :

$$\mathcal{V}_j = \int |\mathbf{f}_j(\mathbf{r})|^2 d^3\mathbf{r}. \quad (\text{I.1.7})$$

$\mathcal{E}_{0,j}$ représente l'amplitude quadratique des fluctuations quantiques du vide du mode. On voit que cette valeur dépend de la géométrie : elle est d'autant plus élevée que le volume du mode est petit. En confinant la lumière, par exemple à l'aide de miroirs, il est donc possible de modifier les propriétés des fluctuations quantiques du vide.

I.1.2 Couplage du mode à un courant classique

Considérons un courant classique $\mathbf{j}(\mathbf{r}, t) = \mathbf{J}_0(\mathbf{r})e^{-i\omega t} + \mathbf{J}_0^*(\mathbf{r})e^{i\omega t}$, oscillant à la pulsation ω . L'interaction de ce courant avec le mode j du champ est décrite par le hamiltonien :

$$\hat{H}_{inj}^I = - \int d^3\mathbf{r} \mathbf{j}(\mathbf{r}, t) \cdot \mathbf{A}_j(\mathbf{r}), \quad (\text{I.1.8a})$$

$$\mathbf{A}_j(\mathbf{r}) = \frac{\mathcal{E}_{0,j}}{\omega_{cj}} \left[\mathbf{f}_j(\mathbf{r}) \hat{a}_j + \mathbf{f}_j^*(\mathbf{r}) \hat{a}_j^\dagger \right]. \quad (\text{I.1.8b})$$

Plaçons-nous en représentation d'interaction par rapport au hamiltonien libre $\hat{H}_{cj}$. Les opérateurs d'annihilation et de création sont alors remplacés par :

$$\hat{a}_{I,j} = \hat{a}_j e^{-i\omega_{cj}t}, \quad (\text{I.1.9a})$$

$$\hat{a}_{I,j}^\dagger = \hat{a}_j^\dagger e^{i\omega_{cj}t}. \quad (\text{I.1.9b})$$

2. Dans toute la suite, on utilisera des caractères gras pour les vecteurs.

Le hamiltonien d'interaction fait intervenir des termes oscillant à la pulsation $\omega + \omega_{cj}$. Ces termes oscillent très rapidement et leur effet sur l'évolution est négligeable. On fera l'approximation de l'onde tournante qui consiste à les supprimer. Le hamiltonien peut alors se mettre sous la forme :

$$\hat{H}_{inj}^I = i\hbar(\gamma_j \hat{a}_j^\dagger - \gamma_j^* \hat{a}_j), \quad (\text{I.1.10})$$

où le taux d'injection γ_j se calcule par l'intégrale de la formule I.1.8a. Le couplage du courant au mode du champ pendant une durée t donne donc lieu à une évolution :

$$\hat{U}_{inj,j}(t) = e^{\gamma_j t \hat{a}_j^\dagger - \gamma_j^* t \hat{a}_j} = \hat{D}_j(\gamma t), \quad (\text{I.1.11})$$

où l'opérateur $\hat{D}_j(\alpha) \equiv e^{\alpha \hat{a}_j^\dagger - \alpha^* \hat{a}_j}$ est l'opérateur déplacement.

I.1.3 La fonction de Wigner

En physique statistique classique, pour décrire un oscillateur harmonique, on peut définir une densité de probabilité sur l'espace des phases représentant la probabilité que le système se trouve au voisinage d'un point de position et de vitesse définies. En physique quantique, il n'est pas possible de définir une telle densité. On dispose en revanche d'un ensemble de quasi-distributions de probabilités permettant de décrire l'état du système. La fonction de Wigner est une de ces fonctions. Elle possède des propriétés intéressantes et peut être mesurée facilement dans nos expériences. Pour un mode du champ électromagnétique décrit par la matrice densité $\hat{\rho}$, elle est définie par³ :

$$W^{[\hat{\rho}]}(\alpha_j) = \text{Tr} \left(\hat{D}_j(-\alpha_j) \hat{\rho} \hat{D}_j(\alpha_j) \hat{P}_j \right), \quad (\text{I.1.12})$$

où $\hat{P}_j$ est l'opérateur parité, défini par :

$$\hat{P}_j = \sum_n (-1)^n |n\rangle \langle n|. \quad (\text{I.1.13})$$

L'amplitude complexe α correspond à une position dans l'espace de Fresnel, qui est l'équivalent, pour un mode du champ, de l'espace des phases pour une particule dans un oscillateur harmonique à une dimension. Les parties réelle et imaginaire sont respectivement les analogues de la position et de la vitesse de la particule. La fonction de Wigner est donc égale à la valeur moyenne de la parité, dans l'état ayant subi un déplacement $\hat{D}(\alpha_j)$. Il y a une bijection entre les fonctions de Wigner et les matrices densité. La connaissance de la fonction de Wigner est donc équivalente à la connaissance de $\hat{\rho}$.

Fonction de Wigner et transformations

Considérons un état $\hat{\rho}$ auquel on fait subir un déplacement $\hat{D}_j(\beta_j)$:

$$\hat{\rho}_d = \hat{D}_j(\beta_j) \hat{\rho} \hat{D}_j(-\beta_j). \quad (\text{I.1.14})$$

3. La définition usuelle de la fonction de Wigner fait intervenir un facteur de normalisation $2/\pi$, que nous omettons ici pour rendre plus direct le lien entre la fonction de Wigner et la parité. Elle n'est alors plus normalisée et on ne peut plus en tirer des distributions de probabilité par intégration.

En utilisant la formule

$$\hat{D}_j(\alpha_j)\hat{D}_j(\beta_j) = \hat{D}_j(\alpha_j + \beta_j)e^{(\alpha_j\beta_j^* - \alpha_j^*\beta_j)/2}, \quad (\text{I.1.15})$$

on peut récrire la fonction de Wigner comme :

$$W^{\hat{\rho}^d}(\alpha_j) = \text{Tr} \left(\hat{D}_j(-\alpha_j + \beta_j) \hat{\rho} \hat{D}_j(\alpha_j - \beta_j) \hat{P}_j \right), \quad (\text{I.1.16a})$$

$$W^{\hat{\rho}^d}(\alpha_j) = W^{\hat{\rho}}(\alpha_j - \beta_j). \quad (\text{I.1.16b})$$

La fonction de Wigner de l'état déplacé est égale à la fonction de Wigner de l'état, déplacée dans l'espace des phase de β_j . À titre d'exemple, considérons l'état cohérent $\hat{D}_j(\alpha_j) |0\rangle = |\alpha_j\rangle$. En utilisant l'identité de Glauber, on montre que

$$|\alpha_j\rangle = \sum_n e^{-|\alpha|^2} \alpha^n / \sqrt{n!} |n\rangle. \quad (\text{I.1.17})$$

La fonction de Wigner de l'état vide vaut

$$W^{[0]}(\alpha_j) = \langle 0 | \hat{D}_j(\alpha_j) \hat{P}_j \hat{D}_j(-\alpha_j) | 0 \rangle, \quad (\text{I.1.18a})$$

$$= \langle 0 | \hat{D}_j(\alpha_j) \hat{P}_j \hat{D}_j(-\alpha_j) \hat{P}_j | 0 \rangle, \quad (\text{I.1.18b})$$

$$= \langle 0 | \hat{D}_j(\alpha_j) \hat{D}_j(+\alpha_j) | 0 \rangle, \quad (\text{I.1.18c})$$

$$= \langle 0 | \hat{D}_j(2\alpha_j) | 0 \rangle, \quad (\text{I.1.18d})$$

$$= \langle 0 | 2\alpha_j \rangle, \quad (\text{I.1.18e})$$

$$= e^{-2|\alpha_j|^2}, \quad (\text{I.1.18f})$$

où on a utilisé l'identité :

$$\hat{P}_j \hat{D}_j(\alpha_j) \hat{P}_j = \hat{D}_j(-\alpha_j), \quad (\text{I.1.19})$$

qui traduit le fait que l'opérateur de parité consiste à faire une symétrie de l'espace des phases par rapport à l'origine. La fonction de Wigner de l'état vide est donc une gaussienne, centrée sur l'origine. On en déduit que celle de l'état cohérent $|\alpha_j\rangle = \hat{D}_j(\alpha_j) |0\rangle$ est une gaussienne centrée sur α_j .

Considérons maintenant le cas d'une rotation $\hat{R}_j(\phi) \equiv e^{i\phi\hat{N}_j}$. On peut montrer, en utilisant la formule de Baker-Campbell-Hausdorff, que

$$\hat{R}_j(\phi) \hat{D}_j(\alpha_j) \hat{R}_j(-\phi) = \hat{D}_j(\alpha_j e^{i\phi}). \quad (\text{I.1.20})$$

On en déduit :

$$W^{\hat{R}_j(\phi) \hat{\rho} \hat{R}_j(-\phi)}(\alpha_j) = W^{\hat{\rho}}(\alpha_j e^{-i\phi}). \quad (\text{I.1.21})$$

La fonction de Wigner de l'état après rotation est la fonction de Wigner de l'état, tournée d'un angle ϕ .

I.1.4 État du champ à deux modes fonction de Wigner

Champ à deux modes

Dans nos expériences, nous utiliserons deux modes du champ électromagnétique. Les états du système constitué de ces deux modes seront donc des éléments de l'espace produit tensoriel :

$$\mathcal{H} = \mathcal{H}_1 \otimes \mathcal{H}_2. \quad (\text{I.1.22})$$

Pour simplifier les notations, on écrira les états de Fock sous la forme $|n_1, n_2\rangle \equiv |n_1\rangle \otimes |n_2\rangle$. Ces états forment une base de $\mathcal{H}$ ⁴. Si $\hat{O}_1$ est un opérateur agissant sur $\mathcal{H}_1$, on utilisera abusivement la même notation pour $\hat{O}_1$ et sa version tensorisée $\hat{O}_1 \otimes \mathbb{1}_2$ agissant sur $\mathcal{H}_1 \otimes \mathcal{H}_2$ et on fera de même pour les opérateurs de $\mathcal{H}_2$.

On définit l'opérateur nombre :

$$\hat{N} \equiv \hat{N}_1 \otimes \mathbb{1}_2 + \mathbb{1}_1 \otimes \hat{N}_2. \quad (\text{I.1.23})$$

Le hamiltonien décrivant l'évolution du système des deux cavités s'écrit :

$$\hat{H}_c = \hat{H}_{c1} + \hat{H}_{c2} = \hbar\omega_{c1} \left(\hat{N}_1 + \frac{1}{2} \right) + \hbar\omega_{c2} \left(\hat{N}_2 + \frac{1}{2} \right). \quad (\text{I.1.24})$$

Un système préparé à l'origine des temps dans l'état $|\psi_0\rangle = \sum_{n_1, n_2} c_{n_1, n_2} |n_1, n_2\rangle$ et évoluant librement est donc à l'instant t dans l'état :

$$|\psi\rangle(t) = e^{-\frac{i\hat{H}_c t}{\hbar}} |\psi_0\rangle = \sum_{n_1, n_2} c_{n_1, n_2} e^{-i(n_1\omega_{c1} + n_2\omega_{c2})t} |n_1, n_2\rangle, \quad (\text{I.1.25})$$

en ne prenant pas en compte le déplacement de Lamb.

Les composantes de Fock tournent donc avec des pulsations dépendant du nombre de photons dans chaque cavité et de leur fréquence respective. Notons $\delta = \omega_{c2} - \omega_{c1}$ la différence de fréquence entre les cavités. Afin de simplifier le problème, on se place dans un référentiel tournant à la pulsation moyenne $(\omega_{c1} + \omega_{c2})/2$ donné par :

$$|n_1, n_2\rangle_t = e^{-i\frac{\omega_{c1} + \omega_{c2}}{2}(n_1 + n_2)t} |n_1, n_2\rangle. \quad (\text{I.1.26})$$

L'évolution est alors décrite par :

$$|\psi(t)\rangle = \sum_{n_1, n_2} e^{-i\delta\frac{n_2 - n_1}{2}} |n_1, n_2\rangle_t. \quad (\text{I.1.27})$$

L'état évolue en effectuant une rotation :

$$\hat{R}(t) = \exp(i\delta t(\hat{N}_2 - \hat{N}_1)/2). \quad (\text{I.1.28})$$

Dans la suite, on se placera dans ce référentiel, sauf mention contraire, et on simplifiera les notations en omettant de noter la dépendance temporelle des états de Fock.

Définition de la fonction de Wigner à deux modes

La fonction de Wigner se généralise à deux modes avec la définition suivante⁵ :

$$W(\alpha_1, \alpha_2) = \text{Tr} \left(\hat{D}_1(-\alpha_1) \hat{D}_2(-\alpha_2) \hat{\rho} \hat{D}_1(\alpha_1) \hat{D}_2(\alpha_2) \hat{P} \right). \quad (\text{I.1.29})$$

Cette fonction dépend désormais de deux paramètres complexes α_1 et α_2 qui correspondent aux déplacements effectués dans les deux modes, avant d'effectuer une mesure de parité globale :

$$\hat{P} = \hat{P}_1 \otimes \hat{P}_2 = e^{i\pi(\hat{N}_1 + \hat{N}_2)}. \quad (\text{I.1.30})$$

4. Lorsque le contexte permet d'éviter les ambiguïtés, on omettra la virgule et on écrira par exemple $|10\rangle$ pour $|1, 0\rangle$.

5. Comme pour la fonction de Wigner à un mode, la définition usuelle fait intervenir un facteur $4/\pi^2$, que nous omettons.

Représentation graphique

Comme la fonction de Wigner à deux modes dépend de deux paramètres complexes, il est difficile de la représenter. On peut cependant en avoir un bon aperçu de la manière suivante. On discrétise l'espace pour le premier paramètre α_1 . Pour chaque valeur de α_1 , on trace la fonction obtenue lorsqu'on fait varier α_2 . La figure I.1 montre la représentation ainsi obtenue de la fonction de Wigner de l'état cohérent à deux modes $|\alpha_1 = 1, \alpha_2 = i\rangle$. Il s'agit d'une gaussienne dans l'espace des phases à 4 dimensions, centrée sur $\alpha_1 = 1, \alpha_2 = i$ et de largeur $1/2$.

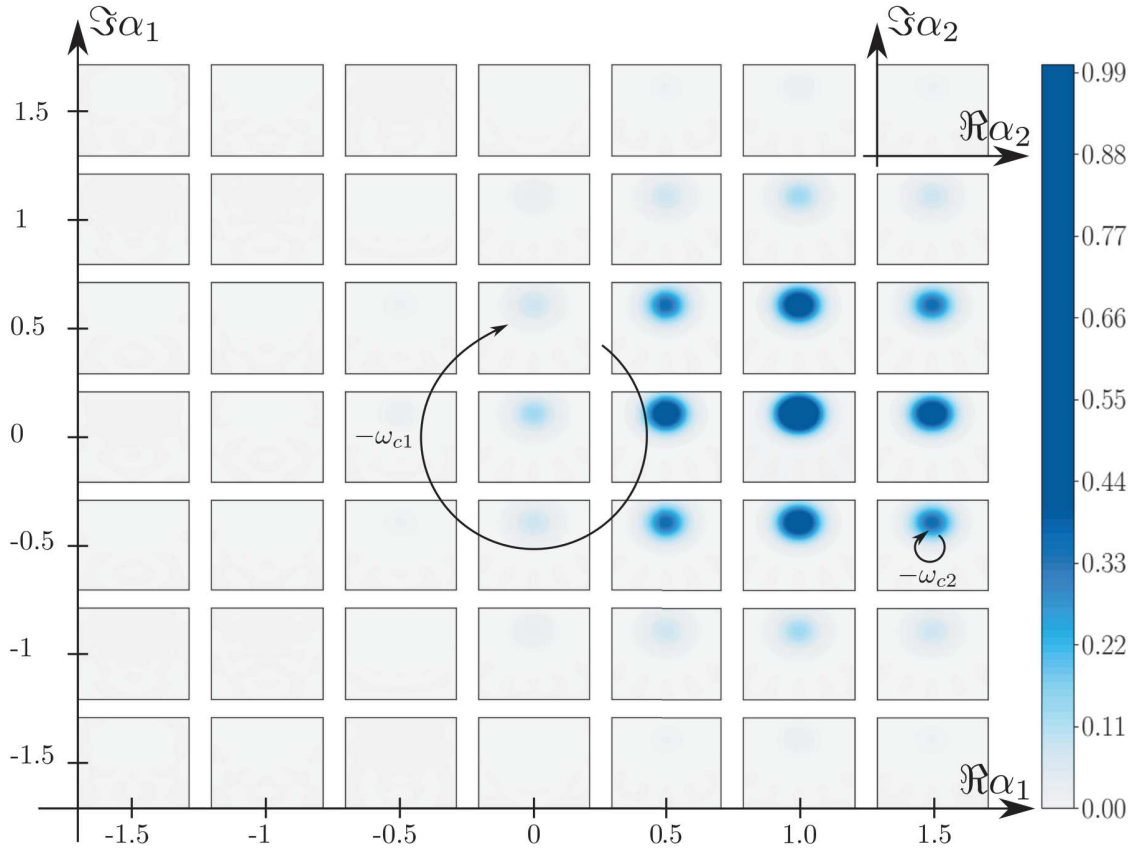


Fig. I.1 Fonction de Wigner de l'état cohérent $|\alpha_1 = 1, \alpha_2 = i\rangle$. α_1 est discrétisé selon une grille. Ses parties réelle et imaginaire prennent des valeurs identiquement espacées entre -1,5 et 1,5. En chaque point de cette grille, on trace la fonction obtenue en balayant les parties réelle et imaginaire de α_2 entre -2 et 2, pour la valeur de α_1 correspondante. Les flèches figurent l'évolution de l'état qui est constitué d'une rotation du paramètre α_1 , à la pulsation $-\omega_{c1}$ et d'une rotation de α_2 , à la pulsation $-\omega_{c2}$.

Évolution temporelle

Revenons pour ce paragraphe dans le référentiel fixe pour les deux modes. Lorsque le champ est isolé, son hamiltonien s'écrit à une constante près :

$$\hat{H}_c = \hbar\omega_{c1}\hat{N}_1 + \hbar\omega_{c2}\hat{N}_2. \quad (\text{I.1.31})$$

L'évolution peut donc se mettre sous la forme :

$$\hat{U}_c(t) = e^{-i\hat{N}_1\omega_{c1}t - i\hat{N}_2\omega_{c2}t} = \hat{R}_1(-\omega_{c1}t) \otimes \hat{R}_2(-\omega_{c2}t). \quad (\text{I.1.32})$$

En utilisant l'équation I.1.20, on en déduit que la fonction de Wigner se met à l'instant t sous la forme :

$$W^{\hat{U}_c(t)\hat{\rho}\hat{U}_c(t)}(\alpha_1, \alpha_2) = W^{\hat{\rho}}(\alpha_1 e^{i\omega_{c1}t}, \alpha_2 e^{i\omega_{c2}t}). \quad (\text{I.1.33})$$

Ce qu'on peut traduire de la manière suivante : les vignettes de la fonction de Wigner représentée figure I.1 tournent en fonction du temps à la pulsation $-\omega_{c1}$, tandis que l'intérieur de chaque vignette tourne à la pulsation $-\omega_{c2}$. Si on se place désormais dans le référentiel tournant défini par I.1.26, il ne reste que la rotation différentielle entre les deux modes, à la pulsation δ .

I.1.5 Décohérence du champ

Jusqu'à présent, on n'a considéré que des évolutions unitaires pour le champ, ce qui est valide pour un système isolé. Cependant, à cause du couplage avec les degrés de liberté de l'environnement, le mode relaxe vers un état d'équilibre thermique. Si l'environnement est à la température T , l'état d'équilibre thermique dans le mode i est décrit par la statistique de Bose-Einstein :

$$P_{th}(n) = e^{-n\beta\hbar\omega_{cj}} (1 - e^{-\beta\hbar\omega_{cj}}), \quad (\text{I.1.34})$$

où $\beta = (k_b T)^{-1}$ avec k_b la constante de Boltzmann. Le nombre moyen de photons est alors :

$$n_{th} = \frac{1}{e^{\beta\hbar\omega_{cj}} - 1}. \quad (\text{I.1.35})$$

Lorsque le champ est mis hors d'équilibre, le retour à l'état précédent, sous l'effet de l'interaction avec l'environnement, peut être décrit par l'équation de Lindblad [27] :

$$\frac{d\hat{\rho}_j}{dt} = -\frac{i}{\hbar} [\hat{H}_j, \hat{\rho}_j] + \sum_{\mu} \left(\hat{L}_{j,\mu} \hat{\rho}_j \hat{L}_{j,\mu}^\dagger - \frac{1}{2} \hat{\rho}_j \hat{L}_{j,\mu}^\dagger \hat{L}_{j,\mu} - \frac{1}{2} \hat{L}_{j,\mu}^\dagger \hat{L}_{j,\mu} \hat{\rho}_j \right) \quad (\text{I.1.36})$$

avec les opérateurs de saut :

$$\hat{L}_{j,\downarrow} = \sqrt{\kappa_j (1 + n_{th})} \hat{a}_j, \quad \hat{L}_{j,\uparrow} = \sqrt{\kappa_j (n_{th})} \hat{a}_j^\dagger. \quad (\text{I.1.37})$$

On verra dans la section III.3.1 l'origine de cette équation. Les opérateurs de saut $\hat{L}_{j,\downarrow}$ et $\hat{L}_{j,\uparrow}$ correspondent respectivement à la perte d'un photon vers l'environnement et à l'introduction d'un photon thermique dans le mode j .

Lorsqu'on s'intéresse uniquement au nombre de photons, ces équations permettent d'établir un ensemble fermé d'équations pour les probabilités d'occupation des états de Fock :

$$\begin{aligned} \frac{dp(n)}{dt} = & \kappa_j(1 + n_{th})(n + 1)p(n + 1) + \kappa_j n_{th}np(n - 1) - \\ & [\kappa_j(1 + n_{th})n + \kappa_j n_{th}(n + 1)]p(n). \end{aligned} \quad (\text{I.1.38})$$

En dérivant la relation $\langle \hat{N}_j \rangle = \sum_n np(n)$ par rapport au temps et en injectant la relation précédente, il vient :

$$\frac{d\langle \hat{N}_j \rangle}{dt} = -\kappa_j (\langle \hat{N}_j \rangle - n_{th}). \quad (\text{I.1.39})$$

L'énergie contenue dans la cavité relaxe donc vers n_{th} . Le temps caractéristique de cette relaxation est le temps de vie

$$T_{cj} = \frac{1}{\kappa_j}. \quad (\text{I.1.40})$$

T_{cj} peut aussi s'interpréter comme le temps de vie d'un photon. Pour un champ préparé initialement dans l'état de Fock $|n\rangle$, en l'absence de champ thermique, il apparaît sur [I.1.36](#) que la population $p(n)$ suit une décroissance exponentielle de temps caractéristique T_{cj}/n .

I.2 L'atome à deux niveaux

L'atome à deux niveaux est un système idéal, permettant de mettre en évidence de nombreux aspects de la physique quantique. Il s'agit d'un atome pouvant se trouver dans deux états : un état fondamental $|g\rangle$ et un état excité $|e\rangle$. Ces deux niveaux sont couplés par une transition dipolaire électrique. Dans notre expérience, les atomes peuvent être raisonnablement décrits par des atomes à deux niveaux. Dans cette section, nous allons introduire le formalisme permettant de décrire un atome à deux niveaux, ainsi que son interaction avec un champ classique.

I.2.1 Le pseudo-spin atomique

Le traitement de l'atome à deux niveaux est formellement identique à celui d'un spin $1/2$ plongé dans un champ magnétique orienté selon Oz . Notons $|+_z\rangle$ l'état *up* du spin et $|-_z\rangle$ l'état *down*. Par analogie avec le cas du spin, on peut écrire $|+_z\rangle \equiv |e\rangle$ et $|-_z\rangle \equiv |g\rangle$ et on définit pour l'atome un pseudo-spin $\hat{\mathbf{S}} = \hbar\hat{\boldsymbol{\sigma}}$, où $\hat{\boldsymbol{\sigma}}$ est le vecteur dont les composantes sont les matrices de Pauli $\hat{\sigma}_x, \hat{\sigma}_y$ et $\hat{\sigma}_z$, qui s'écrivent dans la base $\{|e\rangle, |g\rangle\}$:

$$\hat{\sigma}_x = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \hat{\sigma}_y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad \hat{\sigma}_z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (\text{I.2.1})$$

En choisissant l'origine des énergies à la moyenne des énergies des deux états, le hamiltonien s'écrit :

$$\hat{H}_{at} = \frac{\hbar\omega_{at}}{2}\hat{\sigma}_z, \quad (\text{I.2.2})$$

où ω_{at} est la pulsation de la transition atomique. Un état pur quelconque $|\psi\rangle$ du système peut s'écrire, à une phase globale près, sous la forme :

$$|\psi\rangle = \cos\frac{\theta}{2}|e\rangle + e^{i\phi}\sin\frac{\theta}{2}|g\rangle, \quad (\text{I.2.3})$$

où $\theta \in [0, \pi]$ et $\phi \in [0, 2\pi]$. Un état peut donc être paramétré par un vecteur unitaire $\mathbf{u}$, de coordonnées sphériques θ et ϕ . L'ensemble de ces vecteurs forme une sphère appelée *sphère de Bloch*, représentée sur la figure I.2. L'état $|+\mathbf{u}\rangle$ associé au vecteur de Bloch $\mathbf{u}$ est état propre de l'opérateur $\mathbf{u} \cdot \hat{\boldsymbol{\sigma}}$ avec la valeur propre $+1$.

I.2.2 Évolution temporelle

Le hamiltonien $\hat{H}_{at}$ a pour états propres $|e\rangle$ et $|g\rangle$. Un atome préparé à l'instant initial dans l'état $|\psi\rangle$ de la formule I.2.3 se trouve donc à l'instant t dans l'état :

$$|\psi(t)\rangle = e^{-i\omega_{at}t/2} \cos\frac{\theta}{2}|e\rangle + e^{i(\phi+\omega_{at}t/2)} \sin\frac{\theta}{2}|g\rangle. \quad (\text{I.2.4})$$

Cette évolution prend une forme très simple sur la sphère de Bloch. Il s'agit d'une précession autour de l'axe Oz , dans le sens direct, à la pulsation ω_{at} .

De manière plus générale, tout hamiltonien $\hat{H}_0$ de trace nulle peut s'écrire sous la forme $\hat{H}_0 = \frac{\hbar\omega_0}{2}\mathbf{h} \cdot \hat{\boldsymbol{\sigma}}$. L'évolution est alors une précession autour de l'axe défini par $\mathbf{h}$.

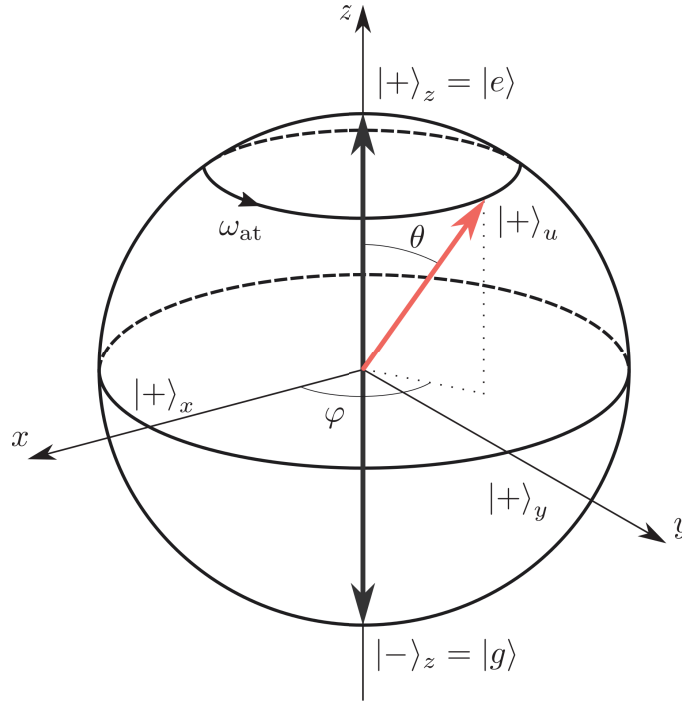


Fig. I.2 Sphère de Bloch. Le parallèle passant par $|+\rangle_u$ figure l'évolution de l'état. Il s'agit d'une rotation dans le sens direct autour de l'état de plus grande énergie $|+\rangle_z$, à la pulsation de la transition $|e\rangle \leftrightarrow |g\rangle$.

I.2.3 Interaction de l'atome à deux niveaux avec un champ classique

Comme en physique classique, l'atome à deux niveaux peut interagir avec un champ électrique par le biais de son dipôle électrique. L'opérateur dipôle électrique $\hat{\mathbf{D}} \equiv q\hat{\mathbf{R}}$, où q est la charge de l'électron et $\hat{\mathbf{R}}$ l'opérateur position, est nul dans les états $|e\rangle$ et $|g\rangle$, de parité définies. Ses seuls éléments non nuls sont donc les termes non diagonaux. On peut définir les opérateurs

$$\hat{\sigma}_+ \equiv |e\rangle \langle g|, \quad \hat{\sigma}_- \equiv |g\rangle \langle e|. \quad (\text{I.2.5})$$

Ces opérateurs correspondent respectivement à la création et à la destruction d'une excitation atomique et vont nous servir à écrire l'opérateur dipôle électrique. En effet, celui-ci a des valeurs nulles dans les états $|e\rangle$ et $|g\rangle$, qui ont des parités bien définies. Le dipôle a donc des valeurs nulles sur la diagonale et peut s'écrire :

$$\hat{\mathbf{D}} = d(\boldsymbol{\epsilon}_a \hat{\sigma}_- + \boldsymbol{\epsilon}_a^* \hat{\sigma}_+), \quad (\text{I.2.6})$$

où l'on a introduit le dipôle d et la polarisation $\boldsymbol{\epsilon}_a$ de la transition $|g\rangle \leftrightarrow |e\rangle$.

Considérons un champ électrique classique, prenant à la position de l'atome⁶ la valeur

$$\mathbf{E}_r(t) = i\mathcal{E}_r \left(\boldsymbol{\epsilon}_r e^{-i(\omega_r t + \phi)} - \boldsymbol{\epsilon}_r^* e^{i(\omega_r t + \phi)} \right), \quad (\text{I.2.7})$$

6. On se place ici dans l'approximation dipolaire en considérant que le champ a une amplitude complexe uniforme à l'échelle de l'atome. Comme les champs utilisés dans nos expériences ont des longueurs d'onde voisines de 6 millimètres, cette approximation est parfaitement justifiée.

où la pulsation ω_r est proche de celle de la transition atomique. L'interaction de l'atome avec ce champ est décrite par le hamiltonien dipolaire :

$$\hat{H}_r = -\hat{\mathbf{D}} \cdot \mathbf{E}_r(t). \quad (\text{I.2.8})$$

On se place dans le référentiel tournant, en posant $|\tilde{\psi}\rangle = e^{-i\omega_r t \hat{\sigma}_z/2} |\psi\rangle$, où $|\psi\rangle$ est le vecteur d'état dans le référentiel initial. Dans ce référentiel, le hamiltonien fait intervenir des termes constants et des termes oscillants à une pulsation de l'ordre de $2\omega_r$. Comme dans le cas de l'interaction d'un mode du champ avec une source classique, nous allons faire l'approximation de l'onde tournante, qui consiste à négliger les termes oscillants. Le hamiltonien s'écrit alors :

$$\tilde{H}_r = \frac{\hbar\Delta_c}{2}\hat{\sigma}_z - \frac{i\hbar}{2}\left(\Omega_c e^{-i\phi}\hat{\sigma}_+ - \Omega_c^* e^{i\phi}\hat{\sigma}_-\right), \quad (\text{I.2.9a})$$

$$= \frac{\hbar}{2} \begin{pmatrix} \Delta_c & -i\Omega_c e^{-i\phi} \\ i\Omega_c^* e^{i\phi} & -\Delta_c \end{pmatrix} \quad (\text{I.2.9b})$$

où on a introduit le désaccord entre l'atome et le champ $\Delta_c = \omega_{at} - \omega_r$ et la pulsation de Rabi classique

$$\Omega_c \equiv |\Omega_c| e^{i\psi} = \frac{2d}{\hbar} E_r \boldsymbol{\epsilon}_a^* \cdot \boldsymbol{\epsilon}_r. \quad (\text{I.2.10})$$

Le hamiltonien peut s'écrire

$$\tilde{H}_r = \frac{\hbar\Omega'_c}{2} \mathbf{n} \cdot \hat{\boldsymbol{\sigma}}, \quad (\text{I.2.11})$$

avec

$$\Omega'_c = \sqrt{|\Omega_c|^2 + \Delta_c^2}, \quad (\text{I.2.12a})$$

$$\mathbf{n} = \frac{1}{\Omega'_c} \begin{bmatrix} |\Omega_c| \sin(\psi - \phi) \\ |\Omega_c| \cos(\psi - \phi) \\ \Delta_c \end{bmatrix}. \quad (\text{I.2.12b})$$

L'évolution a alors l'interprétation géométrique suivante : le vecteur de Bloch effectue une précession autour du vecteur $\mathbf{n}$, à la pulsation effective Ω'_c . Lorsque le champ est fortement désaccordé, $\mathbf{n} \simeq \mathbf{e}_z$ et $\Omega'_c \simeq \Delta_c$: on retrouve la précession de Larmor dans le référentiel tournant. Pour $\Delta_c = 0$, le champ est résonnant et $\mathbf{n}$ est un vecteur du plan équatorial de la sphère de Bloch, dont l'azimut dépend des polarisations relatives de l'atome et du champ et de la phase du champ classique. L'atome effectue des oscillations entre $|e\rangle$ et $|g\rangle$ à la pulsation de Rabi. On les appelle *oscillations de Rabi classiques*. Lorsque la pulsation de Rabi est comparable au désaccord, les oscillations ont lieu autour d'un axe incliné par rapport au plan équatorial. Si on part de l'état $|g\rangle$, la probabilité d'atteindre l'état $|e\rangle$ n'atteint jamais l'unité. Par ailleurs les oscillations ont lieu plus rapidement que pour l'oscillation à résonance.

En contrôlant la fréquence et la phase de l'impulsion, on peut donc réaliser une rotation autour d'un axe quelconque de la sphère de Bloch. En contrôlant par ailleurs la durée de l'interaction, on peut donc réaliser n'importe quelle opération unitaire sur l'atome à deux niveaux et donc contrôler son état. Plaçons-nous dans le cas d'un atome soumis à une

impulsion à résonance, avec une pulsation de Rabi réelle ($\psi = 0$). L'évolution des états atomiques est donnée par :

$$|g\rangle \mapsto \cos \frac{\Omega_c t}{2} |g\rangle - e^{-i\phi} \sin \frac{\Omega_c t}{2} |e\rangle, \quad (\text{I.2.13a})$$

$$|e\rangle \mapsto e^{i\phi} \sin \frac{\Omega_c t}{2} |g\rangle + \cos \frac{\Omega_c t}{2} |e\rangle. \quad (\text{I.2.13b})$$

Pour une durée d'interaction t_π telle que $\Omega_c t_\pi = \pi$, les états atomiques $|e\rangle$ et $|g\rangle$ sont échangés. On parle d'*impulsion* π . Ces impulsions nous permettent de choisir l'état initial de l'atome pour nos expériences.

Si l'impulsion a une durée $t_{\pi/2}$ telle que $\Omega_c t_{\pi/2} = \pi/2$, les états $|e\rangle$ et $|g\rangle$ sont transformés en des superpositions à poids égaux :

$$|g\rangle \mapsto \frac{1}{\sqrt{2}}(|g\rangle - e^{-i\phi} |e\rangle), \quad |e\rangle \mapsto \frac{1}{\sqrt{2}}(e^{i\phi} |g\rangle + |e\rangle). \quad (\text{I.2.14})$$

Ces impulsions sont appelées *impulsions* $\pi/2$. Elles permettent de réaliser des mesures d'interférométrie de Ramsey, que nous présenterons en détail dans la suite.

I.3 Interaction entre atomes et lumière

Dans cette section, nous développons la théorie traitant l'interaction entre le champ quantifié et l'atome à deux niveaux, qui est au cœur de nos expériences. Dans un premier temps, nous rappelons le formalisme de l'atome habillé, qui consiste à décrire le système dans une base qui prend en compte l'effet de la lumière sur les états atomiques. Nous utilisons ce formalisme pour étudier l'évolution du système. Nous appliquons ensuite ces résultats à l'interaction résonnante, dans laquelle un quantum d'énergie est échangé entre l'atome et le champ. Enfin, nous traitons l'interaction dispersive, lorsque le désaccord entre l'atome et le champ est trop important pour permettre d'échanger de l'énergie. Dans ce cas, il y a quand même un effet du champ sur l'atome. En effet, le déplacement lumineux induit par les photons contenus dans la cavité modifie les niveaux d'énergie de l'atome et peut donc changer la phase d'une superposition atomique. Réciproquement, l'atome agit comme un diélectrique qui modifie la phase du champ.

I.3.1 Modèle de Jaynes et Cummings et théorie de l'atome habillé

Le modèle de Jaynes-Cummings est un modèle théorique qui permet de traiter l'interaction entre un atome à deux niveaux et un mode quantifié du champ électromagnétique. Le système est décrit par le hamiltonien

$$\hat{H}_{JC} = \hat{H}_{at} + \hat{H}_{cj} + \hat{H}_{int,j}, \quad (\text{I.3.1})$$

où $\hat{H}_{at}$ est le hamiltonien atomique (I.2.2), $\hat{H}_{cj}$ est le hamiltonien du mode j (I.1.3) et $\hat{H}_{int,j}$ est le hamiltonien d'interaction, analogue du hamiltonien dipolaire (I.2.8), avec le champ quantifié $\hat{\mathbf{E}}_{c,j}$ du mode j (I.1.5) :

$$\hat{H}_{int,j} = -\hat{\mathbf{D}} \cdot \hat{\mathbf{E}}_{c,j} = -id\mathcal{E}_{0,j}(\boldsymbol{\epsilon}_a \hat{\sigma}_- + \boldsymbol{\epsilon}_a^* \hat{\sigma}_+) \left(f_j(\mathbf{r}) \boldsymbol{\epsilon}_{cj} \hat{a}_j - f_j^*(\mathbf{r}) \boldsymbol{\epsilon}_{cj}^* \hat{a}_j^\dagger \right). \quad (\text{I.3.2})$$

Ce hamiltonien fait intervenir deux termes proportionnels à $\hat{\sigma}_- \hat{a}_j$ et $\hat{\sigma}_+ \hat{a}_j^\dagger$. Ces termes correspondent à des processus où on crée ou absorbe simultanément une excitation atomique et une excitation du champ. Ces termes sont largement hors résonance lorsque la fréquence de transition atomique et celle du mode sont proches. On va donc les négliger. Cette approximation correspond, dans le cas quantique, à l'approximation de l'onde tournante déjà effectuée pour le champ classique. Dans le cadre de cette approximation, le hamiltonien d'interaction s'écrit :

$$\hat{H}_{int,j} = -\frac{i\hbar}{2} \left(\Omega_{0,j} \hat{\sigma}_+ \hat{a}_j - \Omega_{0,j}^* \hat{\sigma}_- \hat{a}_j^\dagger \right), \quad (\text{I.3.3})$$

où

$$\Omega_{0,j} = \frac{2\mathcal{E}_{0,j}d}{\hbar} f_j(\mathbf{r}) \boldsymbol{\epsilon}_a^* \cdot \boldsymbol{\epsilon}_{cj}. \quad (\text{I.3.4})$$

est appelée pulsation de Rabi du vide. Cette pulsation dépend de la position : elle est plus élevée là où le champ électrique du mode est plus intense. Elle est d'autant plus grande que le dipôle est élevé et que $\mathcal{E}_{0,j}$ est grand, ce qui signifie, d'après la formule I.1.6, que le volume du mode est faible. La pulsation de Rabi dépend aussi du recouvrement entre

la polarisation du mode et celle de la transition atomique. Le hamiltonien de Jaynes-Cummings peut finalement s'écrire sous la forme [38] :

$$\hat{H}_{JC} = \frac{\hbar\omega_{at}}{2}\hat{\sigma}_z + \hbar\omega_{cj}\left(\hat{a}_j^\dagger\hat{a}_j + \frac{1}{2}\right) - \frac{i\hbar}{2}\left(\Omega_{0,j}\hat{\sigma}_+\hat{a}_j - \Omega_{0,j}^*\hat{\sigma}_-\hat{a}_j^\dagger\right). \quad (\text{I.3.5})$$

Le terme d'interaction, obtenu dans l'approximation de l'onde tournante, ne fait intervenir que deux termes, dans lesquels un quantum d'énergie est échangé entre l'atome et le champ. Il préserve donc le nombre d'excitations $\hat{a}_j^\dagger\hat{a}_j + |e\rangle\langle e|$. Notons $\mathcal{S}_n = \{|e, n\rangle, |g, n+1\rangle\}$, pour $n \geq 0$ et $\mathcal{S}_{-1} = \{|g, 0\rangle\}$, les sous-espaces de l'espace de Hilbert de nombre d'excitations fixé. À cause de cette séparation, si le système est initialement préparé dans un état de $\mathcal{S}_n$, il reste dans ce sous-espace. On peut alors décrire toute l'évolution comme celle d'un système à deux niveaux. En particulier, on peut construire une sphère de Bloch associée à la multiplicité n , telle que représentée sur la figure I.3.

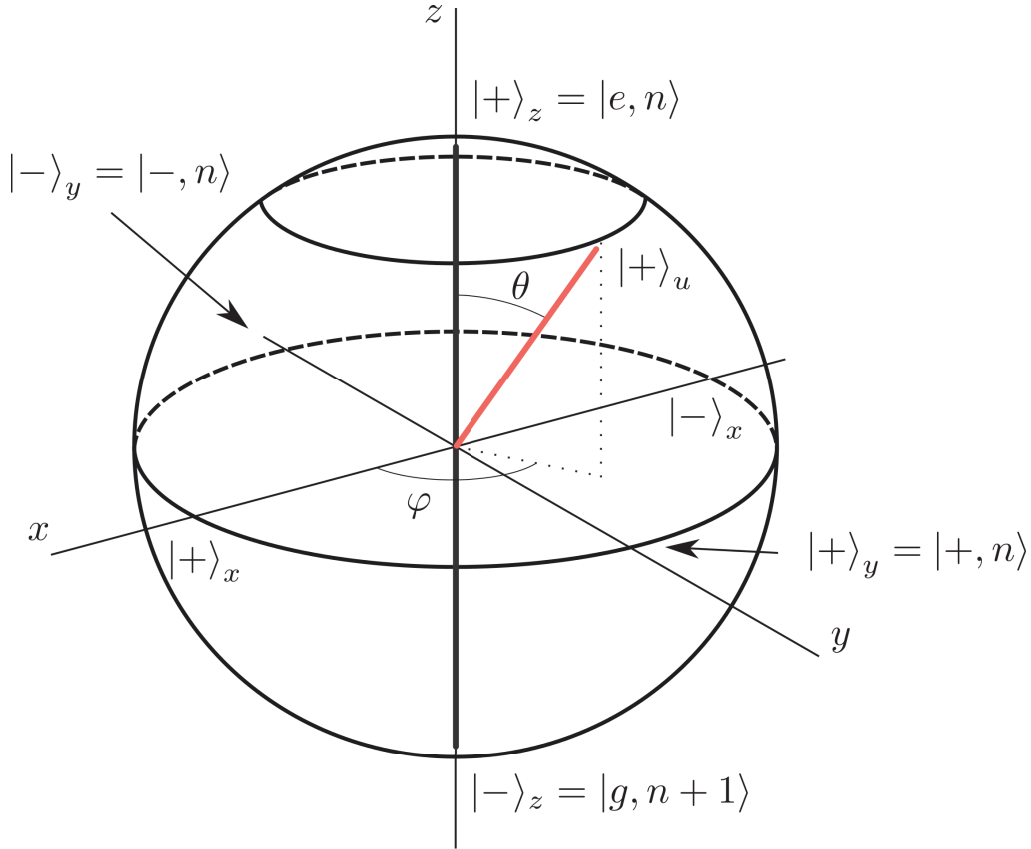


Fig. I.3 Sphère de Bloch pour la multiplicité $\mathcal{S}_n$. Cette sphère permet de représenter l'évolution du système lorsque celui-ci est initialement préparé dans un état de $\mathcal{S}_n$.

Comme le hamiltonien ne couple pas les états de nombres d'excitations différents, il peut être diagonalisé dans chaque $\mathcal{S}_n$ séparément. Dans l'état $|g, 0\rangle$, il n'y a aucune excitation à échanger. Cet état est donc état propre de $\hat{H}_{JC}$. La restriction de $\hat{H}_{JC}$ au

sous-espace $\mathcal{S}_n$ s'écrit :

$$\hat{H}_n = \omega_{cj}(n+1) + \frac{\hbar}{2} \begin{pmatrix} \Delta_j & -i\Omega_{n,j}^* \\ i\Omega_{n,j} & -\Delta_j \end{pmatrix}, \quad (\text{I.3.6})$$

où on a noté $\Delta_j = \omega_{at} - \omega_{cj}$ le désaccord atome-champ et $\Omega_{n,j} = \Omega_{0,j}\sqrt{n+1}$.

Ce hamiltonien est très similaire à celui qu'on avait obtenu pour le champ classique (I.2.9b). Il possède en plus un terme constant qui ne modifie pas la dynamique dans le sous-espace $\mathcal{S}_n$ mais ajoute une phase entre les états de nombre d'excitations différents. Il a par ailleurs une fréquence de Rabi d'autant plus élevée que le nombre de photons est important. L'évolution, dans le sous-espace $\mathcal{S}_n$ est donc une rotation autour de l'axe $\mathbf{n}$, à la pulsation $\Omega'_{n,j}$, avec :

$$\Omega'_{n,j} = \sqrt{|\Omega_{n,j}|^2 + \Delta_j^2}, \quad (\text{I.3.7a})$$

$$\mathbf{n} = \frac{1}{\Omega'_{n,j}} \begin{bmatrix} |\Omega_{n,j}| \sin(\psi) \\ |\Omega_{n,j}| \cos(\psi) \\ \Delta_j \end{bmatrix}, \quad (\text{I.3.7b})$$

où ψ est l'argument complexe de $\Omega_{n,j}$. Contrairement au cas classique, le champ n'a pas de phase. L'azimut de l'axe autour duquel la rotation a lieu ne dépend que des polarisations de l'atome et du champ. On va supposer par la suite que la fréquence de Rabi est réelle positive et que l'atome se situe au maximum du mode i . e. $f_j(\mathbf{r}) = 1$. Les états propres du hamiltonien sont alors :

$$|+, n\rangle_{\theta_n} = \cos \frac{\theta_n}{2} |e, n\rangle + i \sin \frac{\theta_n}{2} |g, n+1\rangle, \quad (\text{I.3.8a})$$

$$|-, n\rangle_{\theta_n} = \sin \frac{\theta_n}{2} |e, n\rangle - i \cos \frac{\theta_n}{2} |g, n+1\rangle, \quad (\text{I.3.8b})$$

où l'angle $\theta_n \in \{0, \pi\}$ est défini par :

$$\tan \theta_n = \frac{\Omega_{n,j}}{\Delta_j}. \quad (\text{I.3.9})$$

Il s'agit d'états intriqués du système atome-champ. On les appelle les *états habillés* de l'atome. Les valeurs propres associées sont données par :

$$E_{\pm, n} = \hbar\omega_{cj}(n+1) \pm \frac{\hbar}{2} \sqrt{\Omega_{n,j}^2 + \Delta_j^2}. \quad (\text{I.3.10})$$

Elles sont représentées sur la figure I.4 en fonction du rapport entre le désaccord et la pulsation de Rabi. Trois régimes d'interaction sont particulièrement utilisés dans nos expériences : l'interaction résonnante, lorsque $\Delta_j = 0$, l'interaction dispersive, pour laquelle $|\Delta_j| \gg \Omega_n$ pour les valeurs de n atteintes par le système et le passage adiabatique. Ce dernier consiste à balayer la fréquence de la transition atomique de manière à passer d'un désaccord positif $\Delta_j \gg \Omega_n$ à un désaccord négatif $-\Delta_j \gg \Omega_n$.

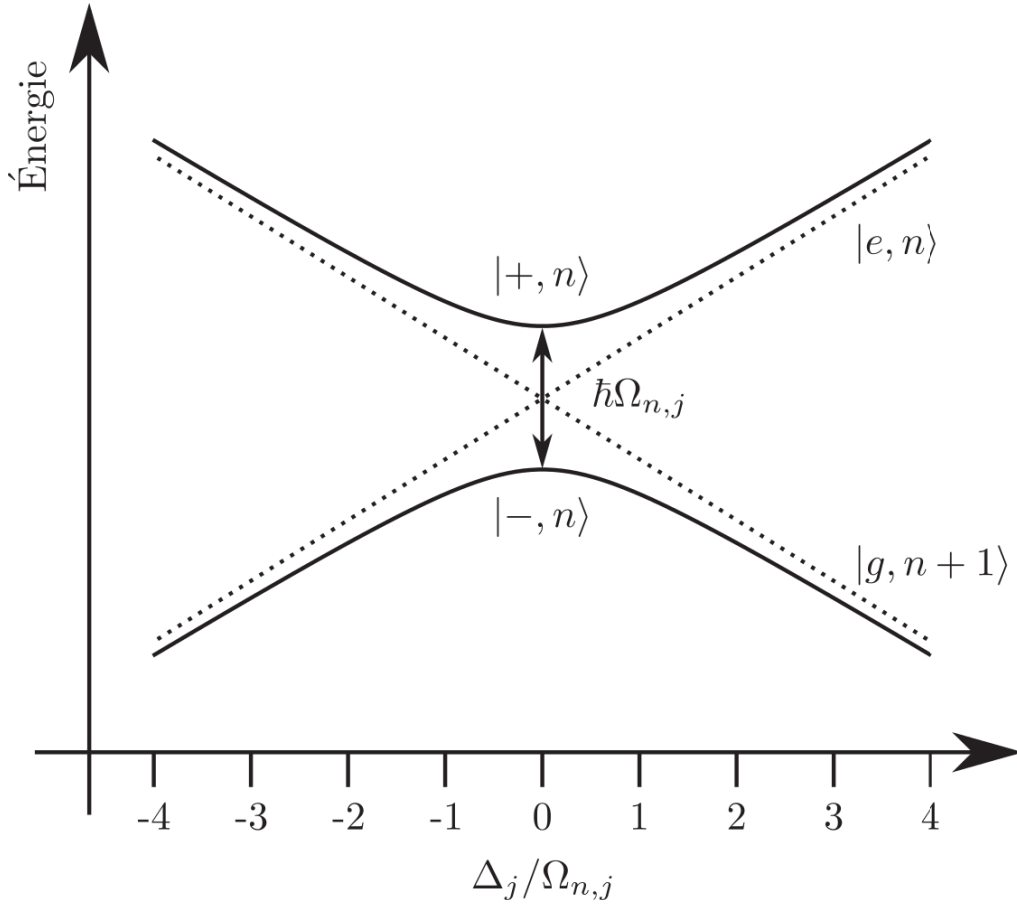


Fig. I.4 Branches d'énergies des états habillés obtenues en variant ω_{at} . La branche supérieure (resp. inférieure) correspond à l'énergie de l'état $|+, n\rangle_{\theta_n}$ (resp. $|-, n\rangle_{\theta_n}$). À résonance ($\Delta_j/\Omega_{n,j} = 0$), les états propres sont les états $|+, n\rangle$ et $|-, n\rangle$. Leurs énergies sont séparées de $\hbar\Omega_{n,j}$. Cet écart croît avec le nombre de photons comme $\sqrt{n}$. Pour les grands désaccords ($|\Delta_j| \gg \Omega_{n,j}$), les états tendent vers les états découplés $|e, n\rangle$ et $|g, n+1\rangle$.

I.3.2 Interaction résonnante

Lorsque l'atome et le champ sont à résonance ($\Delta_j = 0$), les états propres du hamiltonien sont les états maximalement intriqués

$$|+, n\rangle = \frac{1}{\sqrt{2}} (|e, n\rangle + i |g, n+1\rangle), \quad |-, n\rangle = \frac{1}{\sqrt{2}} (|e, n\rangle - i |g, n+1\rangle). \quad (\text{I.3.11})$$

Ces deux états sont des états équatoriaux de la sphère de Bloch de la multiplicité n . Un atome préparé initialement dans l'état $|e, n\rangle$ va donc effectuer des oscillations entre cet état et l'état $|g, n+1\rangle$ appelées *oscillations de Rabi quantiques*. Dans le cas de l'interaction avec un champ classique, les oscillations avaient une pulsation proportionnelle au champ électrique (I.2.10). Ici, la pulsation est proportionnelle à $\sqrt{n+1}$. Il y a donc des oscillations même lorsque le champ est vide. On parle d'oscillations de Rabi du vide et

la pulsation $\Omega_{0,j}$ est appelée *pulsation de Rabi du vide*. Par rapport au cas classique, une autre différence importante est la modification de l'état du champ lorsque l'atome passe de son état excité à l'état fondamental. Lorsque l'oscillation s'effectue dans un champ contenant plusieurs nombres de photons différents, comme un champ cohérent, les différentes oscillations ont lieu à des fréquences différentes, ce qui donne lieu à un effondrement des oscillations. Lorsque le champ a une grande amplitude, les fluctuations du nombre de photons deviennent négligeables devant sa moyenne et le champ est très peu modifié par l'interaction avec l'atome. Il n'y a alors plus d'effondrement et on retrouve les oscillations de Rabi classiques. On ne décrira pas en détail l'interaction résonnante avec un champ cohérent, qui n'est pas utilisée dans les expériences présentées ici et qui a été étudiée en détail dans d'autres travaux de l'équipe [39–41].

I.3.3 Interaction dispersive

Dans le cas de grands désaccords $|\Delta_j| \gg \Omega_{n,j}$,⁷ les états habillés rejoignent asymptotiquement les états découplés $|g, n\rangle$ et $|e, n+1\rangle$. Supposons que $\Delta_j > 0$. On a alors $|+, n\rangle \simeq |e, n\rangle$ et $|- , n\rangle \simeq |g, n+1\rangle$. Effectuons un développement asymptotique au premier ordre en $\Omega_{n,j}/\Delta_j$ des énergies des états habillés (I.4) :

$$E_{e,n} = E_{e,n}^0 + \frac{\hbar\Omega_{0,j}^2}{4\Delta_j}(n+1), \quad (\text{I.3.12a})$$

$$E_{g,n} = E_{g,n}^0 - \frac{\hbar\Omega_{0,j}^2}{4\Delta_j}(n), \quad (\text{I.3.12b})$$

où l'exposant 0 dénote la valeur en l'absence de couplage. Dans ce régime, le hamiltonien peut se mettre, à une constante près, sous la forme approchée

$$\hat{H}_{JC}^{\text{disp}} = \hat{H}_{cj} + \frac{\hbar}{2} \left[\omega_{at} + \frac{\hbar\Omega_{0,j}^2}{2\Delta_j} \left(\hat{a}_j^\dagger \hat{a}_j + \frac{1}{2} \right) \right] \hat{\sigma}_z, \quad (\text{I.3.13})$$

dont les états propres sont les états découplés. L'interaction entre les atomes et le champ se réduit alors à des déplacements lumineux : la présence de n photons induit un décalage de la fréquence de la transition atomique :

$$\Delta\omega_{at}(n) = \frac{\Omega_{0,j}^2}{2\Delta_j} \left(n + \frac{1}{2} \right). \quad (\text{I.3.14})$$

Ce déplacement a lieu même en l'absence de photon : il s'agit alors du *déplacement de Lamb* [42].

Du fait de ce déplacement, l'interaction dispersive d'un atome avec le champ pendant une durée t_{int} conduit à l'accumulation d'une phase relative entre les états $|e\rangle$ et $|g\rangle$ qui vaut

$$\Delta\phi(n) = \phi_{0,j} \left(n + \frac{1}{2} \right), \quad (\text{I.3.15})$$

7. Cette condition ne peut être satisfaite pour tout n , puisque $\Omega_{n,j} \propto \sqrt{n+1}$. On obtient cependant des résultats corrects lorsqu'elle est vérifiée pour tous les n ayant un poids non négligeable dans les états considérés.

avec

$$\phi_{0,j} = \frac{\Omega_{0,j}^2 t_{\text{int}}}{2\Delta_j}. \quad (\text{I.3.16})$$

Le déphasage $\phi_{0,j}$ est appelé *déphasage par photon*. Si l'atome est initialement préparé dans la superposition à poids égaux $(|e\rangle + |g\rangle)/\sqrt{2}$ et le champ contient n photons, l'évolution de l'atome dans le référentiel tournant à ω_{at} est, à une phase près :

$$\frac{1}{\sqrt{2}} (|e\rangle + |g\rangle) \mapsto \frac{1}{\sqrt{2}} (|e\rangle + e^{i(\phi_{0,j}+1/2)} |g\rangle). \quad (\text{I.3.17})$$

Cette évolution est représentée sur la sphère de Bloch de la figure I.5. En plus du déplacement de Lamb, l'atome acquiert une phase proportionnelle au nombre de photons. Cette linéarité est due au fait que nous avons effectué le développement des énergies uniquement au premier ordre. Cette approximation est valide lorsqu'on considère de faibles nombre d'atomes et un grand désaccord. Lorsque le nombre de photons est plus élevé, on doit calculer le déphasage pour chaque valeur du nombre de photons, ce qu'on fera dans la section I.5.4.

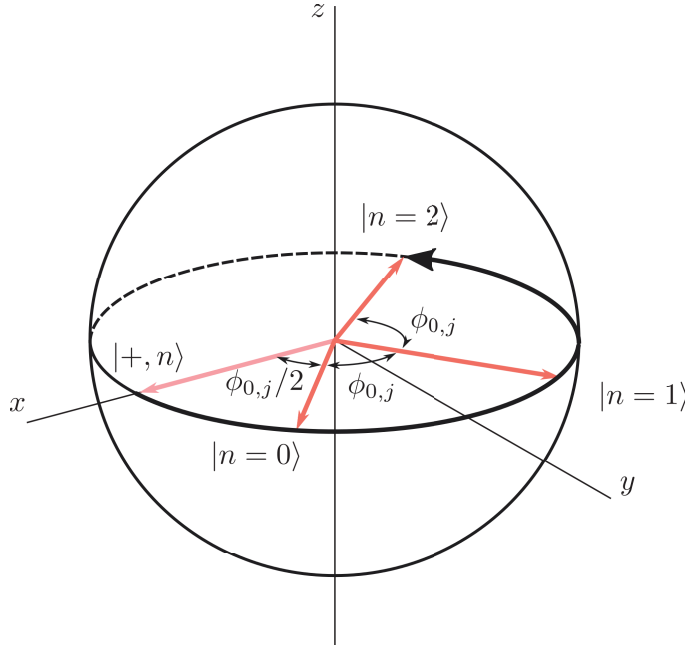


Fig. I.5 Évolution, dans le référentiel tournant à ω_{at} , de l'état de l'atome, initialement préparé dans l'état $|+, n\rangle$, lors de l'interaction avec le champ. Avant l'interaction, l'état est dirigé selon Ox . Les différents états en rouge correspondent aux états de Fock $|n\rangle$, pour n allant de $|0\rangle$ à $|2\rangle$. En plus de la phase $\phi_{0,j}/2$, due au déplacement de Lamb, l'atome acquiert une phase $\phi_{0,j}$ pour chaque photon présent dans le mode j .

I.3.4 Le passage adiabatique

Le passage adiabatique est une technique permettant de transférer l'atome efficacement d'un état vers l'autre, lorsqu'on ne connaît pas le nombre de photons présents dans

le mode. Le principe de la méthode est présenté sur la figure I.6. L'atome est initialement préparé dans l'état $|g\rangle$, en présence de $n + 1$ photons, avec un désaccord $\Delta_0 \gg \Omega_{n,j}$. Ce désaccord est suffisamment important pour que l'état $|g, n + 1\rangle$ du système se confonde avec l'état habillé $|-, n\rangle_{\theta_n^{(0)}}$. On modifie le désaccord de manière à arriver à un désaccord $|\Delta_f| \gg \Omega_{n,j}$, $\Delta_f < 0$. Si cette modification se fait suffisamment lentement, l'état du système suit l'état habillé, de sorte que l'état final est $|-, n\rangle_{\theta_n^{(f)}} \simeq |e, n\rangle$. On est donc passé de l'état $|g, n + 1\rangle$ à l'état $|e, n\rangle$. Nous utilisons cette méthode pour enlever des photons dans la cavité. Elle a l'avantage de fonctionner quel que soit le nombre de photons initial, tant que les conditions sur le désaccord sont vérifiées. L'absorption d'un photon par oscillations de Rabi, quant à elle, ne fonctionne que pour un état de Fock donné et requiert une calibration du temps d'interaction.

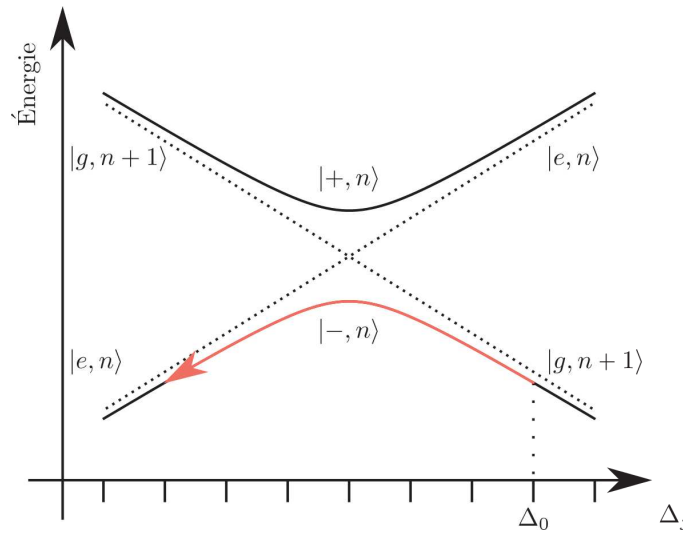


Fig. I.6 Principe du passage adiabatique. Le système est initialement préparé dans l'état $|g, n + 1\rangle$, qui se confond, pour un désaccord Δ_0 suffisamment grand, à l'état habillé de la branche inférieure. Si le désaccord est modifié suffisamment lentement jusqu'à atteindre le désaccord Δ_f , le système suit l'état habillé (flèche rouge). Il termine donc son évolution dans un état habillé proche de $|e, n\rangle$.

I.3.5 Pertinence du modèle d'un atome couplé à un mode du champ

Dans les sections précédentes, nous avons introduit les descriptions quantiques d'un atome à deux niveaux et d'un mode du champ électromagnétique puis nous avons présenté certains phénomènes physiques apparaissant lorsqu'on les couple l'un à l'autre. En général, un atome n'interagit pas avec un seul mode, mais avec un ensemble de modes, qui peut être un continuum dans le cas de l'espace vide. Dans ce cas, l'interaction ne donne pas lieu à des oscillations de Rabi, mais au phénomène d'émission spontanée [43]. Pour que l'interaction soit cohérente, il est nécessaire de découpler l'atome et le champ des degrés de liberté de l'environnement. Par ailleurs, pour que les phénomènes présentés soient mesurables, il faut que les fréquences d'interaction soient dominantes par rapport aux taux de décohérence. L'électrodynamique quantique en cavité (CQED) est une discipline qui

utilise des cavités résonnantes pour augmenter le couplage entre atomes et lumière, afin de mettre en évidence des effets quantiques. Les expériences d'électrodynamique quantique en cavité du laboratoire Kastler-Brossel s'inscrivent dans cette perspective. Elles ont fait l'objet de nombreux développements technologiques afin de parvenir au régime dans lequel un atome interagit de façon considérable avec un photon, avec une fréquence de Rabi largement supérieure aux largeurs des états de l'atome et de la résonance de la cavité. Dans la section suivante, nous allons présenter le dispositif permettant de se mettre dans ce régime.

I.4 Le dispositif expérimental

Nos expériences reposent sur deux éléments essentiels, qui se combinent pour atteindre le régime de couplage fort. En guise d'atome à deux niveaux, on utilise des atomes de Rydberg, préparés dans un état circulaire qui sera défini dans la section I.4.2 [44]. Ces atomes ont pour propriété d'avoir des transitions micro-ondes isolées, qui permettent de les traiter comme des atomes à deux niveaux, dont le dipôle est très élevé, ce qui les fait interagir très fortement avec la lumière. Ils ont par ailleurs de très grands temps de vie, de l'ordre de quelques dizaines de millisecondes, ce qui permet de leur faire parcourir plusieurs dizaines de centimètres dans un jet atomique, sans qu'une proportion notable d'entre eux ne se désexcite. Ils nécessitent cependant d'être manipulés dans des conditions cryogéniques, à cause de leur grande sensibilité au rayonnement thermique micro-onde.

Le rôle de l'oscillateur harmonique est quant à lui joué par un mode d'une cavité micro-onde supraconductrice de très grande finesse, dans une configuration Fabry-Perot ouverte. Ces cavités sont fabriquées au laboratoire grâce à une technique mise au point au CEA [45]. Entre 2006 et 2014, deux paires de miroirs ont été utilisées. Les temps de vie des cavités formées de ces miroirs étaient respectivement de 130 ms et 65 ms. Comme les atomes, ces cavités nécessitent un environnement cryogénique.

Le dispositif sur lequel j'ai travaillé a été conçu pendant le travail de thèse de Sébastien Gleyzes [46]. Il a par la suite été au cœur de nombreux travaux de thèse portant sur les sauts quantiques de la lumière [47], l'effet Zénon quantique [48], la préparation et la reconstruction d'états exotiques du champ et l'étude de leur décohérence [29], la rétroaction quantique [49], la mesure adaptative [50], la métrologie quantique [51], *etc.* Avant mon arrivée au laboratoire, il a subi un déménagement et une reconstruction qui ont permis de rénover plusieurs parties du montage. Nous avons en particulier passé beaucoup de temps à mettre en place deux cavités, ce qui était l'objectif initial du montage, mais qui n'avait jamais pu être réalisé. Ces cavités ont rendu possible l'étude d'états délocalisés du champ électromagnétique, présentée dans ce mémoire.

De nombreuses techniques ont été mises au point pour ces expériences de CQED. On se contentera dans cette partie de présenter le dispositif et d'exposer les principaux éléments nécessaires à la compréhension du présent travail. Les détails du montage et des méthodes utilisées sur le dispositif peuvent être trouvés dans les travaux de thèse des doctorants précédents [30, 32, 34, 46, 52–56].

Le dispositif

Les principaux éléments de notre dispositif expérimental sont présentés sur la figure I.7. Un four chauffé à 200 °C fournit un jet atomique de ^{85}Rb . Les atomes sont sélectionnés en vitesse à 250 m/s par des lasers pulsés, avant d'être excités dans un état de Rydberg circulaire dans la *boîte à circulariser* B. Ils entrent ensuite dans une boîte d'isolation du champ magnétique, dans laquelle se déroulent les phénomènes d'interaction avec le champ qui nous intéressent. Dans cette boîte, ils interagissent successivement avec les deux cavités C_1 et C_2 . Leur état atomique peut être contrôlé avant et après chaque cavité à l'aide des *zones de Ramsey* R_1 , R_2 et R_3 , dans lesquelles on peut leur appliquer des impulsions micro-ondes fournies par la source S_R . Deux autres sources micro-ondes S_1 et S_2 permettent d'injecter des champs cohérents dans C_1 et C_2 respectivement. Enfin, les

atomes sont détectés à la sortie de la boîte d'écrantage par ionisation dans le Channeltron D.

Toute la portion de trajet de B à D a lieu dans un cryostat. Celui-ci comporte un premier étage de refroidissement à l'azote liquide à la température de 77 K, qui écrante tout le dispositif. Un second étage contenant un bain d' ^4He atteint quant à lui 4.2 K. En pompant sur un réservoir d' ^4He alimenté par un capillaire, on atteint la température de 1.6 K. Enfin, un dernier étage cryogénique, connecté au premier grâce à un échangeur contenant de l' ^3He , écrante le dispositif et définit sa température. Celle-ci peut être maintenue à 1.6 K par le pompage de l' ^4He , ou atteindre 0.8 K par pompage de l' ^3He . Tout ce dispositif est pompé par des pompes turbomoléculaires et par cryopompage à une pression inférieure à 10^{-7} mbar.

Dans la suite de cette section, on va expliquer plus en détail la géométrie des cavités, ainsi que la manière dont on mesure et contrôle leur fréquence de résonance. On présentera enfin les propriétés des atomes de Rydberg circulaires et les techniques permettant de les préparer et de les détecter.

1.4.1 Des cavités de grande finesse

Les cavités utilisées dans nos expériences sont des cavités micro-ondes supraconductrices. Les miroirs les constituant sont fabriqués par pulvérisation cathodique d'une couche de niobium de 12 μm sur un substrat en cuivre. Le niobium a une température critique $T_C^{\text{Nb}} = 9.3$ K. Une fois placé dans un environnement cryogénique sous T_C^{Nb} , son absorption s'annule à fréquence nulle. Pour des fréquences finies, l'absorption n'est pas nulle à T_C^{Nb} et continue de décroître jusqu'à environ 1 K. De ce fait, nos cavités ont un temps de vie deux fois plus élevé à 0.8 K qu'à 1.6 K. Il semble donc intéressant de toujours travailler à la température la plus faible. Cependant, à cause d'un incident arrivé au cours de la thèse, le capillaire remplissant le réservoir d' ^3He a été partiellement bouché. Il n'était donc plus possible de pomper sur ce dernier pendant plus d'une heure, ce qui limite fortement la prise de données. Les expériences présentées ici ont donc été réalisées à l'une ou l'autre des températures, selon la contrainte la plus importante entre le temps de vie des cavités et la quantité de données à acquérir.

Les miroirs sont montés dans une configuration de Fabry-Pérot ouverte, nécessaire pour pouvoir appliquer un champ électrique permettant de maintenir les états de Rydberg des atomes. Les miroirs, représentés sur la figure I.8, sont séparés de 2.8 cm de sorte que la fréquence du mode utilisé (TEM_{900}) soit de 51.1 GHz, proche de la fréquence de résonance des atomes. Dans cette configuration, deux modes, de polarisations orthogonales, coexistent dans la cavité. Afin de lever la dégénérescence entre ces deux modes, les miroirs sont toroïdaux : ils ont deux rayons de courbure distincts 39.4 mm et 40.6 mm qui permettent d'obtenir une séparation de 1.26 MHz.

Afin de contrôler finement cette distance, les atomes sont installés dans des blocs cavités (Fig. I.9), dont la longueur est définie par l'épaisseur de cales piézoélectriques. Par application d'une différence de potentiel comprise entre -2500 V et 2500 V, la fréquence de la cavité peut être accordée sur une gamme de ± 5 MHz, ce qui correspond à ± 2.5 μm . La distance de la cavité, une fois installée dans le cryostat, après pompage et refroidissement à température cryogénique, doit donc être correcte à quelques microns près pour tomber dans la gamme d'accord des cales piézoélectriques. Dans ce but, on place contre les

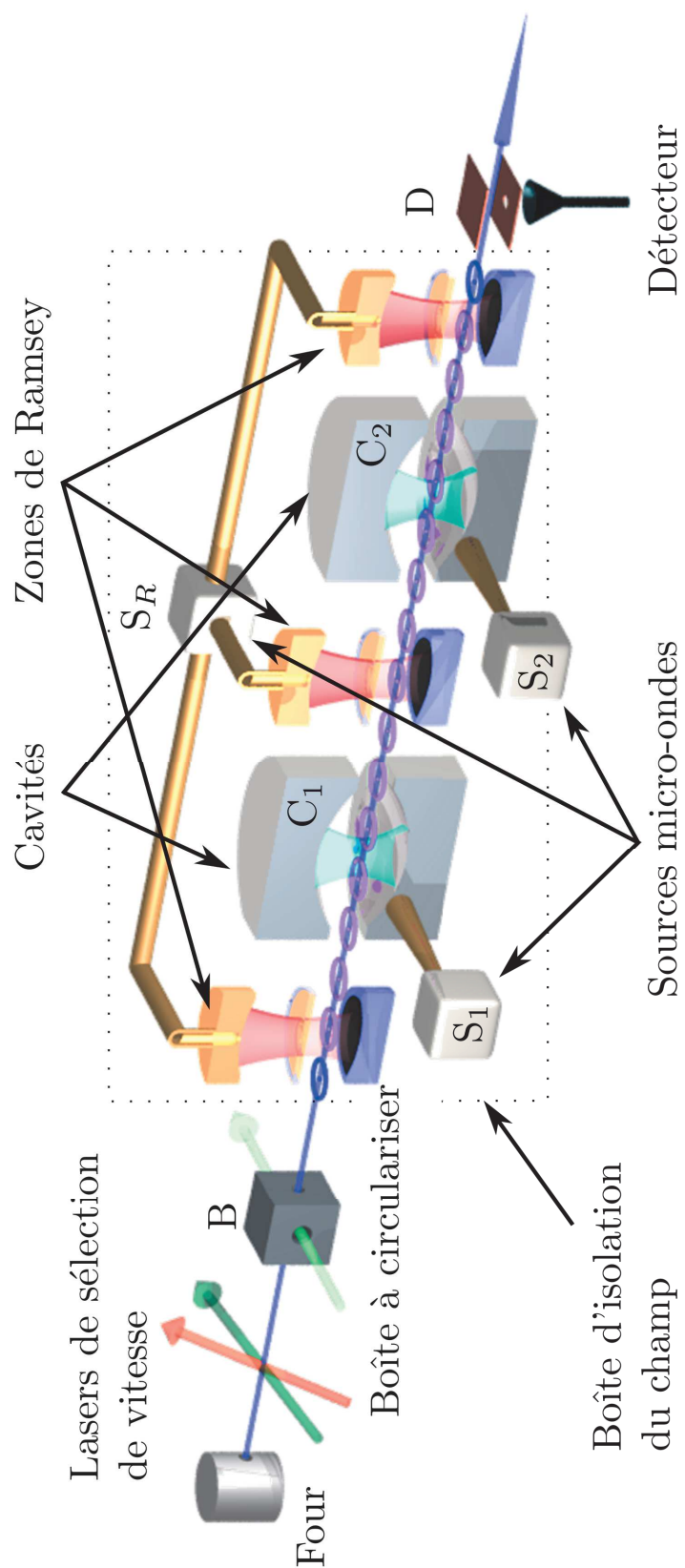


Fig. I.7 Schéma du dispositif expérimental (description dans le texte)

cales piézoélectriques deux cales en cuivres, qui participent à la définition de la longueur de la cavité. Si la fréquence de la cavité, qu'on ne peut mesurer qu'à froid à l'aide des atomes de Rydberg, est trop faible, on ronge ces cales à l'acide, afin d'enlever les quelques microns de trop [46]. Cette phase d'accord nécessite donc de mesurer l'épaisseur des cales à l'aide d'une cavité test par des mesures micro-ondes à chaud, d'installer les cavités dans le cryostat, de mesurer la fréquence de résonance du bloc cavité, puis de réchauffer le cryostat, de démonter les cavités, de ronger les cales et de refaire une mesure de leur épaisseur pour contrôler qu'on a enlevé la bonne quantité. Une telle opération est pénible et nécessite au moins 1 mois de travail. À cause de l'extrême précision mécanique nécessaire et du vieillissement du dispositif, dont de nombreuses pièces ont dû être changées, 8 cycles d'accord ont été nécessaires avant de parvenir à mettre les cavités à résonance.

Le dispositif actuel comprend deux cavités C_1 et C_2 . Leur temps de vie ont été mesurés grâce à une méthode expliquée dans [46] et valent respectivement $T_{c1} = 20 \pm 2$ ms et $T_{c2} = 50 \pm 2$ ms à $T = 0.8$ K.

Forme du mode

Le mode soutenu par les miroirs est un mode gaussien TEM_{900} , représenté schématiquement sur la figure I.8. Les atomes passent entre les miroirs dans le ventre central, de façon à voir le champ d'amplitude maximale. La forme du mode peut s'écrire [27] :

$$f(r, z) = f_z(z) e^{-\frac{r^2}{w(z)^2}}, \quad (I.4.1)$$

où r et z sont les coordonnées cylindriques par rapport à l'axe de la cavité et où $f_z(z)$ est la forme de l'onde stationnaire dans l'axe de la cavité. Cette fonction varie selon une direction orthogonale à l'axe du jet atomique. Elle reste donc constante pendant le trajet des atomes, égale à 1 lorsqu'ils passent par le centre de la cavité. $w(z)$ est la largeur du mode gaussien :

$$w(z) = w_0 \sqrt{1 + \left(\frac{\lambda z}{\pi w_0^2} \right)^2}, \quad (I.4.2)$$

avec $w_0 \simeq 6$ mm la largeur du mode au centre de la cavité et $\lambda = 5.87$ mm la longueur d'onde.

Le volume du mode vaut :

$$\mathcal{V} = \frac{\pi w_0^2 l}{4} = 770 \text{ mm}^3, \quad (I.4.3)$$

où l est la distance entre les miroirs. Cette valeur est assez faible et proche de λ^3 . Il en résulte une amplitude du champ importante :

$$\mathcal{E}_{0,j} = \sqrt{\frac{\hbar \omega_{cj}}{2\epsilon_0 \mathcal{V}_j}} = 1.5 \cdot 10^{-3} \text{ V/m}. \quad (I.4.4)$$

Sources micro-ondes

Afin de réaliser des déplacements du champ contenu dans les cavités, on envoie un champ qui se couple aux modes des cavités par diffraction sur les bords des miroirs.

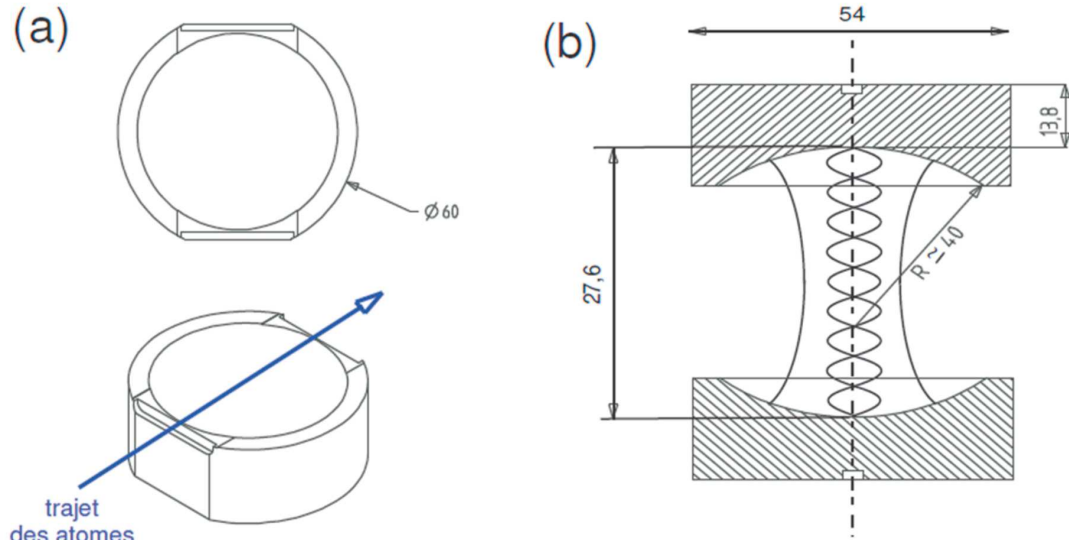


Fig. I.8 (a) Vue de dessus et de trois quarts des miroirs utilisés. (b) Configuration de la cavité, vue de profil. Pour la lisibilité, la distance entre les miroirs est exagérée par rapport à leur épaisseur. La largeur du mode $w(z)$ et la fonction axiale du mode $f_z(z)$ sont indiquées.

Le signal émis par une source micro-onde à 12.8 GHz est envoyé sur une diode PIN contrôlée par un signal TTL permettant de laisser passer ou de bloquer la micro-onde, afin de contrôler la durée et l'instant d'injection. Il est ensuite guidé vers un quadrupleur de fréquence situé au sommet du cryostat. Un atténuateur calibré situé juste après ce dernier permet de contrôler finement la puissance. Le signal sortant à 51.1 GHz descend ensuite dans le cryostat à travers des guides d'onde et sort par un trou situé sur le côté de la cavité. On dispose d'une voie d'injection de ce type pour chaque cavité. On note respectivement S_1 et S_2 les deux sources associées à C_1 et C_2 . Pour la stabilité, les sources sont verrouillées en fréquence sur une horloge atomique.

Mesure de la fréquence de résonance

Les fréquences de résonance des cavités sont deux des paramètres les plus importants de notre expérience. Elles déterminent le désaccord des atomes avec les champs, les fréquences à choisir pour injecter un champ classique dans les cavités et la fréquence d'évolution des états préparés. Pour mesurer la fréquence de résonance d'une cavité, on allume la source S_j pendant quelques millisecondes. On envoie ensuite des atomes après quelques centaines de microsecondes. Si la source est désaccordée de la cavité, il n'y a plus de micro-onde dans le dispositif quand les atomes passent et ceux-ci ne ressentent aucun effet. En revanche, si la source est à résonance, le champ injecté dans la cavité reste suffisamment longtemps pour interagir avec les atomes. On répète cette opération en variant la fréquence d'injection dans le but d'observer une variation de la probabilité de trouver les atomes dans l'état initial, signe qu'on a trouvé la fréquence de résonance.

Lorsqu'on ne connaît pas initialement la fréquence de la cavité, le choix de la durée

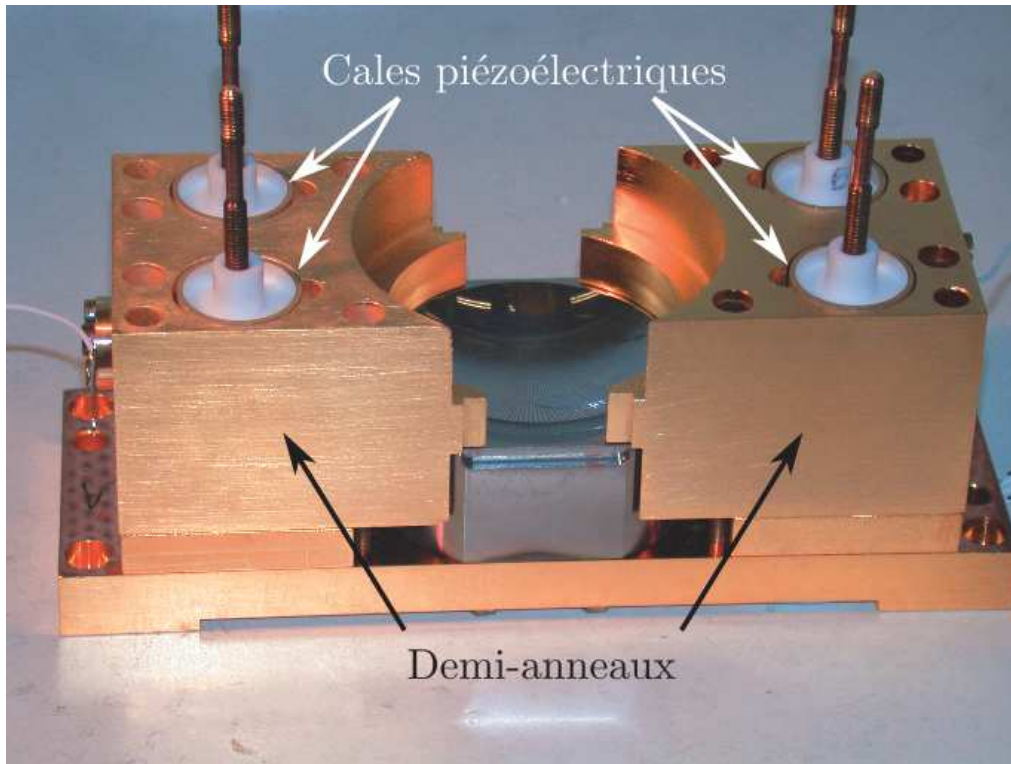
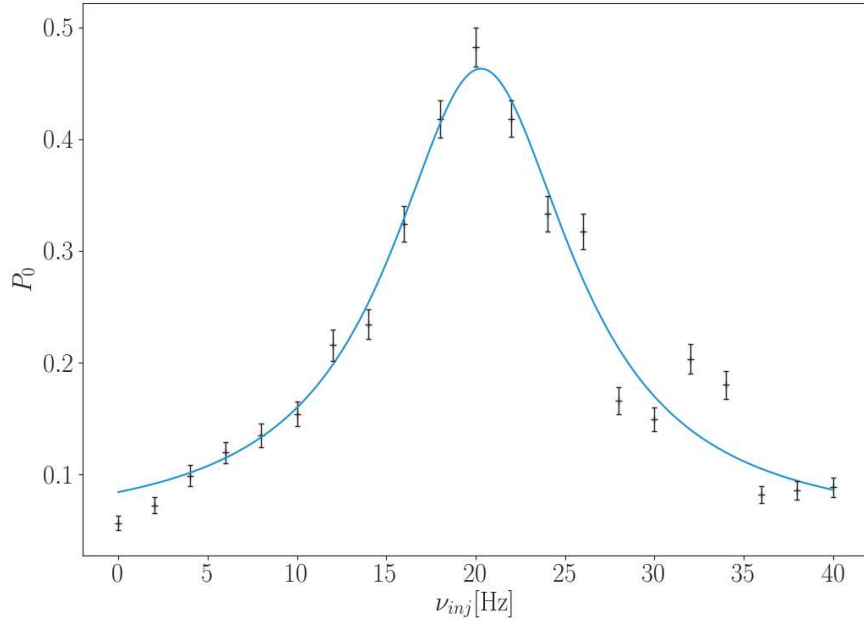
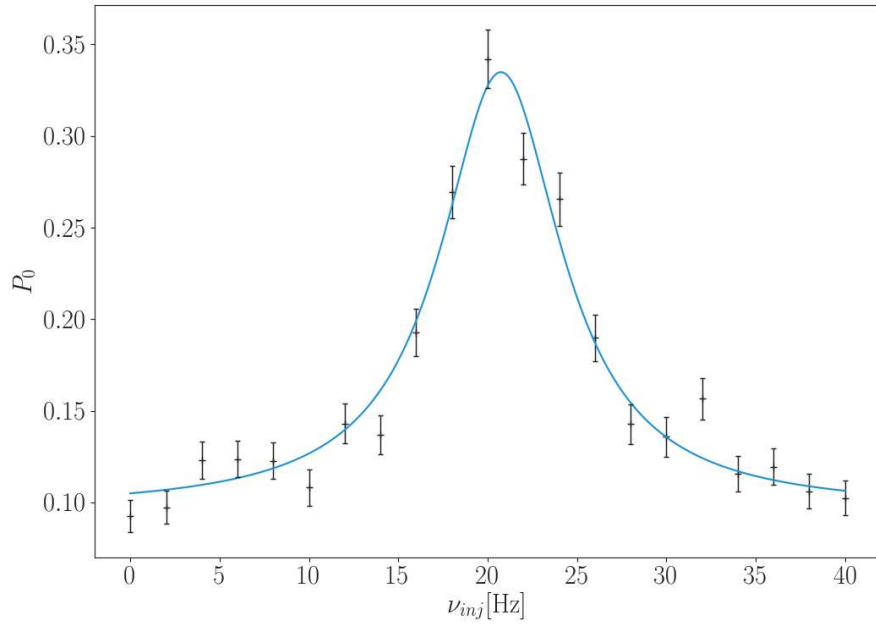


Fig. I.9 Photographie de la partie supérieure du bloc cavité, qui vient se fixer en bas du cryostat, dans la boîte d'écrantage du champ. On voit en particulier un des miroirs de la cavité et les demi-anneaux permettant de recycler une partie des photons sortant du mode de la cavité et de rendre le champ électrique plus homogène au niveau de la cavité. De ces demi-anneaux affleurent les cales piézoélectrique permettant de régler la distance entre les miroirs.

de l'impulsion nécessite de faire un compromis entre, d'une part, le fait qu'elle doit être suffisamment longue pour injecter assez de champ dans la cavité et, d'autre part, le fait qu'elle doit être assez large en fréquence pour pouvoir varier la fréquence de la source en faisant des grands pas, afin de scanner de larges plages de fréquence. Une fois qu'on a déterminé la fréquence de la cavité, on peut effectuer des injections plus longues afin de réduire la largeur spectrale de Fourier des impulsions et d'affiner la mesure de fréquence jusqu'à être limité par la finesse de la cavité. Les figures I.10a et I.10b présentent les spectres des cavités ainsi obtenus. La largeur des raies est de quelques hertz seulement, ce qui correspond à des facteurs de qualité de l'ordre de 10^{10} ! Pour atteindre de telles performances, il faut avoir un dispositif très stable. En particulier, on a travaillé en stabilisant la température des cavités et la pression du bain d'hélium pour une partie des expériences présentées ici. Il arrive cependant que la fréquence des cavités varie entre des expériences, voire pendant une mesure de longue durée. Cette variation peut être abrupte ou contenir des sauts. L'échelle de variation pendant quelques heures est de l'ordre de quelques dizaines de hertz, ce qui correspond à des déplacements des miroirs de quelques dizaines de picohertz.



(a) Probabilité de transition de l'atome en fonction de la fréquence de la source micro-onde pour une injection dans C_1 . La courbe bleue est un ajustement par une lorentzienne de largeur 6.5 ± 0.3 Hz.



(b) Probabilité de transition de l'atome en fonction de la fréquence de la source micro-onde pour une injection dans C_2 . La courbe bleue est un ajustement par une lorentzienne de largeur 4.1 ± 0.3 Hz.

Fig. I.10 Spectres des cavités obtenus en mesurant l'effet du champ contenu dans les cavités sur une transition atomique.

Temps d'interaction effectif

Dans la section I.3, on a fait l'hypothèse que l'atome se trouvait au centre du mode, pour $f_j(\mathbf{r}) = 1$. Nos atomes traversent les cavités à $v = 250$ m/s. La fréquence de Rabi qu'ils voient dépend de la position où ils se trouvent dans le mode et donc de l'instant. Lorsque l'atome est sur les bords du mode, il voit une fréquence de Rabi plus faible et interagit donc moins fortement. On peut introduire un temps effectif pour prendre en compte cet effet.

Lorsque l'atome est à résonance, les hamiltoniens à des temps différents commutent, l'opérateur d'évolution du système peut donc se mettre sous la forme :

$$\hat{U}_\infty = \exp \left[-\frac{i}{\hbar} \int_{-\infty}^{\infty} \hat{H}(t) dt \right] = \exp \left[\frac{i}{\hbar} \hat{H}(0) T_{eff}^r \right], \quad (\text{I.4.5})$$

où

$$T_{eff}^r = \int_{-\infty}^{\infty} f_j(vt) dt = \frac{\sqrt{\pi} w_0}{v} = 42,5 \text{ } \mu\text{s}. \quad (\text{I.4.6})$$

Si l'atome est mis à résonance uniquement pendant une durée T_R , le temps effectif s'écrit :

$$T_{eff} = \int_{-T_R/2}^{T_R/2} f_j(vt) dt. \quad (\text{I.4.7})$$

La figure I.11 montre le temps effectif en fonction de la durée de l'impulsion de Rabi T_R . Ces temps sont approximativement égaux pour des durées d'interaction inférieures à $10 \text{ } \mu\text{s}$.

Lorsque l'atome est très désaccordé, les hamiltoniens à des instants différents commutent aussi et on a alors un temps d'interaction effectif :

$$T_{eff}^d = \int_{-\infty}^{\infty} f_j^2(vt) dt = \sqrt{\frac{\pi}{2}} \frac{w_0}{v} = 30,1 \text{ } \mu\text{s}. \quad (\text{I.4.8})$$

I.4.2 Atomes de Rydberg

Dans nos expériences le rôle de l'atome à deux niveaux est joué par des atomes de ^{85}Rb excités dans des états de Rydberg circulaires. On va voir dans un premier temps comment leurs propriétés en font des candidats idéals pour l'électrodynamique quantique en cavité. On verra ensuite comment on peut préparer des échantillons d'atomes de Rydberg et les détecter.

a) Propriétés des atomes de Rydberg circulaires

Un *état de Rydberg* est un état électronique de grand nombre quantique principal n , typiquement de l'ordre de quelques dizaines. Cet état est dit *circulaire* si par ailleurs ses nombres quantiques orbital et magnétique sont maximaux, c'est-à-dire si $l = |m| = n - 1$. Dans la suite, on notera $|nC\rangle$ l'état circulaire de nombre quantique principal n et de moment angulaire $m = n - 1$. Les autres états de même n mais ne vérifiant pas la condition pour être circulaires sont dits *elliptiques*. Les états circulaires ont des fonctions d'onde toroïdales dont le rayon a pour ordre de grandeur $r_n = n^2 a_0$, où a_0 est le rayon de Bohr. La distance entre le noyau et l'électron excité peut donc être plusieurs ordres de grandeur

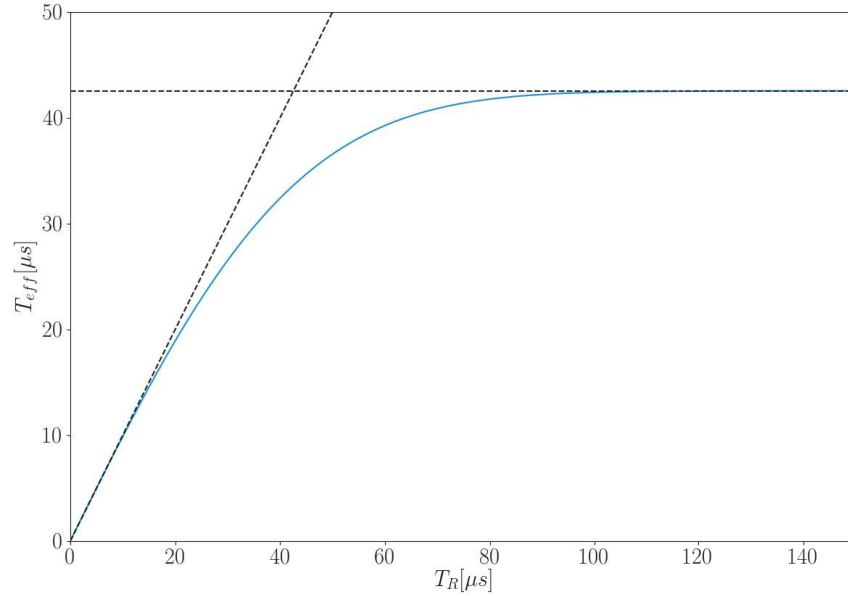


Fig. I.11 Temps effectif d'interaction de l'atome avec la cavité en fonction de la durée T_R de l'impulsion de Rabi. La ligne pointillée horizontale indique le temps effectif final. L'autre ligne pointillée est la droite d'équation $T_{eff} = T_R$. Pour $T_R \leq 10 \mu s$, le temps effectif est égal à T_R à 1.5% près.

plus élevée que dans un atome standard, ce qui leur confère des dipôles de transition très élevés. Les atomes de Rydberg peuvent être raisonnablement décrits comme des atomes hydrogénoïdes et l'énergie de l'état nC peut donc être approximée par :

$$\omega_n = \frac{R_{Rb}}{\hbar} \left(\frac{1}{n^2} - \frac{1}{(n+1)^2} \right), \quad (I.4.9)$$

où $R_{Rb} = R_y(85m_p)/(m_e + 85m_p)$. Pour les valeurs $n \simeq 50$ utilisées dans nos expériences, les transitions entre atomes sont des transitions micro-ondes. L'élément de matrice d_n du dipôle atomique entre les deux états atomiques $|nC\rangle$ et $|(n-1)C\rangle$ vaut, dans le modèle semi-classique :

$$d_n = \frac{a_0 e (n-1)^2}{\sqrt{2}}, \quad (I.4.10)$$

où e est la charge élémentaire.

On peut lever la dégénérescence entre états de Rydberg de même n en appliquant un champ électrique directeur. Dans ce cas, m est toujours un bon nombre quantique, mais l ne l'est plus. Il est remplacé par le nombre quantique parabolique n_1 , qui prend des valeurs entières entre 0 et $n - |m| - 1$. La figure I.12 montre les diagrammes d'énergie obtenus dans ce cas pour les états $|g\rangle \equiv |50C\rangle$ et $|e\rangle \equiv |51C\rangle$ utilisés dans nos expériences. La levée de dégénérescence entre $|nC\rangle$ et les deux états elliptiques les plus proches est due à l'effet Stark linéaire et vaut de l'ordre de 100 MHz/(V/cm). On peut donc très bien séparer les

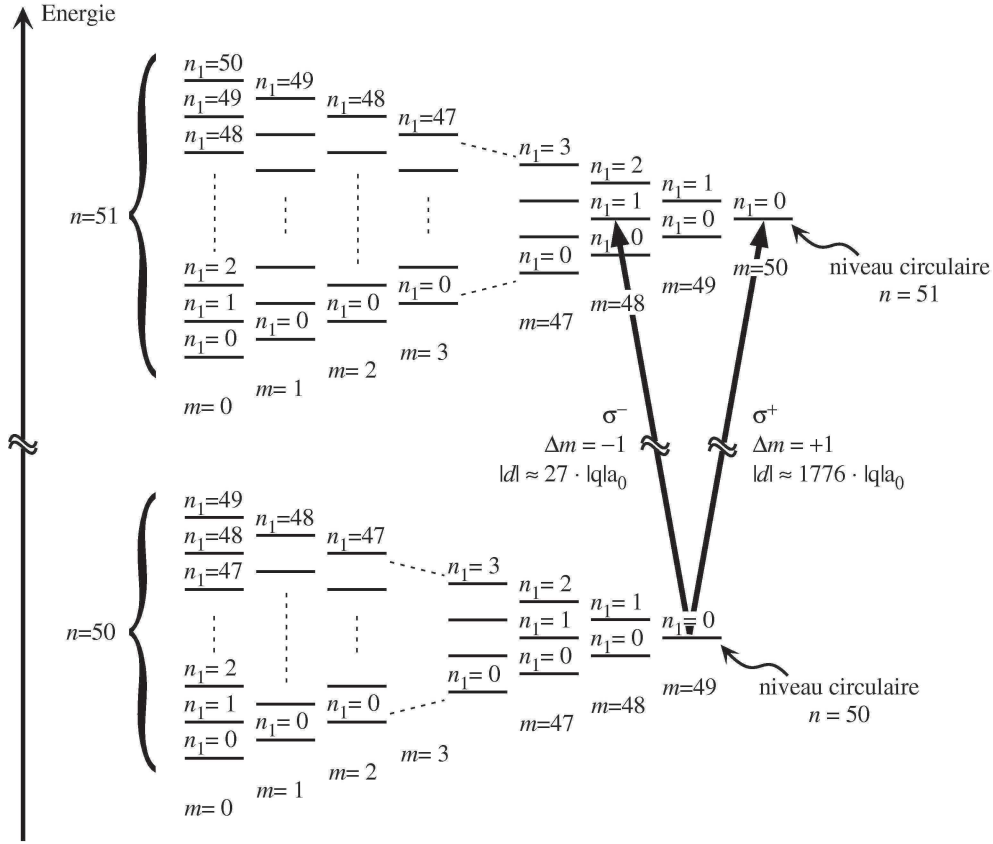


Fig. I.12 Diagramme d'énergie des niveaux $|e\rangle \equiv |51C\rangle$ et $|g\rangle \equiv |50C\rangle$, dans un champ électrique statique. Seuls les états de m positifs sont représentés.

transitions. Par contre il existe une transition de l'état $|g\rangle$ vers l'état $|n_1 = 1, m = 48\rangle$, qui n'est pas soumise à l'effet Stark linéaire. Cette transition a un dipôle bien plus faible que celui de la transition qui nous intéresse. On peut donc considérer la transition $|e\rangle \leftrightarrow |g\rangle$ comme fermée.

Les états circulaires sont insensibles à l'effet Stark linéaire. Ils sont cependant soumis à l'effet Stark quadratique. La transition $|e\rangle \leftrightarrow |g\rangle$ subit à cause de cet effet un décalage de $-255 \text{ kHz}/(\text{V}/\text{cm})^2$. Ce décalage nous permet de régler la fréquence de la transition atomique de manière à mettre l'atome à résonance ou hors résonance à l'aide d'un champ électrique.

Temps de vie

Comme on peut le voir sur le diagramme I.12, le seul état vers lequel $|nC\rangle$ peut relaxer par une transition dipolaire électrique est l'état $(n-1)C$. À température nulle, le taux d'émission spontanée vers cet état est :

$$\Gamma_{0,n} = \frac{\omega_n^3 |d_n|^2}{3\pi\epsilon_0 \hbar c}. \quad (\text{I.4.11})$$

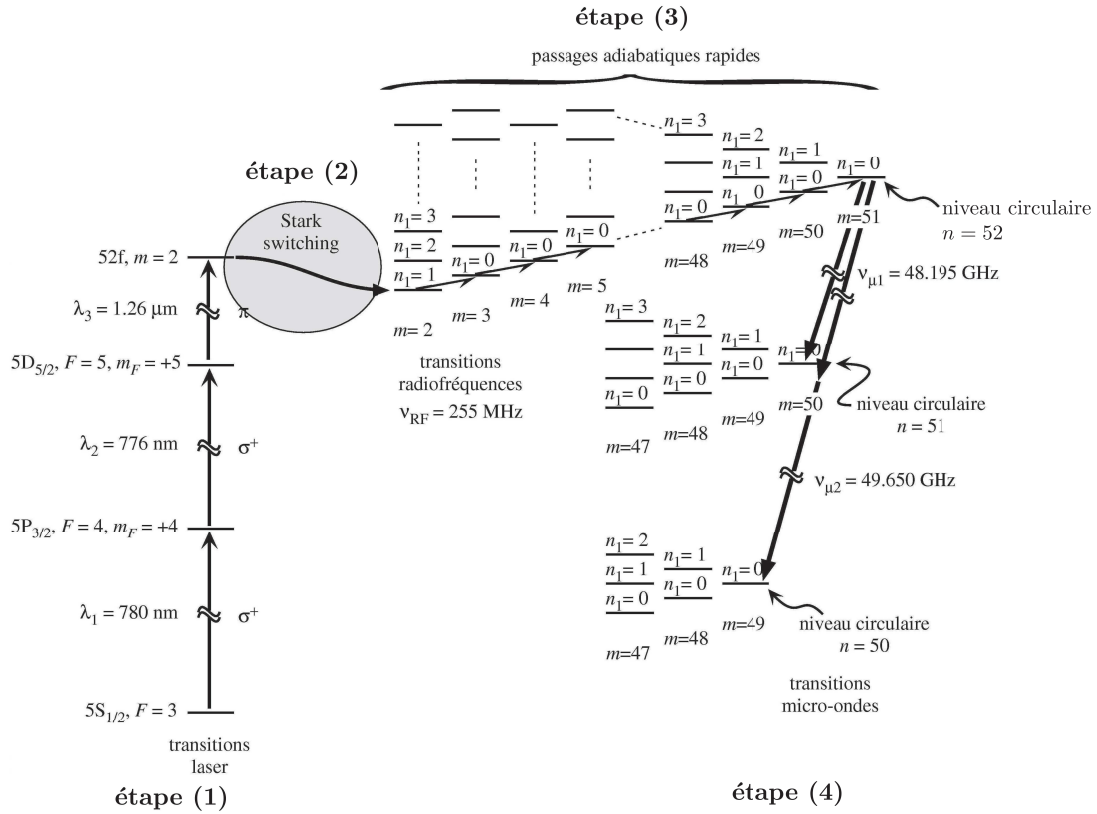


Fig. I.13 Schéma d'excitation dans les états $|g\rangle$ ($n = 50$) et $|e\rangle$ ($n = 51$). (1) Excitation optique vers un niveau de grand n . (2) Glissement Stark. (3) Passage adiabatique vers un niveau circulaire. (4) Purification vers l'état $|e\rangle$ ou $|g\rangle$.

Pour $n = 51$ et $n = 50$ on en déduit que les temps de vie sont respectivement 31,5 ms et 28,5 ms. Ces niveaux sont extrêmement stables. Ceci peut s'interpréter dans une vision semi-classique comme une conséquence du fait que l'électron a une accélération faible dans l'état circulaire, d'autant plus que le rayon de l'orbitale est grand. Dans nos expériences, les atomes mettent environ 1100 μs à parcourir la distance séparant la boîte à circulariser, où ils sont préparés, du détecteur. La probabilité de relaxer est donc de l'ordre de 3,5%, en l'absence de champ thermique.

b) Préparation sélective d'échantillons atomiques

Les atomes utilisés dans nos expériences sont des atomes de ^{85}Rb . Dans un premier temps, ces atomes doivent être excités dans l'état $|e\rangle$ ou $|g\rangle$. On souhaite par ailleurs sélectionner la vitesse des atomes et faire en sorte d'avoir des *échantillons atomiques* de durée fixée à des instants bien précis. Le schéma d'excitation atomique est présenté sur la figure I.13.

La première étape de la préparation est une série de trois excitations laser. L'atome, initialement pompé dans l'état $|5S_{1/2}, F=3\rangle$ absorbe deux photons optiques et un photon

infrarouge qui l'amènent dans l'état de Rydberg $|52f, m = 2\rangle$. L'atome doit encore absorber 49 photons de polarisation σ^+ pour être dans un état circulaire. On commence par effectuer un passage adiabatique en amenant le champ électrique de 0 V/cm à 2.5 V/cm, en présence d'un champ magnétique. L'état $|51f, m = 2\rangle$ est alors branché sur l'état $|n = 51, m = 2, n_1 = 1\rangle$ de la multiplicité Stark et la dégénérescence de la multiplicité est levée. Afin d'apporter les quantas de moment angulaire restants, on effectue un passage adiabatique, pendant lequel l'atome absorbe 49 photons radio-fréquence, préparés avec une polarisation σ_+ par quatre électrodes. Cette procédure, inspirée du passage adiabatique présenté dans la section I.6 consiste à balayer le champ électrique en présence d'une impulsion radio-fréquence de manière à traverser la résonance. Elle est présentée en détail dans la thèse de P. Nussenzveig [57]. Après cette étape, l'atome se trouve dans $|52C\rangle$.

Une dernière étape consiste à effectuer une impulsion micro-onde ramenant l'atome dans l'état $|e\rangle$ ou $|g\rangle$. Cette impulsion est appelée purification. Elle est suivie d'une tension élevée appliquée sur une électrode située après la boîte à circulariser de façon à ioniser les atomes restés dans un état de la multiplicité 52. On a donc préparé $|e\rangle$ ou $|g\rangle$ de façon très pure.

Une fois les états circulaires préparés, il est nécessaire de toujours maintenir un champ électrique directeur, sans quoi les états de la multiplicité Stark se couplent les uns aux autres. Le potentiel de tous les éléments situés le long du trajet atomique est donc contrôlé séparément, à l'aide d'une quarantaine d'électrodes, contrôlées depuis l'extérieur du cryostat.

Sélection en vitesse

Afin de contrôler l'interaction entre les atomes et le champ, on veut faire en sorte d'avoir les atomes au centre du mode à un instant bien déterminé. Dans ce but, on les prépare de façon pulsée. On sélectionne par ailleurs les atomes ayant une vitesse déterminée. On prépare ainsi des échantillons atomiques contenant des atomes présents dans un volume d'excitation à un instant donné. En dehors des échantillons, les atomes restent dans leur état fondamental et n'interagissent donc pas du tout avec le champ.

En l'absence de sélection en vitesse, les atomes ont une distribution de vitesse maxwellienne à la sortie du four. La figure I.15 montre la distribution de vitesse mesurée au niveau du détecteur lorsqu'on excite les atomes dans l'état $|52F\rangle$. Celle-ci n'est pas maxwellienne, puisque les atomes les plus lents ont moins de chance d'atteindre le détecteur, d'une part parce qu'ils sont en chute libre, d'autre part à cause de la relaxation.

Le principe de la sélection en vitesse repose sur la structure hyperfine des états $5S_{1/2}$ et $5P_{3/2}$ du rubidium, présentée sur la figure I.14 [58]. L'état $5S_{1/2}$ possède deux sous-niveaux hyperfins : $|5S_{1/2}, F = 2\rangle$ et $|5S_{1/2}, F = 3\rangle$, séparés de 3 GHz. À température ambiante, ces deux niveaux sont identiquement peuplés. Le niveau $5P_{3/2}$ possède 4 sous-niveaux, pour lesquels $F' \in [1, 4]$, séparés de quelques dizaines de MHz. Notre sélection en vitesse consiste à pomper les atomes ayant la bonne vitesse dans l'état $F = 3$, à partir duquel l'excitation Rydberg a lieu, tandis que tous les autres sont pompés dans l'état $F = 2$. Dans ce but, on utilise deux lasers croisant le faisceau atomique au niveau du four atomique. Le premier laser, appelé laser *dépompeur* est orthogonal au jet et accordé sur la transition $F = 3 \leftrightarrow F' = 3$. Ce laser est tout le temps allumé et pompe tous les atomes dans le niveau $F = 2$. Le second laser, appelé *repompeur*, est situé après le

laser dépompeur et fait un angle θ avec le faisceau, de manière à sélectionner la vitesse des atomes par effet Doppler. Il est désaccordé de la transition $F = 2 \leftrightarrow F' = 3$ de $\Delta\omega = 2\pi \cdot 145$ MHz. Les atomes à résonance avec lui sont donc ceux qui ont la vitesse

$$v = \frac{\lambda_{rep}}{2\pi} \frac{\Delta\omega}{\cos \theta} = 250 \text{ m/s.} \quad (\text{I.4.12})$$

Ces atomes sont repompés dans l'état $F' = 3$ et sont donc susceptibles d'être excités dans l'état circulaire.

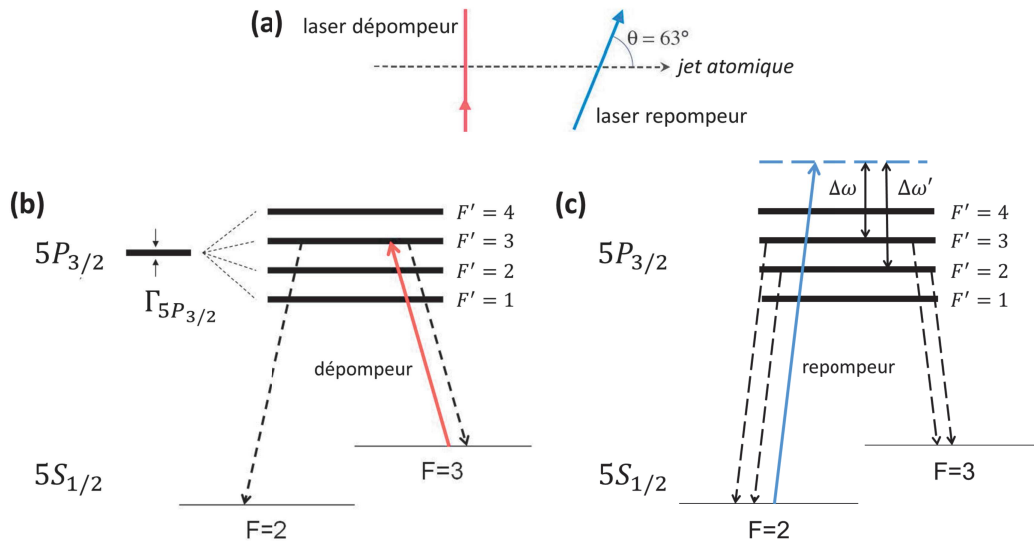


Fig. I.14 Principe de la sélection Doppler. (a) Position des lasers de sélection en vitesse sur le trajet atomique. (b) Transition utilisée pour le dépompeur, en rouge. Comme ce laser est orthogonal au jet, tous les atomes sont pompés dans l'état $F = 2$, quelle que soit leur vitesse. (c) Transition utilisée pour le repompeur, en bleu.

La distribution de vitesse obtenue après la sélection par effet Doppler est présentée sur la figure I.15. Le pic est plus haut que la valeur obtenue pour la distribution initiale de vitesse puisqu'il contient aussi les atomes qui étaient dans l'état $F' = 2$ et qui ont été pompés dans $F = 3$. On observe par ailleurs un deuxième pic autour de 360 m/s. Celui-ci est constitué d'atomes repompés par la transition $F = 2 \leftrightarrow F' = 2$, pour laquelle le laser est désaccordé de $\Delta\omega' = 2\pi \cdot 208$ MHz. On voit par ailleurs que la sélection Doppler permet d'atteindre une largeur de quelques dizaines de mètres par seconde sur la distribution de vitesse. Celle-ci conduirait à des dispersions de l'ordre de quelques centimètres au niveau des cavités. On a donc besoin d'une méthode de sélection plus efficace.

Afin d'affiner la sélection en vitesse, on utilise des excitations pulsées pour pratiquer une sélection en temps de vol. Le repompeur et le premier laser d'excitation à 780 nm, séparés d'environ 36 cm, ne sont allumés que pendant quelques micro-secondes, à 1440 μs d'intervalle. Les seuls atomes susceptibles d'être excités sont donc ceux qui sont sélectionnés par l'effet Doppler et qui se situent dans le faisceau de chaque laser au moment de

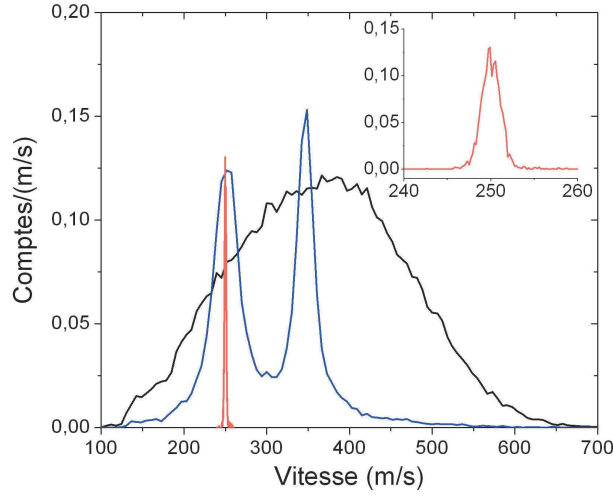


Fig. I.15 Distribution de vitesse des atomes excités dans $|52F\rangle$. Courbe noire : sans sélection de vitesse. Courbe bleue : avec sélection Doppler. Courbe rouge : avec sélection Doppler et en temps de vol. L'insert est un agrandissement de la courbe rouge autour de 250 m/s.

l'impulsion. La distribution de vitesse résultante est montrée en rouge sur la figure I.15. La largeur de la distribution de vitesse n'est plus que d'environ 2 m/s.

Nombre d'atomes par échantillon

L'excitation des atomes vers les états de Rydberg est un processus poissonien. Le nombre d'atomes contenus dans un échantillon atomique est donc une variable aléatoire décrite par la loi :

$$P_a(n_a) = e^{-\mu} \frac{\mu^{n_a}}{n_a!}, \quad (\text{I.4.13})$$

où μ est le nombre moyen d'atomes par échantillon. Ce nombre peut être ajusté en variant la durée des impulsions laser d'excitation. On cherche généralement à avoir $\mu < 1$ de façon à limiter la probabilité d'avoir plusieurs atomes, ce qui perturbe l'interaction de Rabi et réduit le contraste des franges de Ramsey. Typiquement, μ est inférieur à 0.3.

c) Détection des atomes

Après avoir traversé le dispositif et interagi avec les cavités, les atomes sont détectés par ionisation. Les atomes de Rydberg étant fort excités, un champ d'environ 100 V/cm suffit à arracher l'électron de valence pour les atomes que nous utilisons. Le champ d'ionisation dépend par ailleurs de l'énergie de l'électron, ce qui permet de discriminer les différents états. Le montage expérimental comporte deux détecteurs, dont la conception est décrite en détail dans la thèse d'Alexia Auffèves [59]. On se contentera ici de décrire le principe de fonctionnement du détecteur utilisé dans nos expériences.

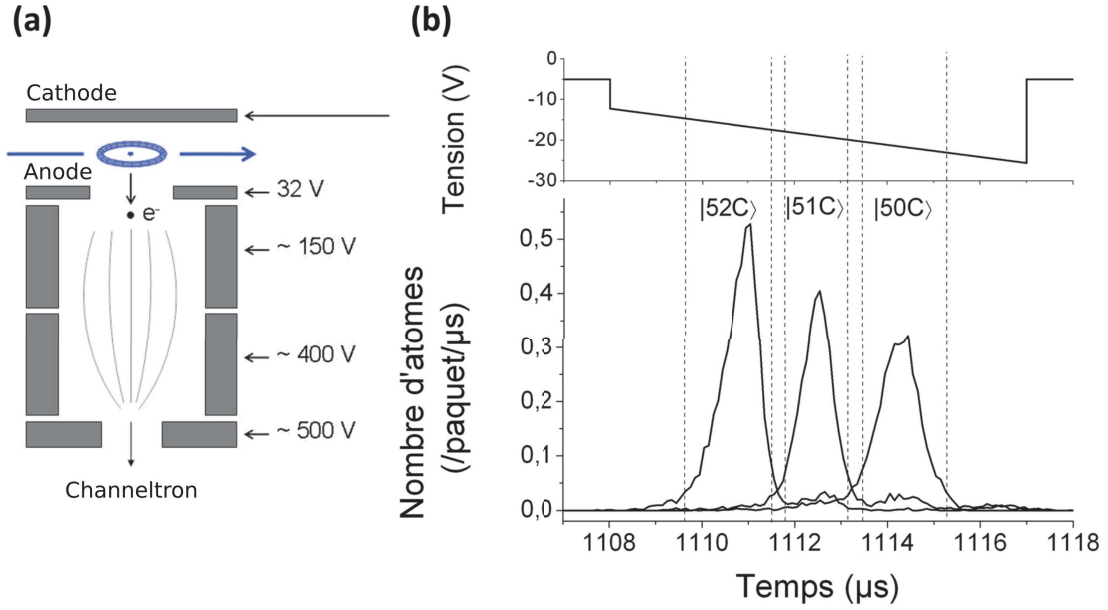


Fig. I.16 *Détection des atomes. (a) Schéma de principe du détecteur. Une rampe de tension, suffisamment courte pour négliger le mouvement de l'atome, est appliquée sur la cathode et arrache l'électron au dessus du trou de l'anode. Celui-ci est ensuite collimaté sur le Channeltron par un ensemble de lentilles électrostatiques. (b) Rampe de potentiel appliqué sur la cathode (en haut) et signal temporel d'ionisation mesuré sur le détecteur. Les pointillés indiquent les fenêtres de détection utilisées.*

Le schéma de principe du détecteur est montré sur la figure I.16. Il comporte une cathode plate située en face d'une anode percée d'un trou de 6 mm. Le potentiel de l'anode est fixé, tandis qu'on applique à la cathode une rampe de potentiel au moment du passage de l'échantillon. L'électron est arraché par la rampe à un instant dépendant de son potentiel d'ionisation, donc de son état. Il est ensuite collimaté par un ensemble de lentilles électrostatiques sur un Channeltron qui génère une cascade électronique. Le signal est amplifié puis envoyé sur un discriminateur. La figure I.16 montre la rampe de tension utilisée et le nombre d'atomes détectés en fonction du temps. On voit qu'on parvient à résoudre les différents états. On peut donc définir des fenêtres de détection, visibles sur la figure, correspondant à chacun d'entre eux.

Le détecteur a une efficacité de détection ϵ . Celle-ci est composée de deux composantes. D'une part, une partie des atomes ne sont pas détectés du tout, parce que les électrons ne sont pas correctement collimatés ou ne parviennent pas à déclencher une cascade électronique. D'autre part, une partie des atomes sont détectés en dehors des fenêtres de détection. En particulier, entre 12 et 14% des atomes sont détectés dans l'intervalle entre les intervalles de détection de $|e\rangle$ et $|g\rangle$. Pour simplifier l'étude, on préfère considérer que ces atomes ne sont pas détectés. Ceci a simplement pour effet de réduire l'efficacité de détection. Grâce à une méthode expliquée dans [60], on a mesuré $\epsilon = 0,5 \pm 0,05\%$.

Par ailleurs, la discrimination entre les états n'est pas parfaite. Il existe une probabilité

η_e^d (resp. η_g^d) de détecter dans $|g\rangle$ (resp. $|e\rangle$) un atome préparé dans $|e\rangle$ (resp. $|g\rangle$). On a mesuré $\eta_e^d = 0.09$ et $\eta_g^d = 0.065$. L'écart entre ces deux valeurs est principalement dû au fait qu'on y a incorporé la probabilité que l'atome relaxe avant d'être détecté, qui est de l'ordre de 3%. Si l'atome relaxe depuis $|e\rangle$, il est alors compté comme $|g\rangle$, tandis que si il relaxe depuis $|g\rangle$, il n'est pas compté.

I.5 Techniques d'électrodynamique quantique

Dans cette section, nous présentons plusieurs techniques d'électrodynamique quantique en cavité qui ont été développées dans l'équipe pour manipuler et mesurer l'état quantique d'un mode du champ électromagnétique à l'aide des atomes de Rydberg. Nous expliquons aussi les nouvelles méthodes mises au point pour faire des oscillations de Rabi dans lesquelles l'atome est initialement proche de résonance avec la cavité. Celle-ci ont été nécessaires pour effectuer des mesures QND de la parité de l'état $|10\rangle + |01\rangle$. Enfin, nous détaillons les techniques permettant d'effectuer des mesures de la fonction de Wigner à deux modes, inédites pour notre dispositif.

I.5.1 Calibration du champ électrique et contrôle de la fréquence atomique

Afin de contrôler l'interaction des atomes avec les cavités, on utilise l'effet Stark quadratique pour contrôler la fréquence de la transition et mettre les atomes à résonance ou hors résonance. On utilise dans ce but une électrode, appelée *killer*, qui permet de changer le champ électrique au sein des cavités. Cette électrode est normalement constituée du miroir inférieur de la cavité, l'autre miroir étant mis à la masse. Cependant, lors du refroidissement du dispositif, nous avons perdu le contact électrique avec le miroir de la cavité 1. Nous avons donc décidé d'utiliser comme électrode killer les deux demi-anneaux du bloc cavité, visibles sur la figure I.9. On applique à l'un d'entre eux un potentiel positif et à l'autre le potentiel opposé.

Ceci présente plusieurs inconvénients. En effet, les demi-anneaux sont situés plus loin du passage des atomes, ce qui implique d'appliquer un potentiel électrique plus élevé. De plus, le champ électrique qui en résulte est moins homogène, à cause de la courbure de ces électrodes et de la présence des miroirs. Enfin, le champ résultant est dirigé d'un demi-anneau vers l'autre. Il est donc horizontal, contrairement au champ directeur qui est vertical dans le reste de la boîte d'écrantage du champ magnétique. Il est donc nécessaire de faire correctement la transition entre le champ directeur vertical et celui horizontal, sans quoi on mélange les états de Rydberg [61].

Afin de s'assurer qu'on peut modifier suffisamment la fréquence de la transition atomique, on calibre la variation de la fréquence de transition atomique par effet Stark en fonction du potentiel appliqué. Pour ce faire, on prépare les atomes dans l'état $|g\rangle$ et on effectue une spectroscopie de la transition $|g\rangle \leftrightarrow |e\rangle$ en envoyant une impulsion de champ micro-onde par le guide d'onde latéral de la cavité, lorsque les atomes sont au milieu de celle-ci. On répète cette séquence en variant le potentiel $V_{K,1}$ appliqué sur le killer. Les résultats de cette mesure sont présentés sur la figure I.17. On voit qu'on peut modifier la fréquence de résonance de la transition atomique d'au moins 400 kHz, ce qui est suffisant pour nos expériences. Pour la cavité 2, on obtient des résultats similaires, à la différence que les tensions à appliquer sont de l'ordre de quelques volts.

I.5.2 Initialisation du mode dans l'état vide

Les séquences expérimentales que nous utilisons nécessitent de partir de l'état vide pour la cavité. Lorsqu'une expérience commence, le champ dans la cavité est généralement dans un état thermique. Il contient alors entre 0.05 et 0.2 photons, selon qu'on travaille à 0.8 K ou 1.6 K. Par ailleurs, les expériences ne sont jamais effectuées une seule fois. Chaque

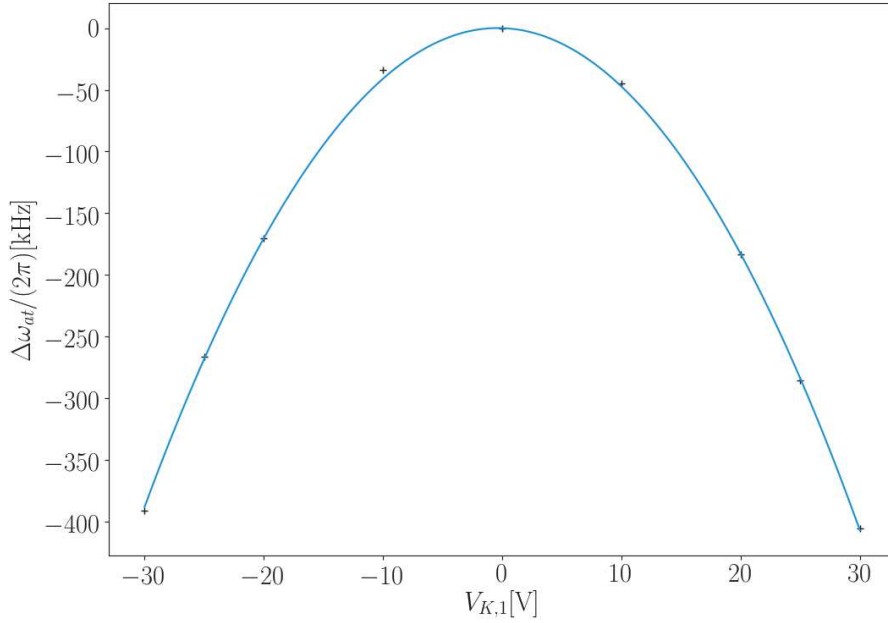


Fig. I.17 Variation de la fréquence de la transition en fonction de $V_{K,1}$. La courbe bleue est un ajustement quadratique par la fonction $V_{K,1} \mapsto aV_{K,1}^2 + bV_{K,1}$ avec $a = -0,44 \text{ kHz/V}^2$ et $b = -0,31 \text{ kHz/V}$.

séquence est répétée des centaines de fois afin de moyennner les résultats de détection pour déterminer les probabilités des différents résultats. À la fin d'une séquence, le champ est dans un état hors équilibre qui peut contenir de nombreux photons si on a injecté un champ dans la cavité. Le temps de retour à l'équilibre est de l'ordre de quelques temps de vie de la cavité, ce qui peut représenter plusieurs centaines de millisecondes. Afin de s'assurer qu'on commence donc toujours dans l'état vide et d'éviter d'attendre le retour à l'équilibre, on envoie donc des échantillons d'atomes, appelés par tradition *atomes serpilières*, qui sont préparés dans $|g\rangle$ et absorbent le champ contenu dans la cavité. Comme on ne sait pas combien de photons le champ contient, on utilise la technique du passage adiabatique, décrite dans la section I.6. Le champ électrique de l'électrode killer est balayé lorsque les atomes sont au centre du mode, de manière à passer du désaccord initial positif, à un désaccord négatif. La figure I.18 montre par exemple la rampe de désaccord utilisée pour les serpilières dans C_2 pour un désaccord initial faible.

I.5.3 Oscillations de Rabi et préparation d'états de Fock à un photon

La pulsation de Rabi du vide $\Omega_{0,j}$ est le paramètre le plus important de nos expériences. Elle caractérise la force du couplage entre les atomes et les photons contenus dans les cavités. Afin de la déterminer, on réalise des expériences d'oscillations de Rabi du vide. Un atome, initialement préparé dans $|e\rangle$, est mis à résonance avec la cavité pendant une durée variable T_R , à l'aide du killer. Avant et après cette interaction à résonance,

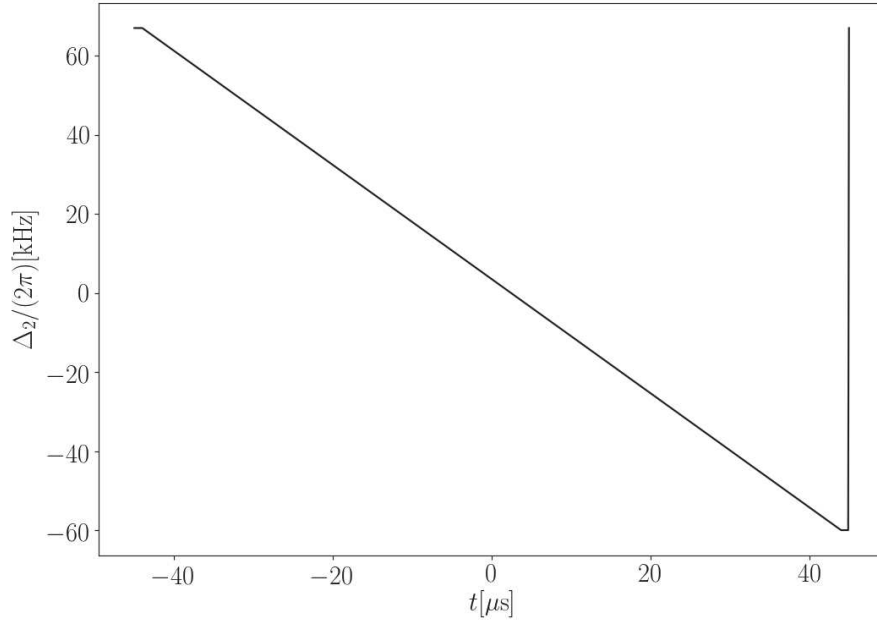


Fig. I.18 Rampe de désaccord utilisée pour les serpillières dans C_2 , exprimée en terme de désaccord entre les atomes et le champ. L'origine des temps est prise au moment où l'atome est au centre de la cavité.

l'atome est fortement désaccordé pour empêcher toute autre interaction. On mesure alors la probabilité de trouver l'atome dans l'état $|e\rangle$ en fonction de T_R .

a) Mesure de la fréquence de Rabi du vide

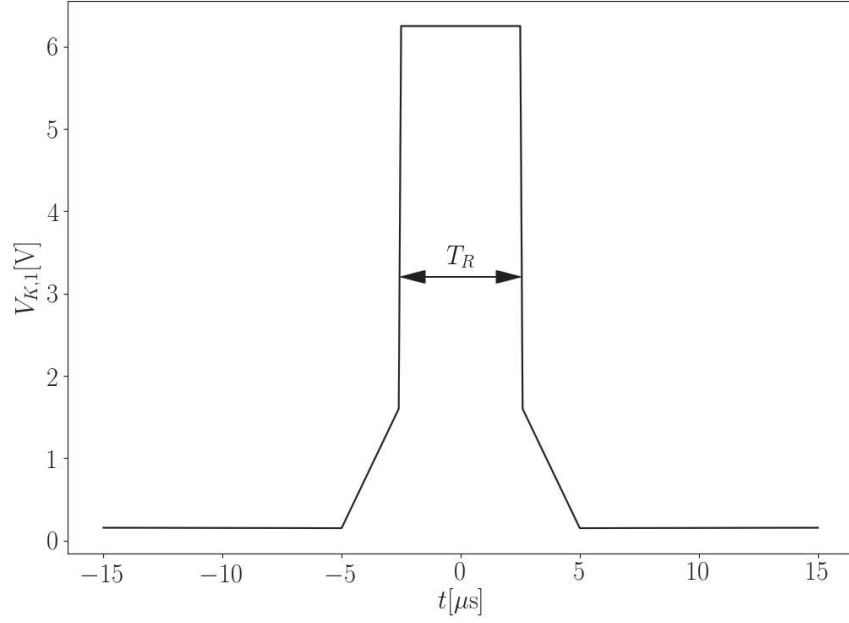
Considérons le cas où le désaccord entre l'atome et la cavité est grand devant la fréquence de Rabi : $|\Delta_j| \gg \Omega_{n,j}$. L'atome, initialement préparé dans $|e\rangle$, entre dans le mode de la cavité, préparée dans l'état $|0\rangle$. Comme le mode a une structure gaussienne, la fréquence de Rabi, proportionnelle à f_j , croît progressivement. Les atomes vont suffisamment lentement pour que le couplage se fasse de manière adiabatique. Initialement, $\Omega_{0,j} = 0$ et $|e, 0\rangle$ est l'état habillé correspondant. L'état du système suit donc l'état habillé pendant la première phase de couplage. Comme le désaccord est grand, cet état reste proche de $|e, 0\rangle$ et il n'y a pas d'échange d'énergie tant que l'atome n'est pas porté à résonance.

À l'instant $t = -T_R/2$, on applique un potentiel sur le killer qui met l'atome à résonance. Comme le système est initialement dans l'état $|e, 0\rangle$, on est dans le cas traité dans la section I.3.2. La seule différence est que l'atome n'est pas au centre du mode, à $f_j = 1$ mais qu'il se déplace pendant l'interaction. Les oscillations ont donc lieu en fonction du temps d'interaction effectif T_{eff} . Comme on l'a vu sur la figure I.11, ce temps peut être confondu avec T_R pour des durées inférieures à 10 μs. À $t = T_R/2$, le killer revient à sa valeur initiale. L'atome est désaccordé et se découple progressivement du champ à mesure qu'il quitte la cavité.

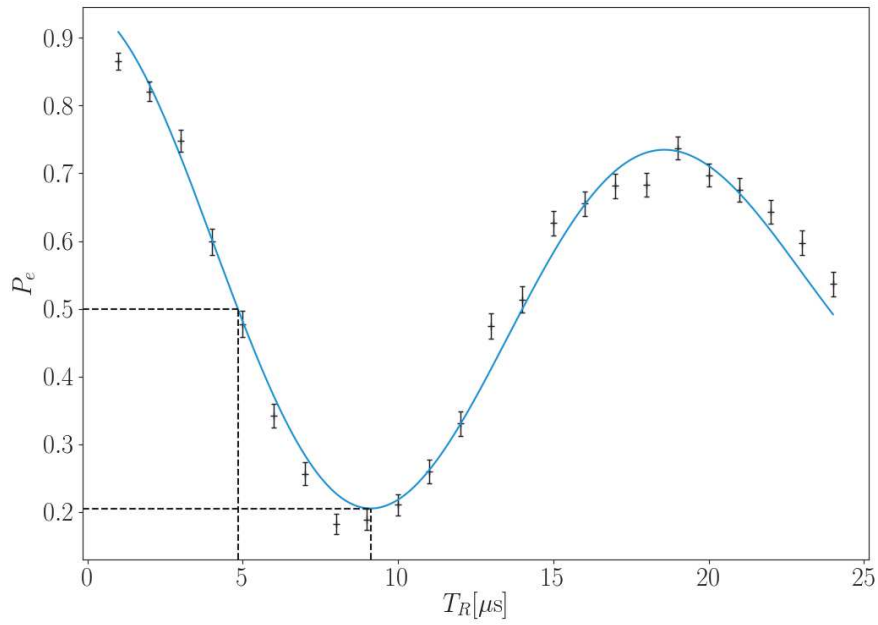
La figure I.20 montre un exemple de rampe de potentiel appliquée dans la cavité 2 et le signal d'oscillations de Rabi obtenu lorsqu'on fait varier la durée de cette rampe. Le désaccord initial vaut $\Delta_2^0 = 2\pi \cdot 272$ kHz. On ajuste ce signal par une fonction prenant en compte le brouillage progressif de l'oscillation. Ceci permet de calibrer la fréquence de Rabi du vide $\Omega_{0,2} = 2\pi \cdot 47$ kHz.

La figure I.19 montre la rampe utilisée et les oscillations résultantes pour C₁, pour un désaccord initial $\Delta_1^0 = 2\pi \cdot 285$ kHz. Dans cette cavité, le killer est constitué de deux potentiels opposés appliqués aux demi-anneaux latéraux. Comme le champ électrique résultant est horizontal, il faut tourner l'axe de quantification de l'atome, initialement vertical. Afin de rendre cette rotation adiabatique, on applique une rampe de champ avant et après l'accord de l'atome à résonance. Les valeurs des potentiels appliqués à résonance ont été optimisées pour maximiser le transfert de l'atome de $|e\rangle$ vers $|g\rangle$. L'ajustement donne la fréquence de Rabi $\Omega_{0,1} = 2\pi \cdot 53$ kHz. Cette valeur légèrement plus élevée est un artefact dû au fait que l'atome se couple un peu au champ pendant la rampe. L'effet du couplage au champ avant l'oscillation de Rabi sera traité dans la suite.

Les oscillations de Rabi présentent un amortissement, dû à la dispersion du couplage entre les atomes et les cavités. En effet, les échantillons atomiques ont une extension spatiale, due aux largeurs du jet atomique et des faisceaux laser d'excitation. Ils ne passent donc pas tous exactement par le centre du mode et ils ne sont pas portés à résonance au même point des cavités. Ils sont donc soumis à des pulsations de Rabi différentes. La moyenne des oscillations sur ces différentes pulsations donne donc un brouillage pour les grands temps d'interaction.

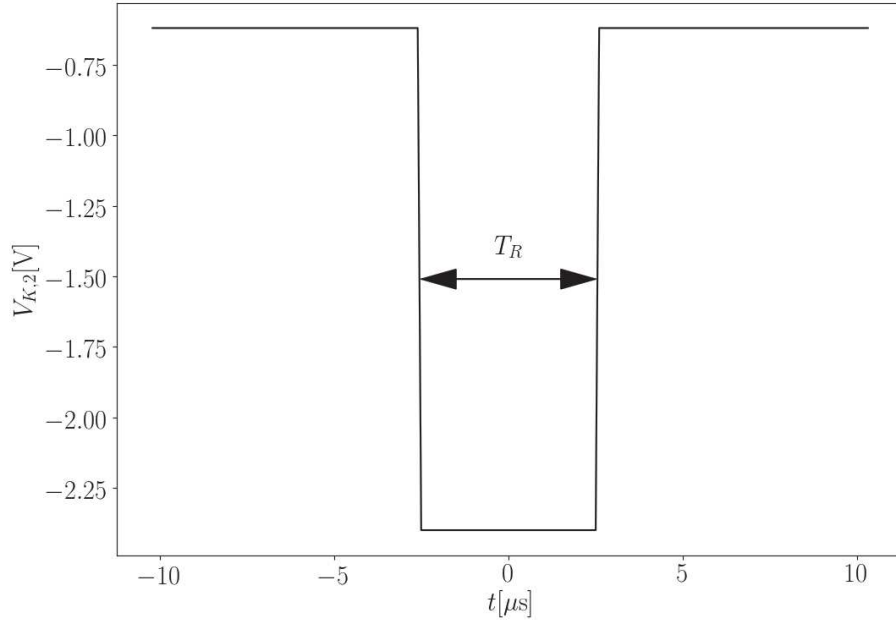


(a) Potentiel appliqué au premier anneau de la cavité 1. Le potentiel appliqué à l'autre anneau est l'opposé de celui-ci.

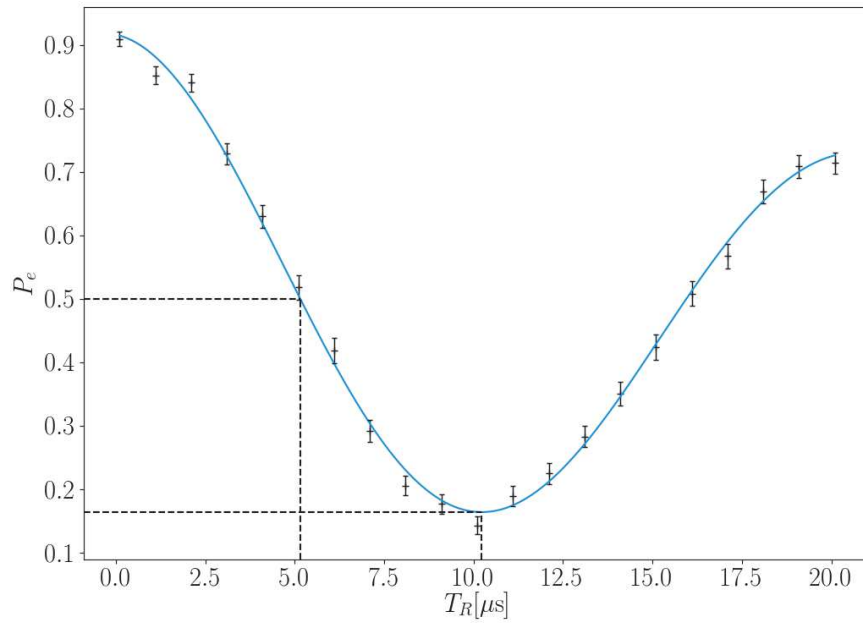


(b) Oscillations de Rabi

Fig. I.19 Oscillations de Rabi du vide dans la cavité 1 et rampe de potentiel utilisée pour mesurer le point $T_R = 5 \mu\text{s}$, pour $\Delta_1^0 = 2\pi \cdot 285 \text{ kHz}$. La courbe bleue est un ajustement par la fonction $P = P_0 + C/2 \cos(\Omega T_R) e^{-t/\tau}$ avec $P_0 = 0.51$, $C = 0.86$, $\Omega = 2\pi \cdot 53 \text{ kHz}$ et $\tau = 28 \mu\text{s}$.



(a) Potentiel appliqué au miroir inférieur de la cavité 2.



(b) Oscillations de Rabi

Fig. I.20 Oscillations de Rabi du vide dans la cavité 2 et rampe de potentiel utilisée pour mesurer le point $T_R = 5 \mu\text{s}$, pour $\Delta_2^0 = 2\pi \cdot 272 \text{ kHz}$. La courbe bleue est un ajustement par la fonction $P = P_0 + C/2 \cos(\Omega T_R) e^{-t/\tau}$ avec $P_0 = 0.49$, $C = 0.86$, $\Omega = 2\pi \cdot 47 \text{ kHz}$ et $\tau = 37 \mu\text{s}$.

b) Oscillations de Rabi à proximité de la résonance : effet de l'habillage initial de l'atome

Pour nos expériences, nous avons eu besoin par moments d'imposer une valeur initiale du désaccord $\Delta_j = 2\pi \cdot 64$ kHz⁸. Pour cette valeur de Δ_j , comparable à $\Omega_{0,j}$, l'état initial $|g, 0\rangle$ n'est pas un état propre du hamiltonien. Il y a donc une évolution avant que le killier ne mette l'atome à résonance. On va voir dans cette section comment les oscillations de Rabi sont affectées par cette évolution initiale et comment on peut néanmoins obtenir des probabilités d'injection convenables.

Lorsque l'atome est peu désaccordé, il se couple au mode de la cavité avant qu'on le mette à résonance. À cause de la forme gaussienne du mode, ce couplage s'effectue de manière adiabatique. Initialement, $\Omega_{0,j} = 0$ puisque l'atome est hors du mode. L'état $|e, 0\rangle$ est donc un état propre du hamiltonien. L'atome suit donc l'état habillé du mode. Si le temps d'interaction est suffisamment faible, la valeur de $\Omega_{0,j}(\mathbf{r})$ atteinte au moment où on met l'atome à résonance est proche de Δ_j . L'état habillé à ce moment-là comporte donc une composante importante de $|g, 1\rangle$ avant que l'atome n'entame les oscillations de Rabi. La figure I.21a montre l'évolution théorique de l'état dans la cavité 2 dans le cas où $\Delta_2 = |\Omega_{0,2}|$. On a pris $\Omega_{0,2}$ imaginaire pur pour cette figure, pour plus de lisibilité. Le temps d'interaction est choisi égal à $\pi/|\Omega_{0,j}|$, de sorte que l'atome effectuerait une impulsion π si il était initialement très désaccordé. Avant d'interagir, l'atome suit l'état habillé positif, qui se rapproche de l'équateur, au fur et à mesure que $\Omega_{0,2}$ devient comparable à Δ_2 . Le vecteur bleu montre l'état habillé juste avant l'interaction $|+\rangle_i$. Lorsque l'atome est mis à résonance, il est donc dans un état qui contient une part non négligeable de $|g, 1\rangle$. Il tourne ensuite pendant T_R autour de l'axe équatorial donné par l'état habillé de plus grande énergie $|+\rangle_R$ (vecteur rouge). Après l'interaction, le système est loin de l'état habillé de basse énergie (vecteur vert), autour duquel il précède. Au fur et à mesure du découplage de l'atome avec le mode du champ, cet état habillé se rapproche de $|g, 1\rangle$. Comme l'atome en est éloigné, il garde un coefficient important sur $|e, 1\rangle$, presque égal à 50%. Les valeurs $\Delta_2 = 0$ et $T_R = \pi/|\Omega_{0,2}|$ utilisées précédemment pour effectuer une impulsion π ne sont donc plus du tout adaptées.

Une stratégie pour obtenir un meilleur transfert est présentée sur la figure I.21b, qui montre une impulsion π dans C_2 , avec $\Delta_2 = 2\pi \cdot 68$ kHz. Au lieu de mettre l'atome à résonance pendant l'impulsion de Rabi, on choisit un désaccord négatif $\Delta_2^R = -2\pi \cdot 32$ kHz, de sorte que l'axe de rotation soit orthogonal à l'état habillé juste avant l'impulsion. L'état précède donc sur un grand cercle de la sphère de Bloch, à une fréquence légèrement supérieure à celle de l'oscillation de Rabi du vide. On choisit la durée de l'oscillation de manière à faire tourner l'état d'un angle proche de π , pour qu'il arrive au voisinage de l'état habillé négatif juste après l'impulsion $|-\rangle_i$ (vecteur vert sur la figure). Pour plus de lisibilité, on a choisi une valeur de la durée d'interaction pour laquelle l'état n'arrive pas exactement sur le vecteur vert. Lors du découplage progressif du mode, l'état suit donc cet état habillé et le système est conduit près de l'état $|g, 1\rangle$.

L'état atomique passe légèrement en dessous de l'état habillé de basse énergie. Ceci est dû au fait que l'angle entre $|+\rangle_i$ et $|+\rangle_R$ est légèrement inférieur à $\pi/2$, mais aussi au fait que $\Omega_{0,2}(\mathbf{r})$ continue d'évoluer pendant l'oscillation de Rabi, puisque l'atome s'approche

8. Comme nous le verrons dans la suite, cette valeur correspond à une mesure de la parité par les atomes dispersifs.

puis s'éloigne du centre du mode. $|+\rangle_R$ a donc un léger mouvement de balancier qui le rapproche de l'axe Ox puis l'en éloigne à nouveau. Pour ces raisons, on optimise les impulsions de Rabi de manière à avoir un transfert maximal vers $|g, 1\rangle$.

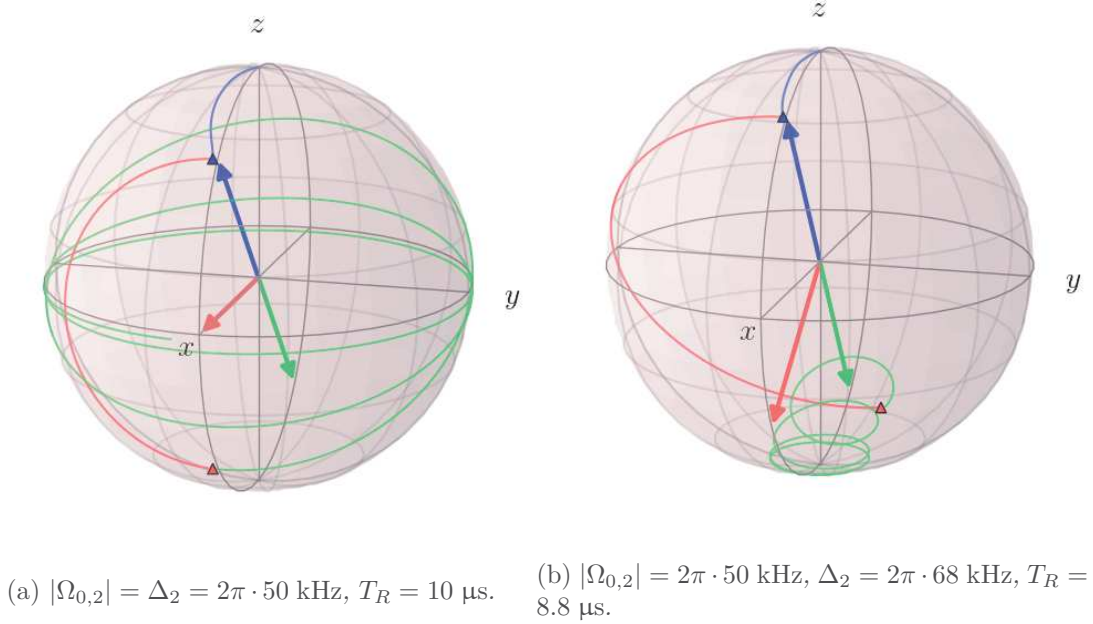
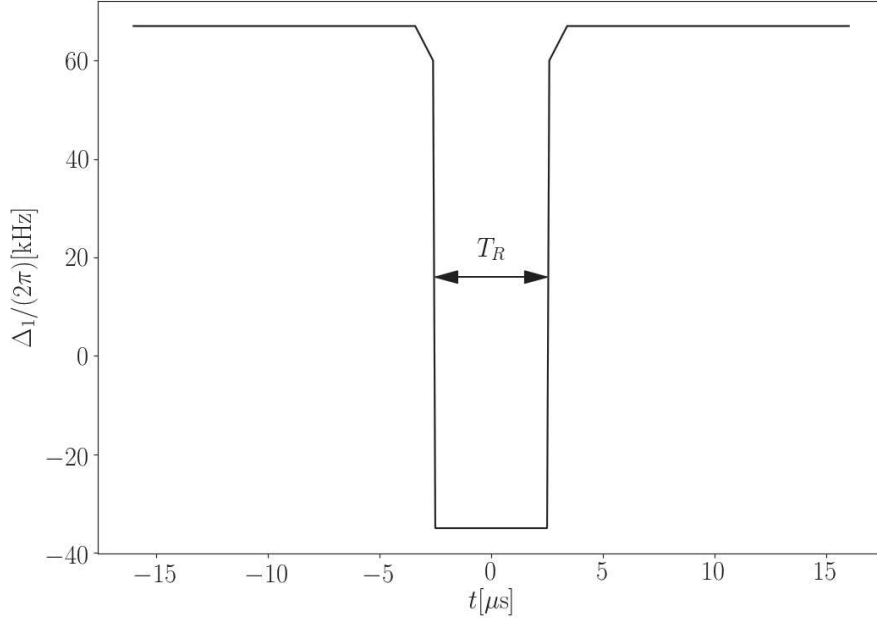


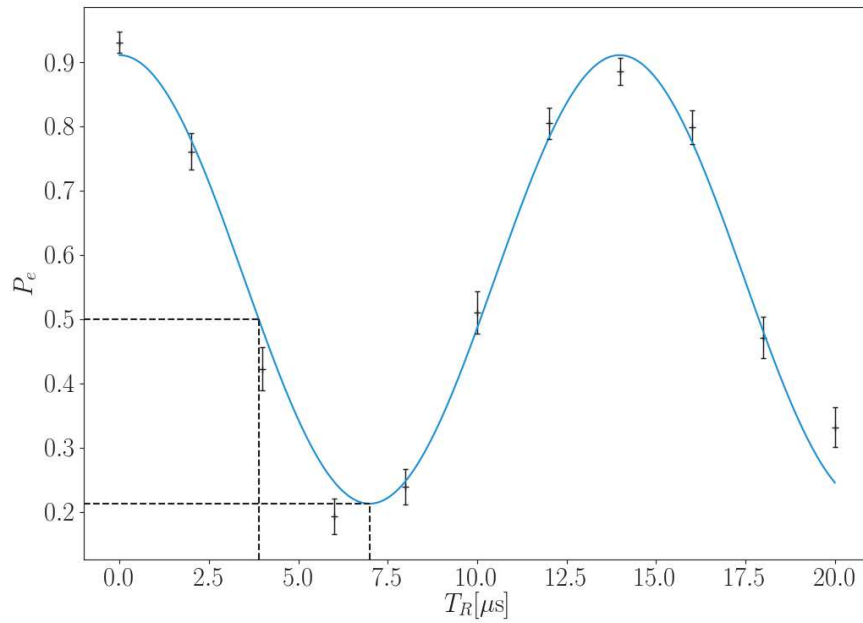
Fig. I.21 Oscillations de Rabi du vide dans la cavité 2 pour deux régimes de paramètres. Le système est initialement dans l'état $|e, 1\rangle$ (pôle nord de la sphère). Durant une première phase, décrite par la courbe bleue, l'atome se couple au mode du champ. La légère courbure de la trajectoire indique que le couplage n'est pas tout à fait adiabatique : l'atome ne suit pas exactement l'état habillé. À $t = -T_R/2$ (triangle bleu), on modifie la fréquence de transition atomique. Sur l'image de gauche, on met l'atome à résonance, tandis que sur celle de droite, on met la transition atomique 32 kHz en dessous de la fréquence de la cavité. Le vecteur bleu (resp. rouge) indique la direction de l'axe de rotation de l'évolution due au hamiltonien juste avant (resp. après) qu'on change le potentiel. Pendant une deuxième phase, de durée T_R , l'atome effectue des oscillations de Rabi autour du vecteur rouge. Enfin, à $t = T_R/2$ (triangle rouge), le kill est remis à sa valeur initiale. La flèche verte indique la direction de l'axe de rotation après la commutation du potentiel. Durant la dernière phase de l'évolution, l'atome se découple progressivement. Il précesse alors autour de l'axe donné par les états propres instantanés du hamiltonien, qui se rapproche progressivement de l'axe Oz . Dans le cas (a), l'état est éloigné de l'état propre du hamiltonien et effectue de grandes oscillations. Finalement, l'atome a environ une chance sur deux d'émettre. Dans le cas (b), l'état suit l'état propre du hamiltonien et l'atome émet presque certainement.

Les figures I.22 et I.23 montrent les impulsions retenues pour C_1 et C_2 respectivement, ainsi que les oscillations de Rabi obtenues. Pour plus de clarté, on a indiqué les désaccords plutôt que les potentiels pour les rampes des électrodes kill, en utilisant les calibrations

des champs électriques effectuées selon la méthode de la section [I.5.1](#). Les durées des impulsions sont choisies de sorte qu'on ait un transfert maximal pour les impulsions π et une probabilité $P_e = 1/2$ pour les impulsions $\pi/2$.

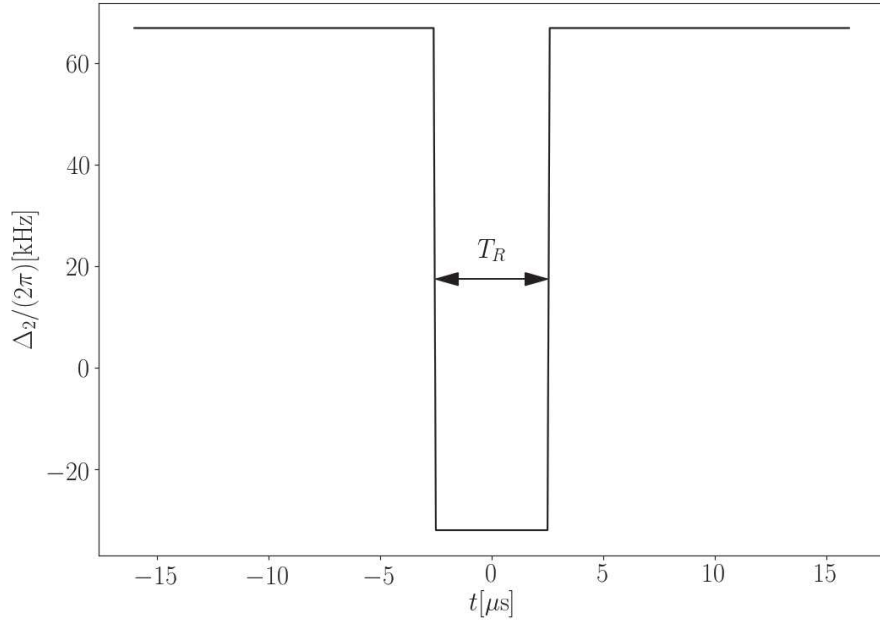


(a) Rampe appliquée à l'atome pour mesurer les oscillations de Rabi dans C_1 . L'échelle verticale indique le désaccord attendu, en kilohertz, entre l'atome et la cavité pour le potentiel appliqué.

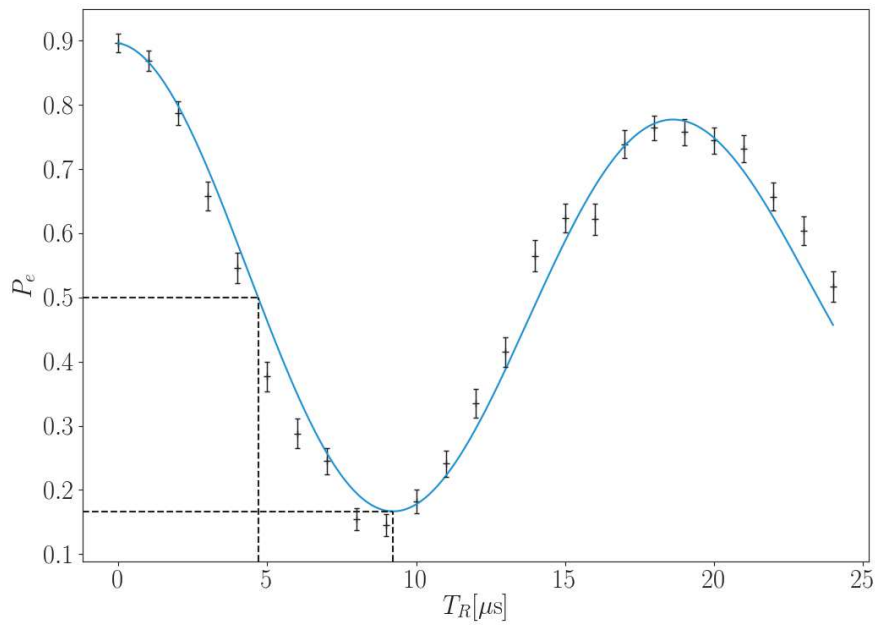


(b) Oscillations de Rabi

Fig. I.22 Oscillations de Rabi du vide dans la cavité 1 et rampe de potentiel utilisée pour mesurer le point $T_R = 5$ μs, pour $\Delta_1^0 = 2\pi \cdot 56$ kHz. La courbe bleue est un ajustement par la fonction $P = P_0 + C/2 \cos(\Omega T_R)$ avec $P_0 = 0.56$, $C = 0.70$, $\Omega = 2\pi \cdot 71.5$ kHz.



(a) Rampe appliquée à l'atome pour mesurer les oscillations de Rabi dans C_2 . L'échelle verticale indique le désaccord attendu, en kilohertz, entre l'atome et la cavité pour le potentiel appliqué.



(b) Oscillations de Rabi

Fig. I.23 Oscillations de Rabi du vide dans la cavité 2 et rampe de potentiel utilisée pour mesurer le point $T_R = 5 \mu\text{s}$, pour $\Delta_2^0 = 2\pi \cdot 68 \text{ kHz}$. La courbe bleue est un ajustement par la fonction $P = P_0 + C/2 \cos(\Omega T_R) e^{-t/\tau}$ avec $P_0 = 0.50$, $C = 0.79$, $\Omega = 2\pi \cdot 53 \text{ kHz}$ et $\tau = 52 \mu\text{s}$.

I.5.4 Interférométrie de Ramsey et mesure de la fonction de Wigner

On a montré dans la section I.3.3 que l'interaction dispersive d'un atome avec un mode du champ donne lieu à un déplacement lumineux qui décale la fréquence de transition atomique d'une valeur environ proportionnelle au nombre de photons contenus dans le mode. Dans cette section, on va expliquer comment on peut utiliser cet effet pour mesurer la parité du nombre de photons contenus dans une cavité. On verra ensuite comment on peut déplacer l'état du champ à l'aide d'une injection micro-onde. Avec ces deux techniques, on peut réaliser une mesure de la fonction de Wigner du champ.

Interférométrie de Ramsey

L'*interférométrie de Ramsey* est une technique de mesure qui utilise des superpositions cohérentes de deux états atomiques pour réaliser des mesures de précision de différences d'énergie [62]. On commence par préparer l'atome dans une superposition :

$$|\psi_0\rangle = \frac{1}{\sqrt{2}} (|e\rangle + |g\rangle). \quad (\text{I.5.1})$$

Dans notre expérience, cet état est préparé en appliquant, dans la zone de Ramsey R_1 , une impulsion résonnante $\pi/2$ à un atome initialement préparé dans $|g\rangle$. On a choisi la référence de phase de sorte que l'état atomique ait une phase nulle après l'impulsion. Si on effectue une seconde impulsion, de phase ϕ_R par rapport à la première impulsion, juste après la première, l'état atomique devient

$$|\psi_f\rangle = \frac{1}{2} \left[(1 - e^{-i\phi_R}) |e\rangle + (1 + e^{i\phi_R}) |g\rangle \right]. \quad (\text{I.5.2})$$

La probabilité de mesurer l'atome dans $|g\rangle$ oscille en fonction de la phase de l'impulsion :

$$P_g = \frac{1}{2} + \frac{1}{2} \cos \phi_R. \quad (\text{I.5.3})$$

Ces oscillations sont connues sous le nom d'*oscillations de Ramsey*.

Choix de la phase de Ramsey

Les impulsions de Ramsey utilisées dans notre expérience sont toutes générées par la même source. Afin de modifier la phase de la deuxième impulsion relativement à la première, on désaccorde la source de la transition atomique d'une quantité $\Delta_r \ll \Omega_c$, où Ω_c est la fréquence de Rabi classique de la transition. De cette manière, la phase relative entre la source et la superposition pendant le temps de vol T_{vol} entre les deux zones de Ramsey vaut $\Delta_r T_{vol}$. Le temps de vol entre deux zones de Ramsey adjacentes vaut 360 μs . Pour obtenir une phase 2π , il faut donc un désaccord $\Delta_r = 2\pi \cdot 2.8 \text{ kHz}$ bien inférieur à la fréquence de Rabi dans les zones de Ramsey, qui est de l'ordre de 100 kHz. On peut donc faire varier la phase en désaccordant très légèrement la source, tout en restant suffisamment proche de la résonance pour éviter de réduire le contraste des oscillations.

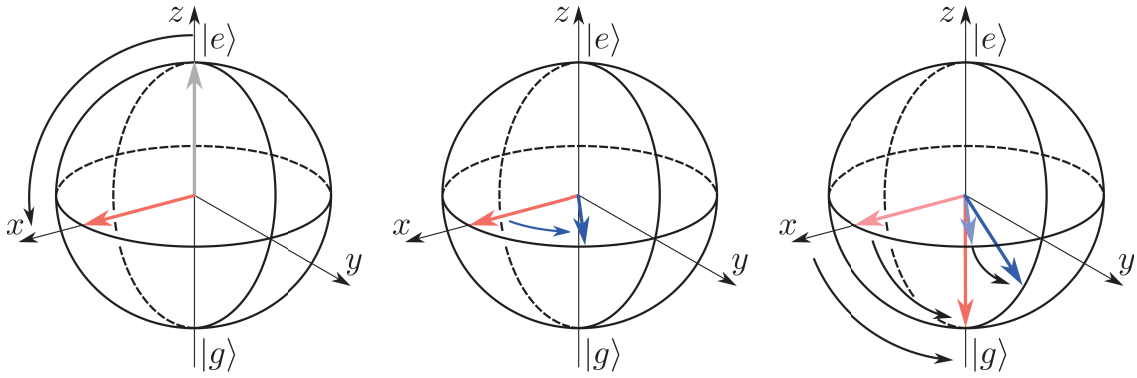


Fig. I.24 Représentation de l'interférométrie de Ramsey sur la sphère de Bloch. À gauche, l'atome subit une première impulsion $\pi/2$ qui l'amène sur l'axe Ox . Au milieu, l'évolution de l'atome lui fait acquérir une phase entre les états $|e\rangle$ et $|g\rangle$. Il tourne donc dans le plan équatorial de la sphère de Bloch. À droite, la deuxième impulsion amène les états équatoriaux dans le plan zOy . L'angle entre les vecteurs bleu et rouge est le même que celui entre les vecteurs bleu pâle et rouge pâle. Les oscillations de phase dans le plan équatorial sont transformées en oscillations de probabilité.

Effet du nombre de photons

Considérons un atome soumis à deux impulsions de Ramsey séparées d'une durée T_{vol} . La figure I.24 montre l'évolution, entre les deux impulsions, de la superposition atomique. Celle-ci acquiert une phase quantique par rapport à la phase de référence en l'absence d'évolution (vecteur rouge). Le vecteur de Bloch tourne donc dans le plan équatorial de la sphère de Bloch (vecteur bleu). L'angle qui décrit sa position dans ce plan est donné par la phase quantique

$$\phi_a = \frac{1}{\hbar} \int_0^{T_{vol}} (E_e - E_g) dt. \quad (I.5.4)$$

La deuxième impulsion reporte cet angle dans le plan zOy de la sphère de Bloch. L'état final vaut donc :

$$|\psi_f\rangle = \frac{1}{2} \left[(1 - e^{-i(\phi_R - \phi_a)}) |e\rangle + (1 + e^{i(\phi_R - \phi_a)}) |g\rangle \right]. \quad (I.5.5)$$

On a donc transformé l'oscillation de phase quantique en oscillation de probabilité qui peut être mesurée directement en détectant l'état de l'atome et en répétant l'expérience un grand nombre de fois.

La phase acquise à cause de l'évolution décale les franges de Ramsey obtenues. En mesurant ce déplacement, on est donc capable de mesurer des modifications des énergies atomiques au cours de l'évolution. En particulier, si l'atome traverse une cavité, il acquiert une phase dépendant du nombre de photons. On a vu précédemment que pour un atome fortement désaccordé, celle-ci est proportionnelle au nombre de photons. On peut utiliser cet effet pour mesurer le nombre de photons contenus dans la cavité. On obtient alors une mesure quantique non destructive (QND) du nombre de photons, puisque celui-ci n'est pas modifié par l'interaction dispersive avec les atomes [63, 64].

Linéarité de la phase pour la mesure de parité

La figure I.25 montre le résultat d'un calcul numérique de la différence de phase $\phi(n+1) - \phi(n)$ entre deux états de Fock successifs, en fonction de Δ_j . Cette différence est calculée à partir de l'équation I.4. Le désaccord est linéaire lorsque les courbes correspondant à des n différents sont confondues. Pour une valeur de $\Phi \equiv \phi(1) - \phi(0)$, les valeurs des phases suivantes sont fixées. Il est donc équivalent de fixer Δ , Φ et $\phi(n)$. On obtient une phase π entre $n = 1$ et $n = 0$ pour $\Delta = 2\pi \cdot 64$ kHz. Pour cette valeur, la phase par photon est loin d'être linéaire. La figure I.26 montre la phase relative $\phi(n) - \phi(0)$ en fonction de n pour ce désaccord. La courbe bleue est une régression quadratique, qui fonctionne correctement pour $n \leq 4$. Pour des nombres plus grands, il faut cependant calculer analytiquement la phase. En toute rigueur, notre mesure n'est pas une mesure de la fonction de Wigner, mais de la pseudo-fonction de Wigner :

$$\tilde{W}^{(\Delta)}(\hat{\rho}, \alpha) = \text{Tr} \left(\hat{D}(-\alpha) \hat{\rho} \hat{D}(\alpha) \cos(\phi(\hat{N})) \right). \quad (\text{I.5.6})$$

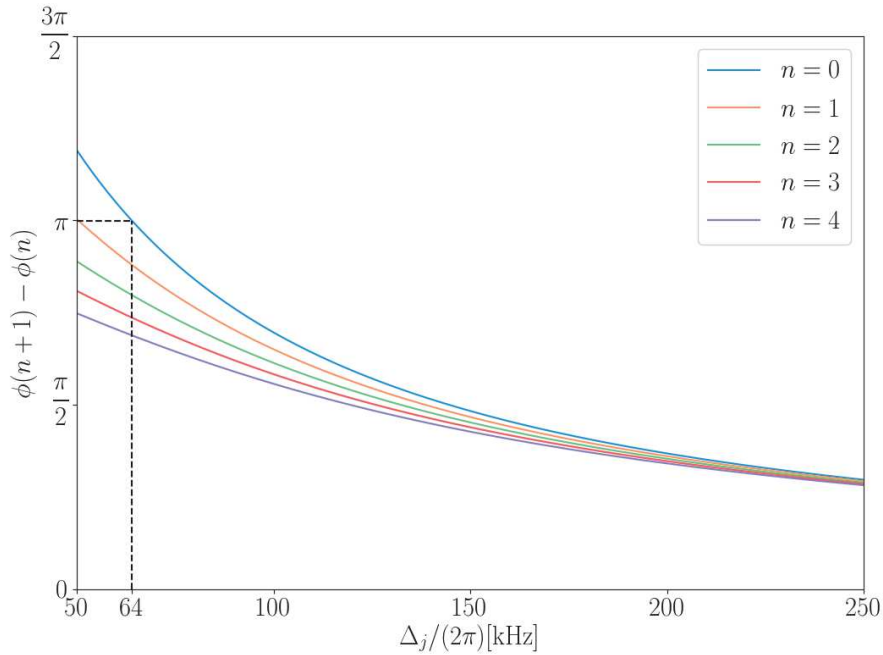


Fig. I.25 Différence entre les phases acquises par l'atome pour deux valeurs de n successives, en fonction du désaccord atome-cavité Δ_j . Les différentes courbes correspondent à des valeurs de n différentes. Les traits en pointillé mettent en évidence la valeur $\Delta_j = 2\pi \cdot 64$ kHz, qui permet d'avoir un déphasage π entre 1 photon et 0 photon. L'écart entre les courbes pour cette valeur est dû à la non linéarité de la phase par rapport au nombre de photons.

Tentons de voir à quel point ceci affecte nos résultats. La pseudo-fonction de Wigner

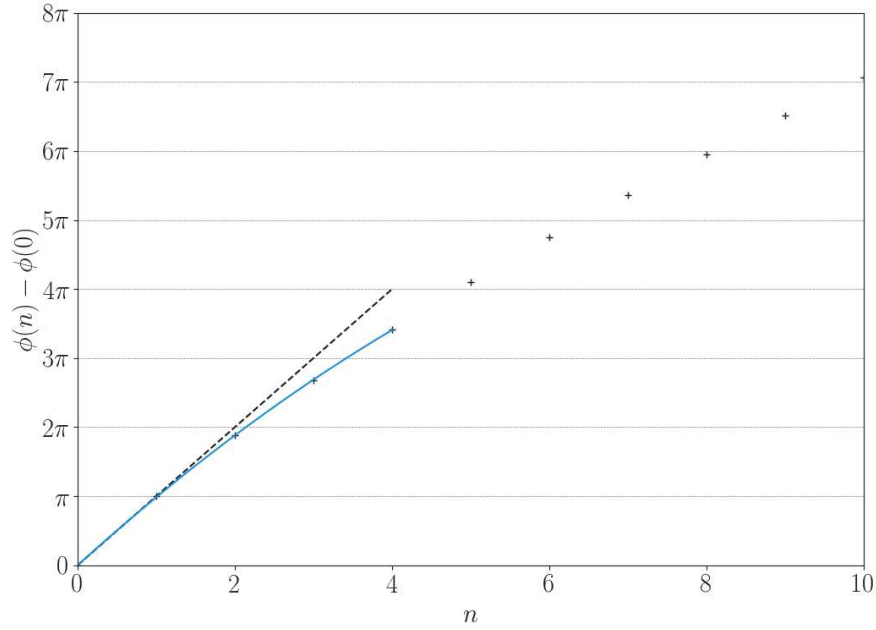


Fig. I.26 Différence entre la phase $\phi(n)$ acquise par l'atome pour n photons présents dans la cavité et celle acquise lorsque la cavité est vide, en fonction de n , pour $\Delta_j = 2\pi \cdot 64$ kHz. La courbe bleue est une régression par la fonction $n \mapsto \phi^{(2)}n^2 + \phi^{(1)}n$, avec $\phi^{(2)} = -0.140$ et $\phi^{(1)} = 3.23$. Elle fournit des résultats corrects pour $n \leq 4$. La courbe en pointillé est celle attendue pour un déphasage linéaire de π par photon.

pour l'état $|0\rangle$ peut s'écrire :

$$\tilde{W}^{(\Delta)}(|0\rangle, \alpha) = \sum_n \cos(\phi(n)) e^{-|\alpha|^2} \frac{|\alpha|^{2n}}{n!}. \quad (\text{I.5.7})$$

En utilisant l'identité

$$\hat{D}(\alpha) \hat{a} \hat{D}(-\alpha) = (\hat{a} - \alpha \mathbb{1}), \quad (\text{I.5.8})$$

on déduit

$$\hat{D}(\alpha) = (\hat{a}^\dagger - \alpha^*) |\alpha\rangle, \quad (\text{I.5.9})$$

ce qui permet de montrer que

$$\begin{aligned} \tilde{W}^{(\Delta)}(|1\rangle, \alpha) = \sum_n \left\{ \cos[\phi(n)] |\alpha|^{2(n+1)} + (n+1) \cos[\phi(n+1)] |\alpha|^{2n} \right. \\ \left. - 2 \cos[\phi(n+1)] |\alpha|^{2(n+1)} \right\} \frac{e^{-|\alpha|^2}}{n!}. \end{aligned} \quad (\text{I.5.10})$$

Pour les phases Φ_1 et Φ_2 déterminées précédemment, ces deux pseudo-fonctions de Wigner sont très proches des fonctions de Wigner correspondantes. L'écart maximal entre

elles est de 0.012 pour un photon et de 0.005 pour zéro photon. En pratique, la fonction de Wigner et la pseudo-fonction de Wigner sont indiscernables pour $n = 0$ ou 1 lorsque Φ est correctement réglée. On peut l'interpréter de la manière suivante : lorsqu'on déplace peu l'état, il contient essentiellement 0 ou 1 photon et on ne voit pas d'écart puisque $\phi(0)$ et $\phi(1)$ ont les valeurs correspondant à la fonction de Wigner. On n'est de plus pas très sensible aux valeurs exactes de ces phases puisqu'on est aux extrema des franges. Lorsqu'on déplace plus l'état, les fonctions de Wigner de $|0\rangle$ et $|1\rangle$ sont brouillées rapidement par l'interférence entre franges correspondant à des nombres de photons différents et tendent vers 0,5. Ce brouillage a lieu tant qu'on a suffisamment de franges et ne dépend que peu de la phase relative entre ces franges. On pourra donc confondre la fonction de Wigner avec la pseudo-fonction de Wigner pour les états à 0 ou 1 photon.

Concluons sur la sensibilité de la pseudo-fonction de Wigner à la phase Φ . La figure I.27 montre les pseudo-fonctions de Wigner des états $|0\rangle$ et $|1\rangle$ pour $\Phi \in [2.5, 3.21, 2.6]$. On voit que la dépendance à Φ n'est pas très marquée pour 1 photon et complètement négligeable pour l'état vide. Pour un photon, l'effet est le plus marqué au voisinage de l'origine. Il apparaît qu'un écart de quelques dixièmes de radians a un effet modéré sur la fonction de Wigner.

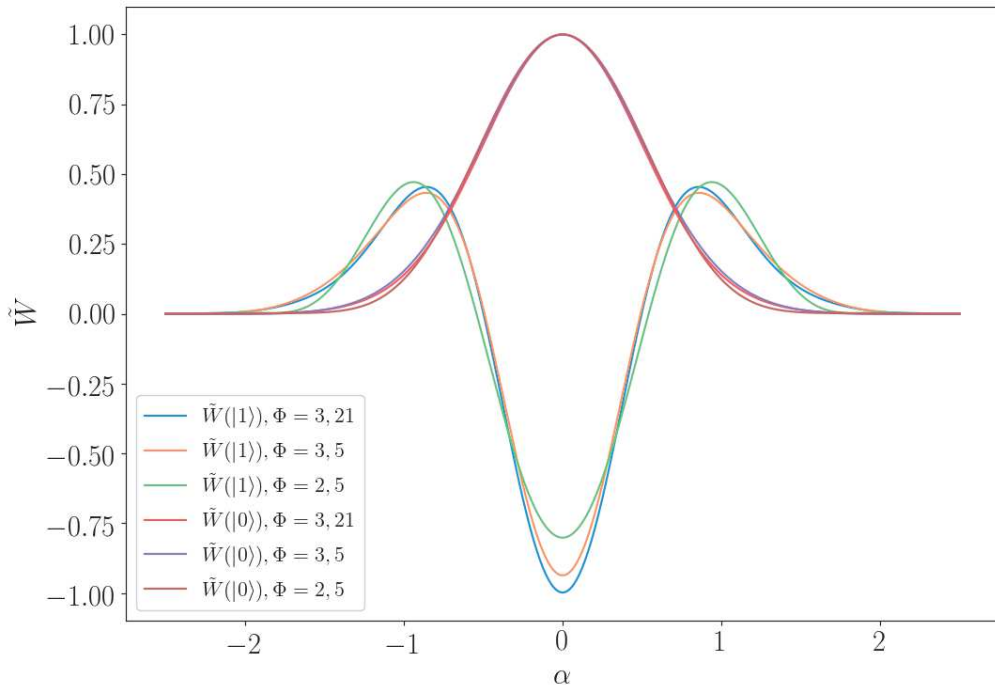


Fig. I.27 Pseudo-fonctions de Wigner des états $|0\rangle$ et $|1\rangle$, pour plusieurs valeurs de Φ .

Calibration du déphasage par photon

Le déphasage par photon dépend du désaccord entre l'atome et la cavité. On le contrôle en réglant le potentiel appliqué au killer de chaque cavité, c'est-à-dire au miroir inférieur de C_2 et aux demi-anneaux de C_1 . Comme les atomes QND sont sensibles aux inhomogénéités de champ électrique, qui brouillent la phase de la superposition atomique, on veut appliquer le champ électrique le plus faible possible au killer. Ce champ doit juste suffire à maintenir l'axe de quantification. On ajuste donc la fréquence de chaque cavité de manière à ce qu'elle soit supérieure à celle de la transition atomique et que le désaccord cherché soit obtenu pour ce champ électrique minimum. Dans les expériences présentées ici, on souhaite effectuer une mesure de parité. Celle-ci correspond, d'après le paragraphe précédent, à un désaccord $\Delta_j^\pi = 64$ kHz. Dans les séquences où on utilise des atomes QND, le désaccord initial pour les oscillations de Rabi sera donc de l'ordre de Δ_j^π , ce qui nécessite de faire des oscillations de Rabi à faible désaccord initial, comme présentées dans la section I.5.3.b.

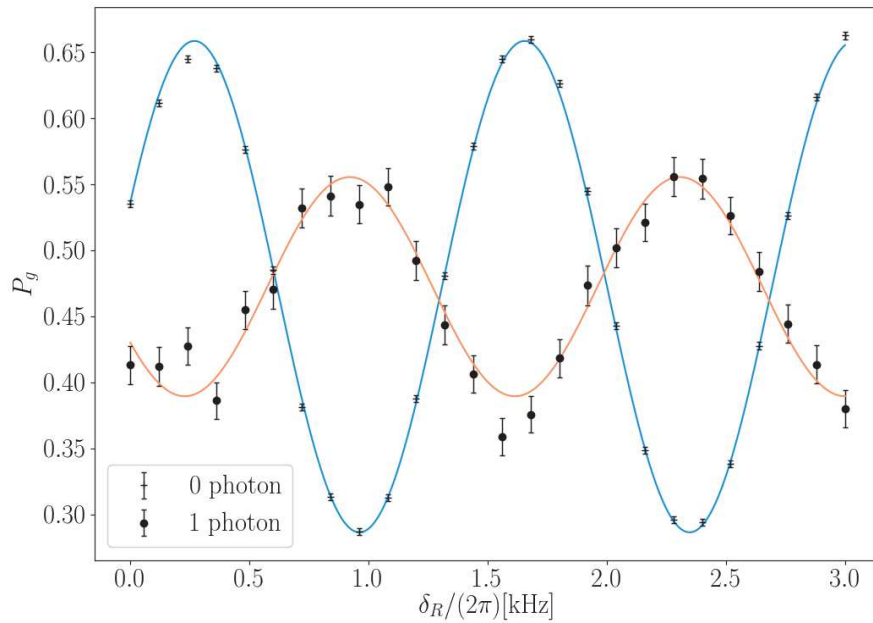


Fig. I.28 Oscillations de Ramsey dans C_1 , pour les champs $|0\rangle$ et $|1\rangle$. Les impulsions $\pi/2$ sont appliquées dans les zones de Ramsey R_1 et R_3 . La courbe bleue est un ajustement de la courbe à zéro photon par la fonction $P = P_0 + C/2 \cos(\delta_R T_R + \phi_0)$ avec $P_0 = 0.472 \pm 0.001$, $C = 0.372 \pm 0.001$, $T_R = 722 \pm 1$ μ s et $\phi_0 = -1.21 \pm 0.01$ rad. La courbe orange est un ajustement de la courbe à un photon par la fonction $P = P_0 + (1 - p_1)C/2 \cos(\delta_R T_R + \phi_0) + p_1 C/2 \cos(\delta_R T_R + \phi_0 + \Phi)$ où tous les paramètres sont pris identiques à ceux obtenus pour la courbe à 0 photon, sauf $\Phi = 3.26 \pm 0.03$ rad et $p_1 = 0.72 \pm 0.01$.

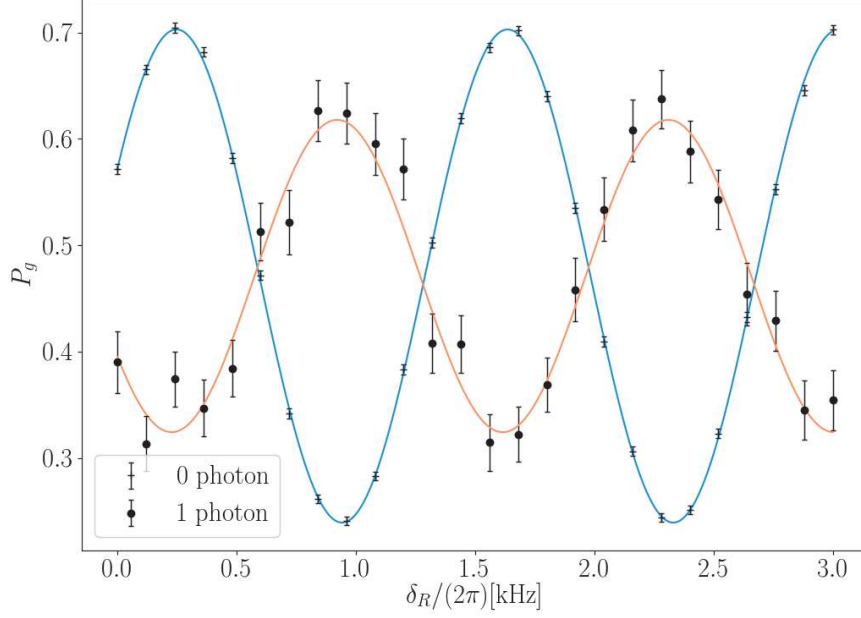


Fig. 1.29 Oscillations de Ramsey dans C_2 , pour les champs $|0\rangle$ et $|1\rangle$. Les impulsions $\pi/2$ sont appliquées dans les zones de Ramsey R_1 et R_3 . La courbe bleue est un ajustement de la courbe à zéro photon par la fonction $P = P_0 + C/2 \cos(\delta_R T_R + \phi_0)$ avec $P_0 = 0.471 \pm 0.001$, $C = 0.464 \pm 0.002$, $T_R = 720 \pm 1 \mu\text{s}$ et $\phi_0 = -1.11 \pm 0.01$ rad. La courbe orange est un ajustement de la courbe à un photon par la fonction $P = P_0 + (1 - p_1)C/2 \cos(\delta_R T_R + \phi_0) + p_1 C/2 \cos(\delta_R T_R + \phi_0 + \Phi)$ où tous les paramètres sont pris identiques à ceux obtenus pour la courbe à 0 photon, sauf $\Phi = 3.21 \pm 0.04$ rad et $p_1 = 0.82 \pm 0.02$.

Pour calibrer expérimentalement le déphasage par photon, on réalise deux mesures d'oscillations de Ramsey. Dans la première, le champ est préparé dans l'état vide, tandis que, dans la deuxième, il est préparé dans l'état $|1\rangle$ par un atome injecteur préparé dans $|e\rangle$ et effectuant une oscillation π dans la cavité. Pour la mesure de l'oscillation avec un photon, on ne garde que les cas où on détecte un seul atome injecteur dans $|g\rangle$. On compare alors la phase des deux oscillations obtenues, qui correspond à l'écart $\Phi \equiv \phi(1) - \phi(0)$. Les figures 1.28 et 1.29 montrent les résultats obtenus. Le contraste de l'oscillation est réduit dans le cas à 1 photon à cause de la probabilité p_0 que le mode soit toujours dans l'état vide parce que l'atome n'est pas parvenu à injecter. Les ajustements prennent en compte cet effet. On en tire les phases $\Phi_1 = 3.26 \pm 0.03$ rad et $\Phi_2 = 3.21 \pm 0.04$ rad.

Calibration de l'injection

Pour mesurer la fonction de Wigner, on doit effectuer des déplacements du champ. On a vu dans la section 1.1.2 qu'on peut déplacer un mode du champ en le couplant

à un courant classique. On va expliquer comment on contrôle l'amplitude complexe du déplacement.

La calibration de l'amplitude d'injection fait appel à l'interférométrie de Ramsey. On choisit la valeur de Δ_r de façon à se placer au sommet des franges de Ramsey lorsque la cavité est vide. On injecte un champ cohérent dans la cavité en allumant la source pendant une durée t_s . Le champ injecté est donc $\alpha_j = \gamma_j t_s$. On envoie ensuite 10 échantillons d'atomes effectuant une mesure QND du nombre de photons. Les multiples composantes de Fock du champ cohérent donnent lieu à des franges d'amplitude et de phase différentes. On observe donc un brouillage progressif à mesure que l'amplitude du champ cohérent augmente.

La probabilité de trouver l'atome dans $|g\rangle$ pour n photons s'écrit :

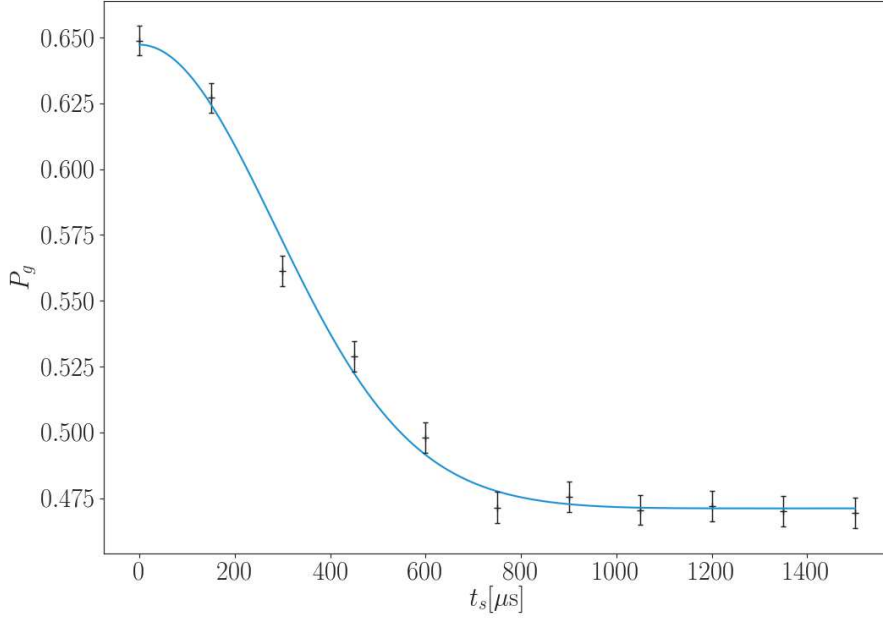
$$P_g^{(n)} = \pi_{0,j} + \frac{C_{F,j}}{2} \cos(\phi_j(n) - \phi(0)), \quad (\text{I.5.11})$$

où $C_{F,j}$ est le contraste des franges de Ramsey. La probabilité de trouver l'atome dans $|g\rangle$ est donc :

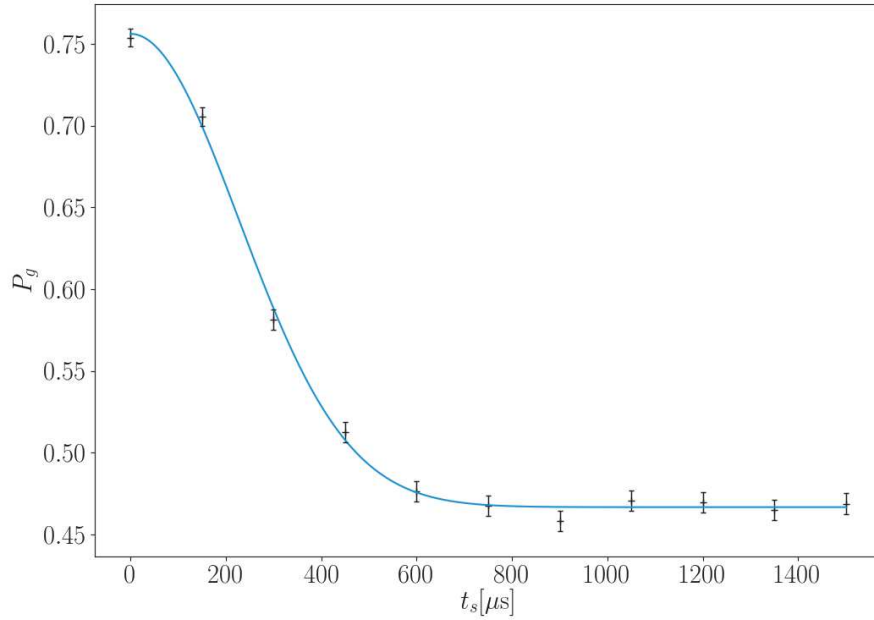
$$P_g = \pi_{0,j} + \frac{C_{F,j}}{2} e^{-|\gamma_j t_s|^2} \sum_n \frac{|\gamma_1 t_s|^{2n}}{n!} \cos(\phi_j(n) - \phi_j(0)). \quad (\text{I.5.12})$$

Les figures [I.30a](#) et [I.30b](#) montrent les résultats obtenus. On en tire les paramètres $\gamma_1 = 1.8$ /ms et $\gamma_2 = 2.2$ /ms.

Pour mesurer la fonction de Wigner, on souhaite varier la valeur de l'amplitude d'injection. On est alors confronté au problème suivant : pour atteindre de faibles valeurs de l'amplitude, on doit faire des injections de courte durée. Ces injections ont alors une grande largeur spectrale et on ne peut plus injecter sélectivement dans une cavité. Si on diminue la puissance de la source pour réduire γ et augmenter la durée des impulsions, la durée d'injection pour balayer les injections plus élevées n'est plus négligeable par rapport au temps de vie des cavités. Pour éviter ces difficultés, nous avons préféré balayer la puissance de la source. Les figures [I.31a](#) et [I.31b](#) montrent les courbes obtenues lorsqu'on effectue une mesure de la fonction de Wigner de l'état vide en balayant la puissance de la source. On effectue des ajustements sur ces courbes en prenant en compte la non-linéarité des quadrupleurs micro-onde. Cette méthode est celle qui sera utilisée pour calibrer les injections. En pratique, la calibration est effectuée en même temps que la mesure en utilisant les réalisations où aucun atome injecteur n'est détecté.

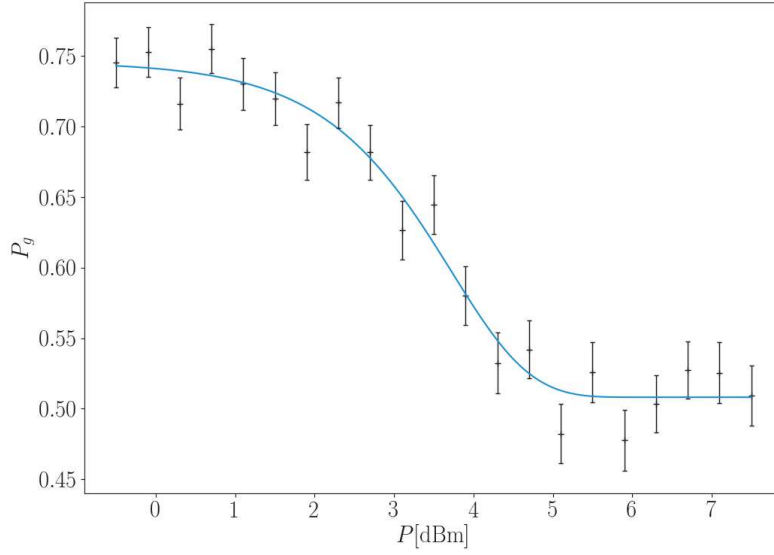


(a) Calibration dans C_1 : $\pi_{0,1} = 0.471 \pm 0.002$, $C_{F,1} = 0.35 \pm 0.01$ et $\gamma_1 = 1.76 \pm 0.07 \text{ ms}^{-1}$.

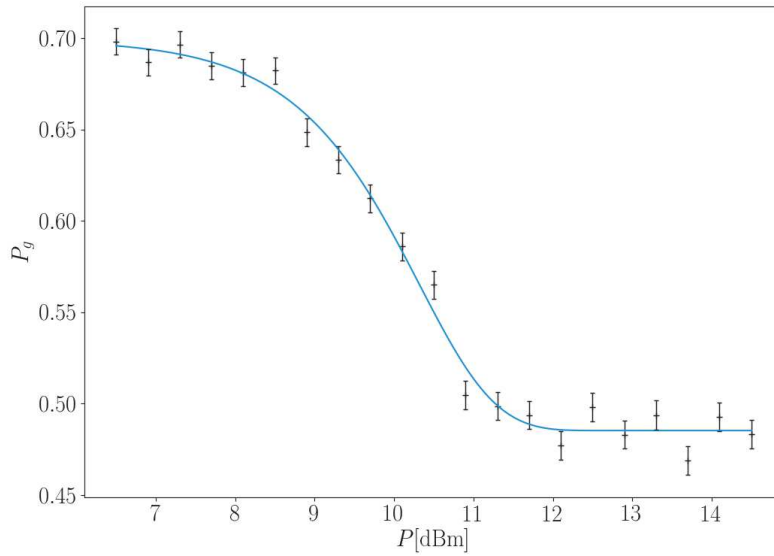


(b) Calibration dans C_1 : $\pi_{0,2} = 0.467 \pm 0.002$, $C_{F,2} = 0.58 \pm 0.01$ et $\gamma_2 = 2.21 \pm 0.06 \text{ ms}^{-1}$.

Fig. I.30 Probabilité de mesurer les atomes QND dans $|g\rangle$ après injection d'un champ dans la cavité C_2 , en fonction de la durée d'injection t_s . Les courbes bleues sont des ajustements par la fonction $P = \pi_{0,j} + C_{F,j}/2e^{-|\gamma_j t_s|^2} \sum_n |\gamma_j t_s|^{2n}/n! \cos[\phi_j(n) - \phi_j(0)]$.



(a) Calibration pour C_1 : $\pi_{0,1} = 0.508 \pm 0.008$, $C_{F,1} = 0.476 \pm 0.03$, $a = 0.45 \pm 0.09$ et $P_0 = -3 \pm 1$ dBm.



(b) Calibration pour C_2 : $\pi_{0,2} = 0.485 \pm 0.003$, $C_{F,2} = 0.43 \pm 0.01$, $a = 0.47 \pm 0.04$ et $P_0 = 3.90 \pm 0.6$ dBm

Fig. I.31 Probabilité de mesurer les atomes QND dans $|g\rangle$ après injection d'un champ dans la cavité C_j , en fonction de la puissance de la source S_i . Les courbes bleues sont des ajustements par la fonction $P = \pi_{0,j} + C_{F,j}/2e^{-|\alpha(P)|^2} \sum_n \alpha(P)^{2n}/n! \cos[\phi_j(n) - \phi_j(0)]$, où $\alpha(P) = \sqrt{1/2000 \cdot 10^{a(P-P_0)}}$.

Calibration de la phase

Les paramètres α_1 et α_2 de la fonction de Wigner sont des nombres complexes dont les arguments correspondent à la phase de l'injection de mesure. Dans nos expériences, nous avons besoin de contrôler précisément la phase relative entre les deux injections. Dans ce but, nous utilisons le circuit de la figure I.32. Les signaux des deux sources micro-onde indépendantes S_1 et S_2 sont divisés en deux. On envoie une partie des signaux vers le circuit d'injection de l'expérience. Les deux autres signaux sont multipliés, puis on filtre la composante haute fréquence. Il en résulte un signal à la différence de fréquence entre les deux signaux, qui vaut δ puisque ces deux signaux sont accordés aux cavités. Ce signal est envoyé sur un comparateur qui envoie à l'ordinateur de contrôle de l'expérience un TTL lors de son passage ascendant par zéro. Ce TTL est utilisé pour déclencher la séquence expérimentale, de sorte que la phase relative entre les deux sources soit toujours la même à un instant donné de l'expérience.

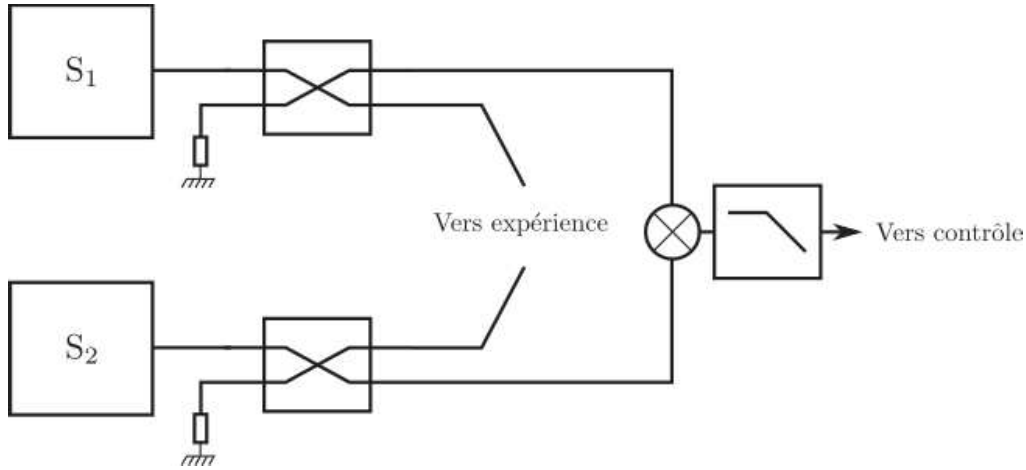


Fig. I.32 Schéma de principe du circuit utilisé pour les injections micro-onde. Les champs produits par les sources S_1 et S_2 sont divisés en deux. Une partie va vers l'expérience, tandis que l'autre va vers un multiplicateur où les deux signaux sont mélangés. Le battement obtenu est filtré pour garder la composante basse fréquence qu'on envoie vers l'ordinateur de contrôle de l'expérience afin de déclencher la séquence expérimentale pour une phase relative donnée des deux injections.

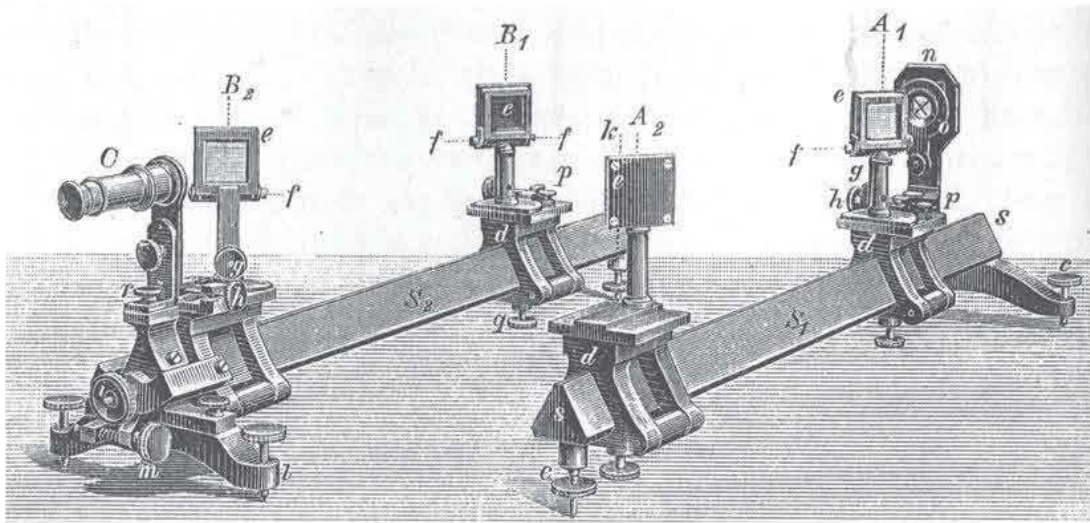
Mesure de la fonction de Wigner

La mesure de la fonction de Wigner est réalisée en combinant la mesure de parité et l'injection micro-onde décrites précédemment. Après la préparation de l'état, on effectue des injections micro-ondes dans les deux cavités. On envoie ensuite des échantillons d'atomes QND. Pour effectuer une mesure de la fonction Wigner à deux modes, il faut mesurer la parité globale. On effectue donc l'interférométrie de Ramsey entre les deux zones de Ramsey externes, de sorte que les atomes acquièrent une phase globale égale à la somme des phases acquises dans chaque cavité. Comme le déphasage par photon n'est pas

exactement linéaire on ne mesure pas la fonction de Wigner, mais une pseudo-fonction de Wigner, qui est très proche de la fonction de Wigner pour des petits nombres de photons. Les résultats obtenus seront présentés dans le chapitre II.

Conclusion du chapitre I : un système bien contrôlé pour étudier des états quantiques délocalisés

Dans ce chapitre, nous avons introduit le formalisme permettant de traiter l'interaction d'un atome à deux niveaux avec un mode d'une cavité micro-onde. Nous avons ensuite présenté le dispositif expérimental permettant de se mettre dans les conditions d'application de ce formalisme. Nous avons ensuite montré que nous sommes en mesure de contrôler l'interaction entre les atomes et les deux cavités de notre dispositif pour réaliser des échanges cohérents de quanta d'énergie entre eux. Nous pouvons par ailleurs effectuer des mesures d'interférométrie de Ramsey permettant de mesurer la parité du champ contenu dans les cavités. Nous disposons donc de tous les éléments nécessaires pour préparer des états non-gaussiens délocalisés du champ à deux modes et effectuer des mesures de sa fonction de Wigner.



Chapitre II

Préparation et mesure d'états intriqués en électrodynamique quantique à deux cavités

Dans ce chapitre, on voit dans un premier temps comment les techniques d'électrodynamique quantique en cavité présentées dans le chapitre I permettent de préparer un état à un photon délocalisé entre les deux cavités. On étudie ses propriétés, notamment en termes d'intrication et de fonction de Wigner. On montre par la suite qu'une manière de sonder les cohérences non locales entre les deux composantes de la superposition est d'effectuer une mesure interférométrique qui est directement sensible à la phase quantique de la superposition. Enfin, on montre les résultats obtenus sur la mesure interférométrique et la mesure de coupes de la fonction de Wigner de l'état préparé, qui mettent en évidence une interférence quantique entre les deux modes et sa disparition progressive à cause de la décohérence.

II.1 Préparation d'un état délocalisé d'un photon

Notre dispositif est particulièrement adapté à l'étude d'états non-gaussiens, puisqu'il permet de réaliser facilement des opérations d'addition et de soustraction de photons, ainsi que des mesures de parité permettant de préparer des états quantiques exotiques par projection de la fonction d'onde [65]. Avec deux cavités, on dispose donc de tous les outils pour préparer des états intriqués non-gaussiens [66].

Dans un premier temps, nous avons utilisé un protocole déjà appliqué à deux modes d'une même cavité micro-onde [35] pour préparer un état intriqué $|10\rangle + |01\rangle$ dans lequel un photon est partagé par les deux cavités. Cet état est le cas le plus simple d'un ensemble d'états appelés *états noon*, présentant un intérêt pour l'interférométrie lorsque le nombre de photons est plus élevé. Nous avons ensuite tenté de préparer l'état *noon* à deux photons. Dans cette section, on va détailler les protocoles permettant de préparer des états intriqués à un et deux photons et présenter les propriétés des états *noon*.

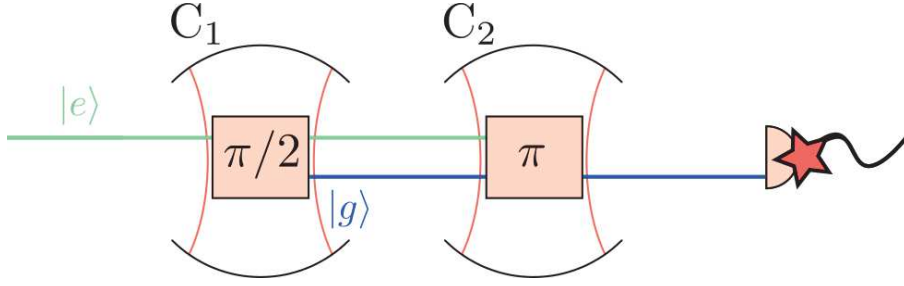


Fig. II.1 Schéma de principe de la préparation de l'état $|10\rangle + |01\rangle$. Les cavités sont initialement préparées dans l'état vide tandis que l'atome est dans l'état excité. Ce dernier effectue d'abord une impulsion $\pi/2$ dans C_1 , qui les intrique, puis une impulsion π dans C_2 , qui transfère l'intrication vers cette dernière. L'atome est alors détecté dans $|g\rangle$ et les cavités sont intriquées.

II.1.1 La séquence d'impulsions

La préparation de l'état $|10\rangle + |01\rangle$ utilise un atome qui sert de médiateur pour transmettre l'intrication d'une cavité à l'autre. L'atome interagit successivement à résonance avec C_1 et C_2 , initialement préparées dans l'état vide, effectuant les impulsions présentées sur la figure II.1. Lors de son passage dans la première cavité, il effectue une première impulsion $\pi/2$, qui les prépare dans un état maximalement intriqué. Dans la deuxième cavité, il effectue une impulsion π , qui réalise une opération SWAP qui échange l'état de l'atome avec celui de la cavité. L'atome est alors dans l'état fondamental et l'état des cavités peut se factoriser :

$$|e, 0, 0\rangle \xrightarrow{\pi/2, C_1} \frac{1}{\sqrt{2}} (|e, 0, 0\rangle + |g, 1, 0\rangle) \xrightarrow{\pi, C_2} \frac{1}{\sqrt{2}} |g\rangle \otimes (|1, 0\rangle + e^{i\psi_p} |0, 1\rangle), \quad (\text{II.1.1})$$

où ψ_p est une phase qui prend en compte la phase acquise par l'atome pendant et entre les interactions avec les deux cavités. Une fois enlevée la composante atomique de l'état, qui se factorise, l'état préparé des cavités est :

$$|\psi\rangle = \frac{1}{\sqrt{2}} (|10\rangle + e^{i\psi_p} |01\rangle). \quad (\text{II.1.2})$$

II.1.2 Les états *noon*

Un état *noon* est un état intriqué dans lequel n photons sont dans une superposition entre deux modes :

$$|\psi_n\rangle = \frac{1}{\sqrt{2}} (|n0\rangle + |0n\rangle). \quad (\text{II.1.3})$$

Ces états présentent un intérêt pour la métrologie quantique. Supposons qu'on dispose d'un interféromètre similaire à l'interféromètre de Mach-Zehnder de la figure II.2, mais dans lequel la lame séparatrice d'entrée prépare l'état $|\psi_n\rangle$. Dans un des deux bras se trouve un échantillon à mesurer qui introduit un déphasage ϕ entre les deux bras. Lors de sa traversée de l'interféromètre l'état *noon* évolue comme :

$$\hat{U}_\phi |\psi_n\rangle = \frac{1}{\sqrt{2}} (|n0\rangle + e^{in\phi} |0n\rangle). \quad (\text{II.1.4})$$

La phase acquise est proportionnelle au nombre de photons. On voit donc que l'état est

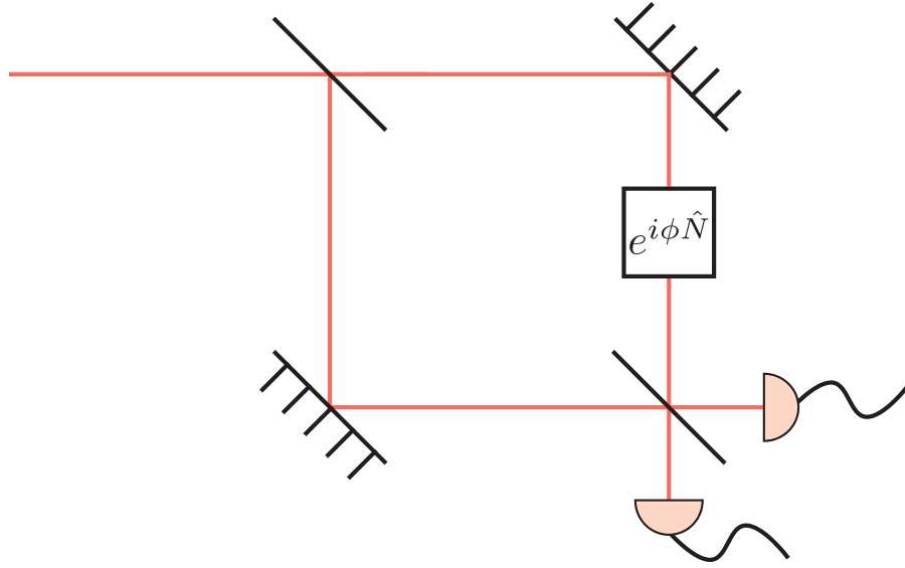


Fig. II.2 Interféromètre de Mach-Zehnder. Dans l'un des deux bras, un échantillon ajoute une phase ϕ , que l'on cherche à mesurer.

d'autant plus sensible que le nombre de photons est grand. Les états *noon* sont particulièrement adaptés pour ce genre de mesures et on peut montrer qu'ils permettent de battre la limite quantique standard sur la mesure de ϕ et d'atteindre la limite de Heisenberg. Cette évolution rapide de la phase quantique se retrouve directement sur la fonction de Wigner. La figure II.3 montre la fonction de Wigner de l'état $|70\rangle + |07\rangle$. Cette fonction est invariante par rotation d'une des coordonnées complexes d'un septième de tour. Si on fixe l'amplitude des injections et qu'on fait varier la phase relative des injections, l'oscillation obtenue a donc une fréquence sept fois plus élevée que celle qu'on aurait pour l'état $|10\rangle + |01\rangle$.

II.1.3 Propriétés de l'état $|10\rangle + |01\rangle$

L'état $|10\rangle + e^{i\phi} |01\rangle$ est un état propre de l'opérateur parité globale, qui contient un seul photon. L'information pertinente pour l'interférométrie dans cet état est la phase quantique. Lorsqu'on laisse évoluer cet état dans les deux cavités, à des échelles de temps suffisamment faibles devant la décohérence, il se transforme comme :

$$\begin{aligned} |10\rangle + e^{i\phi_0} |01\rangle &\mapsto e^{-i\omega_{c1}(t-t_0)} \left(|10\rangle + e^{i\phi_0 - i\delta(t-t_0)} |01\rangle \right), \\ &\mapsto \hat{R}_1(-\omega_{c1}t) \otimes \hat{R}_2(-\omega_{c2}t) \left(|10\rangle + e^{i\phi_0} |01\rangle \right), \end{aligned} \quad (\text{II.1.5})$$

où $\delta = \omega_{c2} - \omega_{c1}$. Il s'agit de l'évolution déjà mentionnée dans la section I.1.4.

La fonction de Wigner de l'état $|10\rangle + |01\rangle$ a pour équation (voir annexe E)

$$W^{|10\rangle+|01\rangle}(\alpha_1, \alpha_2) = (2(\alpha_1 + \alpha_2)(\alpha_1^* + \alpha_2^*) - 1) e^{-2(|\alpha_1|^2 + |\alpha_2|^2)}. \quad (\text{II.1.6})$$

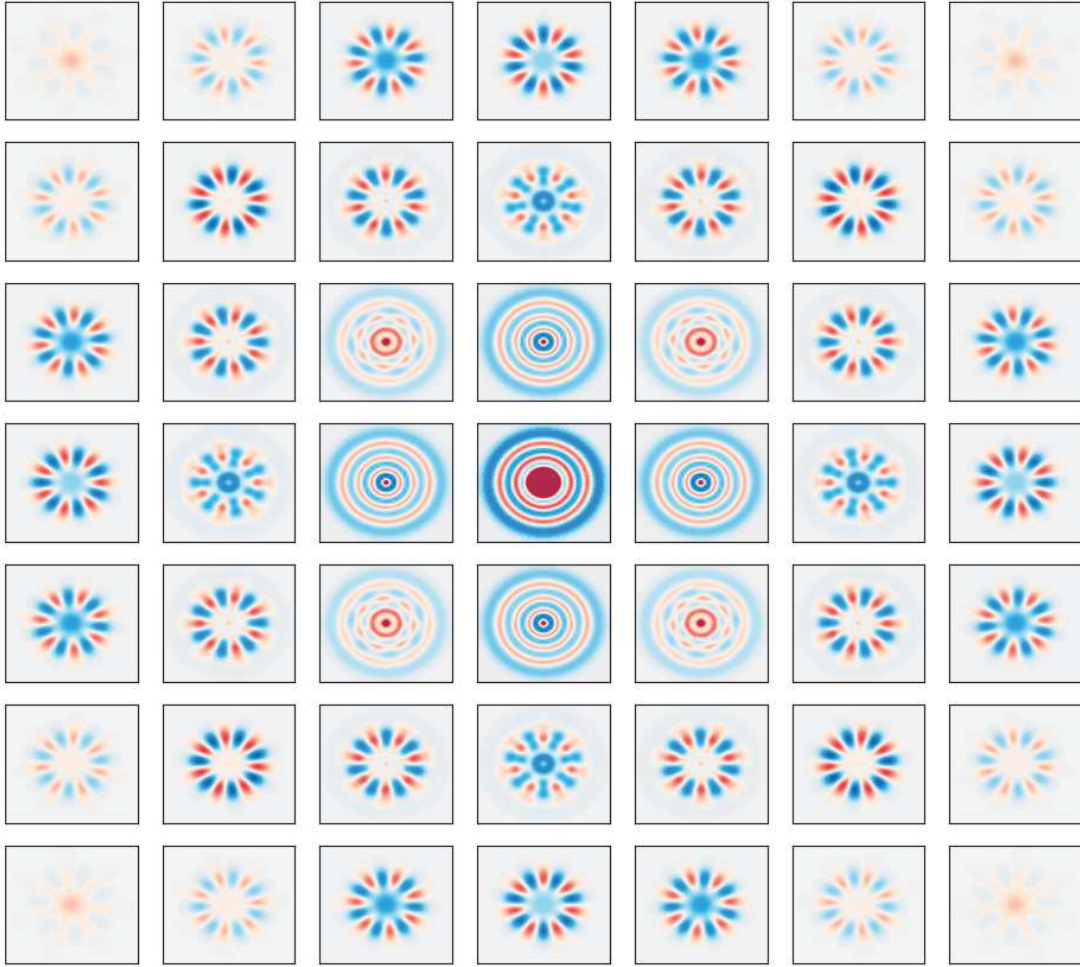


Fig. II.3 — Fonction de Wigner de l'état $|70\rangle + |07\rangle$. $\Re\alpha_1$ et $\Im\alpha_1$ varient entre -1,5 et 1,5. $\Re\alpha_2$ et $\Im\alpha_2$ varient entre -3 et 3. La fonction est invariante par rotation d'un des paramètres complexes d'un septième de tour. Le contraste est exagéré pour plus de visibilité.

(a) *Fonction de Wigner de l'état $|10\rangle + |01\rangle$* (b) *Fonction de Wigner de l'état $|10\rangle - i|01\rangle$*

Fig. II.4 Fonctions de Wigner des états $|10\rangle + |01\rangle$ et $|10\rangle - i|01\rangle$. Après une durée $\pi/(4\delta)$, l'état $|10\rangle + |01\rangle$ est transformé en l'état $|10\rangle - i|01\rangle$. La phase de cet état nous renseigne donc sur le désaccord entre les cavités.

Elle est tracée figure II.4a. Lorsque l'injection est nulle dans les deux cavités, on mesure simplement la parité de l'état qui vaut -1 puisqu'il y a un nombre impair de photons. Si on injecte dans les deux cavités en phase, la fonction de Wigner devient positive lorsque les amplitudes d'injection sont suffisamment grandes, tandis qu'elle est négative si on injecte dans les deux cavités en opposition de phase avec la même amplitude. On verra dans le paragraphe suivant une façon d'interpréter cette propriété. La figure II.4b montre la fonction de Wigner de l'état $|10\rangle - i|01\rangle$. Celle-ci correspond à l'état obtenu après un quart de période d'évolution à la pulsation $-\delta$. Les fonctions obtenues en balayant α_2 ont tourné d'un quart de tour dans le sens anti-trigonométrique.

Fonction de Wigner à un photon et changement de base

Pour mieux comprendre la fonction de Wigner de l'état $|10\rangle + |01\rangle$, on peut effectuer un changement de base de modes pour prendre une base qui prend directement en compte la délocalisation du photon. On va supposer pour simplifier que les deux modes ont la même fréquence et la même pulsation de Rabi Ω_0 au centre du mode, positive. Considérons les opérateurs d'annihilation $\hat{b}_1$ et $\hat{b}_2$ définis par :

$$\hat{b}_1 = \frac{\hat{a}_1 + \hat{a}_2}{\sqrt{2}}, \quad \hat{b}_2 = \frac{\hat{a}_1 - \hat{a}_2}{\sqrt{2}}. \quad (\text{II.1.7})$$

Ces opérateurs vérifient :

$$[\hat{b}_i, \hat{b}_i^\dagger] = 1, \quad (\text{II.1.8a})$$

et, pour $i \neq j$,

$$[\hat{b}_i, \hat{b}_j^\dagger] = 0, \quad (\text{II.1.8b})$$

$$[\hat{b}_i, \hat{b}_j] = 0, \quad (\text{II.1.8c})$$

ce qui en fait une base de modes convenable pour le champ compris dans les cavités. On vérifie par ailleurs que :

$$\hat{b}_1^\dagger \hat{b}_1 + \hat{b}_2^\dagger \hat{b}_2 = \hat{a}_1^\dagger \hat{a}_1 + \hat{a}_2^\dagger \hat{a}_2. \quad (\text{II.1.9})$$

Le terme d'interaction du hamiltonien de Jaynes-Cummings peut se récrire dans cette nouvelle base :

$$\hat{H}_{int} = -\frac{i\hbar\Omega_0}{2} \left(\hat{\sigma}_+ \hat{b}_1 \tilde{f}_1(\mathbf{r}) - \hat{\sigma}_- \hat{b}_1^\dagger \tilde{f}_1(\mathbf{r}) + \hat{\sigma}_+ \hat{b}_2 \tilde{f}_2(\mathbf{r}) - \hat{\sigma}_- \hat{b}_2^\dagger \tilde{f}_2(\mathbf{r}) \right), \quad (\text{II.1.10})$$

avec

$$\tilde{f}_1(\mathbf{r}) = \frac{f_1(\mathbf{r}) + f_2(\mathbf{r})}{\sqrt{2}}, \quad \tilde{f}_2(\mathbf{r}) = \frac{f_1(\mathbf{r}) - f_2(\mathbf{r})}{\sqrt{2}}. \quad (\text{II.1.11})$$

Ces fonctions sont montrées sur la figure II.5b. Il s'agit d'un changement de base non local : la nouvelle base fait intervenir des modes qui ont des composantes sur les deux cavités.

Dans cette nouvelle base, l'état s'écrit simplement :

$$\frac{|10\rangle_a + |01\rangle_a}{\sqrt{2}} = \hat{b}_1^\dagger |00\rangle_b = |10\rangle_b. \quad (\text{II.1.12})$$

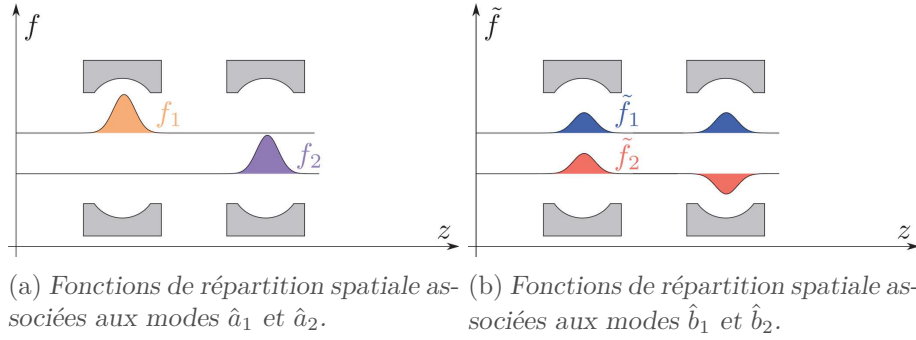


Fig. II.5 Changement de mode $\{\hat{a}_1, \hat{a}_2\} \leftrightarrow \{\hat{b}_1, \hat{b}_2\}$. Les modes de la première base sont localisés chacun sur une cavité, tandis que ceux de la deuxième sont délocalisés sur les cavités et diffèrent par la phase relative des champs associés aux deux cavités.

Il s'agit d'un état dans lequel un photon se trouve dans le mode $\tilde{f}_1$. On peut exprimer la fonction de Wigner de l'état $|10\rangle + |01\rangle$ à partir de la fonction de Wigner dans la base $\{\hat{b}_1, \hat{b}_2\}$. En effet, on montre facilement que :

$$\hat{D}_1^a(\alpha_1)\hat{D}_2^a(\alpha_2) = \hat{D}_1^b(\tilde{\alpha}_1)\hat{D}_2^b(\tilde{\alpha}_2), \quad (\text{II.1.13})$$

où $\hat{D}_i^k$ désigne un déplacement dans le mode i de la base k et où

$$\tilde{\alpha}_1 = \frac{\alpha_1 + \alpha_2}{\sqrt{2}}, \quad \tilde{\alpha}_2 = \frac{\alpha_1 - \alpha_2}{\sqrt{2}}. \quad (\text{II.1.14})$$

Comme l'opérateur parité est le même dans les deux bases, en vertu de l'identité II.1.9, on en déduit qu'on a l'identité :

$$W^{|10\rangle+|01\rangle}(\alpha_1, \alpha_2) = W^{|10\rangle}(\tilde{\alpha}_1, \tilde{\alpha}_2). \quad (\text{II.1.15})$$

Lorsqu'on choisit des amplitudes complexes égales $\alpha_1 = \alpha_2 = \alpha$ de la fonction de Wigner de l'état $|10\rangle + |01\rangle$, on a $\tilde{\alpha}_1 = \sqrt{2}\alpha$ et $\tilde{\alpha}_2 = 0$. La fonction obtenue est donc la fonction de Wigner d'un état de Fock à un photon dans un seul mode :

$$W^{|10\rangle+|01\rangle}(\alpha, \alpha) = W^{|1\rangle}(\sqrt{2}\alpha). \quad (\text{II.1.16})$$

Si on prend en revanche des amplitudes de phases opposées $\alpha_1 = -\alpha_2 = \alpha$, alors $\tilde{\alpha}_1 = 0$ et $\tilde{\alpha}_2 = \sqrt{2}\alpha$. On obtient la fonction de Wigner du vide, qui est multipliée par -1 puisqu'on a toujours 1 photon dans le mode $\hat{b}_2^\dagger$:

$$W^{|10\rangle+|01\rangle}(\alpha, -\alpha) = -W^{|0\rangle}(\sqrt{2}\alpha). \quad (\text{II.1.17})$$

On a donc montré qu'on pouvait voir l'état $|10\rangle + |01\rangle$ comme un état de Fock monomode, dans une base adéquate, au prix d'utiliser des modes qui présentent plusieurs composantes séparées dans l'espace. C'est en fait toujours possible pour un état à un photon ([43], chapitre V). Ce ne serait plus le cas pour un état *noon* à plusieurs photons.

Oscillations de la fonction de Wigner

Lorsqu'on fixe les modules de α_1 et α_2 et qu'on fait varier $\phi_w \equiv \arg \alpha_1 - \arg \alpha_2$, la fonction de Wigner des états *noon* varie sinusoïdalement. L'oscillation qui en résulte effectue une période lorsque ϕ_w varie de 2π . On peut montrer que le contraste maximum est obtenu lorsque $|\alpha_1| = |\alpha_2| \equiv r$.

La figure II.6 montre des coupes de la fonction de Wigner pour $\alpha_1 = \alpha_2 \in \mathbb{R}$, c'est-à-dire pour $\phi_w = 0$, pour différents états. On retrouve les résultats du paragraphe précédent. La fonction de Wigner de l'état $|10\rangle + |01\rangle$ est égale dans ce cas, à une homothétie près, à celle d'un état de Fock à un photon dans un mode. La fonction de Wigner de $|10\rangle - |01\rangle$ est quant à elle égale à l'opposé de celle de l'état vide. Enfin, les coupes des fonctions de Wigner des états $|10\rangle \pm i|01\rangle$, $|10\rangle$, $|01\rangle$ et $|10\rangle\langle 10| + |01\rangle\langle 01|$ sont identiques et situées entre les courbes de $|10\rangle + |01\rangle$ et $|10\rangle - |01\rangle$. Pour distinguer la fonction de Wigner de l'état $|10\rangle + i|01\rangle$ et celle de l'état mélangé, il faut changer ϕ_w . En faisant varier la phase de l'état $|10\rangle + e^{i\phi}|01\rangle$, on obtient un ensemble de courbes qui balaie l'espace compris entre les courbes obtenues pour $|10\rangle + |01\rangle$ et $|10\rangle - |01\rangle$.

Le contraste C_W obtenu lorsque ϕ_w varie, pour r donné, est égal à la différence au point r entre la fonction de Wigner de l'état $|10\rangle + |01\rangle$ et celle de $|10\rangle - |01\rangle$. La figure II.7 montre son évolution en fonction de r . Le maximum de contraste est atteint pour $r = 1/2$ et vaut $C_W = 0.736$. Il s'agit cependant d'un contraste de la fonction de Wigner, qui varie entre -1 et 1. Le contraste sur l'oscillation de probabilité est deux fois plus faible et vaut donc 0.368. Cette valeur est assez faible, alors qu'il s'agit de la plus grande distance sur deux mesures entre deux états purs orthogonaux. La fonction de Wigner ne semble donc pas être l'outil le plus approprié pour discriminer les états *noon*.

État $|10\rangle + |01\rangle$ et intrication

L'état $|10\rangle + |01\rangle$ est un état maximalelement intriqué des deux cavités. Une manière de le montrer serait de mesurer un signal de Bell en transférant l'intrication des cavités vers deux atomes et en mesurant ceux-ci. Cette méthode constitue une preuve définitive de l'intrication mais nécessite de faire des expériences à trois atomes, coûteuses en temps de mesure. Elle est par ailleurs sensible aux imperfections expérimentales et en particulier à la décohérence. En dessous d'un certain seuil, elle ne permet plus de mettre en évidence l'intrication, même si l'état n'est pas séparable. Enfin, la mesure du signal de Bell ne donne que très peu d'information sur l'état quantique du système.

Un autre critère pour détecter l'intrication pour des états à deux modes est la violation d'une inégalité de Bell généralisée, dans la fonction de Wigner [67] [68]. Cette méthode a été proposée pour des états chats de Schrödinger pour notre dispositif [66] et récemment démontrée pour un champ micro-onde couplé à un circuit supraconducteur [69]. Elle est cependant peu adaptée aux états *noon*. En effet, leur fonction de Wigner a une enveloppe gaussienne centrée sur l'origine, tandis que les termes dus aux cohérences non locales sont nuls à l'origine, ce qui réduit la valeur qu'on peut espérer observer sur un signal de Bell dans l'espace des phases. Il devient alors plus difficile de mettre en évidence une violation d'inégalité de Bell, et même impossible pour $n \geq 3$.

Dans la suite, nous allons développer une autre approche, inspirée de l'interférométrie, pour étudier les cohérences non locales responsables de l'intrication de l'état $|10\rangle + |01\rangle$.

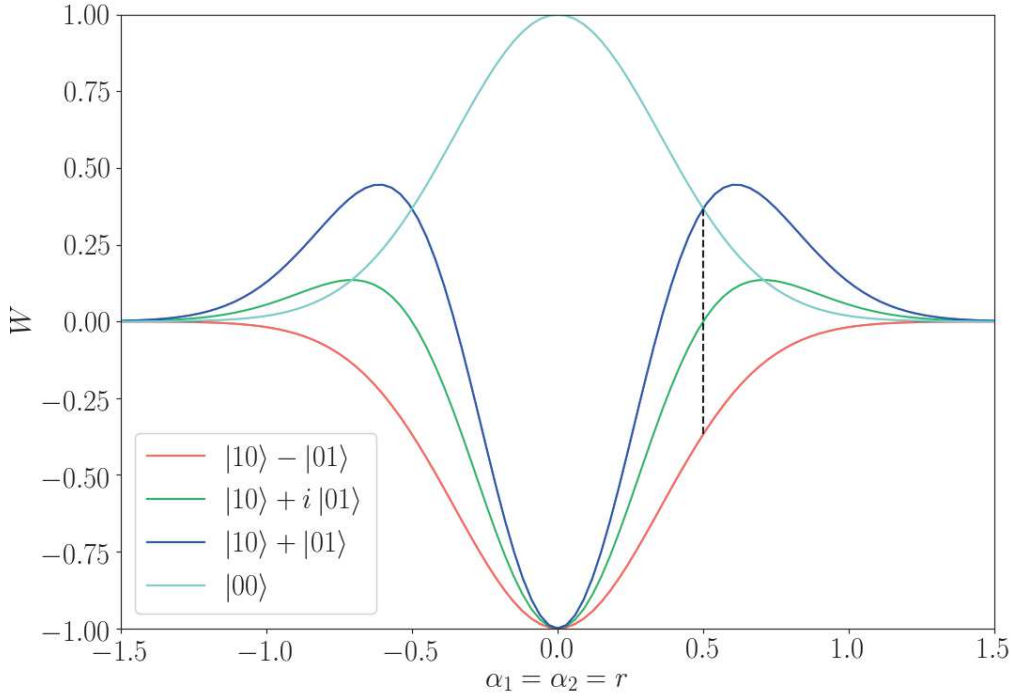


Fig. II.6 Fonctions de Wigner de différents états pour une injection identique dans les deux cavités $\alpha_1 = \alpha_2 = r$. La courbe obtenue pour l'état $|10\rangle + i|01\rangle$ est identique à celles obtenues pour les états $|10\rangle - i|01\rangle$, $|10\rangle$, $|01\rangle$ et $|10\rangle\langle 10| + |01\rangle\langle 01|$.

II.1.4 Préparation d'un état à deux photons

Une fois l'état $|10\rangle + |01\rangle$ préparé, on peut tenter de préparer un état $|20\rangle + |02\rangle$ en injectant un deuxième photon dans les cavités. Dans ce but, on envoie un deuxième atome injecteur, préparé lui aussi dans $|e\rangle$, qui effectue des oscillations de Rabi dans les cavités. Lors du passage de cet atome dans chaque cavité, on veut qu'il injecte un photon si la cavité en contient déjà un et qu'il reste dans l'état excité si elle est vide. Une manière de satisfaire la deuxième condition est d'effectuer une impulsion 2π dans l'état vide. La durée de l'impulsion est alors $t_{2\pi} = 2\pi/\Omega_{0,j}$.

Cette durée d'impulsion permet par ailleurs de satisfaire approximativement la première condition. En effet, dans le champ de 1 photon, la fréquence de l'oscillation de Rabi de l'atome est $\sqrt{2}$ fois plus grande. En utilisant la formule :

$$2\sqrt{2} \simeq \pi, \quad (\text{II.1.18})$$

on en déduit que, pour une durée d'interaction $t_{2\pi}$, l'atome effectue une rotation d'angle presque 3π dans le champ d'un photon. L'état préparé est donc proche de l'état $|20\rangle + |02\rangle$.

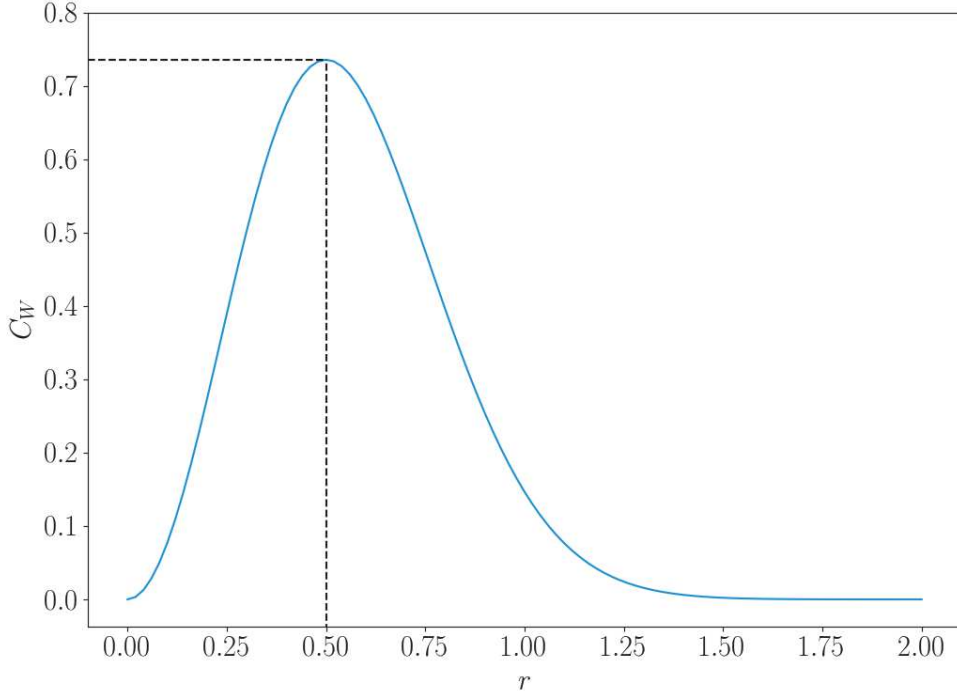


Fig. II.7 *Contraste des oscillations de la fonction de Wigner de $|10\rangle + |01\rangle$ obtenues en faisant varier la phase relative des deux injections, avec des amplitudes d'injection égales : $|\alpha_1| = |\alpha_2| = r$. Le maximum est atteint pour $r = 0.5$ et vaut $W(1/2, 1/2) - W(1/2, -1/2) = 2e^{-1} = 0.736$.*

L'état préparé dans le cas idéal est :

$$|\psi_f\rangle = \frac{c_1}{\sqrt{2}} (|e10\rangle + |e01\rangle) + \frac{c_2}{\sqrt{2}} (|g20\rangle - |g02\rangle), \quad (\text{II.1.19})$$

où $c_1 = \sqrt{p_1}$ et $c_2 = \sqrt{p_2}$ avec

$$p_1 = \frac{1 + \cos(2\sqrt{2}\pi)}{2}, \quad p_2 = \frac{1 - \cos(2\sqrt{2}\pi)}{2}. \quad (\text{II.1.20})$$

On voit qu'il suffit de post-sélectionner les cas où l'atome est dans $|g\rangle$ pour obtenir l'état $|20\rangle + |02\rangle$. La probabilité de réussite de cette post-sélection est $p_2 = 92,9\%$. On parvient donc à préparer l'état dans la plupart des cas.

II.2 La mesure résonnante

L'état $|10\rangle + |01\rangle$ présente des cohérences entre les états $|10\rangle$ et $|01\rangle$. Celles-ci sont responsables de la sensibilité de l'état à la différence de fréquence des deux cavités. Afin de les mettre en évidence, il faut recombinaison les deux composantes de la superposition.

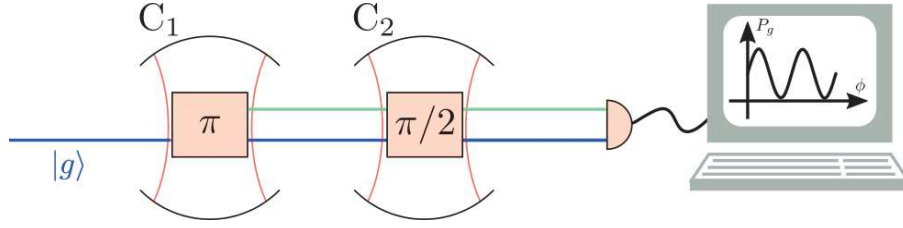


Fig. II.8 Schéma de principe de la mesure de l'état $|10\rangle + e^{i\phi} |01\rangle$ par un atome sonde. L'atome effectue d'abord une impulsion π dans C_1 , qui transfère l'état de la cavité sur l'atome, puis une impulsion $\pi/2$ dans C_2 . On mesure la probabilité de le trouver dans $|g\rangle$, qui dépend de ϕ .

Ceci peut être fait par un atome qui absorbe le photon éventuellement contenu dans la première cavité et le fait interférer avec la partie de la fonction d'onde où il est dans la deuxième cavité.

Dans cette section, nous allons présenter une mesure permettant de réaliser cette interférence. Dans un premier temps, nous décrirons la séquence expérimentale de mesure. Nous montrerons ensuite que cette méthode permet en principe de mettre en évidence l'intrication entre les deux cavités.

II.2.1 Séquence de mesure

Afin de faire battre les deux composantes de la superposition $|10\rangle + |01\rangle$, on utilise un atome résonnant, qu'on appellera dans la suite *atome sonde*, qui copie le champ contenu dans la première cavité et peut ensuite absorber ou émettre dans la deuxième cavité, en fonction de la phase quantique de la superposition.

Le principe de la mesure effectuée par l'atome sonde est présenté sur la figure II.8. Initialement préparé dans $|g\rangle$, il effectue une impulsion π dans la première cavité. L'état de la première cavité est alors transféré à l'atome et celle-ci termine dans l'état vide :

$$|g, 1, 0\rangle + e^{i\phi} |e, 0, 1\rangle \xrightarrow{\pi, C_1} |e, 0, 0\rangle + e^{i\phi} |g, 0, 1\rangle. \quad (\text{II.2.1})$$

Lors de l'interaction avec la deuxième cavité, l'atome effectue une impulsion $\pi/2$. En supposant réelle la pulsation de Rabi $\Omega_{0,2}$ dans la deuxième cavité, l'état évolue comme :

$$|e, 0, 0\rangle + e^{i\phi} |g, 0, 1\rangle \xrightarrow{\pi/2, C_2} (1 - e^{i(\phi - \psi_m)}) |e, 0, 0\rangle + e^{i\phi} (1 + e^{i(\psi_m - \phi)}) |g, 0, 1\rangle. \quad (\text{II.2.2})$$

Comme on l'a fait pour la préparation de l'état, on a introduit une phase ψ_m pour la mesure, qui prend en compte à la fois la phase acquise par l'atome pendant la traversée entre les deux cavités et pendant les interactions avec les cavités.

Afin de sonder les cohérences, on mesure la probabilité de trouver l'atome dans $|g\rangle$, pour différentes valeurs de la phase des cohérences. Si on laisse évoluer l'état, sa phase oscille à la différence de pulsation δ des cavités. Il suffit donc d'attendre entre la préparation et la détection pour observer une oscillation de la probabilité d'absorption du photon. Le contraste de cette oscillation nous donne l'amplitude des cohérences.

La figure II.9 montre la séquence de préparation et mesure utilisée pour sonder les cohérences. Elle consiste en quatre étapes :

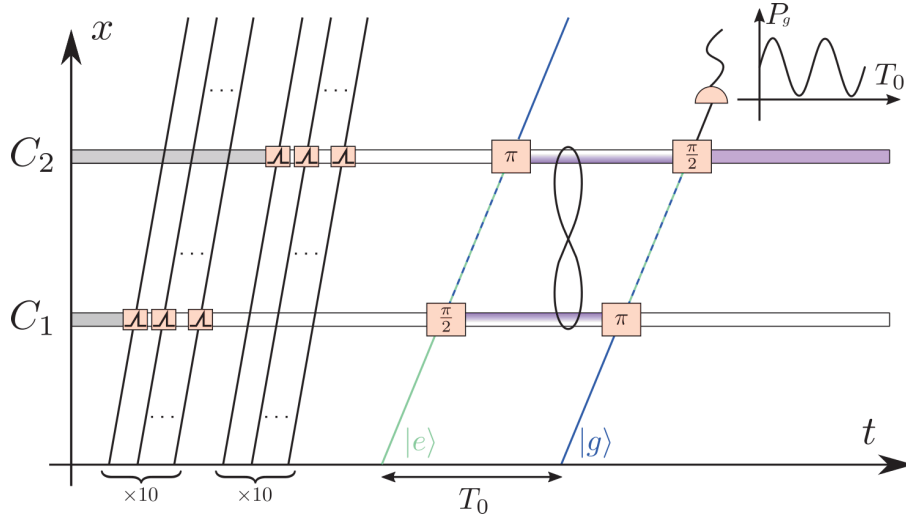


Fig. II.9 Schéma temps-position de la séquence expérimentale utilisée pour étudier les cohérences de l'état $|10\rangle + |01\rangle$. Dans un premier temps, une série d'atome serpillières est envoyée dans chaque cavité pour absorber le champ thermique. On envoie ensuite l'atome injecteur qui prépare l'état intriqué des deux cavités. Cet état évolue pendant une durée T_0 , avant l'arrivée de l'atome sonde. On répète l'expérience pour moyenner la probabilité de trouver l'atome sonde dans $|g\rangle$ et on fait varier T_0 pour observer l'évolution temporelle du signal d'interférence.

1. On initialise les cavités dans l'état vide en envoyant une série de dix atomes serpillières dans chacune d'entre elles ;
2. On prépare l'état $|10\rangle + |01\rangle$ en envoyant un atome injecteur A_I selon le protocole décrit section II.1.1 ;
3. On attend un temps T_0 ;
4. On envoie un atome sonde A_S et on mesure la probabilité de le trouver dans $|g\rangle$.

On obtient alors une courbe de la probabilité de trouver la sonde dans $|g\rangle$ en fonction du temps T_0 . La phase du signal oscillant obtenu est une mesure interférométrique de la différence de fréquence entre les deux cavités. La figure II.10 montre à l'aide de sphères de Bloch comment cette phase est transformée en signal de probabilité.

II.2.2 Optimalité de la mesure résonnante : états propres et information de Fisher

Dans cette section, on va introduire un critère permettant de montrer que l'état est intriqué par la violation d'une inégalité métrologique [70] et on va montrer que la mesure résonnante permet de violer maximalelement cette inégalité.

Le critère repose sur l'information de Fisher, qu'on présentera plus en détails dans le chapitre III, et qui permet de quantifier l'information qu'un signal contient sur un paramètre qu'on souhaite estimer. Dans notre cas, celui-ci est la phase ϕ imprimée dans l'état $|10\rangle + e^{i\phi}|01\rangle$. Le signal de mesure que l'on obtient lorsqu'on fait varier ϕ est tracé

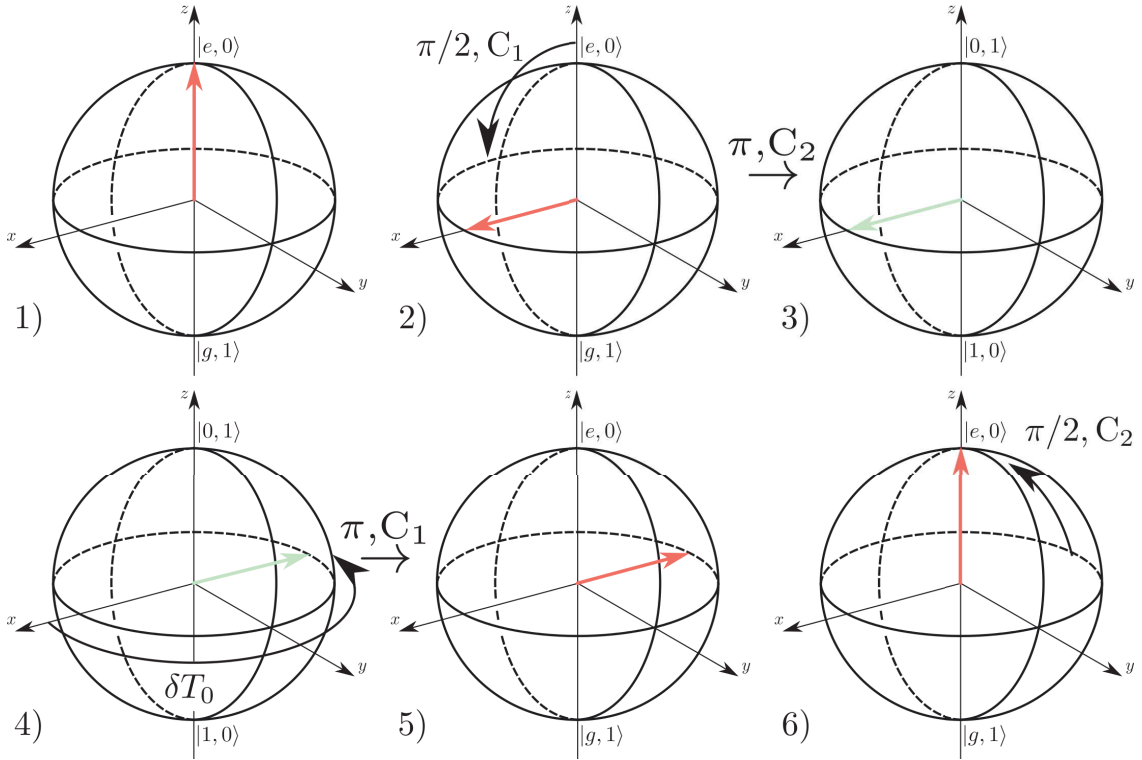


Fig. II.10 Séquence expérimentale vue en termes de sphère de Bloch. Les deux premières vignettes montrent, dans la sphère de Bloch A_I/C_1 , l'oscillation $\pi/2$ qui intrique l'injecteur à la cavité 1. Cette intrication est ensuite transmise à C_2 , ce qu'on voit sur la sphère de Bloch des deux cavités (en haut à droite). L'état des cavités évolue pendant une durée T_0 en tournant dans le plan équatorial de leur sphère de Bloch. L'intrication est ensuite transmise à l'atome sonde par l'impulsion π dans la première cavité, ce qu'on voit sur la sphère de Bloch A_S/C_2 de la cinquième vignette. Enfin, la dernière impulsion transcrit la phase azimutale de l'état après T_0 en probabilité. Pour simplifier, on a pris $\psi_m = \psi_p = 0$.

sur la figure II.11. Il s'agit d'une oscillation décrite par :

$$P_g = \frac{1 + C \cos(\phi)}{2}, \quad (\text{II.2.3})$$

où C est le contraste de l'oscillation mesurée.

L'information de Fisher contenue dans le signal est donnée par :

$$F(\phi) = P_e \left(\frac{d \log P_e}{d\phi} \right)^2 + P_g \left(\frac{d \log P_g}{d\phi} \right)^2, \quad (\text{II.2.4a})$$

$$= \frac{C^2 \sin^2 \phi}{1 - C^2 \cos^2 \phi}, \quad (\text{II.2.4b})$$

et vaut 1, lorsque $C = 1$, quelle que soit la phase.

On prépare initialement l'état $|\psi\rangle = |10\rangle + |01\rangle$ et on le laisse évoluer selon une évolution unitaire générée par $\hat{h}_f = (\hat{N}_1 - \hat{N}_2)/2$, de manière à obtenir l'état $|10\rangle + e^{i\phi}|01\rangle$,

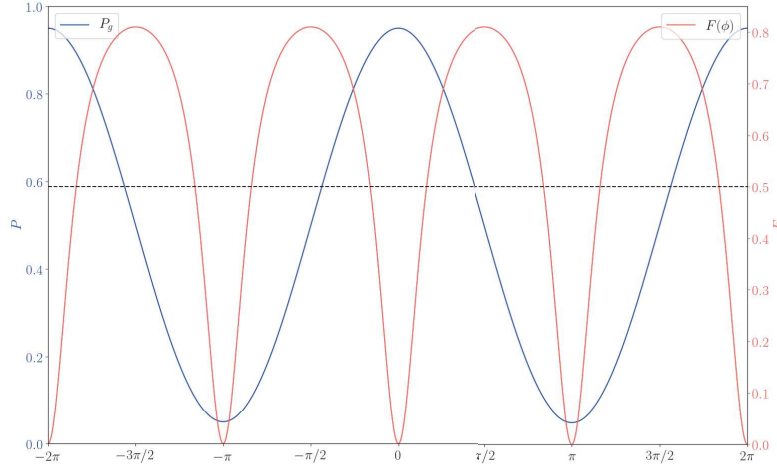


Fig. II.11 En bleu, signal de probabilité obtenu lorsqu'on fait varier la phase ϕ de l'état, en laissant le champ évoluer librement avant la mesure, pour un contraste de la mesure $C = 0.9$. En rouge, information de Fisher sur la phase, associée au signal bleu. Le trait horizontal en pointillé indique l'information de Fisher à dépasser pour mettre en évidence l'intrication avec le critère II.2.6.

qui contient l'information sur ϕ . L'information de Fisher maximale qu'on peut atteindre en effectuant une mesure sur le système est appelée information de Fisher quantique F_Q et peut s'exprimer, pour un état initial pur [71] [72] :

$$F_Q(\phi) = 4 \langle \psi | \Delta \hat{h}_f^2 | \psi \rangle, \quad (\text{II.2.5a})$$

$$= 1. \quad (\text{II.2.5b})$$

Pour un contraste $C = 1$, l'information de Fisher de notre mesure est donc égale à l'information de Fisher quantique, ce qui signifie qu'on extrait le maximum d'information sur la phase.

Le critère pour l'intrication proposé par [70] repose sur l'inégalité :

$$F_Q(\phi, \hat{\rho}_{sep}) \leq 4 \left[(\Delta \hat{h}_1)^2 + (\Delta \hat{h}_2)^2 \right], \quad (\text{II.2.6})$$

valable pour tout état séparable $\hat{\rho}_{sep}$ et pour un paramètre ϕ imprimé sur l'état par une transformation générée par $\hat{h}_f = \hat{h}_1 + \hat{h}_2$, où $\hat{h}_1$ et $\hat{h}_2$ agissent respectivement sur les sous-systèmes 1 et 2. Dans notre cas, $\hat{h}_1 = \hat{N}_1/2$, $\hat{h}_2 = -\hat{N}_2/2$ et on a :

$$4 \left[(\Delta \hat{h}_1)^2 + (\Delta \hat{h}_2)^2 \right] = \frac{1}{2}. \quad (\text{II.2.7})$$

On voit donc que l'inégalité est violée par l'information de Fisher quantique. Plus généralement, si on parvient à obtenir un signal dont l'information de Fisher est supérieure à $1/2$, l'état n'est pas séparable. On voit que notre mesure permet de le faire dès lors que

$$C > \frac{1}{\sqrt{2}} \simeq 0.71. \quad (\text{II.2.8})$$

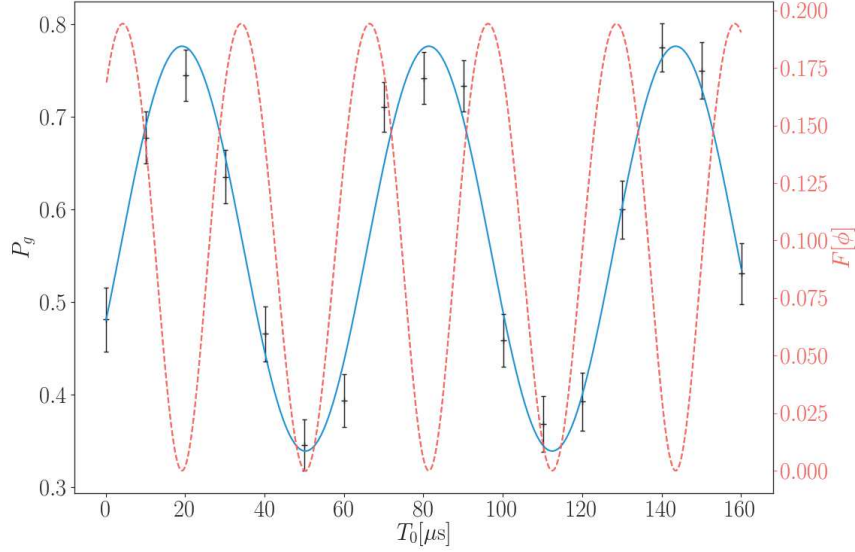


Fig. II.12 Signal d'interférence obtenu pour l'état $|10\rangle + |01\rangle$ à froid. La courbe bleue est un ajustement par la fonction $T_0 \mapsto y_0 + C/2 \cos(2\pi f T_0 + \varphi)$, avec $y_0 = 0.558 \pm 0.007$, $C = 0.44 \pm 0.02$, $f = 16.0 \pm 0.2$ kHz et $\varphi = 4.3 \pm 0.09$. La courbe rouge est l'information de Fisher sur la phase calculée à partir de l'ajustement. Celle-ci atteint 0.19 à mi-frange.

II.3 Résultats

Dans cette section, nous commençons par montrer les signaux interférométriques obtenus, ainsi que l'effet de la décohérence sur ces signaux. Nous présentons ensuite les résultats obtenus sur la mesure de la fonction de Wigner avec deux injections.

II.3.1 Un interféromètre de 6000 km

Signal d'interférence

Le signal obtenu à 0.8 K, pour un délai T_0 allant de 0.5 ms à 0.66 ms est montré sur la figure II.12. Les désaccords initiaux des cavités pour cette mesure valent $\Delta_1^0 = 2\pi \cdot 299.03 \pm 0.05$ kHz et $\Delta_2^0 = 2\pi \cdot 282.99 \pm 0.08$ kHz. Le contraste vaut $C = 0.44 \pm 0.02$. Nous verrons dans la section suivante les principales imperfections qui le limitent. La courbe en rouge est l'information de Fisher sur la phase associée au signal mesuré. On voit que celle-ci est insuffisante pour permettre de conclure à l'intrication par le critère II.2.6. Afin de mettre cette dernière en évidence, il faudra donc réaliser une tomographie complète de l'état, ce qui sera fait dans le chapitre IV.

Le signal permet d'obtenir la différence de fréquence $\delta = 2\pi \cdot 16.0 \pm 0.2$ kHz. La mesure de la fréquence des cavités, par la méthode de la section I.4.1, donne $\delta = 2\pi \cdot 16.0 \pm 0.1$ kHz. Ces deux valeurs sont donc en accord. Notre protocole d'interférométrie permet donc aussi de déterminer la différence de fréquence entre les cavités.

Effet de la décohérence

La décohérence du système à deux cavités peut être décrite à l'aide des outils de la section I.1.5. Afin de mettre en évidence ses effets, on a répété la mesure précédente pour des valeurs différentes de T_0 , pour deux valeurs de la température. Pour $T = 0,8$ K, on a effectué des scans de durée $[T_0, T_0 + 160 \mu\text{s}]$, pour des valeurs de T_0 entre 0.5 ms et 20 ms. Pour $T = 1.6$ K, les valeurs de T_0 ont été prises entre 0.5 ms et 13 ms. La figure II.13 montre l'évolution du contraste en fonction de T_0 . Celui-ci suit le comportement attendu, donné par l'intégration numérique de l'équation de Lindblad, avec pour seul paramètre ajustable le contraste initial. On remarque cependant que le point à 20 ms de la courbe à froid est plus bas qu'attendu. Ceci est dû au fait que les fréquences des cavités varient pendant la mesure, qui prend plusieurs heures lorsqu'on attend 20 ms entre la préparation et la mesure. Cette variation présente des évolutions lentes, dues à des variations de température dans le cryostat, mais aussi des sauts brusques. On peut la limiter en stabilisant la température des cavités, ainsi que la pression du bain d'hélium. Cependant, comme on s'est par la suite limité à des délais faibles entre la préparation et la détection, on n'a pas insisté sur la stabilité des cavités.

La courbe montre qu'on a encore une trace de l'interférence après 20 ms. On a alors réalisé un interféromètre à un photon de longueur effective 6000 km !

II.3.2 Contraste de l'interféromètre

Dans cette section, nous allons détailler les imperfections limitant le contraste C attendu pour le signal d'interférence et donner une estimation grossière de la part que chacune d'entre elle prend dans cette limitation.

a) Erreurs de détection

L'erreur de détection comprend deux composantes (voir section I.4.2.c) :

- Les erreurs de détection à proprement parler, qui font qu'un atome dans $|e\rangle$ (resp. $|g\rangle$) est détecté dans $|g\rangle$ (resp. $|e\rangle$). On note η_e^d (resp. η_g^d) l'erreur associée. On les a estimées à $\eta_e^d = 0.09$ et $\eta_g^d = 0.075$.
- L'efficacité finie $\epsilon = 0.5$ du détecteur.

L'effet des erreurs de détection est de réduire le contraste des franges et de décaler leur valeur moyenne, puisque les erreurs ne sont pas symétriques. L'efficacité finie du détecteur n'est pas en soit un facteur de réduction du contraste. Cependant, elle nous empêche de filtrer correctement tous les événements où deux atomes sont présents dans un échantillon, ce qui réduit le contraste.

b) Erreurs de préparation

L'atome injecteur est préparé dans $|e\rangle$ à partir d'un atome purifié dans $|g\rangle$ auquel on applique une impulsion π dans R_1 . Cette impulsion n'a pas un contraste de 1, à cause de l'étalement du paquet d'onde, qui donne lieu à une dispersion sur les pulsations de Rabi classiques dans la zone de Ramsey. Notons η_e^p la probabilité que l'atome soit dans $|g\rangle$ alors qu'on souhaite le préparer dans $|e\rangle$. La probabilité de détecter l'atome dans $|e\rangle$

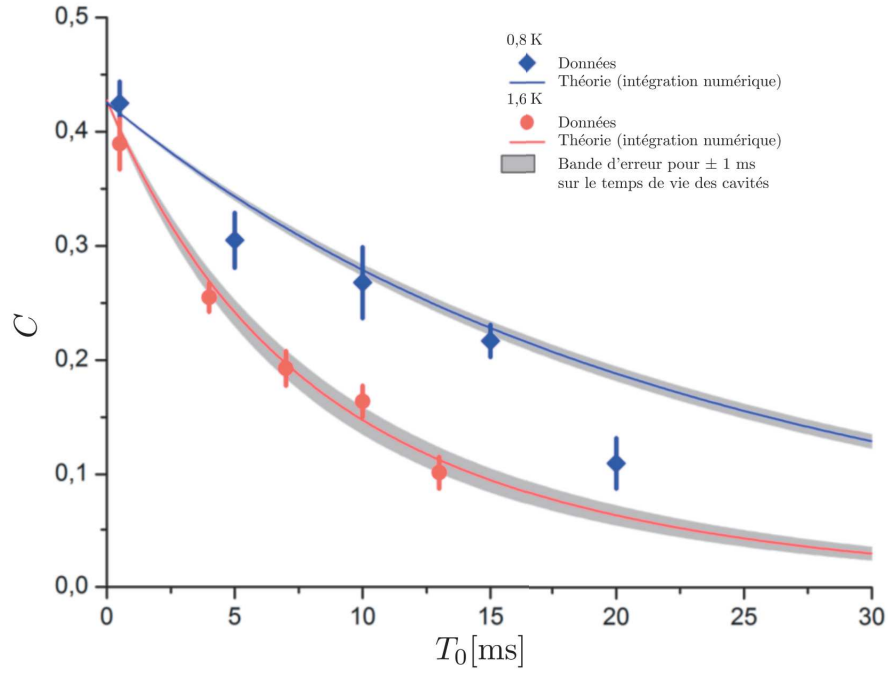


Fig. II.13 Évolution du contraste en fonction du délai entre la préparation et la mesure. Les courbes théoriques sont obtenues par l'intégration numérique de l'équation de Lindblad pour les temps de vie de nos cavités. Le seul paramètre libre de ces courbes est le contraste initial de l'oscillation pris égal à $C = 0.425$. Pour les points à 0,8 K, 20 ms et 1,6 K, 15 ms, le moyennage dure plusieurs heures et on est plus sensible à la fréquence à cause de la durée de la séquence. Pour éliminer la perte de contraste due à la dérive de fréquence, on a d'abord séparé les données en plusieurs jeux, pour lesquels on a estimé la dérive de phase, qu'on a compensée avant d'effectuer l'ajustement.

lorsqu'on le prépare dans $|e\rangle$ est :

$$P_e = (1 - \eta_e^p)(1 - \eta_e^d) + \eta_e^p \eta_g^d = 0.85. \quad (\text{II.3.1})$$

Le second terme correspond à la probabilité que l'atome soit mal préparé et mal détecté, qui est inférieure à 1%. On peut donc le négliger. On a alors :

$$\eta_e^p = 1 - \frac{P_e}{1 - \eta_e^d} = 0.07. \quad (\text{II.3.2})$$

On n'a pas tenu compte de la relaxation des atomes, qui concerne environ 1,5% des atomes avant que ceux-ci n'interagissent avec les cavités. Cet effet est donc pris en compte dans l'efficacité de préparation.

c) Contraste de Rabi

Les oscillations de Rabi dans les cavités ont aussi un contraste limité, à cause de la dispersion des atomes autour de l'axe du jet et dans leur direction de propagation. Afin de

les modéliser, on va supposer que les oscillations du système atome-cavité ont un contraste qui décroît exponentiellement au cours du temps :

$$C(t) = C_R e^{-t/\tau}. \quad (\text{II.3.3})$$

Avec les régressions de la section I.5.3, on obtient les contrastes suivants pour les oscillations de Rabi $\pi/2$ et π , dans le cas des oscillations fort désaccordées et dans le cas des oscillations à faible désaccord initial :

	$C_1, \delta \gg \Omega_{0,1}$	$C_2, \delta \gg \Omega_{0,2}$	$C_1, \delta \gtrsim \Omega_{0,1}$	$C_2, \delta \gtrsim \Omega_{0,2}$
$t_{\pi/2}$	0.73	0.75	0.70	0.73
t_π	0.62	0.65	0.70	0.67

Les contrastes estimés avec les oscillations de Rabi contiennent l'effet de la dispersion des couplages, mais aussi celui des erreurs de préparation et de détection. Pour obtenir le contraste vraiment dû à la dispersion, il faut les diviser par $(1 - \eta_e^p)(1 - \eta_e^d - \eta_g^d) = 0.78$. On trouve alors les contrastes :

	$C_1, \delta \gg \Omega_{0,1}$	$C_2, \delta \gg \Omega_{0,2}$	$C_1, \delta \gtrsim \Omega_{0,1}$	$C_2, \delta \gtrsim \Omega_{0,2}$
$t_{\pi/2}$	0.94	0.96	0.90	0.94
t_π	0.79	0.83	0.90	0.86

d) Échantillons à deux atomes

Les échantillons atomiques ont une répartition poissonnienne du nombre d'atomes. On ne garde que les échantillons dans lesquels on a détecté un seul atome. Pour une efficacité de détection $\epsilon = 0.5$, si un atome contient un échantillon, on a une probabilité $1/2$ de détecter l'atome et de garder l'échantillon. S'il en contient deux, on a aussi une probabilité $1/2$ de ne détecter qu'un seul atome et de garder le paquet. Si il en contient trois, la probabilité de n'en détecter qu'un vaut $3/8$. La probabilité qu'il en contienne plus sans être filtré est négligeable. Une proportion $P_1 = 90.5\%$ des atomes sont donc issus d'échantillons monoatomiques, 9% d'échantillons composés de deux atomes et 0.5% d'échantillons à trois atomes. On va supposer que lorsque les échantillons contiennent plusieurs atomes, il n'y a aucune interférence.

e) Bilan

Dans le cadre des hypothèses précédentes, le contraste de l'interféromètre vaut :

$$C = C_1^{(\pi)} C_1^{(\pi/2)} C_2^{(\pi)} C_2^{(\pi/2)} (1 - \eta_e^p)(1 - \eta_e^d - \eta_g^d) P_1^2, \quad (\text{II.3.4})$$

où $C_j^{(k)}$ est le contraste de l'impulsion k dans la cavité j . C vaut 0.39 dans le cas du grand désaccord initial et 0.42 dans le cas du faible désaccord initial. La principale limitation est liée au contraste des oscillations de Rabi, qui intervient 4 fois. Il faut cependant noter qu'on n'a pas pris en compte le fait qu'on effectue une post-sélection en ne gardant que les événements où l'atome est détecté dans $|g\rangle$. Au premier ordre dans toutes les imperfections, cette post-sélection enlève les événements où l'impulsion π dans la seconde cavité a échoué. Le contraste est alors

$$C = C_1^{(\pi)} C_1^{(\pi/2)} C_2^{(\pi/2)} (1 - \eta_e^p)(1 - \eta_e^d - \eta_g^d) P_1^2. \quad (\text{II.3.5})$$

Le contraste dans le cas du grand désaccord initial vaut alors 0.46. Cet ordre de grandeur est en accord avec le contraste de la figure II.12. Pour le faible désaccord initial, il vaut 0.49.

II.3.3 Mesure de coupes de la fonction de Wigner de l'état $|10\rangle + |01\rangle$

La mesure de la fonction de Wigner s'est avérée difficile à réaliser. En effet, comme on l'a déjà dit, le contraste est assez réduit entre le signal attendu pour l'état $|10\rangle + |01\rangle$ et celui pour $|10\rangle - |01\rangle$. Nous avons effectué des mesures de deux types de coupes de la fonction de Wigner : la première consiste à effectuer des injections de même amplitude dans les deux cavités et de balayer cette amplitude. La deuxième consiste à varier la phase de mesure en laissant évoluer l'état. Cette seconde mesure est peu concluante à cause de la limitation du contraste. On va présenter brièvement les résultats obtenus et les limitations qui rendent cette mesure difficile.

Balayage en amplitude

La première mesure consiste à effectuer un balayage de la puissance des sources micro-onde, en utilisant les méthodes présentées dans la section I.5.4 pour calibrer les amplitudes de manière à avoir $|\alpha_1| = |\alpha_2| \equiv r$. On s'attend alors à obtenir un signal proche de ceux de la figure II.6. En fonction de la phase relative des injections et de la phase de l'état, on peut avoir n'importe quelle courbe intermédiaire entre celle de $|10\rangle + |01\rangle$ et $|10\rangle - |01\rangle$. L'avantage de cette méthode est que, quelle que soit la phase de l'état *noon*, la fonction de Wigner vaut 0 pour une injection nulle et tend vers 1/2 pour des grandes injections. On s'attend donc à avoir un signal contrasté.

Pour réaliser la mesure, on envoie d'abord un atome injecteur, dans un échantillon contenant en moyenne $\mu = 0,32$ atome. On attend 1350 μs après l'instant d'excitation de cet atome, pour qu'il soit détecté, avant d'effectuer simultanément deux injections de 746 μs . La durée de ces injections est choisie de sorte qu'elle soit égale à $1/\delta$, afin que l'injection d'une cavité ait une densité spectrale de puissance nulle à la fréquence de l'autre cavité. On envoie ensuite, 2200 μs après l'atome injecteur, 10 échantillons d'atomes QND, séparés de 100 μs , contenant en moyenne 0,18 atomes. La mesure est répétée 140×100 fois. La figure II.14 montre les signaux obtenus lorsqu'on détecte un atome injecteur dans l'état $|g\rangle$, et lorsqu'aucun atome injecteur n'est détecté. Le second de ces deux signaux est la fonction de Wigner de l'état vide des deux cavités, avec une légère pollution due aux atomes injecteurs non détectés. On en tire l'ajustement en bleu. Cet ajustement nous permet d'effectuer une calibration de l'amplitude d'injection dans les deux cavités. En utilisant cette calibration, on peut effectuer un ajustement de notre signal par la fonction de Wigner d'un état *noon* avec des paramètres d'injection réels identiques¹ :

$$W^{|10\rangle + e^{i\phi}|01\rangle}(r, r) = \left(8r^2 \cos^2 \frac{\phi}{2} - 1\right) e^{-4r^2}, \quad (\text{II.3.6})$$

en prenant pour seuls paramètres libres la probabilité d'avoir 0 photon et la phase de l'état *noon*. Celui-ci fournit la valeur $p_1 = 0,77 \pm 0,02$ pour la probabilité d'avoir 1 photon, qui

1. On suppose ici que les injections sont en phase et on ajuste par rapport à la phase de l'état. Ceci revient à effectuer une rotation des paramètres complexes de la fonction de Wigner.

est à prendre avec précaution, puisque le modèle n'inclue pas la possibilité que les poids de $|10\rangle$ et $|01\rangle$ soient différents. Il fournit par ailleurs la valeur de la phase $\phi = 2.02 \pm 0.2$ rad.

La courbe obtenue est proche des données expérimentales, cependant, elle présente plusieurs limitations pour l'étude de l'état *noon*. En effet, elle est proche de la courbe, verte, pouvant correspondre aux états $|10\rangle + e^{\pm i\phi} |01\rangle$ et $|1\rangle\langle 1| + |0\rangle\langle 0|$. On ne peut donc pas conclure à la présence de cohérences entre $|10\rangle$ et $|01\rangle$. On voit qu'il est nécessaire de connaître la différence entre la phase de l'état et les phases de la mesure pour pouvoir prouver la présence de cohérences. Afin de l'obtenir, on peut faire varier la phase de l'état.

Balayage de la phase

Les phases complexes des paramètres de la fonction de Wigner correspondent aux phases des sources utilisées pour réaliser les injections. La référence de phase à choisir pour ces phase dépend en général du protocole de préparation de l'état. Par exemple, lors de la préparation d'un état *chat de Schrödinger* $|\alpha\rangle + |-\alpha\rangle$ à partir d'un état cohérent, la référence de phase est donnée par la source qui a servi à préparer l'état cohérent [29]. Les états que nous avons mesurés dans ce travail sont préparés sans injection dans les cavités. Par ailleurs, ils sont préparés par un atome dans l'état $|e\rangle$, qui n'a pas de phase quantique. Il n'y a donc pas de référence de phase pour notre état. La phase globale des impulsions utilisées pour les déplacements n'importe pas, ce qui se traduit par les propriétés de symétrie des fonctions de Wigner des états *noon*.

Cependant, la phase relative entre le déplacement dans C_1 et celui dans C_2 importe. On a vu dans la section I.5.4 qu'on la contrôle à l'aide d'un battement entre les signaux des deux sources. On veut balayer la phase relative $\phi_S = \arg \alpha_2 - \arg \alpha_1$ des deux injections. D'après l'équation I.1.33, l'évolution temporelle de l'état revient à changer la phase des injections micro-onde. Plutôt que de changer ϕ_S , il est plus simple de changer la phase relative ϕ_c de l'état entre les deux cavités. Dans ce but, on balaye l'instant t_0 de préparation de l'état en gardant fixe l'instant d'injection. Cette méthode est schématisée sur la figure II.15.

Le déroulé de la séquence est donc le suivant. À l'instant $t = 0$, le signal de battement atteint 0 par valeurs ascendantes. Le signal de S_2 est alors en quadrature et en retard de phase avec celui de S_1 . La phase relative des deux injections évolue donc comme :

$$\phi_S = -\frac{\pi}{2} + \delta t. \quad (\text{II.3.7})$$

À l'instant t_0 , l'atome injecteur prépare l'état, avec une phase ψ_p . La phase de l'état évolue ensuite à δ selon :

$$\phi_c = \psi_p + \delta(t - t_0). \quad (\text{II.3.8})$$

La différence entre ces deux phases, qui est la seule phase à jouer un rôle dans la fonction de Wigner, à cause de la symétrie des états *noon*, vaut donc au moment de la mesure :

$$\Delta\phi = \phi_c - \phi_S = \psi_p + \frac{\pi}{2} - \delta t_0. \quad (\text{II.3.9})$$

Cette phase ne dépend pas de l'instant d'injection t_{mes} . Ceci est dû au fait que les sources sont accordées sur les cavités. ϕ_c et ϕ_S croissent donc toutes les deux linéairement avec un taux δ . En revanche, si l'instant de préparation t_0 varie, cette phase est modifiée. On

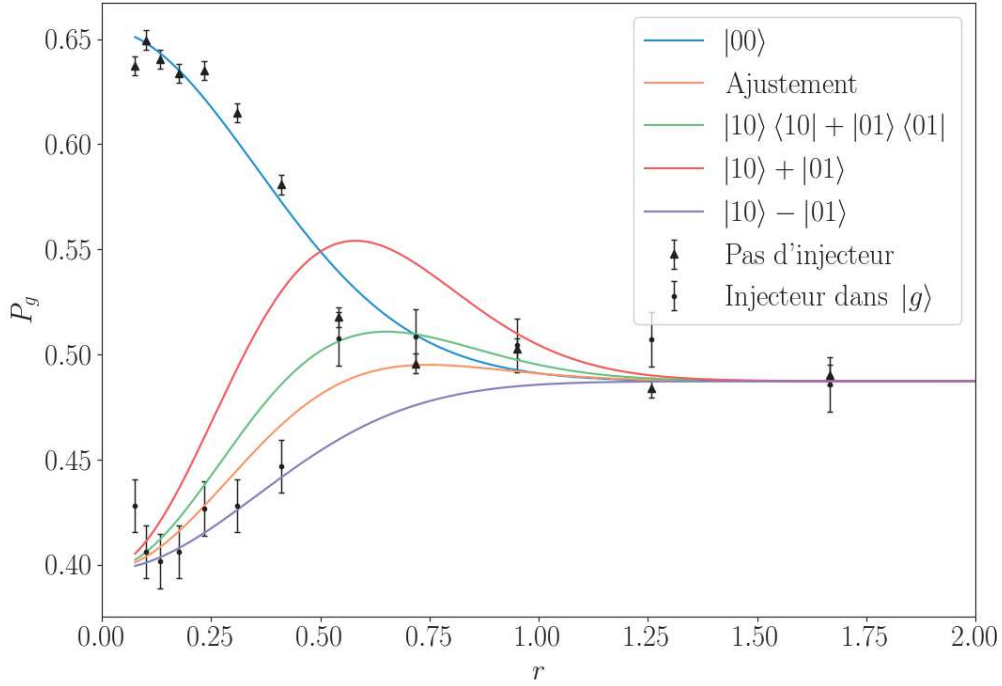


Fig. II.14 Coupes de fonctions de Wigner obtenues en balayant l'amplitude $r = \alpha_1 = \alpha_2$. Les triangles correspondent aux événements où on n'a détecté aucun atome injecteur. Le champ est alors dans l'état vide. On a utilisé ces points pour calibrer l'amplitude d'injection : $\alpha(P) = \sqrt{1/2000 \cdot 10^{a(P-P_0)}}$ avec $a = 0.609 \pm 0.012$ et $P_0 = -2.7 \pm 0.1$ dBm, et où P est la puissance de la source S_1 . La courbe bleue est un ajustement de ces points par la fonction $P = \pi_0 + (C_F/2)W^{(00)}(r, r)$, avec les paramètres $\pi_0 = 0.487 \pm 0.002$ et $C_F = 0.335 \pm 0.006$. Les cercles représentent les données obtenues lorsqu'on détecte l'injecteur dans $|g\rangle$, c'est-à-dire la fonction de Wigner de l'état 1001 . La courbe jaune est un ajustement sur ces données avec la fonction $P = \pi_0 + (C_F/2) \left[(1 - p_1)W^{(0,0)}(r, r) + p_1 W^{(10)+e^{i\phi}|01\rangle}(r, r) \right]$, où $p_1 = 0.77 \pm 0.02$, $\phi = 2.02 \pm 0.2$ rad et les paramètres π_0 et C_F sont ceux obtenus par l'ajustement sur l'état vide.

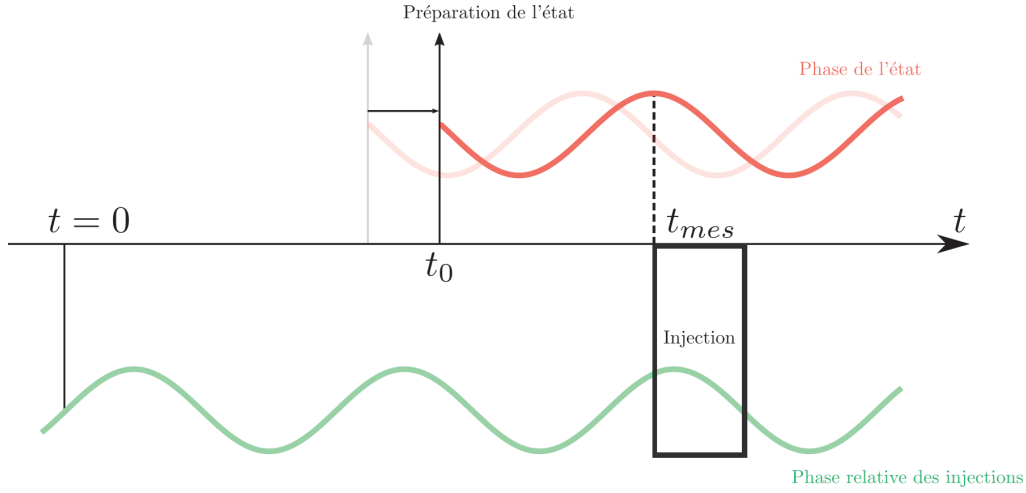


Fig. II.15 Schéma de principe du choix de la phase de mesure. À l'instant $t = 0$, le signal de battement des sources, en vert, permet de lancer la séquence pour une phase relative des injections donnée. L'atome injecteur est envoyé à un instant t_0 variable. La phase initiale de l'état préparé est toujours la même et sa phase ϕ_c évolue à la même pulsation que la phase relative des deux injections. Si on prépare l'atome injecteur à un instant différent, la phase $\phi_c(t_{mes})$ au moment de l'injection est modifiée, ce qui est équivalent à modifier la phase de mesure $\phi_S(t_{mes})$.

s'attend donc à voir une oscillation à δ . Le contraste attendu pour celle-ci est cependant très faible. Le contraste maximal obtenu pour une mesure idéale sur un état $|10\rangle + |01\rangle$ vaut $e^{-1} = 0.37$. À cause du mélange de l'état avec l'état vide, ce contraste est réduit à 0.28. En prenant en compte le contraste de Ramsey $C_F = 0.34$, on s'attend à un contraste $C_W = 0.092$.

La figure II.16 montre les résultats obtenus. Les points de mesure sont en noirs. La courbe bleue est un ajustement dans lequel la phase est le seul paramètre libre. La fréquence de l'oscillation est mesurée par des mesures de fréquence des cavités. Le contraste et la valeur moyenne de l'oscillation sont ceux obtenus par ajustement sur la courbe de balayage de l'amplitude. Le signal semble présenter une oscillation avec les paramètres attendus, à l'exception de la phase. Le point rouge montre la valeur déterminée par l'ajustement de la courbe de balayage en amplitude, qui correspond à la valeur $t_i = 28 \mu s$. Pour cette valeur de t_i , la phase estimée avec le balayage en amplitude est $\phi^{(a)} = 2.02 \pm 0.25$ rad, tandis que celle estimée pour le balayage en phase vaut $\phi^{(p)} = 0.03 \pm 0.13$ rad. Ces deux valeurs sont en désaccord. Dans les séquences utilisées, il s'écoule $7200 \mu s^2$ entre l'instant où on fixe la phase et celui où l'injection commence. La différence de phase précédente pourrait donc s'expliquer par une variation 44 Hz de la fréquence relative des deux injections. Un tel écart est compatible avec les dérives observées à l'échelle de séquences de mesure longues de quelques heures.

En conclusion, on observe l'oscillation attendue mais le contraste est très faible et la

2. Ce délai est en grande partie dû au temps nécessaire pour envoyer les atomes serpillières avant d'envoyer l'atome injecteur.

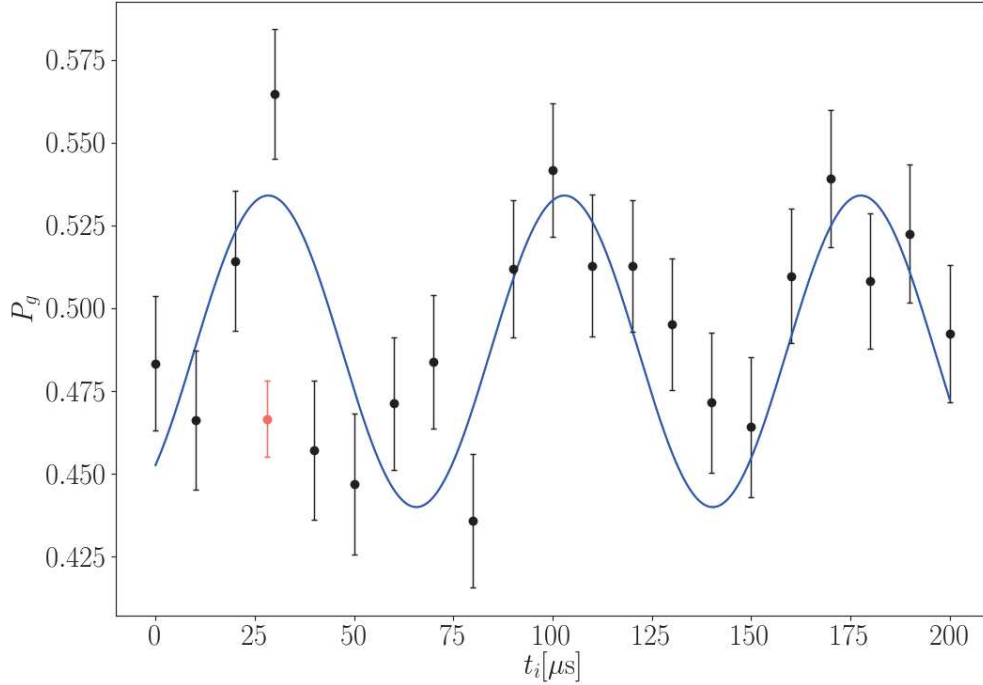


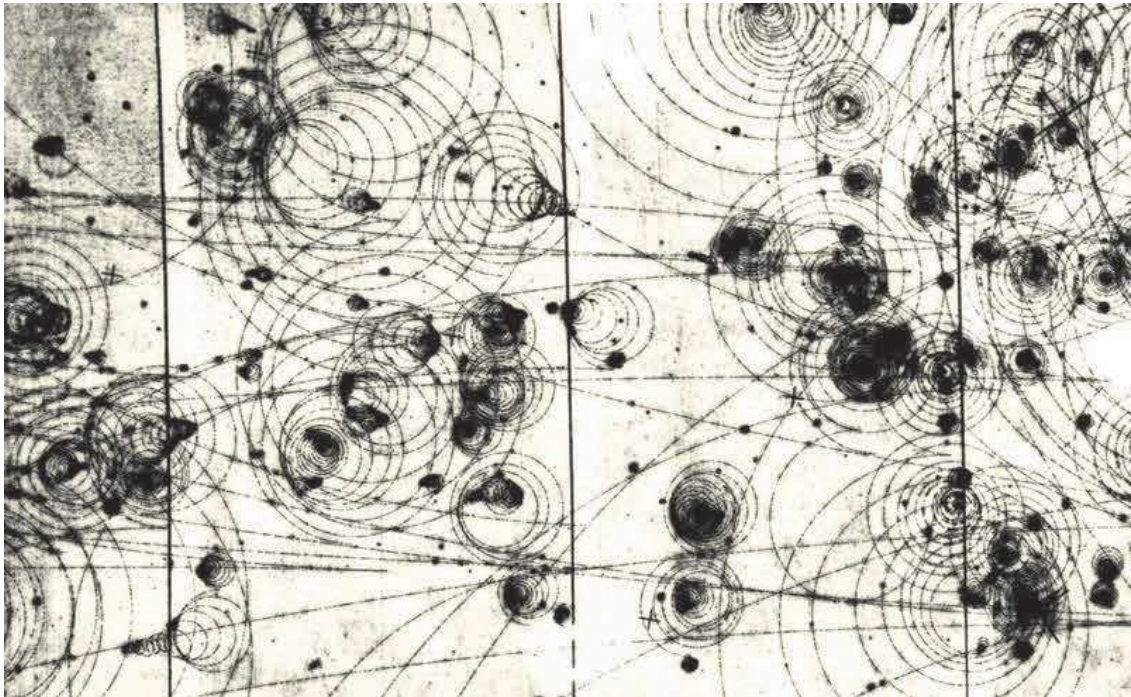
Fig. II.16 Oscillation obtenue en balayant la phase ϕ de l'état $|10\rangle + e^{i\phi}|01\rangle$ lors de la mesure de la fonction de Wigner. On s'est placé à l'amplitude d'injection $r = 1/2$ pour avoir un contraste maximum. La courbe bleue est un ajustement par la fonction $P = \pi_0 + (C_F/2) \cos(\delta t_i + \phi_0)$, où $\phi_0 = 2.38 \pm 0.13$ rad est le seul paramètre libre. $\delta = 2\pi \cdot 13.4$ kHz est obtenu en mesurant les fréquences des cavités et les autres paramètres sont ceux obtenus par l'ajustement sur le balayage en amplitude. Le point rouge est la valeur estimée pour $t_i = 28 \mu\text{s}$ avec l'ajustement du balayage en amplitude.

valeur obtenue pour la phase ne semble pas reproductible d'une mesure à l'autre. Cette méthode est de plus particulièrement lourde à mettre en place puisqu'elle nécessite de calibrer les amplitudes d'injection et de verrouiller les phases.

Conclusion du chapitre II : des méthodes de mesures complémentaires

Dans ce chapitre, nous avons montré que nous sommes capables de préparer un état délocalisé d'un photon présent dans deux cavités à la fois. Nous avons mis en évidence de deux manières différentes les cohérences quantiques de cet état. La première d'entre elles consiste à réaliser une mesure interférométrique en envoyant un atome sonde, dont la probabilité d'absorber dépend de la phase des cohérences entre les deux membres de la superposition quantique. Cette méthode présente un contraste maximal dans le cas idéal et permet donc de discriminer très efficacement les états $|10\rangle + e^{i\phi}|01\rangle$ de phases différentes. Elle nous a permis de mettre en évidence les cohérence quantiques jusqu'à

20 ms après la préparation de l'état. Elle ne permet cependant pas de reconstruire l'état entièrement. Une autre approche consiste à effectuer des mesures de la fonction de Wigner à deux modes de l'état. Celle-ci permet en principe de reconstruire totalement l'état, si on la mesure en suffisamment de points différents. Elle présente aussi des oscillations en fonction de la phase quantique de l'état. Cependant, ces oscillations ont un contraste faible, y compris dans le cas idéal. La fonction de Wigner n'est donc pas un outil adapté pour mesurer les cohérences quantiques des états *noon*. Nous verrons dans le chapitre IV comment on peut combiner ces deux mesures pour effectuer une tomographie de l'état délocalisé.



Chapitre III

Principes de la tomographie d'états quantiques

Dans le chapitre II, nous avons montré des expériences de préparation et de mesure d'états intriqués à deux cavités. Les mesures effectuées mettent en évidence des cohérences quantiques, responsables de l'oscillation observée. Cependant, le contraste obtenu ne permet pas de conclure à l'intrication avec le critère II.2.6 que nous avons envisagé. Dans ce chapitre, nous allons nous intéresser aux techniques de reconstruction. Celles-ci permettent de construire des estimateurs qui, à partir d'un ensemble de mesures, toutes réalisées avec le même état initial, fournissent un état qui s'approche de ce dernier. Nous allons présenter une méthode qui permet de prendre en compte les protocoles de mesure les plus généraux qu'on puisse appliquer au système, faisant intervenir plusieurs mesures successives, ainsi que de la décohérence entre ces mesures. Cette méthode de *tomographie par trajectoires* enrichit la gamme de protocoles de reconstruction existants en permettant de prendre en compte toute l'information contenue dans les corrélations entre mesures effectuées successivement lors d'une même trajectoire quantique suivie par l'état.

Dans la section III.1, nous donnons quelques rappels sur la théorie de l'estimation d'un paramètre réel. Nous introduisons en particulier l'information de Fisher qui décrit l'information apportée par la mesure sur le paramètre à estimer. Celle-ci intervient dans la borne de Cramer-Rao qui donne la limite de précision accessible sur la détermination du paramètre, pour un nombre de répétitions de la mesure donné. Dans la section III.2, le cadre de la tomographie d'états quantiques est introduit. On présente ensuite les principales méthodes de reconstruction qui permettent d'estimer un état quantique à l'aide d'un ensemble de mesures instantanées. La section III.3 détaille la nouvelle méthode de reconstruction utilisée dans ce travail de thèse, qui permet de prendre en compte une série de mesures successives pour chaque réalisation de l'état à déterminer. Elle montre en particulier comment appliquer cette méthode à la reconstruction de l'état d'un système à deux niveaux. Enfin, la section III.4 pose la question de l'estimation des incertitudes associées à la tomographie. Nous y proposons une manière de calculer des barres d'erreur directement sur les coefficients de la matrice densité. Nous comparons cette approche à une des méthodes proposées pour établir des régions de confiance de la reconstruction. Notre protocole de tomographie sera ensuite appliqué à la reconstruction de l'état de nos cavités dans le chapitre IV.

III.1 Quelques rappels de théorie de l'estimation

III.1.1 Cadre de l'estimation d'un paramètre réel

Supposons qu'on veuille estimer la valeur d'un paramètre θ , à l'aide d'un ensemble de résultats de mesure qui dépendent de lui. À cause des incertitudes de mesure, mais aussi des fluctuations statistiques, on n'aura pas, à θ fixé, toujours les mêmes résultats. On modélise le résultat de chaque mesure par une variable aléatoire X qui peut prendre les valeurs $x \in \mathcal{X}$. Les résultats de mesure peuvent être discrets, comme par exemple le nombre de photons impactant un détecteur pendant un intervalle de temps, ou continus, comme la valeur de la température d'un objet macroscopique. Dans le premier cas, on décrit la loi de probabilité par une fonction de masse

$$p_\theta : \mathcal{X} \rightarrow [0, 1], \quad (\text{III.1.1a})$$

$$x \mapsto p(x|\theta), \quad (\text{III.1.1b})$$

qui associe, à chaque résultat possible, la probabilité de l'obtenir pour la valeur θ du paramètre.

Dans le deuxième cas, la loi de probabilité est décrite par une densité de probabilité f_X^θ telle que la probabilité conditionnelle d'obtenir un résultat situé dans l'intervalle $[x_1, x_2]$ s'écrit

$$p(x \in [x_1, x_2]) = \int_{x_1}^{x_2} f_X^\theta(x) dx. \quad (\text{III.1.2})$$

Dans toute la suite, pour alléger les notations, on notera indistinctement $p(x|\theta)$ la probabilité conditionnelle d'obtenir le résultat x lorsque le paramètre vaut θ et on utilisera le symbole $\int dx$ pour les sommations, de sorte que l'espérance d'une fonction $f(X)$ s'écrive :

$$\mathbb{E}[f(X)] = \int dx p(x|\theta) f(x). \quad (\text{III.1.3})$$

En répétant l'expérience un grand nombre de fois, on acquiert de l'information sur la fonction de masse ou la densité de probabilité et on peut utiliser cette information pour estimer la valeur de θ . Un ensemble de n mesures est décrit par un vecteur $\vec{X} = (X_1, \dots, X_n)$. On construit à cet effet un estimateur $\Theta_{\text{est}}(\vec{X})$ ¹. L'estimateur est une variable aléatoire qui dépend de la variable aléatoire $\vec{X}$. Sur un échantillon x tiré, l'estimateur fournit une estimation $\theta_{\text{est}} = \Theta_{\text{est}}(x)$ de θ .

Propriétés importantes d'un estimateur

Définition 1 (Biais) On appelle biais de l'estimateur la quantité

$$B(\Theta_{\text{est}}) = \mathbb{E}[\Theta_{\text{est}}] - \theta, \quad (\text{III.1.4})$$

où $\mathbb{E}[X]$ est l'espérance de la loi X . Lorsque $B(\Theta_{\text{est}}) = 0$, l'estimateur est dit sans biais.

1. Pour éviter la confusion avec les opérateurs de l'espace de Hilbert, on ne suivra pas la convention habituelle de noter l'estimateur avec un chapeau.

Définition 2 (Variance) *On appelle variance de l'estimateur la quantité*

$$\Delta^2 \Theta_{est} = \mathbb{E} \left[(\Theta_{est} - \mathbb{E} [\Theta_{est}])^2 \right]. \quad (\text{III.1.5})$$

Un estimateur est dit efficace lorsque sa variance est faible. On verra plus loin la limite imposée à l'efficacité d'un estimateur.

III.1.2 Efficacité de l'estimateur

a) Information de Fisher et borne de Cramer-Rao

La question de savoir quelle est la limite imposée à la précision d'un estimateur est centrale en statistique. Elle a été abordée par Fisher en termes d'efficacité des estimateurs [73] [74], puis une borne a été donnée indépendamment par Cramer [75] et Rao [76]. Dans la suite, on introduit l'information de Fisher et on présente la borne de Cramer-Rao, qui sera démontrée dans l'annexe A, puis on montre qu'elle est atteinte dans la limite d'un grand nombre de mesures par une estimation par maximum de vraisemblance.

Supposons qu'on réalise une mesure décrite par la variable aléatoire X , dont la loi de probabilité dépend d'un paramètre θ qu'on souhaite estimer. La variance sur un estimateur Θ_{est} de θ est minorée par une borne appelée borne de Cramer-Rao :

$$\Delta^2 \Theta_{est} \geq \frac{1}{F(\theta)}. \quad (\text{III.1.6})$$

Elle fait intervenir l'information de Fisher $F(\theta)$ qui dispose de plusieurs expressions équivalentes :

$$F(\theta) \equiv \mathbb{E} \left[\left(\frac{\partial}{\partial \theta} \log p(x|\theta) \right)^2 \right], \quad (\text{III.1.7a})$$

$$= \mathbb{E} [V(X, \theta)^2], \quad (\text{III.1.7b})$$

$$= -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \log p(x|\theta) \right], \quad (\text{III.1.7c})$$

où la *fonction de score* $V(X, \theta)$ est définie par

$$V(X, \theta) \equiv \partial \log p(x|\theta) / \partial \theta. \quad (\text{III.1.8})$$

L'identité III.1.7c montre que l'information de Fisher est l'espérance de la courbure de la probabilité conditionnelle $p(x|\theta)$. Pour estimer efficacement un paramètre, il faut donc que cette dernière varie beaucoup avec le paramètre à estimer.

Lorsqu'on effectue n mesures identiques et indépendantes, l'information de Fisher vaut :

$$F^{(n)}(\theta) = nF^{(1)}(\theta), \quad (\text{III.1.9})$$

où $F^{(1)}$ est l'information de Fisher pour une mesure. La variance décroît donc proportionnellement à l'inverse du nombre de mesures effectuées, avec pour constante de proportionnalité l'inverse de l'information de Fisher.

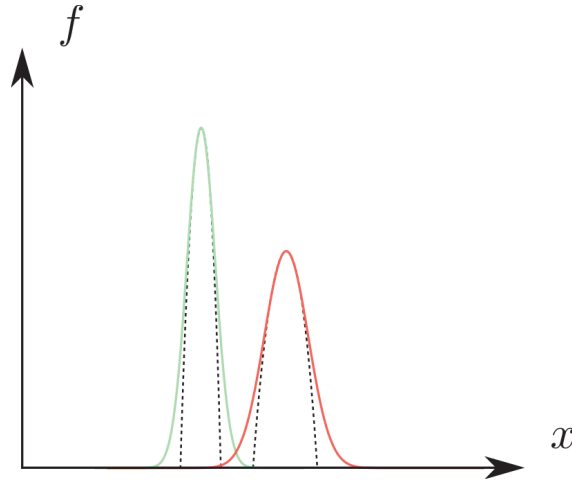


Fig. III.1 Densités de probabilité de deux lois normales f_1 (en vert) et f_2 (en rouge) de largeurs respectives $\mu_1 < \mu_2$. La courbure de la deuxième loi est plus faible que celle de la première. La variance de l'estimateur sera donc plus grande.

Exemple : Centre d'une gaussienne

Considérons pour fixer les idées une mesure dont la loi de probabilité est donnée par une gaussienne

$$f_\mu(x|\theta) = \frac{1}{\sqrt{2\pi\mu^2}} e^{-\frac{(x-\theta)^2}{2\mu^2}}, \quad (\text{III.1.10})$$

dont on veut déterminer le centre θ . La fonction de score est la dérivée par rapport à θ de la fonction de densité de probabilité :

$$\begin{aligned} V(X, \theta) &= \frac{\partial}{\partial \theta} \log f_\mu(X|\theta), \\ &= \frac{X - \theta}{\mu^2}. \end{aligned} \quad (\text{III.1.11})$$

L'information de Fisher est, d'après III.1.7c :

$$\begin{aligned} F(\theta) &= -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \log f_\mu(X|\theta) \right], \\ &= \frac{1}{\mu^2}. \end{aligned} \quad (\text{III.1.12})$$

On en déduit qu'un estimateur du centre de la gaussienne aura une variance au moins égale à μ^2 . L'écart-type sur l'estimation de θ après n mesures sera donc $\mu/\sqrt{n}$. On retrouve ici que l'information de Fisher est liée à la courbure de la densité de probabilité. On voit sur la figure III.1 que plus la courbure de f est grande, plus l'estimation pourra être précise.

b) Estimation par maximum de vraisemblance

On va voir maintenant comment on peut atteindre la limite donnée par la borne de Cramer-Rao. On considère un ensemble de n répétitions d'une mesure, décrites par des variables aléatoires identiques et indépendantes $\vec{X} = (X_1, \dots, X_n)$. On appelle *vraisemblance*

d'une valeur ϑ du paramètre la probabilité conditionnelle $p(\vec{x}|\vartheta)$ d'obtenir l'ensemble de résultats $\vec{x}$. L'estimation par maximum de vraisemblance consiste à choisir l'estimateur Θ_{est}^{ML} qui, à un ensemble de résultats $\vec{x}$, associe la valeur θ_{est}^{ML} qui maximise la vraisemblance :

$$\Theta_{\text{est}}^{ML} : \vec{x} \mapsto \theta_{\text{est}}^{ML} \equiv \arg \max_{\vartheta} p(\vec{x}|\vartheta). \quad (\text{III.1.13})$$

On appelle *log-vraisemblance* le logarithme de la fonction de vraisemblance. La fonction $\log$ étant strictement croissante, il est équivalent de maximiser la vraisemblance et la log-vraisemblance. En pratique, il est plus commode de travailler avec la log-vraisemblance, dès lors que les événements observés ont une faible probabilité à cause du grand nombre de résultats de mesure possibles. Effectuons un développement de la log-vraisemblance autour de son maximum :

$$\log p(\vec{x}|\theta) = \log p(\vec{x}|\theta_{\text{est}}^{ML}) + \frac{1}{2} (\theta - \theta_{\text{est}}^{ML})^2 \left. \frac{\partial^2 \log p(\vec{x}|\theta_{\text{est}}^{ML})}{\partial \theta^2} \right|_{\theta_{\text{est}}^{ML}} + O \left\{ (\theta - \theta_{\text{est}}^{ML})^3 \right\}. \quad (\text{III.1.14})$$

Le coefficient intervenant dans le terme d'ordre 2 vaut :

$$\left. \frac{\partial^2}{\partial \theta^2} \log p(\vec{x}|\theta_{\text{est}}^{ML}) \right|_{\theta_{\text{est}}^{ML}} = \sum_{i=1}^n \left. \frac{\partial^2}{\partial \theta^2} \log p(x_i|\theta) \right|_{\theta_{\text{est}}^{ML}} \underset{n \rightarrow \infty}{\sim} n \int dx p(x|\theta) \frac{\partial^2}{\partial \theta^2} \log p(x|\theta). \quad (\text{III.1.15})$$

On reconnaît, dans le terme de droite, l'expression III.1.7c de l'information de Fisher. On a donc, quand le nombre de mesures tend vers l'infini :

$$\begin{aligned} \log p(\vec{x}|\theta) &\simeq \log p(\vec{x}|\theta_{\text{est}}^{ML}) - \frac{1}{2} (\theta - \theta_{\text{est}}^{ML})^2 nF(\theta), \\ p(\vec{x}|\theta) &\simeq p(\vec{x}|\theta_{\text{est}}^{ML}) e^{-nF(\theta) \frac{(\theta - \theta_{\text{est}}^{ML})^2}{2}}. \end{aligned} \quad (\text{III.1.16})$$

En utilisant la loi de Bayes, on a

$$p(\vec{x}|\theta) = p(\theta|\vec{x}) \frac{p(\vec{x})}{p(\theta)}. \quad (\text{III.1.17})$$

Si on n'a aucune information *a priori* sur θ , on est amené à poser $p(\theta) = 1$. On obtient alors en combinant les deux expressions précédentes et en utilisant la condition de normalisation de la probabilité :

$$p(\theta|\vec{x}) = \frac{1}{\Delta \theta_{\text{est}}^{ML}} e^{-\frac{(\theta - \theta_{\text{est}}^{ML})^2}{2(\Delta \theta_{\text{est}}^{ML})^2}}. \quad (\text{III.1.18})$$

La probabilité *a posteriori* de θ est donc une gaussienne, dont la largeur vaut :

$$\Delta \theta_{\text{est}}^{ML} = \frac{1}{\sqrt{nF(\theta_{\text{est}}^{ML})}}. \quad (\text{III.1.19})$$

L'estimation par maximum de vraisemblance sature donc la borne de Cramer-Rao, dans la limite d'un grand nombre de mesures.

III.2 Principe de la tomographie quantique et des méthodes de reconstruction

Dans cette section, on présente le problème de la reconstruction d'états quantiques. On commence par rappeler quelques propriétés élémentaires de l'ensemble des matrices densités et de leur représentation. On introduit ensuite les principales méthodes de reconstruction quantiques utilisées.

III.2.1 Rappels sur l'opérateur densité

a) Définitions et propriétés

Lorsqu'on décrit un système quantique dans une expérience, on utilise un opérateur densité $\hat{\rho}$. Celui-ci satisfait les conditions

$$\hat{\rho} = \hat{\rho}^\dagger, \quad (\text{III.2.1a})$$

$$\text{Tr}(\hat{\rho}) = 1 \quad (\text{III.2.1b})$$

et pour tout vecteur d'état $|\psi\rangle$,

$$\langle\psi|\hat{\rho}|\psi\rangle \geq 0, \quad (\text{III.2.1c})$$

i. e. $\hat{\rho}$ est positif.

La première condition assure que $\hat{\rho}$ est un opérateur tandis que les deux suivantes permettent d'interpréter les termes diagonaux de la matrice densité comme des probabilités. Dans ce cadre, la probabilité p_ψ de trouver le système dans l'état $|\psi\rangle$ est donnée par la règle de Born :

$$p_\psi = \text{Tr}(\hat{\rho}|\psi\rangle\langle\psi|). \quad (\text{III.2.2})$$

L'opérateur densité peut être décomposé sur une base de l'espace de Hilbert. On peut alors écrire la matrice densité :

$$\hat{\rho} = \begin{matrix} & \begin{matrix} \langle 0| & \langle 1| & \dots & \langle d-1| \end{matrix} \\ \begin{matrix} |0\rangle \\ |1\rangle \\ \vdots \\ |d-1\rangle \end{matrix} & \begin{pmatrix} p_0 & \rho_{0,1} & \dots & \rho_{0,d-1} \\ \rho_{1,0} & p_1 & \dots & \rho_{1,d-1} \\ \vdots & \vdots & \ddots & \vdots \\ \rho_{d-1,0} & \rho_{d-1,1} & \dots & p_{d-1} \end{pmatrix} \end{matrix} \quad (\text{III.2.3})$$

Dans la suite, suivant l'usage habituel en physique quantique, on confondra opérateur densité et matrice densité et on omettra la base de représentation lorsque celle-ci est évidente (e. g. la base de Fock pour un oscillateur harmonique).

Les termes diagonaux dans $\hat{\rho}$ sont appelés *populations* et p_k représente la probabilité de trouver le système dans l'état $|k\rangle$. Les termes non diagonaux sont les *cohérences quantiques* de l'état, qui jouent un rôle crucial en information quantique. La présence de cohérences entre deux états donne lieu à des interférences quantiques susceptibles d'être mises en évidence dans des expériences appropriées. Ces cohérences dépendent de la base choisie

pour représenter l'opérateur densité. Considérons par exemple un bit quantique dans l'état $|+\rangle = (|0\rangle + |1\rangle)/\sqrt{2}$. Sa matrice densité peut s'écrire :

$$\rho_+ = |+\rangle \langle +| \quad (\text{III.2.4a})$$

$$\rho_+ = \frac{1}{2} \cdot \begin{array}{c} |0\rangle \\ |1\rangle \end{array} \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix} = \begin{array}{c} |+\rangle \\ |-\rangle \end{array} \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \quad (\text{III.2.4b})$$

où $|-\rangle = (|0\rangle - |1\rangle)/\sqrt{2}$ est un état orthogonal à $|+\rangle$. Dans la première base, l'état présente des cohérences maximales (si elles étaient plus grandes, la matrice ne serait plus positive), tandis que dans la deuxième base les cohérences, bien que toujours maximales, sont nulles. La bonne propriété pour classer les matrices densité est leur décomposition spectrale, qui est indépendante de la base choisie pour représenter l'opérateur. Ainsi, dans le cas précédent, les deux matrices admettent pour valeurs propres 0 et 1, ce qui correspond à un état pur. Plus généralement, on peut construire des quantités invariantes sous des opérations unitaires locales qui permettent de décrire les propriétés de cohérence des états quantiques [77] [78].

b) Ensemble $\mathcal{D}(\mathcal{H})$ des matrices densité

D'après la condition III.2.1a, l'ensemble $\mathcal{D}(\mathcal{H})$ des matrices densité est inclus dans l'espace vectoriel $\mathcal{O}(\mathcal{H})$ des opérateurs sur l'espace vectoriel $\mathcal{H}$. La condition III.2.1b, quand à elle est une équation linéaire sur coefficients de $\hat{\rho}$. Elle définit donc un hyperplan. $\mathcal{D}(\mathcal{H})$ est donc inclus dans l'hyperplan affine des matrices hermitiennes de trace 1. Enfin, la condition III.2.1c décrit la forme de $\mathcal{D}(\mathcal{H})$ dans cet hyperplan. Il s'agit d'une variété dont le bord est l'ensemble des matrices positives mais pas strictement positives, c'est à dire des matrices positives dont au moins une des valeurs propres est égale à zéro. Il s'agit donc des matrices densité de rang non maximal (le rang d'une matrice densité est le nombre de ses valeurs propres non nulles). On notera $\mathcal{B}(\mathcal{H})$ le bord de $\mathcal{D}(\mathcal{H})$.

Toute matrice densité $\hat{\rho}$ est diagonalisable et on peut écrire

$$\hat{\rho} = \sum_i p_i |\psi_i\rangle \langle \psi_i|, \quad (\text{III.2.5})$$

où les $|\psi_i\rangle \langle \psi_i|$ sont des projecteurs sur des états purs orthogonaux et où les p_i sont des probabilités. On peut donc écrire toute matrice densité comme le barycentre d'un ensemble d'états purs, avec des poids donnés par les probabilités de trouver ces états purs.

Par ailleurs, il est facile de vérifier qu'une somme pondérée de la forme III.2.5 est une matrice densité, y compris lorsque les $|\psi_i\rangle \langle \psi_i|$ ne sont pas orthogonaux. Il en résulte que l'ensemble des matrices densité est un ensemble convexe. On verra plus tard que cette propriété est importante pour résoudre des problèmes d'optimisation sur $\mathcal{D}(\mathcal{H})$.

Dans l'écriture matricielle III.2.3, il apparaît que la matrice densité contient d coefficients réels pour les populations ainsi que $d(d-1)$ coefficients complexes soit $2d(d-1)$ coefficients réels pour les cohérences. La condition III.2.1a impose cependant $d(d-1)$ relations entre les cohérences et leur conjuguée par transposition (égalité des parties réelles

et valeurs opposées pour les parties imaginaires). Par ailleurs la trace impose une relation entre les populations. La matrice densité contient donc $d^2 - 1$ coefficients réels indépendants. Il est intéressant de disposer d'une paramétrisation pratique de $\hat{\rho}$ en terme de nombres réels, qu'il s'agira alors d'estimer afin de réaliser une tomographie de l'état.

c) Paramétrisation de $\mathcal{D}(\mathcal{H})$ pour un qubit

Le bit quantique ou *qubit* est un élément fondamental de l'information quantique dont il constitue la brique élémentaire. C'est aussi une description adaptée pour des systèmes quantiques à deux niveaux. L'espace de Hilbert d'un qubit est $\mathcal{H} = \text{Vect} \{|0\rangle, |1\rangle\}$. Dans la suite, on l'utilisera comme exemple pour illustrer les méthodes présentées et on réservera pour le chapitre IV l'étude des subtilités introduites par la complexité du système à deux cavités auquel nous avons appliqué la méthode de tomographie par trajectoires. Dans le cas d'un bit quantique, il faut 3 paramètres réels pour représenter $\hat{\rho}$, qui appartient à l'hyperplan affine des matrices hermitiennes de trace unité. On peut donc écrire $\hat{\rho} = \frac{1}{2}\mathbb{1} + \hat{\rho}_0$ où $\hat{\rho}_0$ est de trace nulle. L'ensemble $\mathcal{O}(\mathcal{H})_0$ des opérateurs hermitiens de trace nulle est un espace vectoriel de dimension 3 dont une base est donnée par les matrices de Pauli $\hat{\sigma}_x$, $\hat{\sigma}_y$ et $\hat{\sigma}_z$. On peut donc écrire

$$\hat{\rho} = \frac{1}{2}(\mathbb{1} + \mathbf{r} \cdot \hat{\boldsymbol{\sigma}}), \quad (\text{III.2.6})$$

où $\mathbf{r} = (x, y, z)$ est un vecteur réel appelé *vecteur de Bloch*, d'après III.2.1a, et $\hat{\boldsymbol{\sigma}} = (\hat{\sigma}_x, \hat{\sigma}_y, \hat{\sigma}_z)$. On peut calculer facilement $\vec{r}$ à l'aide des relations suivantes :

$$x = \text{Tr}(\hat{\rho}\hat{\sigma}_x), \quad (\text{III.2.7a})$$

$$y = \text{Tr}(\hat{\rho}\hat{\sigma}_y), \quad (\text{III.2.7b})$$

$$z = \text{Tr}(\hat{\rho}\hat{\sigma}_z). \quad (\text{III.2.7c})$$

Notons p_0 et p_1 les valeurs propres de $\hat{\rho}$. Les conditions III.2.1b et III.2.1c imposent

$$p_0 + p_1 = 1, \quad (\text{III.2.8a})$$

$$0 \leq p_0 \leq 1, \quad (\text{III.2.8b})$$

$$0 \leq p_1 \leq 1. \quad (\text{III.2.8c})$$

En mettant III.2.6 au carré, on montre que :

$$\text{Tr}(\hat{\rho}^2) = \frac{1}{2}(1 + \vec{r}^2). \quad (\text{III.2.9})$$

Par ailleurs, les relations III.2.8 permettent de montrer que $0 \leq p_0^2 + p_1^2 \leq 1$ avec égalité à droite si et seulement si une seule des probabilités est non nulle et égale à 1, c'est-à-dire si l'état est pur. En injectant cette inégalité dans la précédente, on trouve que

$$|\vec{r}| \leq 1, \quad (\text{III.2.10})$$

avec égalité si et seulement si $\hat{\rho}$ est pur.

Le résultat précédent nous informe sur la structure de l'ensemble des matrices densité. On peut paramétriser ces matrices par un vecteur réel, qui appartient à la boule unité.

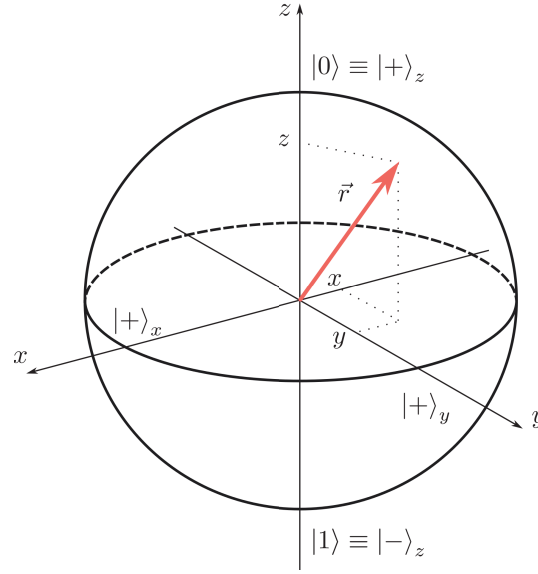


Fig. III.2 La sphère de Bloch permet de représenter l'ensemble des matrices densité. Chaque point de la boule représente un vecteur de Bloch, qui correspond à une matrice densité d'un qubit, d'après la formule III.2.6. Le pôle nord de la sphère est l'état $|0\rangle$ tandis que le pôle sud est l'état $|1\rangle$. Les états du plan équatorial sont des superpositions à poids égaux de $|0\rangle$ et $|1\rangle$. Les états du bord de la sphère sont les états purs, tandis que ceux de l'intérieur sont mélangés. L'état du centre de la sphère est l'état d'entropie maximale, totalement mélangé.

Par ailleurs, le bord de cet ensemble est constitué des états purs, dont le vecteur de Bloch est de norme 1. On peut donc représenter l'ensemble des états quantiques d'un système à deux niveaux par une boule appelée communément sphère de Bloch, visible sur la figure III.2. Dans le cas du qubit, le pôle nord (resp. sud) représente conventionnellement l'état $|0\rangle$ (resp. $|1\rangle$).

La sphère de Bloch présente l'intérêt qu'on peut lire directement les probabilités de mesure dans les différentes bases. En effet, la probabilité de mesurer le spin dans l'état $|+_z\rangle$ (resp. $|-_z\rangle$), lors d'une mesure selon Oz vaut $p_z = (1 + z)/2$ (resp. $(1 - z)/2$). Plus généralement, la probabilité d'observer l'état $|+_u\rangle$ dans une mesure selon l'axe Ou vaut

$$(1 + \vec{r} \cdot \vec{u})/2. \quad (\text{III.2.11})$$

Ceci permet d'illustrer la différence entre une superposition quantique de deux états, qui présente des cohérences dans une base correctement choisie, et un mélange statistique de ces mêmes états. Les cohérences donnent la longueur du vecteur $\vec{r}$, qui est elle-même responsable de la variation des probabilités de mesure lors d'un changement de base. Cette sensibilité fait que l'état change lors d'une évolution hamiltonienne, alors qu'un état complètement mélangé reste identique à lui-même. En particulier, ceci explique qu'on peut envoyer un état pur sur un autre état pur par une évolution hamiltonienne, ce qui sert à faire du contrôle cohérent d'états quantiques. On peut aussi utiliser des états purs pour encoder de l'information sur un paramètre en les faisant évoluer suivant un hamiltonien dépendant de ce paramètre. Les idées précédentes peuvent être visualisées directement

sur la sphère de Bloch dans le cas d'un qubit, mais restent vraies pour des systèmes plus compliqués.

Cette représentation est parfois trompeuse lorsqu'on parle de vecteurs d'états. En effet, il existe alors une liberté dans le choix de la phase quantique d'un état et il y a donc une infinité d'états $e^{i\phi}|\psi\rangle$ qui correspondent à un même point de la surface de la sphère de Bloch.

Revenons au cas général d'un espace de Hilbert de dimension d . Dans ce cas, il y a $d^2 - 1$ coefficients indépendants dans $\hat{\rho}$. On peut définir un vecteur de Bloch généralisé à l'aide des matrices de Gell-Man [78] :

$$\hat{u}_{jk} = |j\rangle\langle k| + |k\rangle\langle j|, \quad (\text{III.2.12a})$$

$$\hat{v}_{jk} = -i(|j\rangle\langle k| - |k\rangle\langle j|), \quad (\text{III.2.12b})$$

$$\hat{w}_l = \sqrt{\frac{2}{l(l+1)}} \sum_{j=1}^l (|j\rangle\langle j| - l|l+1\rangle\langle l+1|), \quad (\text{III.2.12c})$$

définie pour $0 \leq j \leq k \leq d-1$ et $0 \leq l \leq d-2$. On note $(\hat{\lambda}_i)_{i \in \llbracket 1, d^2-1 \rrbracket}$ la famille des matrices de Gell-Mann. On a alors une décomposition de $\hat{\rho}$ qui généralise III.2.6 :

$$\hat{\rho} = \frac{1}{d} \mathbb{1} + \sum_i \frac{1}{2} \lambda_i \hat{\lambda}_i, \quad (\text{III.2.13a})$$

où

$$\lambda_i = \text{Tr}(\hat{\rho} \hat{\lambda}_i). \quad (\text{III.2.13b})$$

III.2.2 Méthodes usuelles de reconstruction

Cadre général

Le cadre général de la tomographie d'états quantiques est le suivant. On dispose d'un système décrit par un espace de Hilbert de dimension d ,² qu'on est capable de préparer de manière répétée dans un état $\hat{\rho}$.

On réalise des mesures sur cet état, afin d'obtenir des informations sur les paramètres décrivant la matrice densité. On dispose de plusieurs configurations de mesure qui correspondent à un ensemble de paramètres expérimentaux associés à une réalisation de l'expérience. Par exemple, dans le cas d'un qubit, une configuration correspond à la direction selon laquelle on choisit de faire la mesure dans la sphère de Bloch. À une configuration de mesure c sont associés des résultats possibles $y \in Y^{(c)}$, avec chacun la probabilité $p^{\hat{\rho}}(y|c)$. Les mesures peuvent être décrites dans le cas général par un POVM (*positive operator-valued measurement*), c'est-à-dire un ensemble d'opérateurs $(\hat{E}_y^{(c)})_{y \in Y^{(c)}}$, hermitiens et positifs, vérifiant :

$$\sum_y \hat{E}_y^{(c)} = \mathbb{1}. \quad (\text{III.2.14})$$

2. On se restreint ici au cas de la dimension finie. La tomographie d'états dans un espace de dimension infinie joue un rôle important en optique quantique lorsqu'on effectue des mesures par détection homodyne. Il existe des techniques adaptées à cette situation mais elles dépassent le cadre de l'étude présente.

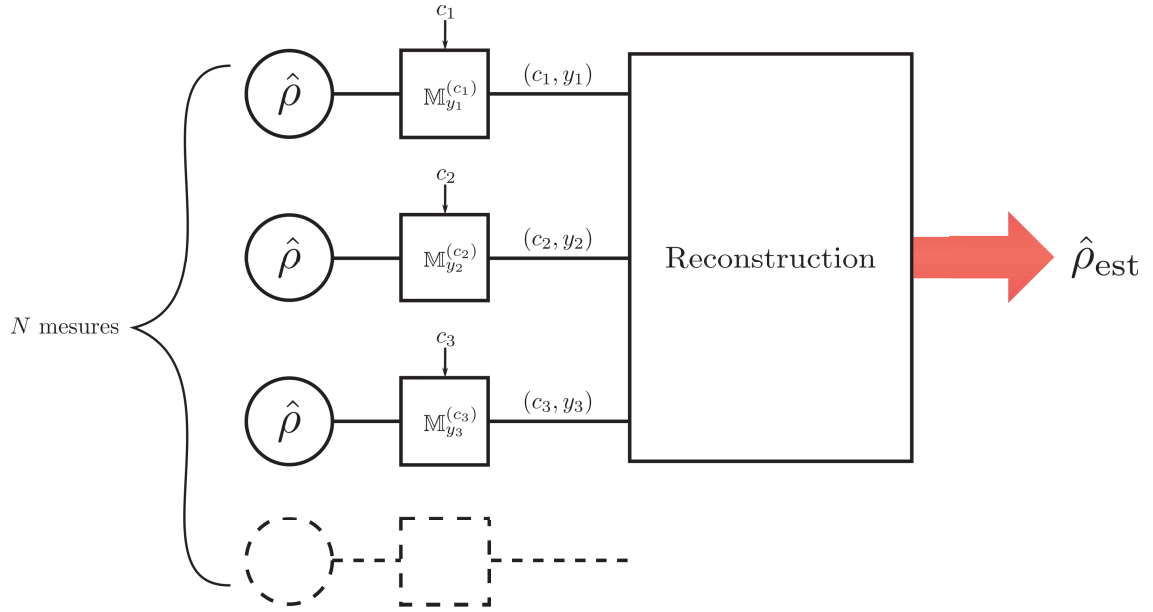


Fig. III.3 Schéma de principe de la tomographie. On prépare, de manière répétée, le système dans l'état $\hat{\rho}$. Pour chaque préparation i , on choisit une configuration de mesure c_i . On obtient alors un résultat y_i , avec une probabilité $\text{Tr}(\hat{\rho}\hat{E}_{y_i}^{(c)})$. On attribue ensuite à l'ensemble des couples (c_i, y_i) un état reconstruit $\hat{\rho}_{\text{est}}$, à l'aide d'une méthode de reconstruction.

On a alors

$$p^{\hat{\rho}}(y|c) = \text{Tr}(\hat{\rho}\hat{E}_y^{(c)}). \quad (\text{III.2.15})$$

Le principe du protocole de tomographie est schématisé sur la figure III.3. Pour chaque préparation de $\hat{\rho}$, on choisit une configuration de mesure et on attribue ensuite à l'ensemble de ces résultats une estimation $\hat{\rho}_{\text{est}}$ construite à partir d'une méthode de reconstruction. On décrit dans la suite plusieurs méthodes couramment utilisées.

La méthode d'inversion

L'équation III.2.15 montre que les probabilités d'occurrence des résultats d'une mesure sont données par un produit scalaire de la matrice $\hat{\rho}$ avec une matrice de mesure, où on utilise le produit scalaire de Frobenius sur l'ensemble des matrices hermitiennes :

$$\langle A, B \rangle \equiv \text{Tr}(AB) = \sum_{ij} A_{ij} B_{ji}. \quad (\text{III.2.16})$$

Il s'agit donc d'une équation linéaire faisant intervenir les éléments de $\hat{\rho}$. Si on dispose de $d^2 - 1$ matrices de mesure linéairement indépendantes, on peut inverser le système et exprimer les éléments de $\hat{\rho}$ en fonction des probabilités d'obtenir des résultats de mesure. On peut ensuite estimer ces probabilités en répétant les mesures un grand nombre de fois.

La méthode standard consiste à effectuer l'inversion précédente en utilisant, au lieu des probabilités auxquelles on n'a pas accès, les fréquences expérimentales d'occurrences

des résultats :

$$f(y|c) = \frac{N_{(y,c)}}{N_c}, \quad (\text{III.2.17})$$

où $N_{(y,c)}$ est le nombre de fois qu'on a obtenu le résultat y dans la configuration de mesure c et N_c le nombre de mesures dans cette configuration. On résout alors le système :

$$f(y_i|c_i) = \text{Tr} \left(\hat{\rho}_{\text{est}} \hat{E}_{y_i}^{(c_i)} \right), i \in \llbracket 1, d^2 - 1 \rrbracket, \quad (\text{III.2.18})$$

pour trouver $\hat{\rho}_{\text{est}}$.

Par exemple, dans le cas du qubit, on peut choisir comme configurations de mesurer selon les axes Ox , Oy et Oz . Supposons que les mesures réalisées sont projectives : $\hat{E}_+^{(i)} \equiv |+\rangle\langle +|$ et $\hat{E}_-^{(i)} \equiv |-\rangle\langle -|$ avec $i \in \{x, y, z\}$. On a alors $f(+|i) = N_{(+,i)}/N_i$ et on n'a pas besoin de $f(-|i)$. Le système III.2.18 prend alors la forme très simple :

$$f(+|x) = \frac{1+x}{2}, \quad (\text{III.2.19a})$$

$$f(+|y) = \frac{1+y}{2}, \quad (\text{III.2.19b})$$

$$f(+|z) = \frac{1+z}{2}. \quad (\text{III.2.19c})$$

La solution de ce système s'interprète simplement : pour chaque axe de mesure, on estime la coordonnée correspondante du vecteur de Bloch à partir des fréquences de mesure. L'état reconstruit est alors donné en injectant le vecteur de Bloch estimé dans III.2.6.

Cette méthode présente l'inconvénient de donner, en général, un résultat qui n'est pas une matrice densité. En effet, le bruit statistique sur les mesures fait que les probabilités estimées ne correspondent pas à une matrice densité. Par exemple, imaginons que l'état $\hat{\rho}$ soit un état totalement mélangé : $\hat{\rho} = \frac{1}{2}\mathbb{1}$. Si on effectue une mesure selon chaque axe, on obtiendra ± 1 pour chaque composante du vecteur de Bloch et donc un vecteur de Bloch de norme 3, qui ne correspond à aucune matrice densité. La figure III.4 illustre une situation dans laquelle le vecteur de Bloch estimé se trouve en dehors de la sphère de Bloch. L'ensemble des résultats possibles de l'inversion est donné par le parallélépipède en pointillé, circonscrit par les plans $x = \pm 1$, $y = \pm 1$ et $z = \pm 1$.

Le maximum de vraisemblance

La difficulté posée par l'inversion linéaire peut être levée en utilisant la méthode du maximum de vraisemblance [79]. Dans celle-ci, on cherche la matrice densité $\hat{\rho}_{\text{est}}$ qui maximise la vraisemblance $\mathcal{V}(\hat{\rho})$, c'est-à-dire la probabilité d'obtenir tous les résultats, étant donné un état initial $\hat{\rho}$. Dans ce cas, $\hat{\rho}_{\text{est}}$ est une matrice densité par construction.

Supposons qu'on ait réalisé N mesures et notons y_i le résultat de la i ème mesure, réalisée dans la configuration c_i . On a alors :

$$\mathcal{V}(\hat{\rho}) = \prod_{i=1}^N p^{\hat{\rho}}(y_i|c_i) = \prod_{i=1}^N \text{Tr} \left(\hat{\rho} \hat{E}_{y_i}^{(c_i)} \right). \quad (\text{III.2.20})$$

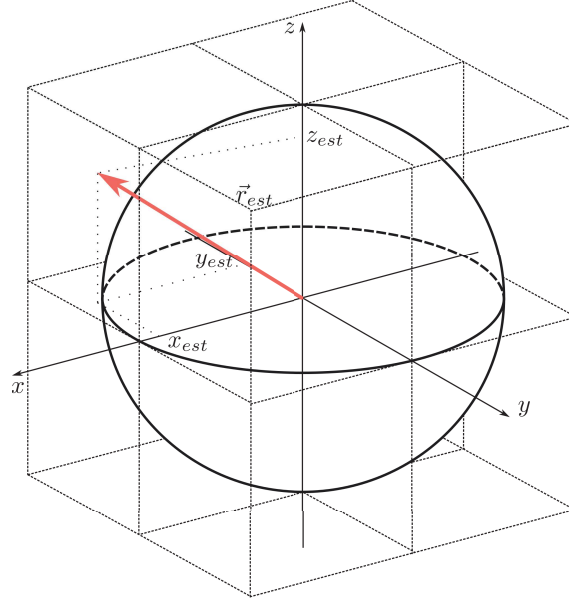


Fig. III.4 Rôle des erreurs dans la reconstruction par inversion. À cause de l'incertitude statistique, l'estimateur $\vec{r}_{est}$ peut avoir une norme supérieure à 1 et sortir de la sphère de Bloch. Le quadrilatère en pointillé donne l'ensemble des résultats possibles pour la reconstruction.

Étant donné un ensemble de résultats, $\mathcal{V}$ est une fonction de $\hat{\rho}$ qu'on cherche à maximiser. En pratique, on utilise plutôt la log-vraisemblance

$$\mathcal{L} = \log \mathcal{V}, \quad (\text{III.2.21})$$

qui est plus adaptée pour traiter numériquement de petites probabilités. Maximiser $\mathcal{L}$ est équivalent à maximiser $\mathcal{V}$ du fait de la stricte croissance de la fonction log. L'estimateur par maximum de vraisemblance est alors défini par :

$$\hat{\rho}_{\text{ML}} \equiv \max_{\hat{\rho}} \mathcal{L}(\hat{\rho}). \quad (\text{III.2.22})$$

Le logarithme étant une fonction concave, la log-vraisemblance l'est aussi. Comme l'ensemble des matrices densité est un ensemble fermé convexe, le problème de maximisation précédent est un problème convexe qui peut être résolu numériquement efficacement (voir [80], chapitre 4 et le chapitre III.3 dans la suite).

Plaçons nous dans le cas élémentaire du paragraphe 3.2 où on mesure un qubit avec des mesures projectives. On a alors :

$$\begin{aligned} \mathcal{L}(\hat{\rho} | N_{+x}, N_{-x}, N_{+y}, N_{-y}, N_{+z}, N_{-z}) = \\ N_{+x} \log \frac{1+x}{2} + N_{+y} \log \frac{1+y}{2} + N_{+z} \log \frac{1+z}{2} \\ + N_{-x} \log \frac{1-x}{2} + N_{-y} \log \frac{1-y}{2} + N_{-z} \log \frac{1-z}{2}. \end{aligned} \quad (\text{III.2.23a})$$

Cette fonction se décompose aisément sous la forme :

$$\mathcal{L}(\hat{\rho}) = \mathcal{L}_x(x) + \mathcal{L}_y(y) + \mathcal{L}_z(z) \quad (\text{III.2.24})$$

et on peut donc effectuer séparément la maximisation pour chaque paramètre. L'estimation d'une des composantes du vecteur de Bloch est alors identique à celle de la probabilité d'une loi de Bernoulli. Le maximum de la fonction de vraisemblance est atteint pour

$$x = \frac{N_{+x} - N_{-x}}{N_{+x} + N_{-x}}, \quad (\text{III.2.25a})$$

$$y = \frac{N_{+y} - N_{-y}}{N_{+y} + N_{-y}}, \quad (\text{III.2.25b})$$

$$z = \frac{N_{+z} - N_{-z}}{N_{+z} + N_{-z}}. \quad (\text{III.2.25c})$$

Ce maximum coïncide avec la solution du système III.2.19, obtenu dans le cas de l'inversion linéaire. Il y a alors deux situations possibles.

Si le maximum est atteint à l'intérieur de la sphère de Bloch, alors la reconstruction par la méthode par inversion et celle par maximum de vraisemblance donnent le même résultat. Si il est atteint en dehors de la sphère de Bloch, le maximum de vraisemblance donne un état situé sur le bord de la sphère de Bloch, c'est à dire un état pur. Cet état est obtenu en renormalisant le vecteur obtenu par III.2.25 de sorte qu'il soit unitaire. La reconstruction par maximum de vraisemblance donne donc toujours une matrice densité et elle coïncide avec la reconstruction par inversion dans le cas où cette dernière donne aussi une matrice densité. Ceci est toujours vrai dans les dimensions supérieures (voir [81], paragraphe 3.3.3).

Cette méthode présente cependant l'inconvénient qu'elle a tendance à surestimer la pureté des états reconstruits. Par exemple, dans le cas de la figure III.4, l'état qui maximise la vraisemblance sera un état pur.

La reconstruction bayésienne

La reconstruction bayésienne vise à prendre en compte les différentes informations *a priori* dont on dispose sur l'état, encodées dans une distribution de probabilités initiale. Cette distribution de probabilités est ensuite mise à jour à l'aide des résultats de mesure obtenus. La méthode bayésienne ne donne pas un état mais une distribution de probabilités sur les états, ce qui permet de calculer naturellement des valeurs moyennes ou des variances sur la valeur estimée par la reconstruction de n'importe quelle observable. La contrepartie de cette simplicité est le recours à des intégrales sur l'espace des matrices densité qui sont très coûteuses à calculer dans des espaces de grandes dimensions.

MaxEnt

La méthode de maximisation d'entropie diffère des méthodes présentées précédemment puisqu'elle vise à estimer un état en présence d'informations incomplètes. Elle a été appliquée avec succès à notre dispositif pour reconstruire des états non classiques du champ micro-onde [29]. Elle n'a cependant pas permis de prendre en compte toute l'information statistique des mesures effectuées à l'époque, puisqu'on a du moyenner les mesures effectuées afin de rendre l'analyse numérique possible.

Minimisation de norme

De nombreuses méthodes reposant sur des minimisations de normes ont été proposées [82]. Elles visent généralement à trouver la matrice de $\mathcal{D}(\mathcal{H})$ la plus proche des données mesurées, pour une certaine métrique. Comme la méthode du maximum de vraisemblance, elles ont l'avantage de donner un état physique. Elles ont par ailleurs une interprétation géométrique. La difficulté soulevée par ces méthodes est généralement qu'il existe de nombreux choix de normes et de projections sur l'ensemble des opérateurs linéaires et qu'on dispose de peu d'arguments pour en privilégier une.

Méthodes d'acquisition comprimée

Les méthodes d'acquisition comprimée (*compressed sensing* en anglais) utilisent le fait que la matrice à reconstruire est généralement de rang faible pour réduire le nombre de mesures à effectuer, en utilisant des mesures aléatoires [6, 19]. Pour une matrice de rang r dans un espace de dimension d , on s'attendrait à pouvoir reconstruire l'état avec $O(rd)$ configurations de mesure différentes, alors que les méthodes de tomographie par inversion en requièrent $d^2 - 1$. Les méthodes d'acquisition comprimée consistent à utiliser des bases de mesure choisies aléatoirement et à appliquer un algorithme de reconstruction. Pour un nombre de bases suffisant, de l'ordre de $O(rd \log^2 d)$, la probabilité d'échec devient très faible. Ces méthodes ont un intérêt évident lorsque la dimension de l'espace de Hilbert augmente et illustrent le fait qu'on peut profiter de la structure de l'ensemble des matrices de rang faible pour simplifier la reconstruction.

Tomographie adaptative

Les méthodes de tomographie adaptative constituent un ensemble de recettes permettant d'accélérer la convergence de la tomographie en choisissant les mesures à effectuer en fonction de l'information déjà acquise sur l'état à reconstruire. Il n'existe pas de cadre général pour la tomographie adaptative, mais de nombreux protocoles ont été proposés et appliqués [83–87].

III.2.3 Limitations des méthodes usuelles

Le cadre présenté jusqu'à présent fait intervenir des mesures uniques. Lorsqu'on est capable de faire des mesures qui ne détruisent pas le système, il est possible d'effectuer des mesures successives qui permettent d'augmenter l'information obtenue. Lorsqu'on peut réaliser une mesure QND, l'intérêt d'effectuer des mesures successives est évident, puisqu'on augmente la statistique de mesure en réalisant plusieurs mesures par préparation de l'état. Par ailleurs, utiliser plusieurs mesures lors de la même réalisation permet d'extraire plus d'information que les mesures prises isolément. Par exemple, lorsqu'on mesure le nombre de photons contenus dans une cavité, l'utilisation de multiples échantillons atomiques permet de réaliser une mesure projective à l'aide d'une série d'échantillons atomiques mesurés, même si la mesure effectuée par un seul atome n'est pas projective (voir [30], III.4.4). Rien n'oblige en principe à se limiter à des mesures QND. Même lorsque la mesure perturbe le système, on peut obtenir des informations précieuses en effectuant

plusieurs mesures successives, à condition de savoir précisément quel est l'effet de la mesure. Imaginons par exemple, dans le cas d'une cavité une mesure qui enlève un photon, tout en indiquant la présence de ce photon par un clic. Pour $n \geq 1$, une telle mesure peut être décrite par l'opérateur $M_n = |n-1\rangle\langle n|$. Lorsque la cavité est vide, il n'y a pas de clic et l'état n'est pas changé, on a alors l'opérateur : $M_0 = |0\rangle\langle 0|$. En effectuant cette mesure une seule fois, on n'obtient pas d'information sur le nombre de photons, à part celle qu'on sait si il est nul ou non. En l'appliquant $n+1$ fois, on réalise une mesure projective pour les états de nombre de photons inférieur à n .

Dans la suite, nous allons développer un cadre permettant de prendre en compte des mesures successives sur le système, ce qui permet d'intégrer dans la reconstruction des informations subtiles sur les corrélations entre plusieurs mesures dans le cas d'une réalisation donnée de l'expérience.

III.3 Tomographie par trajectoires

III.3.1 Trajectoires quantiques, opérateurs de Kraus et matrice d'effet

Avec le progrès des techniques de manipulation de systèmes quantiques isolés, il est devenu possible de contrôler et de mesurer des systèmes quantiques, sans les détruire. Ceci a permis de réaliser des séries de mesures sur des systèmes afin de suivre des trajectoires quantiques individuelles, d'étudier les sauts quantiques de systèmes constitués d'atomes ou de photons et de mettre en œuvre des protocoles de rétroaction quantique.

La relation d'incertitude de Heisenberg implique que la mesure d'une observable, en réduisant l'incertitude sur celle-ci, augmente celle sur l'observable conjuguée. Cet effet est appelé *action en retour*. Cette action en retour a été étudiée dans de nombreux systèmes quantiques et a même été utilisée pour préparer des états non classiques du champ contenu dans une cavité micro-onde [30]. Plus généralement, des mesures du système sont susceptibles de perturber l'état au delà de la limite de Heisenberg. Afin d'extraire correctement l'information contenue dans des séquences de mesures successives sur le système, on a donc besoin d'un cadre général permettant de décrire la transformation de l'état du système lors d'une mesure.

Par ailleurs, plus le nombre de mesures effectuées lors d'une même réalisation de l'état augmente, plus la durée de mesure augmente. Lorsque celle-ci devient comparable au temps d'interaction du système avec son environnement, il est nécessaire de prendre en compte les effets de décohérence induits par cette interaction.

Dans la suite, on mettra en place un cadre très général permettant de tenir compte à la fois de la décohérence et de l'action en retour, ainsi que de l'apport d'information fournie par la mesure. Ce cadre est une description unifiée de toute l'évolution (incluant les mesures et la décohérence) en terme d'opérateurs de Kraus dont les propriétés seront étudiées par la suite.

a) Mesure généralisée

Dans une première approche de la physique quantique, on considère généralement des mesures projectives, ou mesures de von Neumann. Considérant une observable $\hat{A}$, la

probabilité d'obtenir le résultat m_A est donnée par la règle de Born :

$$p(m_A) = \sum_k \langle m_A, k | \hat{\rho} | m_A, k \rangle, \quad (\text{III.3.1})$$

où les $|m_A, k\rangle \langle m_A, k|$ forment une base orthonormée de l'espace propre associé au résultat de mesure m_A . Après la mesure de m_A , le système est alors dans l'état :

$$\mathbb{M}_{m_A}(\hat{\rho}) = \frac{\sum_k |m_A, k\rangle \langle m_A, k| \hat{\rho} | m_A, k\rangle \langle m_A, k|}{p(m_A)}. \quad (\text{III.3.2})$$

Cette expression justifie la dénomination de mesure projective : l'état est projeté sur le sous-espace propre associé à la valeur mesurée, comme illustré figure III.5.

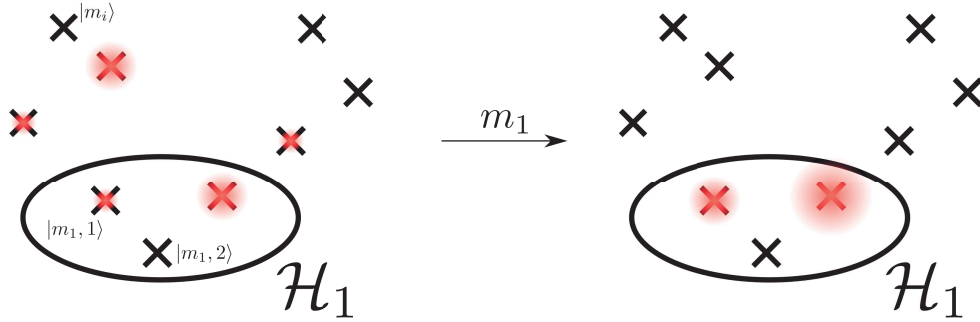


Fig. III.5 Effet d'une mesure projective sur l'état. Les disques rouges symbolisent les poids relatifs des différents états. Après la mesure de m_1 , l'état est projeté sur le sous-espace propre $\mathcal{H}_1$ associé.

Ce type de mesure n'est pas toujours adapté à la description d'un dispositif expérimental. Considérons, par exemple, le cas de la photodétection : on dispose d'un détecteur qui est capable d'absorber un photon et d'en tirer un signal macroscopique mesurable à l'aide d'une chaîne d'amplification. Si le mode dans lequel on détecte est dans l'état vide $|0\rangle$, rien ne se passe et le système reste dans $|0\rangle$. En revanche, si le mode est dans l'état de Fock $|1\rangle$ contenant un photon, le détecteur clique et le mode finit dans l'état $|0\rangle$ puisque le photon a été absorbé. Ces situations ne peuvent être décrites par une mesure de Von Neumann puisque l'état final ne correspond pas au résultat de la mesure.

On appelle *mesure généralisée* une mesure $\mathcal{M}$ définie par un ensemble d'opérateurs M_i de $\mathcal{H}$, associés à un résultat de mesure m_i , dont les probabilités de mesure sont données par

$$p(m_i) = \text{Tr} \left(M_i \hat{\rho} M_i^\dagger \right) \quad (\text{III.3.3})$$

et où l'état final est, après l'obtention du résultat m_i :

$$\mathbb{M}_{m_i}(\hat{\rho}) = \frac{M_i \hat{\rho} M_i^\dagger}{p(m_i)}. \quad (\text{III.3.4})$$

Dans l'exemple de photodétection précédent, la mesure d'un photon serait décrite par l'opérateur $M_1 \equiv |0\rangle \langle 1|$ et l'absence de détection par $M_0 \equiv |0\rangle \langle 0|$.

Les M_i n'ont pas nécessairement besoin d'être hermitiens. Leur nombre est arbitraire et peut dépasser la dimension de l'espace de Hilbert. La seule condition imposée est que la somme des probabilités de tous les résultats soit égale à un, pour tout $\hat{\rho}$, ce qui se traduit par la relation :

$$\sum_i M_i^\dagger M_i = \mathbb{1}. \quad (\text{III.3.5})$$

Les évolutions décrites par les formules III.3.2 et III.3.4 ne sont pas unitaires. Elles peuvent modifier la pureté de l'état. Lorsque la mesure n'est pas lue, on doit prendre en compte tous les résultats potentiels de mesure, et l'état évolue selon

$$\mathbb{M}_u(\hat{\rho}) = \sum_i p(m_i) \mathbb{M}_{m_i}(\hat{\rho}) = \sum_i M_i \hat{\rho} M_i^\dagger. \quad (\text{III.3.6})$$

L'état est alors mélangé, mais l'évolution est de nouveau linéaire.

b) Description de mesures imparfaites

En général la mesure est entachée d'erreur : on a une probabilité $p(r|k)$ d'obtenir le résultat r au lieu de k , à cause de la capacité de discrimination finie du détecteur. Dans ce cas, l'état après mesure de r est :

$$\mathbb{M}_r(\hat{\rho}) = \frac{\sum_k p(r|k) \mathbb{M}_k(\hat{\rho})}{\sum_k p(r|k)}. \quad (\text{III.3.7})$$

Les $(\mathbb{M}_r)_r$ ainsi obtenus définissent alors une nouvelle mesure généralisée.

Le formalisme de la mesure généralisée permet de décrire des mesures plus générales que les mesures projectives, dans lesquelles l'effet de la mesure est plus important que l'action en retour, qui est l'effet minimal traduit par la relation d'incertitude d'Heisenberg. La figure III.6 illustre le fait que la mesure généralisée peut perturber les états qui sont détectés par la mesure. Ce formalisme nous sera utile dans la suite pour décrire des mesures faibles qui ne projettent pas entièrement l'état du système, ainsi que les erreurs de mesure. Dans le paragraphe III.3.1.d, nous verrons qu'il permet d'interpréter la décohérence comme une mesure généralisée non lue réalisée par un environnement fictif.

c) Cartes quantiques et somme de Kraus

Considérons un système S . On peut se demander quelle est la transformation linéaire la plus générale transformant une matrice densité $\hat{\rho}_S$ en une autre matrice densité. Une telle transformation $\mathcal{L}_S$ est appelée une carte quantique. Afin d'être admissible, elle doit remplir les conditions suivantes :

$$\mathcal{L}_S(p\hat{\rho}_S + q\hat{\rho}'_S) = p\mathcal{L}_S(\hat{\rho}_S) + q\mathcal{L}_S(\hat{\rho}'_S), \quad (\text{III.3.8a})$$

$$\mathcal{L}_S(\hat{\rho}_S)^\dagger = \mathcal{L}_S(\hat{\rho}_S), \quad (\text{III.3.8b})$$

$$\text{Tr}(\mathcal{L}_S[\hat{\rho}_S]) = 1, \quad (\text{III.3.8c})$$

$$\langle \psi_S | \mathcal{L}_S(\hat{\rho}_S) | \psi_S \rangle \geq 0, \quad (\text{III.3.8d})$$

$$\langle \psi^{SE} | \mathcal{L}_S \otimes \mathbb{1}_E(\hat{\rho}_{SE}) | \psi^{SE} \rangle \geq 0, \quad (\text{III.3.8e})$$

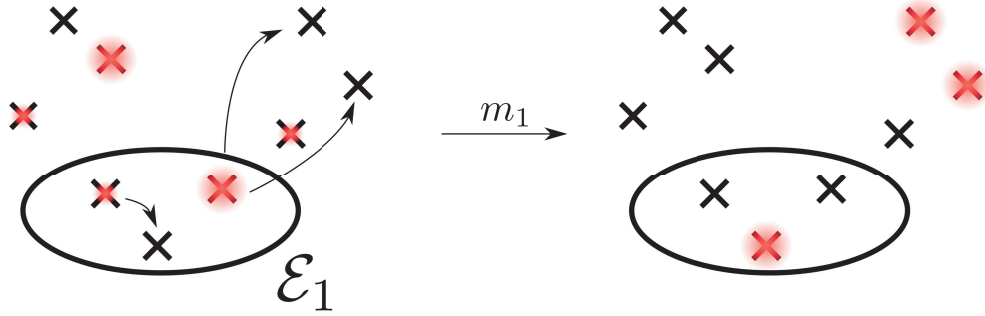


Fig. III.6 Effet d'une mesure généralisée sur l'état. Les disques rouges symbolisent les poids relatifs des différents états. Les flèches symbolisent la transition d'un état vers un autre lors de la mesure. L'ensemble $\mathcal{E}_1$ des états associés au résultat m_1 n'est plus un sous-espace propre de la mesure et le système peut se retrouver dans un état qui n'appartient pas à $\mathcal{E}_1$ après la mesure.

où $p, q \in \mathbb{R}$, $\hat{\rho}_S, \hat{\rho}'_S \in \mathcal{D}(\mathcal{H}_S)$, $|\psi_S\rangle \in S$ et où $|\psi^{SE}\rangle$ et $\hat{\rho}_{SE}$ sont des états du système SE avec E un système quantique quelconque.

La première condition exprime simplement la linéarité de la carte quantique. Les trois suivantes assurent que l'image d'une matrice densité par la carte quantique est une autre matrice densité. La dernière condition est nécessaire pour que la carte quantique soit *complètement positive*, c'est-à-dire qu'elle préserve la positivité des états produits tensoriels lorsqu'on décrit le système conjointement avec un autre système.

Il est possible de montrer que toute carte quantique peut se mettre sous la forme d'une somme de Kraus :

$$\mathcal{L}_S(\hat{\rho}_S) = \sum_{\mu=1}^{N_K} M_\mu \hat{\rho}_S M_\mu^\dagger, \quad (\text{III.3.9})$$

où les M_μ , appelés opérateurs de Kraus sont $N_K \leq d_S^2$, où $d_S \equiv \dim(\mathcal{H}_S)$ opérateurs vérifiant :

$$\sum_{\mu} M_\mu^\dagger M_\mu = \mathbb{1}, \quad (\text{III.3.10})$$

à condition que le système ne soit pas initialement intriqué avec son environnement. Un exemple élémentaire de décomposition de Kraus est celui d'une évolution unitaire $\hat{U}$:

$$\mathcal{L}_S(\hat{\rho}_S) = \hat{U} \hat{\rho}_S \hat{U}^\dagger, \quad (\text{III.3.11})$$

la somme a alors un seul terme et la pureté de l'état est préservée dans ce cas.

Ce résultat montre que, quelque soit la complexité de l'évolution du système, celle-ci peut être décrite par au plus $d_S^2 - 1$ opérateurs, tant qu'elle reste linéaire, c'est-à-dire tant qu'on ne révèle pas de résultat de mesure sur le système. Il montre aussi qu'on peut rendre compte de l'interaction du système avec un environnement arbitrairement compliqué en utilisant un environnement fictif ainsi que des opérateurs de Kraus convenablement choisis. Dans un contexte expérimental, on peut souvent se contenter de trouver les opérateurs de Kraus décrivant convenablement les principaux phénomènes en jeu, sans faire une description complète de l'environnement, ce qui serait impossible.

d) Interaction d'un système avec un environnement

Le résultat précédent est très puissant et permet de décrire de façon simplifiée le système sans entrer en détails dans la description de son environnement. Afin de voir comment on peut déterminer les opérateurs de Kraus, nous allons étudier ce qui se passe lorsqu'on dispose d'une description de l'évolution conjointe du système et de son environnement.

Considérons un système S d'espace de Hilbert $\mathcal{H}_S$, de dimension d_S et un environnement E d'espace de Hilbert $\mathcal{H}_E$ de dimension d_E . On suppose que le système et son environnement sont initialement dans un état séparable :

$$\hat{\rho}_i^{SE} = \hat{\rho}^S \otimes \hat{\rho}^E \quad (\text{III.3.12})$$

et que leur évolution conjointe est unitaire :

$$\hat{\rho}_f^{SE} = \hat{U}_{SE} \left(\hat{\rho}^S \otimes \hat{\rho}^E \right) \hat{U}_{SE}^\dagger. \quad (\text{III.3.13})$$

Si on considère S uniquement, l'état final est donné par :

$$\hat{\rho}_f^S = \text{Tr}_E \left(\hat{\rho}_f^{SE} \right). \quad (\text{III.3.14})$$

Afin de calculer la trace partielle, considérons une base $(|k\rangle)_{1 \leq k \leq d_E}$ de $\mathcal{H}_E$ qui diagonalise $\hat{\rho}^E$:

$$\hat{\rho}^E = \sum_{k=1}^{d_E} p_k |k^E\rangle \langle k^E|. \quad (\text{III.3.15})$$

On a alors :

$$\begin{aligned} \hat{\rho}_f^S &= \sum_{k,k'=1}^{d_E} p_{k'} \langle k^E | \hat{U}_{SE} | k'^E \rangle \hat{\rho}_S \langle k'^E | \hat{U}_{SE}^\dagger | k^E \rangle \\ &\equiv \sum_{k,k'=1}^{d_E} M_{k,k'} \hat{\rho}_S M_{k,k'}^\dagger. \end{aligned} \quad (\text{III.3.16a})$$

La notation $\langle k^E | \hat{U}_{SE} | k'^E \rangle$ est une notation abrégée pour un opérateur agissant sur $\mathcal{H}_S$. Les $M_{k,k'}$ sont donc des opérateurs de $\mathcal{H}_S$. Si on décompose l'opérateur d'évolution $\hat{U}_{SE} = \sum_{m,n} \sum_{p,q} u_{m,n,p,q} |m^S\rangle \langle n^S| \otimes |p^E\rangle \langle q^E|$, on a :

$$M_{k,k'} = \sqrt{p_{k'}} \sum_{m,n} u_{m,n,k,k'} |m^S\rangle \langle n^S|. \quad (\text{III.3.17})$$

On vérifie que les $M_{k,k'}$ vérifient la condition III.3.5 :

$$\sum_{k,k'} M_{k,k'}^\dagger M_{k,k'} = \sum_{k,k'} p_{k'} \langle k'^E | \hat{U}_{SE}^\dagger | k^E \rangle \langle k^E | \hat{U}_{SE} | k'^E \rangle, \quad (\text{III.3.18a})$$

$$= \sum_{k'} p_{k'} \langle k'^E | \hat{U}_{SE}^\dagger \hat{U}_{SE} | k'^E \rangle, \quad (\text{III.3.18b})$$

$$= \sum_{k'} p_{k'} \mathbb{1}_S, \quad (\text{III.3.18c})$$

$$= \mathbb{1}_S. \quad (\text{III.3.18d})$$

On a utilisé la relation de fermeture pour passer de la première ligne à la deuxième, on a utilisé la relation de fermeture, puis l'unitarité de $\hat{U}_{SE}$ pour passer à la troisième ligne.

Dans certains cas, l'environnement peut être un « mètre », c'est-à-dire un système ancillaire qu'on va mesurer par la suite pour obtenir de l'information sur l'état de S . La relation précédente montre qu'on peut en effet interpréter les opérateurs de Kraus comme des éléments d'une mesure généralisée. Imaginons qu'on fasse interagir S et E suivant l'évolution III.3.13 et qu'on décide de décrire l'état de S en tenant compte cette fois de l'état de E . Pour ceci, on prépare initialement E dans l'état $\hat{\rho}^E = |k'^E\rangle\langle k'^E|$ et on réalise une mesure projective dans la base $(|k\rangle)_{1 \leq k \leq d_E}$, qui donne le résultat $|k^E\rangle$. L'état global est alors projeté sur l'état :

$$\hat{\rho}_{|k}^{SE} = \frac{1}{p_k} \langle k^E | \hat{U}_{SE} | k'^E \rangle \hat{\rho}_S \langle k'^E | \hat{U}_{SE}^\dagger | k^E \rangle \otimes |k^E\rangle\langle k^E|, \quad (\text{III.3.19})$$

où p_k est la probabilité d'obtenir le résultat k . On obtient donc, comme état réduit pour S :

$$\hat{\rho}_f^S = \frac{\hat{M}_{k,k'} \hat{\rho}_S \hat{M}_{k,k'}^\dagger}{p_k}. \quad (\text{III.3.20})$$

On retrouve donc bien une mesure généralisée et on peut interpréter la formule III.3.16, dans le cas où on ne mesure pas l'environnement, de la manière suivante : l'interaction fait fuir de l'information sur le système vers l'environnement. On peut voir cette information comme une mesure généralisée non lue. Cette information permet en principe de savoir par quel chemin le système est passé. Il y a donc une suppression des cohérences entre les composantes de $\hat{\rho}_S$ correspondant à des résultats de mesure de l'environnement différents. Cette interprétation fournit un mécanisme permettant d'interpréter la décohérence. Lors de l'interaction entre un système et son environnement, de l'information sur le premier est encodée dans le second. Si on ne tient pas compte de celle-ci, il faut mélanger les états en proportion des probabilités de mesure des différents résultats correspondant aux opérateurs de Kraus entrant dans la décomposition en terme d'opérateurs de Kraus.

La relation III.3.16 permet de déterminer les opérateurs de Kraus, connaissant l'état initial de l'environnement, et l'opérateur décrivant l'interaction du système avec son environnement. Nous en ferons usage par la suite pour rendre compte de l'effet de mesures successives et du couplage à l'environnement sur un système, sans avoir à décrire le détail de ces effets, ce qui augmenterait très fortement la taille de l'espace de Hilbert et rendrait les calculs impossibles à réaliser en pratique.

e) Équation pilote de Lindblad

L'équation III.3.9 montre qu'on peut mettre toute évolution linéaire d'un état, décrit par une carte quantique, sous la forme d'une somme de Kraus. Les opérateurs de cette somme peuvent être interprétés comme des sauts quantiques correspondant à des mesures d'observables bien choisies d'un système de dimension supérieure. On a vu dans la section précédente comment calculer ces opérateurs. Suivant [27], on va désormais esquisser la dérivation d'une équation différentielle décrivant l'évolution temporelle de $\hat{\rho}_S$ lorsque le système interagit avec un environnement qui ne garde pas mémoire des corrélations avec S .

Considérons un système S interagissant avec l'environnement E pendant un temps τ , court devant le temps caractéristique d'interaction τ_I . Si l'état conjoint du système et de l'environnement est initialement séparable, c'est-à-dire si $\hat{\rho}^{SE}(t) = \hat{\rho}_S(t) \otimes \hat{\rho}^E(t)$, on peut écrire l'état après une durée τ sous la forme d'une somme de Kraus :

$$\hat{\rho}_S(t + \tau) = \mathcal{L}_S^\tau(\hat{\rho}_S(t)). \quad (\text{III.3.21})$$

On peut alors définir la dérivée temporelle de $\hat{\rho}_S$ comme :

$$\frac{d\hat{\rho}_S}{dt} = \frac{\mathcal{L}_S^\tau(\hat{\rho}_S(t)) - \hat{\rho}_S(t)}{\tau}. \quad (\text{III.3.22})$$

Si $\mathcal{L}_S$ est une fonction suffisamment régulière de τ [88], on peut écrire $\mathcal{L}_S^\tau(\hat{\rho}_S(t + \tau)) = \hat{\rho}_S(t) + O(\tau)$. Un des opérateurs de Kraus doit donc être proche de $\mathbb{1}$, ce qu'on peut écrire :

$$M_0 = \mathbb{1} - iK\tau + O(\tau^2), \quad (\text{III.3.23})$$

où K est un opérateur indépendant de τ . On peut isoler les parties hermitienne et anti-hermitienne de K :

$$H = \hbar \frac{K + K^\dagger}{2}, \quad (\text{III.3.24a})$$

$$J = i \frac{K - K^\dagger}{2}. \quad (\text{III.3.24b})$$

On a alors, au premier en τ :

$$M_0 \hat{\rho}_S M_0 = \hat{\rho}_S - \frac{i\tau}{\hbar} [H, \hat{\rho}_S] - \tau (J \hat{\rho}_S + \hat{\rho}_S J). \quad (\text{III.3.25})$$

Les autres termes de la somme de Kraus sont d'ordre τ et on peut les écrire $M_\mu = \sqrt{\tau} L_\mu$. La condition de normalisation des opérateurs de Kraus donne alors :

$$\sum_{\mu=0}^{N_K-1} M_\mu^\dagger M_\mu = \mathbb{1} - 2J\tau + \sum_{\mu=1}^{N_K-1} \tau L_\mu^\dagger L_\mu + O(\tau^2) = \mathbb{1}, \quad (\text{III.3.26})$$

dont on déduit :

$$J = \frac{1}{2} \sum_{\mu=1}^{N_K-1} L_\mu^\dagger L_\mu. \quad (\text{III.3.27})$$

On peut alors mettre l'évolution sous la forme de l'équation pilote, dite de Lindblad :

$$\frac{d\hat{\rho}_S}{dt} = -\frac{i}{\hbar} [H, \hat{\rho}_S] + \sum_{\mu=1}^{N_K-1} \left(L_\mu \hat{\rho}_S \tau L_\mu^\dagger - \frac{1}{2} \hat{\rho}_S L_\mu^\dagger L_\mu - \frac{1}{2} L_\mu^\dagger L_\mu \hat{\rho}_S \right). \quad (\text{III.3.28})$$

Cette équation permet de décrire l'évolution du système lors de son couplage à un environnement. Pour pouvoir l'écrire, on a cependant fait l'hypothèse que l'état du système et de son environnement reste séparable. Ceci est valide en pratique lorsque l'environnement peut être considéré comme un réservoir dont le temps d'autocorrélation τ_E est petit

devant le pas de temps τ choisi pour établir l'équation de Lindblad. On doit donc avoir la hiérarchie d'échelles de temps

$$\tau_E \ll \tau \ll \tau_I. \quad (\text{III.3.29})$$

En pratique, cette condition est généralement vérifiée. Par exemple, dans le cas de la relaxation du champ contenu dans une de nos cavités, l'environnement est l'ensemble des modes micro-ondes à la fréquence de la cavité, dont le temps d'autocorrélation est de l'ordre de grandeur de la période du champ micro-onde, égale à une fraction de nano-seconde et le temps d'interaction entre le mode de la cavité et cet environnement est de l'ordre du temps de vie de la cavité, qui vaut quelques dizaines de millisecondes. On peut donc facilement trouver une échelle de temps intermédiaire entre les deux.

Connaissant le détail de l'interaction entre le système et son environnement, on peut calculer les opérateurs de Kraus qui en résultent. En pratique, il suffit souvent d'identifier des opérateurs de Kraus décrivant correctement les principaux mécanismes d'interaction du système avec son environnement, en gardant en tête qu'il en suffit de d_S^2 pour décrire l'évolution la plus générale.

f) Matrices d'effet

Dans ce paragraphe, nous allons montrer comment on peut associer à une évolution du système une matrice d'effet qui décrit l'information qu'on acquiert sur lui par la mesure. On considère une réalisation r de l'expérience. Le système est préparé dans l'état $\hat{\rho}$ et on effectue T_r mesures successives dessus, entre lesquelles il évolue suivant sa dynamique propre, tout en interagissant avec son environnement. On appellera trajectoire de mesure l'ensemble des résultats de mesure pour une telle réalisation. On note $y_t^{(r)}$ le résultat de la t -ième mesure de la trajectoire r et $Y^{(r)} \equiv (y_t^{(r)})_{t \in \llbracket 1, T_r \rrbracket}$. On appelle $\hat{\rho}_t^{(r)}$ l'état du système après la mesure t où $t \in \llbracket 1, T_r \rrbracket$. La probabilité d'observer le résultat y lors de la t -ième mesure s'écrit alors

$$p(y_t^{(r)} = y | \hat{\rho}) = \text{Tr} \left(\mathbb{K}_t^{(r)}(\hat{\rho}_{t-1}^{(r)}) \right), \quad (\text{III.3.30})$$

et l'état résultant

$$\hat{\rho}_t^{(r)} = \frac{\mathbb{K}_t^{(r)}(\hat{\rho}_{t-1}^{(r)})}{\text{Tr} \left(\mathbb{K}_t^{(r)}(\hat{\rho}_{t-1}^{(r)}) \right)}, \quad (\text{III.3.31})$$

où $\mathbb{K}_t^{(r)}$ est le super-opérateur correspondant au résultat de mesure $y_t^{(r)}$ lors de la mesure t dans la trajectoire r . Cette formulation est très générale, comme on l'a vu dans la section [III.3.1.c](#) et permet d'introduire dans les super-opérateurs $\mathbb{K}_t^{(r)}$ l'évolution du système entre deux mesures, y compris en présence de décohérence, sous la forme :

$$\mathbb{K}_t^{(r)} = \mathbb{M}_{y_t^{(r)}} \circ \mathbb{U}_t \circ \mathbb{D}_t, \quad (\text{III.3.32})$$

où $\mathbb{U}_t$ est la carte quantique associée à une opération unitaire effectuée avant la mesure t et où $\mathbb{D}_t$ est la carte quantique décrivant la décohérence entre les mesures $t-1$ et t , qui peut être calculée à partir de l'équation de Lindblad [III.3.28](#).

Des exemples détaillés de super-opérateurs utilisés pour nos tomographies sont donnés dans le chapitre [IV](#).

La probabilité d'obtenir la trajectoire r s'écrit :

$$p(Y^{(r)}|\hat{\rho}) = \prod_{t=1}^{T_r} p(y_t^{(r)}|\hat{\rho}) \quad (\text{III.3.33a})$$

$$= \prod_{t=1}^{T_r} \text{Tr} \left(\mathbb{K}_t^{(r)}(\hat{\rho}_t^{(r)}) \right) \quad (\text{III.3.33b})$$

$$= \text{Tr} \left(\mathbb{K}_t^{(r)} \circ \mathbb{K}_{t-1}^{(r)} \circ \dots \circ \mathbb{K}_1^{(r)}(\hat{\rho}_0^{(r)}) \right) \quad (\text{III.3.33c})$$

$$= \text{Tr} \left(\mathbb{K}_t^{(r)} \circ \mathbb{K}_{t-1}^{(r)} \circ \dots \circ \mathbb{K}_1^{(r)}(\hat{\rho}) \right) \quad (\text{III.3.33d})$$

Pour la troisième égalité, on a procédé par récurrence en utilisant [III.3.31](#). La probabilité d'obtenir les résultats de la trajectoire r peut donc se calculer simplement comme la trace d'un opérateur obtenu en faisant agir une série de super-opérateurs sur l'état initial. Ce calcul est cependant très coûteux. Chaque super-opérateur peut être représenté par une matrice de taille $(d^2)^2$. Or pour reconstruire $\hat{\rho}$, on va avoir besoin de calculer $p(Y^{(r)}|\hat{\rho})$ pour de nombreuses matrices tests $\hat{\rho}$.

Afin de simplifier le calcul, on va utiliser l'adjoint $\tilde{\mathbb{K}}$ de $\mathbb{K}$, qui est un super-opérateur tel que :

$$\text{Tr} (A\mathbb{K}(B)) = \text{Tr} \left(\tilde{\mathbb{K}}(A)B \right). \quad (\text{III.3.34})$$

Pour un super-opérateur défini par

$$\mathbb{K}(\hat{\rho}) = \sum_{\mu} M_{\mu} \hat{\rho} M_{\mu}^{\dagger}, \quad (\text{III.3.35})$$

l'adjoint est donné par :

$$\tilde{\mathbb{K}}(\hat{\rho}) = \sum_{\mu} M_{\mu}^{\dagger} \hat{\rho} M_{\mu}. \quad (\text{III.3.36})$$

De plus, on a

$$(\mathbb{K}_2 \circ \mathbb{K}_1) \sim \tilde{\mathbb{K}}_1 \circ \tilde{\mathbb{K}}_2. \quad (\text{III.3.37})$$

En utilisant ces relations et [III.3.33](#), on a donc :

$$p(Y^{(r)}|\hat{\rho}) = \text{Tr} \left(\mathbb{K}_t^{(r)} \circ \mathbb{K}_{t-1}^{(r)} \circ \dots \circ \mathbb{K}_1^{(r)}(\hat{\rho}) \mathbb{1} \right), \quad (\text{III.3.38a})$$

$$= \text{Tr} \left(\hat{\rho} \tilde{\mathbb{K}}_1^{(r)} \circ \tilde{\mathbb{K}}_2^{(r)} \circ \dots \circ \tilde{\mathbb{K}}_t^{(r)}(\mathbb{1}) \right). \quad (\text{III.3.38b})$$

On appelle matrice d'effet $\hat{E}^{(r)}$ associée à la trajectoire r l'opérateur

$$\hat{E}^{(r)} \equiv \tilde{\mathbb{K}}_1^{(r)} \circ \tilde{\mathbb{K}}_2^{(r)} \circ \dots \circ \tilde{\mathbb{K}}_t^{(r)}(\mathbb{1}), \quad (\text{III.3.39})$$

de sorte que

$$p(Y^{(r)}|\hat{\rho}) = \text{Tr} \left(\hat{\rho} \hat{E}^{(r)} \right). \quad (\text{III.3.40})$$

Cette expression montre qu'on peut calculer très facilement la probabilité d'obtenir la trajectoire r comme le produit scalaire de Frobenius de $\hat{\rho}$ avec $\hat{E}^{(r)}$. On montre facilement, à partir de l'expression [III.3.36](#) que la matrice d'effet est hermitienne et positive.

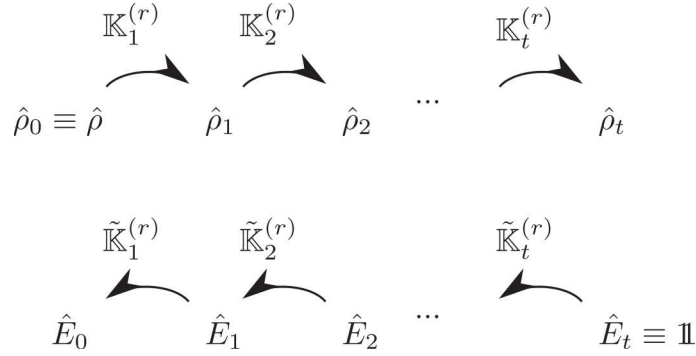


Fig. III.7 Représentation schématique du calcul de la matrice d'effet. On part de l'identité à laquelle on fait suivre une évolution à rebrousse-temps à l'aide des opérateurs de Kraus adjoints.

Attardons-nous un peu sur la forme de la matrice d'effet. On peut la voir comme l'application, à l'opérateur identité, d'une évolution à rebrousse-temps, comme schématisé sur la figure III.7. La matrice d'effet à l'instant t contient toutes les mesures effectuées après t , comme illustré sur la figure III.8. Elle a été utilisée comme un moyen d'améliorer l'estimation du résultat d'une mesure non lue faite à l'instant t en prenant en compte les mesures effectuées ultérieurement et en construisant un *état passé* contenant davantage d'information que la matrice densité [89]. Cette technique a été utilisée sur notre expérience à l'analyse du nombre de photons contenus dans une cavité [33] et a fait l'objet de la thèse de Théo Rybarczyk [34].

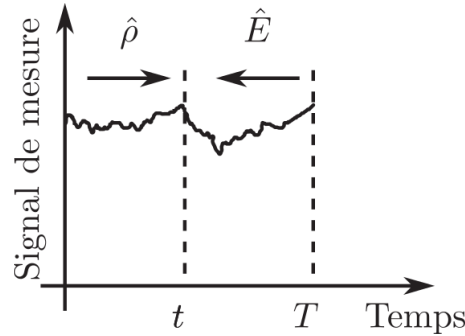


Fig. III.8 La matrice densité décrit l'information dont on dispose sur un système au temps t à l'aide des mesures effectuées précédemment. La matrice d'effet décrit l'information supplémentaire qu'on acquiert par des mesures après t sur les valeurs de mesures qu'on a effectuées à t . (figure reprise de [89])

On s'est contenté ici de donner la méthode de calcul de $\hat{E}$ pour une évolution décrite par un super-opérateur, ce qui suffira pour la suite. Dans le cas d'une évolution continue, on peut montrer que $\hat{E}$ suit une équation similaire à l'équation de Lindblad III.3.28 où le sens du temps est inversé [89].

g) Conclusion de la section

Dans cette section, on a vu que l'évolution d'un système quantique peut être décrite de façon unifiée à l'aide de super-opérateurs faisant intervenir des opérateurs de Kraus. On a vu par la suite comment on peut résumer l'information contenue dans une trajectoire expérimentale compliquée, faisant intervenir de multiples mesures dans un opérateur de l'espace de Hilbert appelé matrice d'effet, à partir duquel on peut calculer facilement la probabilité d'obtenir les résultats. Dans la suite, on va voir comment utiliser ces outils pour effectuer une reconstruction par maximum de vraisemblance plus générale que celle introduite précédemment.

III.3.2 Tomographie par trajectoires : principe et mise en application

Dans cette partie, on introduit une méthode de tomographie d'états quantiques par maximisation de vraisemblance qui utilise les matrices d'effet introduites précédemment pour prendre en compte l'information résultant de mesures multiples. L'accent est mis sur le nombre de mesures nécessaires ainsi que sur leur choix. Afin d'illustrer ces idées, le cas du qubit est traité.

a) Principe de la méthode et construction des matrices d'effet

La reconstruction que nous utilisons est une reconstruction par maximum de vraisemblance, mais où contrairement au cadre de la section III.2.2, on effectue des séries de mesures successives. Dans les méthodes présentées jusqu'ici, on effectuait une seule mesure par réalisation de l'état $\hat{\rho}$. La prise en compte de mesures multiples rend l'analyse beaucoup plus riche, mais aussi beaucoup plus complexe, comme on peut le voir en comparant les figures III.3 et III.9.

Commençons par poser le cadre de la reconstruction. On effectue N préparations successives du système dans l'état $\hat{\rho}$. Reprenons les notations de la section III.3.1.f. Les résultats de la trajectoire r sont notés

$$Y^{(r)} \equiv \left(y_t^{(r)} \right)_{t \in \llbracket 1, T_r \rrbracket}, \quad (\text{III.3.41})$$

et les résultats de l'ensemble des trajectoires

$$Y \equiv \left(Y^{(r)} \right)_{r \in \llbracket 1, N \rrbracket}. \quad (\text{III.3.42})$$

La méthode de tomographie que nous utilisons procède par maximisation de la vraisemblance sur l'ensemble Y des résultats de mesure, c'est à dire qu'on cherche la matrice $\hat{\rho}_{\text{ML}}$ qui maximise la log-vraisemblance

$$\mathcal{L}(\hat{\rho}) = \log p(Y|\hat{\rho}). \quad (\text{III.3.43})$$

Cette méthode prend donc en compte toutes les trajectoires individuellement sans moyenner d'observables sur un ensemble de mesures. On l'appellera donc *reconstruction par trajectoires*.

Afin de calculer $\mathcal{L}$, on utilise les outils de la section III.3.1, en particulier les représentations des diverses évolutions du système sous la forme de cartes quantiques et les

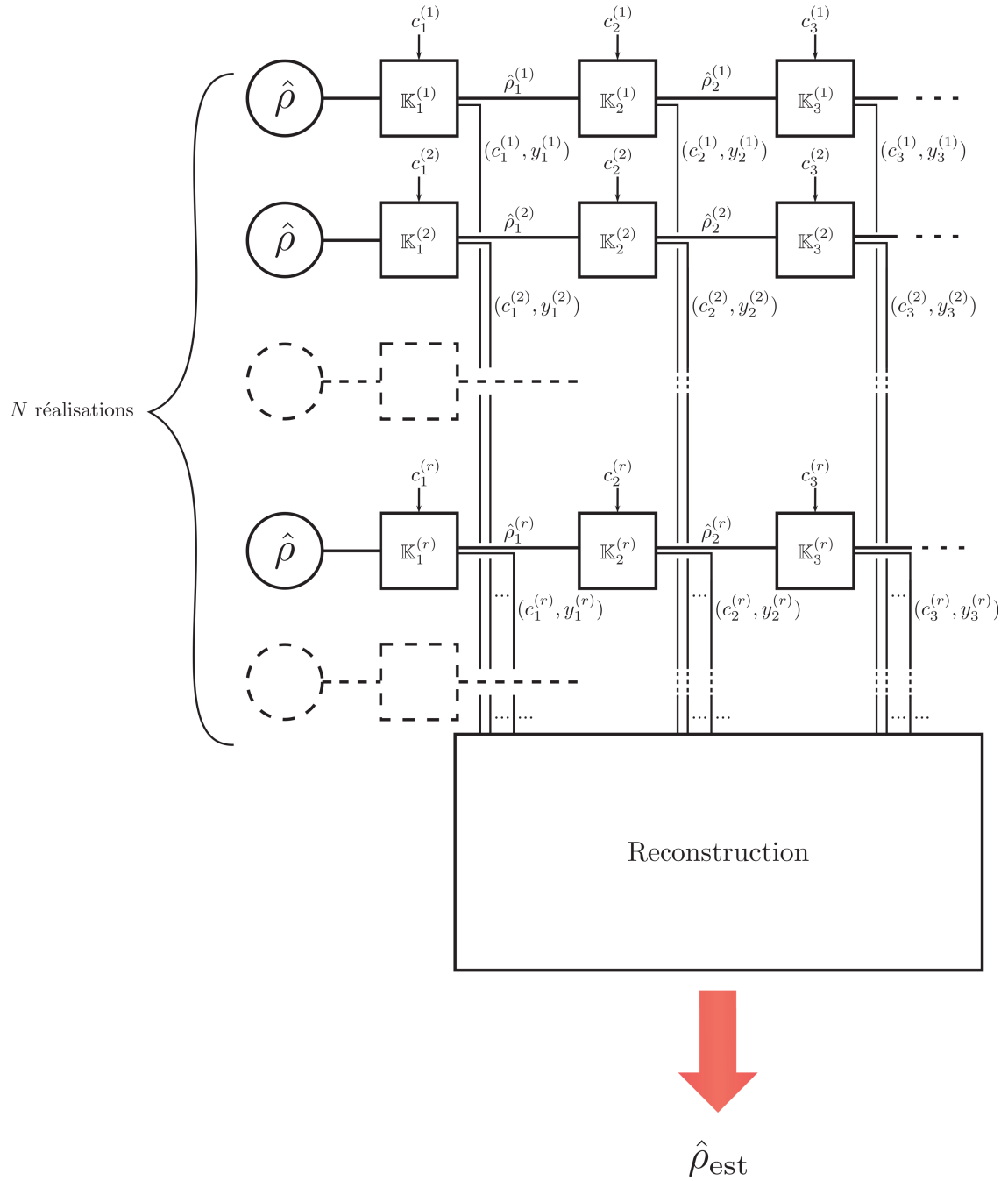


Fig. III.9 Schéma de principe de notre reconstruction. Pour chaque réalisation r , on choisit une série de mesures différentes décrites par les configurations $(c_t^{(r)})_{t \in \llbracket 1, T_r \rrbracket}$. À chaque pas de temps t , on obtient le résultat $y_t^{(r)}$. On a alors la trajectoire de mesures $(y_t^{(r)})_{t \in \llbracket 1, T_r \rrbracket}$. Les mesures perturbent l'état, qui passe de $\hat{\rho}_{t-1}^{(r)}$ à $\hat{\rho}_t^{(r)} \propto \mathbb{K}_t^{(r)}(\hat{\rho}_{t-1}^{(r)})$. On construit un estimateur en utilisant la connaissance de l'ensemble des configurations et résultats.

matrices d'effet. Remarquons que dans la section III.3.1.f, on avait calculé la probabilité d'obtenir la trajectoire effectivement suivie par le système, en utilisant l'état réel $\hat{\rho}$ préparé. Ceci n'est pas possible en pratique. En effet, l'état $\hat{\rho}$ est l'état que nous souhaitons reconstruire, auquel nous n'avons pas accès par principe. Dans toute la suite, on calculera donc des probabilités en fonction d'un état quelconque $\hat{\varrho}$ qui servira uniquement de paramètre aux fonctions qu'on cherchera à maximiser.

Comme les différentes réalisations sont indépendantes, les probabilités se factorisent et on a :

$$\mathcal{L}(\hat{\varrho}) = \log p(Y|\hat{\varrho}) = \log \prod_{r=1}^N p(Y^{(r)}|\hat{\varrho}), \quad (\text{III.3.44a})$$

$$= \sum_{r=1}^N \log p(Y^{(r)}|\hat{\varrho}), \quad (\text{III.3.44b})$$

$$= \sum_{r=1}^N \log \text{Tr} \left(\hat{\varrho} \hat{E}^{(r)} \right). \quad (\text{III.3.44c})$$

Le problème est donc ramené à maximiser une fonction convexe de termes qui sont des produits scalaires de Frobenius entre l'état considéré et une matrice d'effet. L'état $\hat{\rho}_{\text{ML}}$ est donc celui qui maximise, d'une certaine manière, le recouvrement avec les matrices d'effet qui décrivent les différentes mesures. Cette simplification des termes du problème constitue un gain énorme en complexité du calcul. En effet, on peut précalculer les matrices d'effet et il n'y a donc plus à effectuer le calcul de l'évolution pour chaque itération de l'algorithme utilisé pour trouver le maximum de la vraisemblance. Nous allons donner dans la suite plus de détails sur les matrices d'effet et sur la manière de les calculer.

b) Le choix des mesures et leur nombre

Pour reconstruire de manière non ambiguë l'état, il faut $d^2 - 1$ mesures indépendantes. Si on ne dispose pas d'assez de mesures indépendantes, il existe des directions de l'espace des opérateurs hermitiens selon lesquelles la vraisemblance ne varie pas. En extrayant une base du sous-espace vectoriel de $\mathcal{H}$ engendré par l'ensemble des mesures effectuées, on peut déterminer les directions selon lesquelles on a de l'information. Dans le cas du qubit, la mesure selon un axe, si elle est suffisamment répétée, permet de confiner les candidats pour la reconstruction à un plan dans la boule de Bloch. Si on ajoute la mesure selon un autre axe, on les confine sur un axe. Enfin, la mesure selon trois axes permet de situer précisément la matrice densité.

c) Calcul du gradient et détermination de l'estimateur

La fonction de log-vraisemblance est une fonction de d^2 variables réelles à valeurs dans $\mathbb{R}$. Son gradient $\nabla \mathcal{L}$ est le vecteur de $\mathcal{O}(\mathcal{H})$ vérifiant

$$\mathcal{L}(\hat{\varrho} + \delta \hat{h}) \simeq \mathcal{L}(\hat{\varrho}) + \langle \nabla \mathcal{L}(\hat{\varrho}), \delta \hat{h} \rangle, \quad (\text{III.3.45})$$

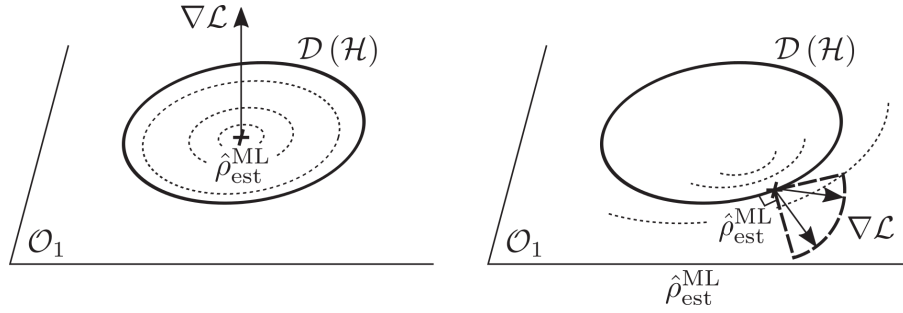


Fig. III.10 Gradient de la vraisemblance en $\hat{\rho}_{ML}$ dans deux situations : à gauche, la matrice densité est à l'intérieur de $\mathcal{D}(\mathcal{H})$. Le gradient est orthogonal à l'hyperplan des opérateurs de trace 1, donc proportionnel à l'identité. À droite, $\hat{\rho}_{ML}$ est sur le bord de $\mathcal{D}(\mathcal{H})$. Le gradient peut donc pointer dans une direction arbitraire vers l'extérieur. Il y a donc moins de conditions sur ses composantes, ce qui se apparaît sur le nombre d'égalités imposées par la condition III.3.48b.

où $\delta \hat{h}$ est un petit vecteur de $\mathcal{O}(\mathcal{H})$ et où le produit scalaire est le produit de Frobenius. Calculons :

$$\mathcal{L}(\hat{\rho} + \delta \hat{h}) = \sum_{r=1}^N \log \langle \hat{\rho} + \delta \hat{h}, \hat{E}^{(r)} \rangle, \quad (\text{III.3.46a})$$

$$= \sum_{r=1}^N \log \langle \hat{\rho}, \hat{E}^{(r)} \rangle + \sum_{r=1}^N \log \left(1 + \frac{\langle \delta \hat{h}, \hat{E}^{(r)} \rangle}{\langle \hat{\rho}, \hat{E}^{(r)} \rangle} \right), \quad (\text{III.3.46b})$$

$$\simeq \mathcal{L}(\hat{\rho}) + \sum_{r=1}^N \frac{\langle \delta \hat{h}, \hat{E}^{(r)} \rangle}{\langle \hat{\rho}, \hat{E}^{(r)} \rangle}. \quad (\text{III.3.46c})$$

En identifiant avec III.3.45, il vient :

$$\nabla \mathcal{L}(\hat{\rho}) = \sum_{r=1}^N \frac{\hat{E}^{(r)}}{\text{Tr}(\hat{\rho} \hat{E}^{(r)})}. \quad (\text{III.3.47})$$

Le gradient donne la direction dans laquelle la vraisemblance varie le plus. Comme la vraisemblance est maximale en $\hat{\rho}_{ML}$, il ne peut pointer que dans une direction extérieure à la variété des matrices densité, comme illustré sur la figure III.10.

Le critère d'optimalité pour $\hat{\rho}_{ML}$ est donné par [90] :

$$[\hat{\rho}_{ML}, \nabla \mathcal{L}_{ML}] = 0, \quad (\text{III.3.48a})$$

$$\lambda_{ML} P_{ML} \leq \nabla \mathcal{L}_{ML} \leq \lambda_{ML} \mathbb{1}, \quad (\text{III.3.48b})$$

où $\nabla \mathcal{L}_{ML} \equiv \nabla \mathcal{L}(\hat{\rho}_{ML})$, λ_{ML} est un entier positif et P_{ML} est le projecteur sur l'image de $\hat{\rho}_{ML}$. Dans III.3.48b, les inégalités entre opérateurs signifient que la différence des deux opérateurs est positive. Si on l'exprime dans une base $\mathcal{B} = \{k_1, \dots, k_d\}$ où P_{ML} est diagonal :

$$P_{ML} = \text{diag}(1, 1, \dots, \frac{1}{i_0}, 0, \dots, 0), \quad (\text{III.3.49})$$

ce critère prend une forme plus simple :

$$(\nabla \mathcal{L}_{\text{ML}})_i = \lambda_{\text{ML}}, \quad \text{si } i \leq i_0, \quad (\text{III.3.50a})$$

$$0 \leq (\nabla \mathcal{L}_{\text{ML}})_i \leq \lambda_{\text{ML}}, \quad \text{si } i \geq i_0. \quad (\text{III.3.50b})$$

Si P_{ML} n'est pas de rang plein, le nombre de contraintes est donc réduit, ce qui est illustré sur la figure III.10.

Dans le cas où on a effectué suffisamment de mesures indépendantes et où la matrice est de rang plein, le gradient au maximum de vraisemblance est proportionnel à l'identité. Considérons qu'on a effectué des mesures décrites par un ensemble de POVM ($\hat{E}_{i_k}^{(k)}$). Dans ce cas, les opérateurs de mesure associés à l'ensemble des POVM décomposent l'identité. La formule III.3.47 montre que le gradient est la somme de termes proportionnels aux opérateurs de mesure. Pour respecter la condition III.3.48b, il faut que le nombre de fois que chaque opérateur $\hat{E}_{i_d}^{(d)}$ apparaît dans la somme du gradient soit proportionnel à $\text{Tr}(\hat{\rho}_{\text{ML}} \hat{E}_{i_d}^{(d)})$. Dit autrement, la fréquence d'obtention d'un résultat de mesure est égale à sa probabilité pour l'état $\hat{\rho}_{\text{ML}}$. La matrice $\hat{\rho}_{\text{ML}}$ peut donc être vue comme une moyenne des matrices d'effet, pondérées par l'inverse des probabilités des trajectoires de mesures. On peut généraliser cette interprétation à des ensembles de mesures non complets [91].

d) Algorithme de montée de gradient

La fonction $\mathcal{L}$ est une fonction concave de plusieurs variables qu'on cherche à maximiser sur l'ensemble $\mathcal{D}(\mathcal{H})$ qui est convexe. Une manière de trouver l'unique maximum $\hat{\rho}^*$ est d'utiliser une technique de montée de gradient³. Le principe est illustré sur la figure III.11. Pour aller au sommet de $\mathcal{L}$, il suffit de toujours aller dans la direction qui monte le plus. Cette direction est donnée par le gradient et cette méthode est assurée de converger vers $\hat{\rho}^*$ par convexité.

En pratique, l'algorithme est le suivant :

- on part d'une matrice densité arbitraire $\hat{\rho}_0$, en pratique l'identité ;
- on fait un pas proportionnel au gradient, suffisamment petit pour que $\mathcal{L}$ augmente ;
- on projette l'opérateur obtenu sur l'ensemble des matrices densité.

On répète ces étapes jusqu'à ce que certaines conditions d'arrêt soient réalisées.

Choix de la longueur du pas

La fonction $\mathcal{L}$ est une fonction de plusieurs variables à valeurs dans $\mathbb{R}^-$. Ses variations autour d'un point $\hat{\rho}_0$ sont données par

$$\mathcal{L}(\hat{\rho}_0 + \hat{h}) = \mathcal{L}(\hat{\rho}_0) + \text{Tr}(\nabla \mathcal{L}(\hat{\rho}_0) \hat{h}) + \frac{1}{2} \hat{h} \cdot \nabla^2 \mathcal{L}(\hat{\rho}_0) \cdot \hat{h} + O(\|\hat{h}\|^3), \quad (\text{III.3.51})$$

où $\hat{h} \in \mathcal{O}(\mathcal{H})$. Le hessien $\nabla^2 \mathcal{L}(\hat{\rho}_0)$ est un tenseur d'ordre 2 sur $\mathcal{O}(\mathcal{H})$, de sorte que $\hat{h} \cdot \nabla^2 \mathcal{L}(\hat{\rho}_0) \cdot \hat{h}$ est un scalaire, qui est négatif puisque $\mathcal{L}$ est une fonction concave. On montrera par la suite comment on peut calculer son effet et on verra qu'il joue un rôle

3. En général, les problèmes de maximisation de fonctions concaves sont ramenés à des minimisations de fonctions convexes, ce qui est équivalent et on parle alors de descente de gradient, mais pour garder à l'esprit qu'on cherche à maximiser la log-vraisemblance, on ne posera pas ici de problème équivalent.

important pour quantifier l'incertitude sur nos résultats de reconstruction. Supposons qu'on se place en $\hat{\rho}_0$ et qu'on fasse un petit pas $\hat{h} = \epsilon \nabla \mathcal{L}(\hat{\rho}_0)$ dans la direction du gradient. On a alors :

$$\mathcal{L}(\hat{\rho}_0 + \epsilon \hat{h}) = \mathcal{L}(\hat{\rho}_0) + \epsilon \text{Tr}(\nabla \mathcal{L}(\hat{\rho}_0)^2) + \frac{\epsilon^2}{2} \nabla \mathcal{L}(\hat{\rho}_0) \cdot \nabla^2 \mathcal{L}(\hat{\rho}_0) \cdot \nabla \mathcal{L}(\hat{\rho}_0) + O(\epsilon^3). \quad (\text{III.3.52})$$

À l'ordre 2 en ϵ , la condition pour que la log-vraisemblance augmente est alors :

$$\epsilon \leq - \frac{2 \text{Tr}(\nabla \mathcal{L}(\hat{\rho}_0)^2)}{\nabla \mathcal{L}(\hat{\rho}_0) \cdot \nabla^2 \mathcal{L}(\hat{\rho}_0) \cdot \nabla \mathcal{L}(\hat{\rho}_0)} \equiv \epsilon_m. \quad (\text{III.3.53})$$

En pratique, on choisira $\epsilon = \epsilon_m/2$, qui donne le sommet de la parabole correspondant au développement à l'ordre deux de la vraisemblance. Le principe de la méthode est résumé sur la figure III.11.

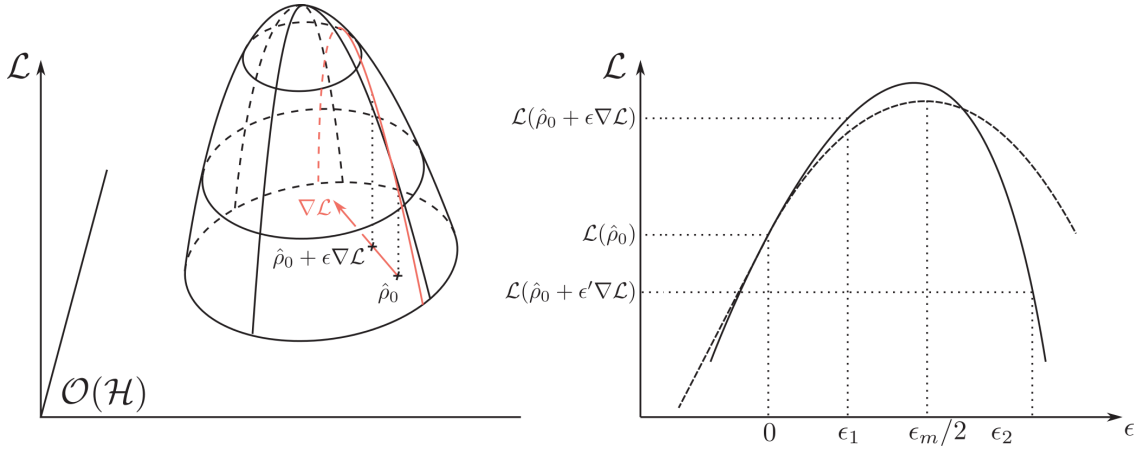


Fig. III.11 Illustration de la montée de gradient. À gauche, log-vraisemblance en fonction de l'opérateur $\hat{\rho}$. On représente schématiquement l'ensemble des opérateurs par un plan. Lorsqu'on se déplace selon le gradient, $\mathcal{L}$ augmente jusqu'à une certaine valeur puis décroît. L'ensemble des valeurs de $\mathcal{L}$ accessibles en faisant varier l'amplitude ϵ du pas est représenté en rouge. À droite, vraisemblance en fonction de ϵ . Le pas ϵ_1 est suffisamment petit et $\mathcal{L}$ augmente pendant ce pas. En revanche ϵ_2 est trop grand et la vraisemblance décroît. Les pointillés correspondent à l'interpolation quadratique de la vraisemblance qui donne le critère III.3.52.

Projection sur $\mathcal{D}(\mathcal{H})$

La méthode présentée précédemment permet de trouver le maximum de $\mathcal{L}$ sur $\mathcal{O}(\mathcal{H})$ mais on doit aussi tenir compte des contraintes sur la matrice densité. En effet, l'opérateur $\hat{\sigma} \equiv \hat{\rho}_0 + \epsilon \nabla \mathcal{L}$ n'est pas en général un opérateur densité. En effet, l'opérateur gradient est un opérateur positif et en général non nul. La trace de $\hat{\sigma}$ n'est donc pas nécessairement

égale à un. Par ailleurs, lorsque $\hat{\rho}_0$ est proche du bord de $\mathcal{D}(\mathcal{H})$, il se peut que $\hat{\sigma}$ ne soit pas positif. On projette donc $\hat{\sigma}$ sur $\mathcal{D}(\mathcal{H})$. Le détail de la façon dont on réalise cette projection est donné dans [90].

e) Interprétation de la méthode : cas du qubit

Afin d'illustrer quelques aspects de notre méthode de reconstruction, considérons le cas d'un qubit dont on veut réaliser la tomographie. Commençons par voir comment on peut calculer $\text{Tr}(\varrho \hat{E})$, où $\hat{E}$ est une matrice d'effet quelconque. Comme $\hat{E}$ est positive et hermitienne⁴, $\hat{E}/\text{Tr}(\hat{E})$ est une matrice densité et on peut écrire :

$$\hat{E} = \text{Tr}(\hat{E}) \frac{\mathbb{1} + \mathbf{e} \cdot \hat{\sigma}}{2}. \quad (\text{III.3.54})$$

On a donc :

$$\text{Tr}(\varrho \hat{E}) = \text{Tr}(\hat{E}) \frac{(1 + \mathbf{r} \cdot \mathbf{e})}{2}. \quad (\text{III.3.55})$$

Ce résultat s'interprète facilement. La probabilité d'obtenir $\hat{E}$ est maximale si $\mathbf{r}$ et $\mathbf{e}$ sont colinéaires et de même sens. Elle est nulle si les deux vecteurs de Bloch sont opposés et de norme 1, c'est à dire si $\hat{\varrho}$ et $\hat{E}$ sont des états purs orthogonaux.

À $\text{Tr}(\hat{E})$ fixée, cette probabilité varie d'autant plus que la norme de $\mathbf{e}$ est importante. Ceci montre que les mesures les plus informatives sont celles dont les matrices d'effet ont un vecteur de Bloch unité, c'est-à-dire des mesures projectives.

Mesures successives

Supposons qu'on effectue sur le qubit une mesure projective, par exemple selon Oz . Après celle-ci, quelque soit l'état initial, le qubit est dans un des deux états propres de la mesure. Il ne reste plus aucune information sur l'état initial. Effectuer une seconde mesure n'apporte donc pas plus d'information sur l'état. Si par exemple on a obtenu le résultat de mesure $+_z$, la première mesure est décrite par $M_1 = |+_z\rangle \langle +_z|$. Notons M_2 l'opérateur de Kraus décrivant la mesure suivante. La matrice d'effet est :

$$\hat{E} = M_1^\dagger M_2^\dagger M_2 M_1 = p_2 |+_z\rangle \langle +_z|, \quad (\text{III.3.56})$$

où $p_2 = \langle +_z | M_2^\dagger M_2 | +_z \rangle$ est la probabilité de la deuxième mesure, indépendante de $\hat{\rho}$. E est proportionnel à la matrice d'effet obtenue avec la première mesure seulement. Multiplier une matrice d'effet par une constante ne change rien à la reconstruction, puisque le seul effet est d'ajouter un terme constant dans $\mathcal{L}$. Imaginons que la deuxième mesure soit une mesure projective selon Ox . Dans ce cas, on a une chance sur deux d'obtenir le résultat $+_x$ et une chance sur deux d'obtenir $-_x$. La probabilité d'obtenir l'un ou l'autre est donc donnée par la moitié de la probabilité d'obtenir $+_z$ initialement, mais ceci est indépendant de l'état. On voit ici que la valeur absolue de la log-vraisemblance n'a aucune importance. Elle peut être très négative si on a une mesure qui a beaucoup de résultats possibles. Seules les variations de $\mathcal{L}$ jouent un rôle dans la reconstruction.

4. On ne tient pas compte du cas où $\hat{E}$ est nulle, qui correspond à une trajectoire de mesure qui ne peut pas arriver.

Considérons maintenant le cas où les mesures ne sont pas projectives. On pose :

$$M_1 = \sqrt{\frac{1+z_m}{2}} |+_z\rangle \langle+_z| + \sqrt{\frac{1-z_m}{2}} |-_z\rangle \langle-_z|, \quad (\text{III.3.57a})$$

$$M_2 = \sqrt{\frac{1+x_m}{2}} |+_x\rangle \langle+_x| + \sqrt{\frac{1-x_m}{2}} |-_x\rangle \langle-_x|. \quad (\text{III.3.57b})$$

On obtient alors

$$\hat{E} = M_1^\dagger M_2^\dagger M_2 M_1, \quad (\text{III.3.58a})$$

$$= \frac{1}{2} \frac{\mathbb{1} + z_m \hat{\sigma}_z + \sqrt{1-z_m^2} x_m \hat{\sigma}_x}{2}. \quad (\text{III.3.58b})$$

Ce résultat s'interprète très simplement. Le mesure selon Ox est d'autant plus brouillée que la mesure selon Oz est projective. On remarque cependant que la longueur du vecteur de Bloch décrivant la mesure est toujours supérieure ou égale à celle obtenue pour la première mesure seule, avec égalité si et seulement si $z_m = \pm 1$, *i. e.* la première mesure est projective. Il est donc presque toujours intéressant d'ajouter des mesures supplémentaires.

Dans le cas où les deux mesures ont lieu selon Oz , on obtient

$$\hat{E} = \frac{1+z_m^2}{4} \left(\mathbb{1} + \frac{2z_m}{1+z_m^2} \hat{\sigma}_z \right) \quad (\text{III.3.59})$$

et on trouve aussi que la norme du vecteur de Bloch décrivant la matrice d'effet augmente.

III.3.3 Prise en compte des imperfections et pertinence des résultats

a) Les résultats sont-ils physiques ?

Comme on l'a vu dans la section III.2.2, l'incertitude statistique de la mesure fait que les fréquences d'occurrence calculées à partir d'un jeu de données ne correspondent pas nécessairement à des probabilités de mesure qui décrivent un état physique. Dans ce cas, la reconstruction par inversion échoue. La vraisemblance est alors maximale en dehors de l'ensemble des matrices densité. La maximisation de vraisemblance donne un état physique, qui a la particularité d'être situé sur le bord de l'ensemble des matrices densité, ce qui signifie qu'il a un rang non maximal.

Ce n'est pas la seule situation dans laquelle les résultats sont susceptibles de ne pas correspondre à un état physique. Considérons le cas d'une mesure imparfaite, décrite par un POVM contenant deux mesures associées aux vecteurs de Bloch $\vec{r}_+ = (e_x, 0, 0)$ et $\vec{r}_- = (-e_x, 0, 0)$. $e_x \in [0, 1]$ est un paramètre permettant de prendre en compte une diminution du contraste de la mesure selon Ox . On note n_+ (resp. n_-) le nombre de résultats $\vec{r}_+$ (resp. $\vec{r}_-$) obtenus. La vraisemblance est alors donnée par :

$$\mathcal{L}(x) = n_+ \log \left(\frac{1+e_x x}{2} \right) + n_- \log \left(\frac{1-e_x x}{2} \right) \quad (\text{III.3.60})$$

et elle atteint son maximum pour $e_x x = (n_+ - n_-)/(n_+ + n_-)$. Dans le cas idéal, elle atteint son maximum pour la valeur de x donnant une probabilité de mesure de $+x$ égale à la valeur mesurée. Lorsque le contraste de la mesure est diminué, le maximum est atteint

pour une valeur renormalisée par e_x . Ce maximum peut alors être atteint pour une valeur de $|x|$ supérieure à 1. Ceci est particulièrement susceptible d'arriver si on sous-estime le contraste. Il est donc nécessaire de bien calibrer les paramètres décrivant la mesure.

La question de savoir si il est justifié de donner comme résultat d'une reconstruction un état de rang non maximal est à l'origine d'un débat [92]. En effet, lorsqu'un état est de rang non maximal, on peut construire des mesures dont certains résultats ont une probabilité nulle d'advenir. Il a été avancé que faire une telle assertion sur la base d'un ensemble fini de mesures n'est pas légitime, si l'on considère la matrice densité comme une prédiction concernant les observations futures. Cependant, dans notre cas, s'interdire de considérer des matrices de rang déficient pose aussi problème. En effet, pour avoir des matrices de rang plein, il faut assigner à toutes les populations une valeur non nulle. Une manière de procéder, proposée dans [92] consiste à donner à chaque population une valeur $p_{min} \simeq 1/(N + d)$ où N est le nombre de mesures effectuées et d le nombre de résultats de mesure possibles.

Ceci pose un problème dans notre cas puisque la dimension de l'espace de Hilbert est en principe infinie. On doit donc en pratique tronquer l'espace de Hilbert. À cause de la condition sur la trace de la matrice densité, les populations dans les états de poids important dépendraient alors de la taille de l'espace tronqué choisi. Une manière de circonvenir ce problème serait de donner une probabilité minimum qui tend asymptotiquement vers 0 lorsque le nombre de photons tend vers l'infini, mais une telle méthode n'aurait pas d'interprétation opérationnelle en terme de probabilité, en plus d'être arbitraire.

On acceptera donc d'obtenir des états purs comme résultats de la reconstruction, en sachant que ceux-ci fournissent un résumé de l'ensemble des résultats obtenus et non un engagement sur des résultats futurs. Ceci sera suffisant dans de nombreux cas, compte tenu du fait qu'on prépare des états presque purs, dans lesquels un état qui nous intéresse est associé à une population proche de 1 et pollué par des états faiblement peuplés qu'on ne souhaite pas connaître précisément. Dans la sous-section suivante, on va voir qu'on peut simplifier grandement la reconstruction dans ce cadre, ce qui présente un avantage certain pour de nombreuses applications en information quantique.

b) Le cas des états purs ou presque

Considérons un système dans un état pur $|\psi\rangle$. Si on peut faire une mesure ayant une valeur propre non dégénérée associée à $|\psi\rangle$, on obtiendra un résultat certain. En moyennant suffisamment, on peut alors se convaincre que l'état est $|\psi\rangle$. On est alors capable de reconstruire l'état à l'aide d'un seul opérateur de mesure. Par exemple, dans le cas du qubit, si on effectue 100 mesures selon Oz et qu'on obtient toujours le résultat $+$, z est vraisemblablement très proche 1. Les contraintes sur l'ensemble des matrices font que dans ce cas il n'est pas nécessaire de mesurer dans les autres directions. En effet, on a $x \leq 1 - z^2$ et $y \leq 1 - z^2$.

Cet effet est aussi présent en plus grande dimension lorsqu'on est au bord de l'ensemble des matrices densité, qui est alors constitué des matrices de rang non maximal. Ceci a pour conséquence de réduire le nombre d'opérateurs de mesure nécessaires. Cette propriété est utilisée dans le cas du *compressed sensing* pour simplifier la mesure d'états de rang non maximal.

III.4 Incertitudes de la tomographie

Dans cette partie, on s'intéresse à la question de la pertinence des résultats de la reconstruction et à la confiance qu'on peut leur attribuer. Celle-ci est, en toute rigueur, donnée par des volumes de confiance dans l'espace des matrices densité. Cependant, la détermination et l'interprétation de ces volumes est un problème compliqué pour des systèmes de grande dimension. On montre qu'on peut calculer des barres d'erreur sur la variance de la valeur moyenne d'un opérateur lorsqu'on répète la reconstruction. Ces barres d'erreur permettent en particulier de déterminer le rôle de l'incertitude statistique sur chaque terme de la matrice densité reconstruite et constitue, dans les bonnes conditions, un indicateur de la fiabilité de la mesure.

III.4.1 Surfaces d'isovraisemblance et volumes de confiance

Lorsqu'on construit un estimateur pour un paramètre qu'on veut déterminer expérimentalement, on souhaite généralement disposer d'un intervalle de confiance, dans lequel la valeur vraie se situe avec une probabilité donnée. Dans notre cas, on veut estimer $d^2 - 1$ paramètres et on construit un estimateur qui maximise une fonction de vraisemblance faisant intervenir ces paramètres. Plus la fonction de vraisemblance est piquée, plus l'ambiguïté sur l'estimateur sera faible. Une manière de traduire la sélectivité de notre reconstruction est alors de regarder le volume compris dans une surface d'isovraisemblance. Une autre approche consiste à construire des volumes de confiance [93–95].

Intervalles et régions de confiance

On considère un problème d'estimation, tel que décrit dans le cadre de la section III.1. Dans cette section, on a vu qu'on pouvait construire un estimateur ponctuel θ_{est} pour le paramètre θ . En général, on souhaite avoir une information sur la pertinence de cet estimateur, qui nous dise dans quelle mesure cet estimateur est proche de θ . On pourrait vouloir donner un intervalle dans lequel il est probable de trouver θ . Cependant, cela n'a pas de sens de parler de la probabilité que θ soit dans l'intervalle $[\theta_1, \theta_2]$, puisque θ n'est pas une variable aléatoire. On peut en revanche construire un intervalle aléatoire $I(X) \equiv [\theta_1(X), \theta_2(X)]$ tel que la probabilité que $I(X)$ ne recouvre pas θ soit inférieure ou égale à un seuil d'erreur α . À titre d'exemple, la figure B.1 montre l'intervalle de confiance de Clopper-Pearson pour une loi binomiale [96], dont la construction est donnée dans l'annexe B. Cet intervalle est construit de telle sorte que la probabilité qu'il recouvre la valeur réelle est égale à 0,95.

Dans le cas de l'estimation d'une matrice densité, on ne cherche pas un intervalle de confiance, mais une région de confiance, c'est-à-dire un hypervolume de $\mathcal{D}(\mathcal{H})$ tel que la probabilité que l'hypervolume associé à une estimation recouvre $\hat{\rho}$ soit supérieure à $1 - \alpha$. Remarquons que cette définition ne caractérise pas une région de confiance unique. Il y a une infinité de choix possibles. Une région qui la vérifie est dite *conservative*. Elle est dite optimale si elle a le plus petit volume parmi les régions conservatives.

La construction de régions de confiance présente deux défis qui les rendent peu utilisées en pratique. Le premier défi est de donner une méthode efficace et applicable aux données expérimentales pour construire des régions de confiance optimales. Il a été montré que ce

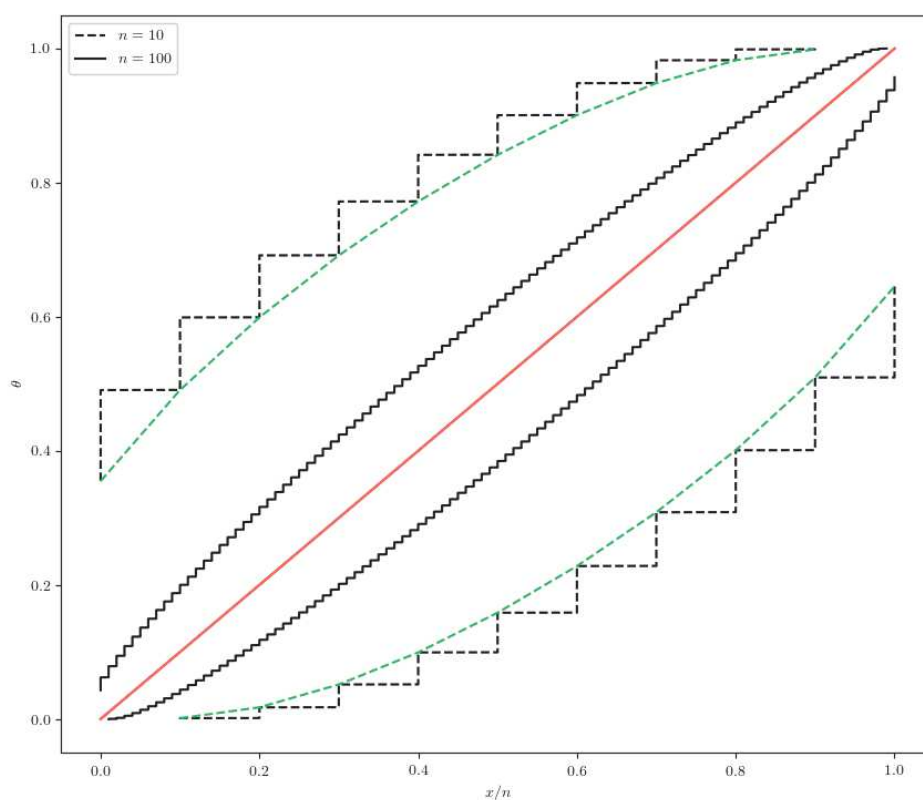


Fig. III.12 Bornes de Clopper-Pearson sur l'estimation du paramètre de deux loi binomiales avec $n = 10$ (en pointillé) et 100 (trait plein) tirages. En abscisse on trace la valeur x/n du rapport entre le nombre de résultats positifs et le nombre de tirages. En ordonnée, on trace la valeur du paramètre θ de la loi binomiale. Pour un résultat x/n , la ligne rouge correspond à la valeur θ_{est} estimée tandis que l'intervalle de confiance est l'intervalle compris entre les deux courbes noires. Les courbes vertes servent à guider l'œil sur cet intervalle.

problème est NP-difficile [97]. La construction de régions de confiance fait donc intervenir un compromis entre le coût algorithmique de la méthode et le resserrement de la région de confiance autour de la valeur estimée.

Le deuxième défi est simplement de visualiser l'objet mathématique obtenu. En effet, il est très compliqué de représenter de manière pratique un hypervolume ne présentant pas de symétrie dès lors que la dimension est supérieure à 3. En pratique, on peut représenter les régions de confiance dans le cas d'un qubit en traçant des portions de la sphère de Bloch mais le problème devient très compliqué dès lors qu'on passe à des espaces de dimensions supérieures. Plusieurs méthodes de construction de régions de confiance pour la tomographie quantique ont été proposées [94] [95].

On va présenter une méthode qui permet de déterminer les régions de confiance sous la forme d'un polytope dont les facettes sont données par les éléments d'un POVM [93].

Polytope

Supposons qu'on veuille reconstruire une matrice $\hat{\rho}$ à l'aide d'un ensemble de mesures effectuées à l'aide d'un POVM $(\hat{E}_i)_{i \in I}$. On note $\vec{n} = (n_i)_{i \in I}$ les résultats de mesure correspondants. Pour une matrice densité $\hat{\rho}$, la probabilité d'obtenir le résultat i s'écrit $p_i = \text{Tr}(\hat{E}_i \hat{\rho}) = \langle \hat{E}_i, \hat{\rho} \rangle$, où le produit scalaire est le produit de Frobenius III.2.16. À une valeur donnée de p_i correspond donc un plan dans $\mathcal{O}(\mathcal{H})$.

L'idée de la construction de [93] est d'associer à chaque élément du POVM un intervalle de confiance en considérant la mesure i comme une loi binomiale de paramètre p_i , qu'on estime en lui donnant une borne supérieur de Clopper-Pearson, associée à une probabilité d'erreur α_i . De cette manière, chaque mesure définit un hyperplan d'équation

$$\text{Tr}(\hat{E}_i \hat{\rho}) \leq p_i, \quad (\text{III.4.1})$$

qui coupe $\mathcal{D}(\mathcal{H})$ en deux régions dont l'une est un intervalle de confiance pour l'estimation de p_i . Si le POVM est à information complète, l'intersection de toutes ces régions forme un polytope de $\mathcal{D}(\mathcal{H})$ dont l'intersection avec $\mathcal{D}(\mathcal{H})$ est une région de confiance avec une probabilité d'erreur $\alpha \equiv \sum_i \alpha_i$.

Afin de simplifier le calcul des régions de confiance ainsi obtenues, [93] montre qu'elles sont contenues dans l'ensemble $\Gamma(\vec{n}) \equiv \cap_i \Gamma_i(n_i)$ où

$$\Gamma_i(n_i) = \left\{ \hat{\rho} \in \mathcal{D}(\mathcal{H}) \mid \text{Tr}(\hat{E}_i \hat{\rho}) \leq \frac{n_i}{n} + \delta(n_i, \alpha_i) \right\}, \quad (\text{III.4.2})$$

avec $\delta(n_i, \alpha_i)$ le réel positif vérifiant

$$D\left(\frac{n_i}{n} \parallel \frac{n_i}{n} + \delta(n_i, \alpha_i)\right) = -\frac{1}{n} \log(\alpha_i), \quad (\text{III.4.3})$$

où $D(x \parallel y)$ est la distance de Kullback-Leibler entre x et y .

Afin d'illustrer la méthode, considérons le cas de l'estimation de l'état d'un qubit par des mesures idéales selon Ox , Oy et Oz . Dans ce cas, les facettes du polytope sont des plans de x , y ou z constant. Considérons qu'on a réalisé n mesures de $\hat{\sigma}_z$ et qu'on

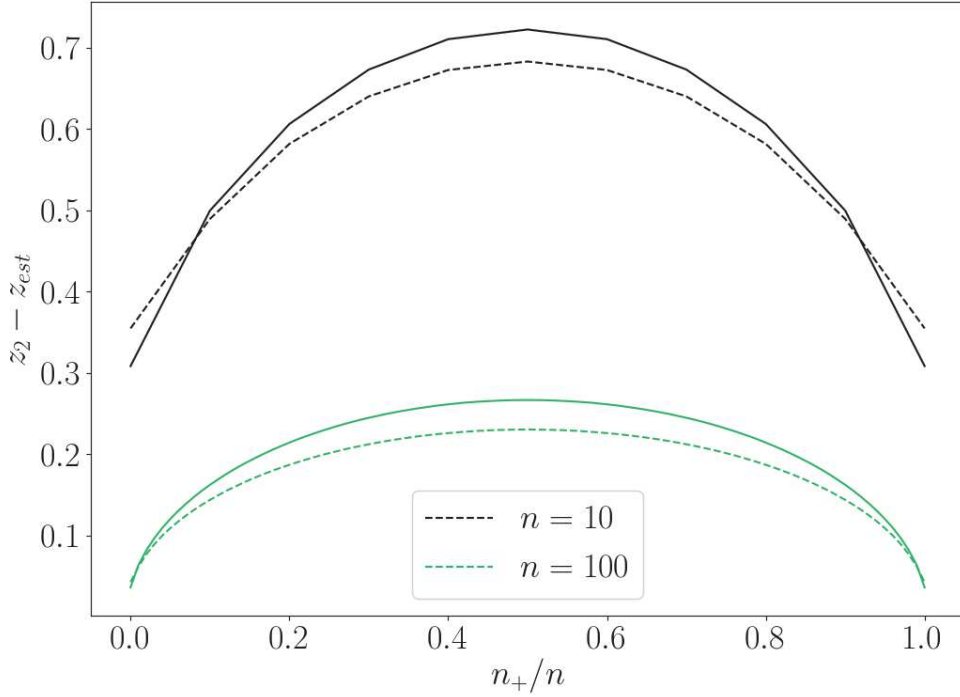


Fig. III.13 Largeur de l'intervalle de confiance en fonction du résultat pour l'intervalle de Clopper-Pearson (courbes en pointillé) et pour la formule III.4.2 (courbes pleines)

a obtenu n_+ résultats positifs. On choisit de répartir équitablement α entre toutes les mesures ($\alpha_i = \alpha/6$). Alors III.4.2 donne :

$$z \leq \frac{2n_+ - n}{n} + 2\delta(n_+, \alpha/6), \quad (\text{III.4.4a})$$

$$z \geq \frac{2n_+ - n}{n} - 2\delta(n - n_+, \alpha/6). \quad (\text{III.4.4b})$$

La figure III.13 montre la largeur de l'intervalle de confiance pour deux valeurs du nombre de tirages en fonction du nombre de résultats positifs. Les courbes en pointillé sont obtenues pour l'intervalle de Clopper-Pearson, tandis que celles en traits pleins sont obtenues à l'aide de la formule simplifiée III.4.2, qui a tendance à donner un intervalle légèrement plus grand. L'intervalle de confiance est bien plus étroit lorsque la fréquence mesurée est proche de 0 ou 1. On retrouve l'idée qu'il est plus facile d'estimer le paramètre d'une loi binomiale très biaisée. La figure III.14 montre la décroissance de la largeur des intervalles de confiance avec le nombre de tirages effectués.

Cette méthode présente plusieurs avantages. Tout d'abord, les régions de confiance construites sont plus petites que celles obtenues par d'autres méthodes proposées auparavant, notamment [94]. Ensuite, les polytopes sont plus faciles à visualiser et à manipuler que des formes arbitraires en grande dimension. En grande dimension, les auteurs

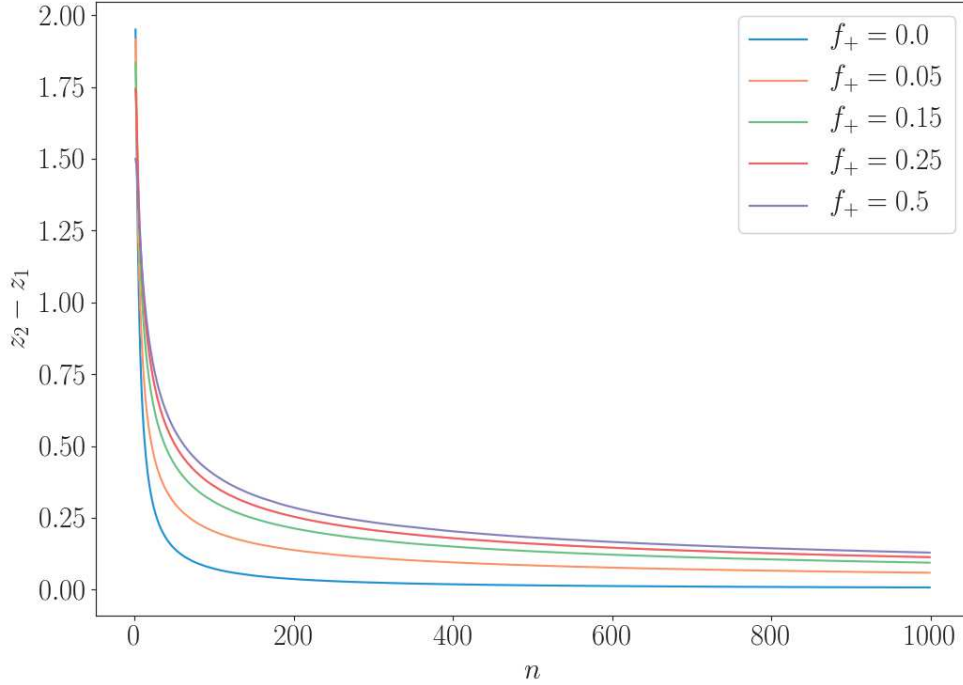


Fig. III.14 Largeur de l'intervalle de confiance en fonction du nombre de mesures, pour plusieurs fréquences d'occurrence du résultat positif

proposent d'extraire l'information contenue dans le polytope en donnant des figures de mérites, comme par exemple la fidélité par rapport à un état, dont on peut déterminer l'intervalle de confiance en trouvant son maximum et son minimum sur le polytope. Cette technique nécessite cependant d'échantillonner le polytope en tirant des états au hasard à l'intérieur, ce qui pose à nouveau des questions sur la manière de définir une mesure sur l'ensemble des matrices densité.

Une des difficultés de cette méthode est qu'elle privilégie les directions dans lesquelles on a mesuré, même quand ce n'est pas justifié. Par exemple, si on a mesuré un qubit n fois dans chaque direction et qu'on a obtenu à chaque fois exactement la moitié de résultats positifs et de résultats négatifs, la largeur de la région de confiance dans la direction $\vec{u} = (1, 1, 1) / \sqrt{3}$ est $\sqrt{3}$ fois plus grande que dans les directions Ox , Oy et Oz , alors qu'on s'attendrait à avoir la même information dans toutes les directions. Une manière d'affiner la région de confiance est de construire des facettes d'ordre supérieur en combinant des mesures, mais il y a alors un certain arbitraire dans le choix des combinaisons et cette construction demande beaucoup d'opérations en grande dimension.

Une deuxième limite de cette méthode est qu'elle nécessite d'avoir un grand nombre de résultats pour chaque base de mesure. Si on dispose de nombreuses mesures dans des bases toutes différentes, la borne de Clopper-Pearson pour chacune de ces mesures est assez lâche et on a donc peu d'information. Par exemple, on voit sur la figure III.14 que pour dix mesures idéales d'un qubit dans la même direction, il faut au moins 9 résultats

dans la même direction pour pouvoir exclure le centre de la sphère de Bloch de la région de confiance.

III.4.2 Information de Fisher et calcul des barres d'erreur

Les outils présentés précédemment pour caractériser l'incertitude sur l'estimation de l'état quantique présentent plusieurs difficultés en pratique. Tout d'abord, ils font intervenir des volumes dans des espaces de grande dimension, qui sont difficiles à représenter. Ensuite, ils nécessitent parfois des calculs coûteux pour obtenir des incertitudes dérivées sur des grandeurs d'intérêt, comme la fidélité, un terme spécifique de la matrice densité ou la valeur moyenne d'une observable, à cause de la nécessité d'intégrer sur tout le volume. Enfin, ils font parfois appel à des choix arbitraires sur une distribution de probabilité *a priori* ou sur la façon de prendre en compte les mesures.

Dans la suite, nous allons développer une approche permettant de caractériser l'incertitude sur notre méthode de tomographie. Le principe est d'appliquer le concept d'information de Fisher introduite dans la section III.1.2.a à notre reconstruction. Cette méthode permet de donner les incertitudes sur les termes de la matrice densité reconstruite $\hat{\rho}_{\text{est}}$, liées au caractère aléatoire des mesures utilisées pour la tomographie. Elle permet aussi de donner des incertitudes sur les valeurs moyennes d'observables calculées à partir de $\hat{\rho}_{\text{est}}$.

L'avantage de cette approche est qu'elle permet, comme notre méthode de reconstruction, de prendre en compte des mesures toutes différentes, sans faire l'hypothèse qu'on a suffisamment moyenné chaque mesure. Nous verrons aussi qu'elle permet de mettre en évidence l'absence d'information sur certains termes de la matrice densité, ce qui en fait un outil particulièrement approprié pour effectuer des tomographies partielles.

Calcul de la hessienne et information de Fisher

Dans la section III.1, on a vu que dans le cadre de l'estimation d'un paramètre, on pouvait quantifier l'information acquise sur ce paramètre lors d'une mesure à l'aide de l'information de Fisher. On a aussi vu que cette information donne une limite inférieure à la variance qu'on peut atteindre, la borne de Cramer-Rao. On a vu que l'information de Fisher peut s'exprimer comme l'espérance de la dérivée seconde de la vraisemblance par rapport au paramètre qu'on souhaite estimer.

Dans notre cas, on ne cherche plus à estimer un paramètre, mais un ensemble de paramètres. On a déjà introduit dans la section III.3.2.d le hessien de la vraisemblance $\nabla^2 \mathcal{L}$, tel que les variations de $\mathcal{L}$ autour de $\hat{\rho}_0$ s'écrivent :

$$\mathcal{L}(\hat{\rho} + \hat{h}) = \mathcal{L}(\hat{\rho}) + \text{Tr}(\nabla \mathcal{L}(\hat{\rho}) \hat{h}) + \frac{1}{2} \hat{h} \cdot \nabla^2 \mathcal{L}(\hat{\rho}) \cdot \hat{h} + O(\|\hat{h}\|^3). \quad (\text{III.4.5})$$

On voit donc que cet opérateur joue un rôle analogue à la courbure pour une fonction à plusieurs dimensions. On peut en donner une expression explicite en utilisant l'expression III.3.44c pour la vraisemblance. Pour tous opérateurs A et B , on a :

$$A \cdot (\nabla^2 \mathcal{L}(\hat{\rho})) \cdot B = - \sum_{r=1}^N \frac{\text{Tr}(A \hat{E}^{(r)}) \text{Tr}(B \hat{E}^{(r)})}{\text{Tr}(\hat{\rho} \hat{E}^{(r)})^2}. \quad (\text{III.4.6})$$

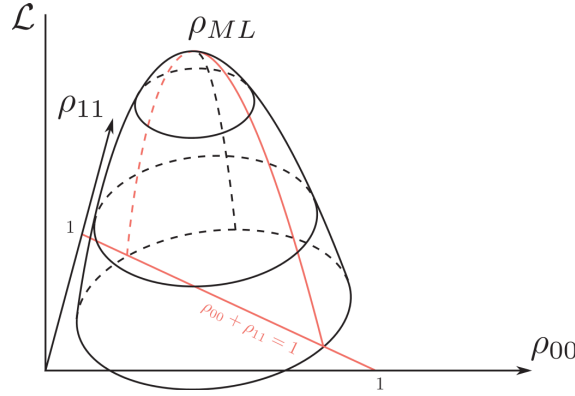


Fig. III.15 Rôle de la contrainte $\text{Tr}(\hat{\rho}) = 1$ dans l'estimation des barres d'erreur. La ligne rouge représente cette contrainte. Les variations de $\mathcal{L}$ dans la direction orthogonale à cette contrainte ne jouent aucun rôle dans l'incertitude de la détermination de p_{00} et p_{11} .

Cette expression permet d'étudier les variations de $\mathcal{L}$ pour des petits déplacements autour d'une valeur $\hat{\rho}$. Intuitivement, si $\mathcal{L}$ varie beaucoup lorsqu'un paramètre de la matrice densité varie, les résultats qui donnent une valeur de ce paramètre différente de celle qui maximise la vraisemblance sont très peu probables.

Il faut cependant faire attention au fait qu'il existe des contraintes liant les paramètres entre eux. Par exemple, la trace de la matrice densité est fixée et on ne doit donc pas tenir compte de la variance obtenue lorsqu'on effectue des déplacements parallèles à $\mathbb{1}$. Par ailleurs, lorsque le rang de $\hat{\rho}_{\text{ML}}$ n'est pas plein, elle est située au bord de $\mathcal{D}(\mathcal{H})$, dont la forme est compliquée et dont il faut aussi tenir compte pour estimer les variances. La référence [28] fournit une méthode qui prend en compte ces difficultés pour calculer la variance de la valeur moyenne $\sigma^2(A)$ d'un opérateur A lorsqu'on effectue plusieurs reconstructions. Celle-ci s'exprime comme :

$$\sigma^2(A) \equiv \text{Tr} \left(A_{\parallel} \mathbf{R}^{-1} (A_{\parallel}) \right), \quad (\text{III.4.7})$$

où le super-opérateur $\mathbf{R} : \mathcal{O}(\mathcal{H}) \rightarrow \mathcal{O}(\mathcal{H})$ est défini, pour tout opérateur hermitien X de trace nulle :

$$\mathbf{R}(X) \equiv \sum_{r=1}^N \frac{\text{Tr} \left(X \hat{E}_{\parallel}^{(r)} \right)}{\text{Tr} \left(\hat{\rho}_{\text{ML}} \hat{E}^{(r)} \right)^2} \hat{E}_{\parallel}^{(r)} + \frac{1}{2} (\lambda_{\text{ML}} \mathbb{1} - \nabla \mathcal{L}_{\text{ML}}) X \hat{\rho}_{\text{ML}}^+ + \frac{1}{2} \hat{\rho}_{\text{ML}}^+ X (\lambda_{\text{ML}} \mathbb{1} - \nabla \mathcal{L}_{\text{ML}}) \quad (\text{III.4.8})$$

où on a utilisé la projection orthogonale $A_{\parallel}$ de l'opérateur A sur l'espace tangent en $\hat{\rho}_{\text{ML}}$ à la sous-variété des opérateurs hermitiens de trace nulle et de même rang que $\hat{\rho}_{\text{ML}}$:

$$A_{\parallel} = A - \frac{\text{Tr}(A P_{\text{ML}})}{\text{Tr}(P_{\text{ML}})} P_{\text{ML}} - (\mathbb{1} - P_{\text{ML}}) A (\mathbb{1} - P_{\text{ML}}), \quad (\text{III.4.9})$$

et où $\hat{\rho}_{\text{ML}}^+$ est le pseudo-inverse de Moore-Penrose de $\hat{\rho}_{\text{ML}}$:

$$\hat{\rho}_{\text{ML}}^+ = \hat{P} \text{diag}(p_1^{-1}, \dots, p_r^{-1}, 0, \dots, 0) \hat{P}^{-1}, \quad (\text{III.4.10a})$$

où $\hat{P}$ est une matrice unitaire telle que

$$\hat{\rho}_{\text{ML}} = \hat{P} \text{diag}(p_1, \dots, p_r, 0, \dots, 0) \hat{P}^{-1}. \quad (\text{III.4.10b})$$

Les deux derniers termes dans l'expression de $\mathbf{R}$ tiennent compte de la structure de $\mathcal{D}(\mathcal{H})$ au voisinage d'une matrice de rang déficient. Lorsque $\hat{\rho}_{\text{ML}}$ est de rang maximal, d'après le critère III.3.48b d'optimalité pour $\hat{\rho}_{\text{ML}}$, ces deux termes s'annulent. Dans ce cas, $\mathbf{R}(X)$ est la projection, pour le produit de Frobenius, de l'opposé de la hessienne $\nabla^2 \mathcal{L}(X)$ sur l'ensemble des opérateurs hermitiens de trace nulle.

L'expression III.4.7 permet de calculer des incertitudes sur des coefficients de la matrice densité. Posons $x_{m,n} \equiv \Re \varrho_{mn}$ et $y_{m,n} \equiv \Im \varrho_{mn}$. En notant $X^{(mn)} \equiv (|m\rangle \langle n| + |n\rangle \langle m|)/2$ et $Y^{(mn)} \equiv i(|m\rangle \langle n| - |n\rangle \langle m|)/2$, on a, pour une matrice densité quelconque $\hat{\varrho}$:

$$\langle \hat{\varrho} X^{(mn)} \rangle = x_{m,n}, \quad (\text{III.4.11})$$

$$\langle \hat{\varrho} Y^{(mn)} \rangle = y_{m,n}. \quad (\text{III.4.12})$$

On a donc :

$$\sigma^2(x_{m,n}) = \sigma^2(X^{(mn)}), \quad (\text{III.4.13})$$

$$\sigma^2(y_{m,n}) = \sigma^2(Y^{(mn)}). \quad (\text{III.4.14})$$

Lorsque les erreurs sur les parties réelles et imaginaires des composantes de la matrice densité ne sont pas corrélées, on peut les utiliser pour obtenir des barres d'erreur sur l'amplitude $r_{m,n} \equiv |\varrho_{mn}|$ et la phase $\phi_{m,n} \equiv \arg(\varrho_{mn})$:

$$\sigma^2(r_{m,n}) = \frac{\sqrt{(x_{m,n}\sigma(x_{m,n}))^2 + (y_{m,n}\sigma(y_{m,n}))^2}}{r_{m,n}}, \quad (\text{III.4.15})$$

$$\sigma^2(\phi_{m,n}) = \frac{\sqrt{(y_{m,n}\sigma(x_{m,n}))^2 + (x_{m,n}\sigma(y_{m,n}))^2}}{r_{m,n}^2}. \quad (\text{III.4.16})$$

Plusieurs remarques importantes s'imposent :

- $\sigma^2(A)$ n'est pas la variance quantique associée à A . Il s'agit de la variance qu'on obtiendrait en réalisant plusieurs estimations de $\langle A \rangle$ avec des matrices densité estimées par notre méthode en utilisant des jeux de données faisant intervenir les mêmes mesures.
- L'expression de $\sigma^2(A)$ n'est valide, en toute rigueur, que dans la limite d'un grand nombre de mesures. En effet, la hessienne est calculée en $\hat{\rho}_{\text{ML}}$, qui n'approche $\hat{\rho}$ qu'asymptotiquement.
- L'expression précédente permet de calculer des barres d'erreur dans le cas où $\hat{\rho}_{\text{ML}}$ n'est pas de rang plein. Elle est cependant à prendre avec précautions dans ce cas. On verra dans la suite un exemple d'illustration.
- Les incertitudes données par III.4.7 contiennent une part statistique, liée à la variance des mesures et une part « systématique », liée au manque d'information sur certains paramètres de la matrice densité. Cette part systématique ne tient pas

compte des contraintes sur la matrice densité et elle peut donc être surestimée. Dans ce cas, les incertitudes calculées sont supérieures à la variance effectivement obtenue en répétant la reconstruction.

Afin d'illustrer ces aspects, on va présenter dans la suite quelques exemples de calcul de $\sigma^2(A)$ dans le cas du qubit.

III.4.3 Incertitudes sur la reconstruction d'un qubit

Supposons qu'on a réalisé un ensemble de mesures à information complète, donné par les mesures projectives $M_i = |i\rangle\langle i|$ où $i \in \{+x, -x, +y, -y, +z, -z\}$ donnant les résultats $\{n_{+x}, n_{-x}, n_{+y}, n_{-y}, n_{+z}, n_{-z}\}$. On note

$$\mathbf{r} = (x, y, z) = \left(\frac{n_{+x} - n_{-x}}{n_{+x} + n_{-x}}, \frac{n_{+y} - n_{-y}}{n_{+y} + n_{-y}}, \frac{n_{+z} - n_{-z}}{n_{+z} + n_{-z}} \right). \quad (\text{III.4.17})$$

Si $|\mathbf{r}|^2 < 1$, alors $\mathbf{r}$ est le vecteur de Bloch correspondant à l'estimation par maximum de vraisemblance, qui est de rang plein.

Les projections des opérateurs de mesure sur l'ensemble des matrices de traces nulles valent :

$$|+x\rangle\langle +x|_{\parallel} = \hat{\sigma}_x/2, \quad | +y\rangle\langle +y|_{\parallel} = \hat{\sigma}_y/2, \quad | +z\rangle\langle +z|_{\parallel} = \hat{\sigma}_z/2, \quad (\text{III.4.18})$$

$$|-x\rangle\langle -x|_{\parallel} = -\hat{\sigma}_x/2, \quad | -y\rangle\langle -y|_{\parallel} = -\hat{\sigma}_y/2, \quad | -z\rangle\langle -z|_{\parallel} = -\hat{\sigma}_z/2. \quad (\text{III.4.19})$$

Pour calculer $\mathbf{R}(\hat{\sigma}_x)$ en utilisant III.4.8, on doit calculer des termes du type $\text{Tr}(\hat{E}_i X)$. Ces termes ne sont non nuls que pour $\hat{E}_i = |+x\rangle\langle +x|$ ou $\hat{E}_i = |-x\rangle\langle -x|$. On a par ailleurs $\text{Tr}(\hat{\rho}_{\text{ML}} |+x\rangle\langle +x|) = (1+x)/2$ et $\text{Tr}(\hat{\rho}_{\text{ML}} |-x\rangle\langle -x|) = (1-x)/2$. Comme $\hat{\sigma}_x$ est de trace nulle, on a $\hat{\sigma}_x|_{\parallel} = \hat{\sigma}_x$. On a donc :

$$\begin{aligned} \mathbf{R}(\hat{\sigma}_x) &= n_{+x} \frac{\text{Tr}(\hat{\sigma}_x |+x\rangle\langle +x|_{\parallel})}{\text{Tr}(\hat{\rho}_{\text{ML}} |+x\rangle\langle +x|)^2} |+x\rangle\langle +x|_{\parallel} + n_{-x} \frac{\text{Tr}(\hat{\sigma}_x |-x\rangle\langle -x|_{\parallel})}{\text{Tr}(\hat{\rho}_{\text{ML}} |-x\rangle\langle -x|)^2} |-x\rangle\langle -x|_{\parallel}, \\ &= \frac{n_{+x}}{\left(\frac{1+x}{2}\right)^2} \frac{\hat{\sigma}_x}{2} + \frac{n_{-x}}{\left(\frac{1-x}{2}\right)^2} \frac{-\hat{\sigma}_x}{2}, \\ &= \frac{2(n_{+x} + n_{-x})}{(1+x)(1-x)} \hat{\sigma}_x. \end{aligned} \quad (\text{III.4.20})$$

On en déduit l'incertitude :

$$\sigma^2(\hat{\sigma}_x) = \frac{(1+x)(1-x)}{n_x}. \quad (\text{III.4.21})$$

Celle-ci correspond à la variance de la loi binomiale associée à $n_x \equiv n_{+x} + n_{-x}$ mesures selon Ox . On trouve des expressions similaires pour des mesures selon Oy et Oz .

Considérons désormais une mesure imparfaite selon Ox , décrite par les opérateurs associés aux vecteurs de Bloch $(e_x, 0, 0)$ et $(-e_x, 0, 0)$. Dans ce cas, l'expression de x

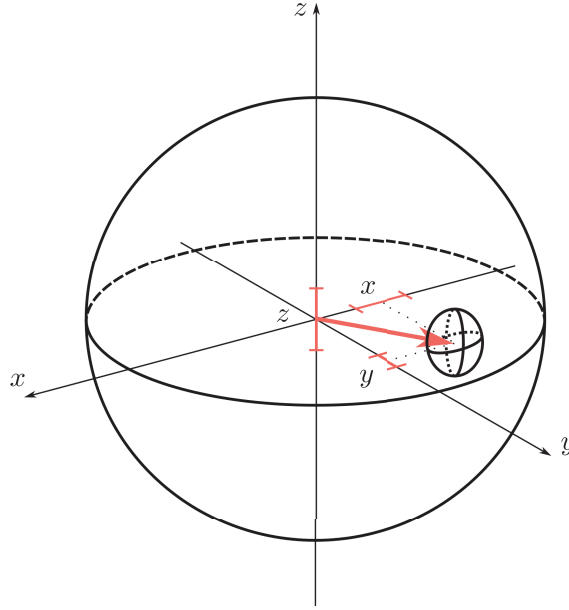


Fig. III.16 Ellipsoïde d'incertitude donné par les barres d'erreur pour le qubit. Les axes de l'ellipsoïde correspondent aux axes de mesure, ici Ox , Oy et Oz . Pour un même nombre de mesures, les barres d'erreur sont plus étroites selon Oy puisque y est la plus grande coordonnée du vecteur de Bloch.

donnée dans III.4.17 est remplacée par $x = (n_{+x} - n_{-x})/e(n_{+x} + n_{-x})$ et on calcule de la même manière que :

$$\sigma^2(\hat{\sigma}_x) = \frac{(1 + e_x x)(1 - e_x x)}{e_x^2 n_x}. \quad (\text{III.4.22})$$

D'une part, la variance de la loi binomiale augmente, d'autre part on a un nombre effectif de mesures $e_x^2 n_x$ plus faible.

Pour une mesure selon $\vec{u} = (u_x, u_y, u_z)$, on a l'incertitude :

$$\sigma^2(\vec{u} \cdot \hat{\sigma}) = \frac{(1+x)(1-x)}{n_x} u_x^2 + \frac{(1+y)(1-y)}{n_y} u_y^2 + \frac{(1+z)(1-z)}{n_z} u_z^2. \quad (\text{III.4.23})$$

L'incertitude en fonction de la direction de mesure forme donc un ellipsoïde dont les axes sont les axes de mesure, comme illustré sur la figure III.16.

Les barres d'erreur sont-elles surestimées ? Rôle des contraintes sur la matrice densité

Les barres d'erreur calculées précédemment constituent un développement limité de la fonction de log-vraisemblance. Lorsque la matrice densité est proche du bord de $\mathcal{D}(\mathcal{H})$, ce développement ne rend pas correctement compte des variations de $\hat{\rho}$. Considérons le cas d'un qubit mesuré selon Ox , étudié dans la section III.4.3. Commençons par supposer la mesure idéale. On peut montrer, à l'aide de l'expression III.4.21 de l'incertitude sur

$\sigma^2(\hat{\sigma}_x)$ que $x + \Delta x \leq 1$, avec $\Delta x \equiv \sqrt{\sigma^2(\hat{\sigma}_x)}$ si et seulement si

$$x \leq 1 - \frac{2}{n_x + 1}. \quad (\text{III.4.24})$$

Cependant, pour $x = 1$, qui correspond à avoir mesuré toujours la valeur $+1$, la variance s'annule. Le plus grand x strictement inférieur à 1 accessible avec n_x mesures est $1 - 2/n_x$. On en déduit qu'il n'y a aucune valeur de x pour laquelle l'intervalle d'incertitude dépasse 1.

Si on considère maintenant le cas d'une mesure avec un contraste limité e_x , on a :

$$x = \frac{1}{e_x} \frac{n_{+x} - n_{-x}}{n_{+x} + n_{-x}}. \quad (\text{III.4.25})$$

Si par exemple

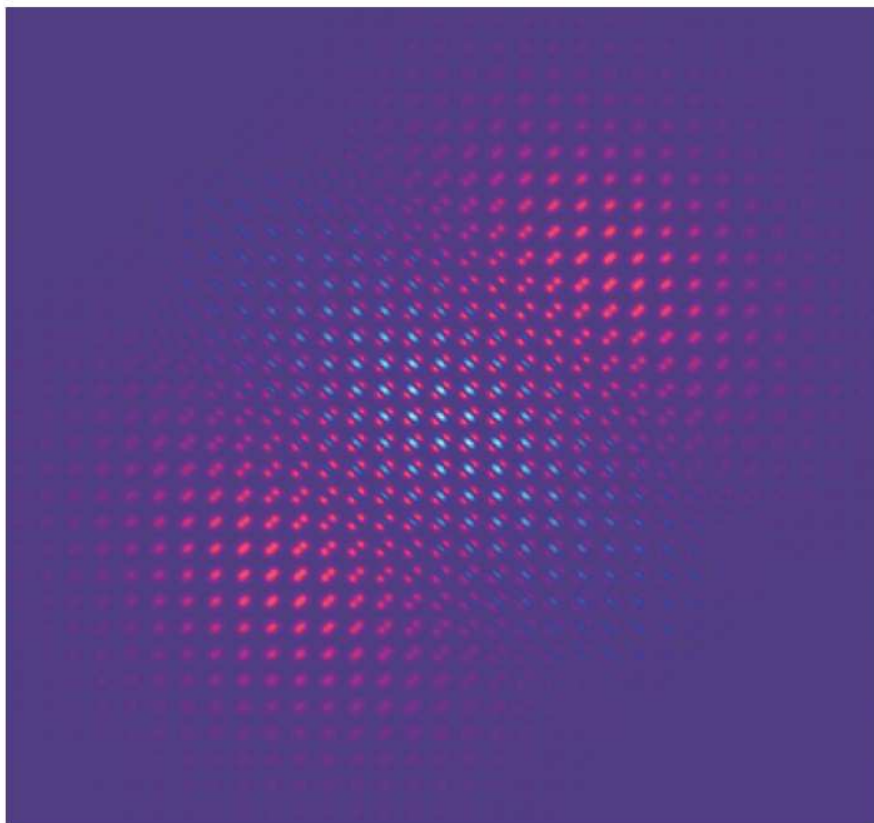
$$\frac{n_{+x} - n_{-x}}{n_{+x} + n_{-x}} = e, \quad (\text{III.4.26})$$

alors $x = 1$ mais la barre d'erreur sur x ne s'annule pas. Les barres d'erreur incluent donc des valeurs qui ne sont pas dans la sphère de Bloch. Les barres d'erreur ne tiennent donc pas toujours compte de la structure de $\mathcal{D}(\mathcal{H})$.

Un autre exemple est le cas d'un qubit proche de $|+_z\rangle$. Pour faire varier x , on est tenu de réduire z , ce qui est très défavorable du point de vue de la vraisemblance. Même si on n'a fait que peu de mesures selon Ox , la valeur de x doit donc être proche de $1/2$. Cet effet n'est pas pris en compte dans l'information de Fisher qui ne tient compte que des variations de $\mathcal{V}$ au voisinage de l'état reconstruit.

Conclusion du chapitre III : une méthode de tomographie reposant sur les trajectoires quantiques

Dans ce chapitre, nous avons introduit les concepts de base de l'estimation de paramètres classiques et de la tomographie d'états quantiques. Nous avons vu que l'information obtenue sur un paramètre estimé peut être quantifiée par l'information de Fisher, qui intervient dans la borne de Cramer-Rao sur la variance de l'estimateur. On peut atteindre cette borne dans la limite d'un grand nombre de mesures en effectuant une estimation par maximum de vraisemblance. Nous avons ensuite introduit les principes de l'estimation d'états quantiques. Celle-ci consiste en général à effectuer, sur un ensemble de préparations de l'état, suffisamment de mesures indépendantes pour pouvoir déterminer la matrice densité, qui doivent être de plus assez répétées pour réduire l'incertitude statistique. Nous avons introduit une nouvelle méthode de tomographie pour laquelle on réalise sur chaque copie de l'état une série de mesures afin d'obtenir de l'information sur la trajectoire quantique suivie par le système. Lorsque la mesure n'est pas projective, on peut la rendre plus pure en la répétant. Même lorsqu'elle est projective, si elle est dégénérée, on peut extraire de l'information supplémentaire en ajoutant une mesure après elle. Il y a donc un intérêt évident à effectuer plusieurs mesures sur une même réalisation de l'état. Finalement, nous avons vu qu'on peut calculer des barres d'erreur sur n'importe quelle fonction linéaire de l'opérateur densité, en utilisant une forme généralisée de l'information de Fisher. Tous ces aspects seront illustrés dans le prochain chapitre, dans lequel nous appliquerons la tomographie par trajectoires à la reconstruction d'un état dans un espace de Hilbert de grande dimension.



Chapitre IV

Reconstruction d'états intriqués

Dans le chapitre II, on a montré qu'on peut mettre en évidence les cohérences de l'état $|10\rangle + |01\rangle$ à l'aide d'un signal d'interférence obtenu en mesurant la probabilité de trouver l'atome sonde dans l'état $|g\rangle$. Ce signal est cependant compatible avec différents états des deux cavités. Par ailleurs, on a vu que certains critères pour mettre en évidence l'intrication entre les deux cavités échouent dans notre cas. Grâce aux outils présentés dans le chapitre III, on peut essayer de reconstruire l'état quantique des deux cavités.

Dans ce chapitre, nous allons voir comment les différentes méthodes de mesure de l'état permettent d'acquérir différentes informations sur la matrice densité et, à terme, de la reconstruire. Dans un premier temps, on donne les matrices d'effet théorique correspondant à trois protocoles de mesure (avec un atome résonnant, avec des échantillons d'atomes dispersifs et avec des échantillons d'atomes résonnants) et on explique les principales caractéristiques de chacune d'entre elles. Ensuite, on montre les résultats obtenus lors de l'application de ces différents protocoles à l'état $|10\rangle + |01\rangle$. Enfin, la section IV.3 montre la possibilité d'utiliser cette méthode pour faire des reconstruction partielles, qui visent à acquérir de l'information sur une partie restreinte de la matrice densité. Dans notre cas on reconstruit les cohérences de la superposition que l'on utilise pour déterminer sa phase quantique.

IV.1 Matrices d'effet théoriques pour nos mesures

Dans cette section, on présente les méthodes de calcul permettant de déterminer les matrices d'effet pour les trois protocoles de mesure utilisés dans cette thèse. On commence par présenter la mesure résonnante à un atome, qui consiste à appliquer la méthode de mesure utilisée dans le chapitre II, constituée d'un atome interagissant successivement avec les deux cavités, à résonance. On étudie ensuite une autre stratégie de mesure dans laquelle on envoie plusieurs échantillons atomiques qui effectuent tous la même interaction que précédemment. On montre que cette mesure permet d'améliorer la discrimination des populations des cavités. Ensuite, on montre les matrices d'effet obtenues pour plusieurs échantillons d'atomes interagissant de manière dispersive avec les cavités, de manière à effectuer une mesure de parité globale. On étudie en particulier la pureté des matrices et l'effet de la décohérence. On se sert de ces résultats pour étudier les matrices d'effet associées à des mesures de fonction de Wigner.

IV.1.1 La mesure résonnante

a) Mesure résonnante à un atome

Considérons un échantillon atomique contenant un atome, qu'on fait interagir à résonance avec la cavité C_j pendant une durée d'interaction effective t . En utilisant les expressions I.3.8 des états propres du hamiltonien et I.4 de leurs énergies, on montre que l'évolution est donnée, en représentation d'interaction, par :

$$|g, n+1\rangle \mapsto \cos \frac{\Omega_{n,j}t}{2} |g, n+1\rangle - \sin \frac{\Omega_{n,j}t}{2} |e, n\rangle, \quad (\text{IV.1.1a})$$

$$|e, n\rangle \mapsto \sin \frac{\Omega_{n,j}t}{2} |g, n+1\rangle + \cos \frac{\Omega_{n,j}t}{2} |e, n\rangle. \quad (\text{IV.1.1b})$$

Les opérateurs de Kraus $W_{k,l}$ décrivant l'effet sur le champ de la mesure dans l d'un atome envoyé dans k sont donc :

$$W_{g,g}(t) = \sum_n \cos(\chi_{n-1}t) |n\rangle \langle n|, \quad (\text{IV.1.2a})$$

$$W_{g,e}(t) = - \sum_n \sin(\chi_{n-1}t) |n-1\rangle \langle n|, \quad (\text{IV.1.2b})$$

$$W_{e,g}(t) = \sum_n \sin(\chi_n t) |n+1\rangle \langle n|, \quad (\text{IV.1.2c})$$

$$W_{e,e}(t) = \sum_n \cos(\chi_n t) |n\rangle \langle n|, \quad (\text{IV.1.2d})$$

où $\chi_n = \Omega_{0,j}\sqrt{n+1}/2$. Dans la séquence de mesure, l'atome, initialement préparé dans $|g\rangle$ interagit successivement avec C_1 et C_2 pendant des durées respectives t_1 et t_2 . L'effet de la mesure sur l'état global des deux cavités est alors décrit par les opérateurs de Kraus :

$$M_e^{[res]}(t_1, t_2) = \sum_k W_{g,k}(t_1) \otimes W_{k,e}(t_2), \quad (\text{IV.1.3a})$$

$$M_g^{[res]}(t_1, t_2) = \sum_k W_{g,k}(t_1) \otimes W_{k,g}(t_2), \quad (\text{IV.1.3b})$$

correspondant respectivement à une mesure dans $|e\rangle$ et dans $|g\rangle$. Si on effectue une unique mesure résonnante, les matrices d'effet correspondant à la détection de $|e\rangle$ et $|g\rangle$ sont :

$$\hat{E}_e^{[res]}(t_1, t_2) = M_e^{[res]\dagger}(t_1, t_2) M_e^{[res]}(t_1, t_2), \quad (\text{IV.1.4a})$$

$$\hat{E}_g^{[res]}(t_1, t_2) = M_g^{[res]\dagger}(t_1, t_2) M_g^{[res]}(t_1, t_2). \quad (\text{IV.1.4b})$$

Dans la section II.2, on a utilisé t_1 et t_2 correspondant respectivement à une impulsion π et une impulsion $\pi/2$, ce qui donnait lieu à un signal oscillant suivant la phase de l'état préparé. Les matrices $\hat{E}_e^{[res]}$ et $\hat{E}_g^{[res]}$ correspondant à ces temps d'interaction sont données sur la figure IV.1. Elles présentent des termes non diagonaux qui sont responsables de la sensibilité de la mesure aux cohérences $|n_1 - 1, n_2 + 1\rangle \langle n_1, n_2|$. Dans la suite, on utilisera

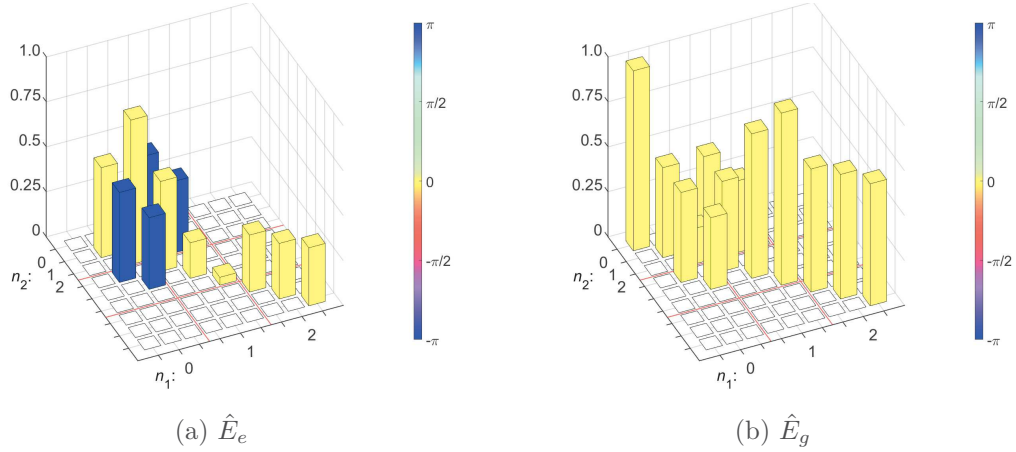


Fig. IV.1 À gauche, matrice d'effet associée à la mesure de l'atome résonnant dans l'état $|e\rangle$. À droite, matrice d'effet associé à sa mesure dans $|g\rangle$. Les termes non diagonaux indiquent la sensibilité de la mesure aux cohérences $|n_1 - 1, n_2 + 1\rangle \langle n_1, n_2|$. On peut par ailleurs remarquer que lorsque le champ est dans l'état $|0\rangle$, il n'y a aucun photon à absorber et l'atome est systématiquement détecté dans $|g\rangle$. L'état $|10\rangle - |01\rangle$ est état propre de $\hat{E}_e$, avec la probabilité 1 tandis que l'état $|10\rangle + |01\rangle$ est état propre de $\hat{E}_g$, avec la probabilité 1.

toujours les mêmes temps d'interaction et on omettra d'expliciter la dépendance en t_1 et t_2 des opérateurs, afin d'alléger les notations.

Tentons de comprendre l'origine des termes non diagonaux en considérant le cas de la détection dans $|g\rangle$. Il existe deux chemins dans l'espace des configurations atomiques menant à ce résultat. Ou bien l'atome n'absorbe de photon dans aucune des cavités, ou bien il absorbe un photon dans C_1 et l'émet dans C_2 . Ces deux chemins correspondent respectivement aux termes $M_1 \equiv W_{g,g} \otimes W_{g,g}$ et $M_2 \equiv W_{g,e} \otimes W_{e,g}$ de la somme IV.1.3b. Le premier terme fait intervenir uniquement des termes du type $|n_1, n_2\rangle \langle n_1, n_2|$ qui préservent le nombre de photons dans chaque cavité. Le second contient des termes du type $|n_1 - 1, n_2 + 1\rangle \langle n_1, n_2|$ qui correspondent au transfert d'un photon de C_1 à C_2 . Lorsqu'on calcule la matrice d'effet, les termes $M_1^\dagger M_1$ et $M_2^\dagger M_2$ sont diagonaux dans la base de Fock. Ce sont donc les termes croisés $M_1^\dagger M_2$ et $M_2^\dagger M_1$ qui sont responsables de la sensibilité de la mesure aux cohérences. Il y a une interférence entre les deux chemins pris par l'atome lorsque le champ est dans une superposition d'état.

C'est cette interférence qui est responsable de la sensibilité de la mesure à la phase du champ. Pour le voir, commençons par chercher les vecteurs propres de $\hat{E}_e^{[res]}$ et $\hat{E}_g^{[res]}$. Un premier vecteur propre évident de $\hat{E}_g^{[res]}$ est $|0, 0\rangle$, associé à la valeur propre 1 : lorsque les deux cavités sont vides, il n'y a aucune excitation dans le système et l'atome ressort dans son état fondamental avec certitude. On remarque aussi aisément que l'état $|10\rangle - |01\rangle$ est vecteur propre de $\hat{E}_g^{[res]}$, avec la valeur propre 0 et de $\hat{E}_e^{[res]}$ avec la valeur propre 1. Lorsque le champ est dans cet état, l'atome sort donc systématiquement des cavités dans l'état excité. En revanche, l'état $|10\rangle + |01\rangle$ est vecteur propre de $\hat{E}_g^{[res]}$ avec la valeur 1 et de $\hat{E}_e^{[res]}$ avec la valeur propre 0 et ne donne donc jamais lieu à l'absorption du photon.

Cette mesure est donc particulièrement adaptée pour discriminer les états $|10\rangle + |01\rangle$ et $|10\rangle - |01\rangle$. On retrouve les résultats du chapitre II.

Si on veut, de façon plus générale, être sensible à n'importe quelle superposition du type $|10\rangle + e^{i\phi}|01\rangle$, il suffit de laisser le système évoluer avant la mesure. En effet, on a vu dans le chapitre II que l'état évolue au cours du temps à cause de la différence de fréquence entre les deux cavités, de sorte qu'en l'absence d'interaction avec les atomes et pour une durée τ suffisamment faible pour négliger la décohérence :

$$\hat{\rho}(t + \tau) = \mathbb{R}_\tau(\hat{\rho}(t)) = \hat{R}(\tau)\hat{\rho}(t)\hat{R}^\dagger(\tau), \quad (\text{IV.1.5})$$

où l'opérateur de rotation dans l'espace des phases $\hat{R}(\tau)$ est donné par :

$$\hat{R}(\tau) = \exp(i\delta\tau(\hat{N}_2 - \hat{N}_1)/2), \quad (\text{IV.1.6})$$

où δ est la différence de fréquence entre les deux cavités et $\hat{N}_i$ l'opérateur nombre de photons de la cavité i . Si on attend une durée T_0 avant la mesure, les matrices d'effet précédentes sont remplacées par :

$$\hat{E}_e^{[res]} = \hat{R}(T_0)^\dagger M_e^{[res]\dagger} M_e^{[res]} \hat{R}(T_0), \quad (\text{IV.1.7a})$$

$$\hat{E}_g^{[res]} = \hat{R}(T_0)^\dagger M_g^{[res]\dagger} M_g^{[res]} \hat{R}(T_0). \quad (\text{IV.1.7b})$$

Cette rotation a pour effet de changer la phase des vecteurs propres des matrices d'effet. De cette manière, on peut générer des matrices d'effet sensibles à n'importe quel état $|10\rangle + e^{i\phi}|01\rangle$.

On peut enfin remarquer que la mesure résonnante est légèrement sensible aux nombres de photons. Par exemple, l'état $|00\rangle$ donnera toujours une détection dans $|g\rangle$; l'état $|02\rangle$ donne 80 % de $|g\rangle$; l'état $|1,1\rangle$ donne 80 % de $|e\rangle$. Ces informations ne suffisent pas à discriminer totalement le nombre de photons. On pourrait faire varier ces différentes probabilités en changeant la durée des interactions avec les cavités. Cependant, l'information sur la distribution des nombres de photons est plus facilement obtenue à l'aide d'une mesure non destructive qui sera introduite par la suite. Dans le prochain paragraphe, on montre qu'on peut aussi renforcer la sensibilité de la mesure résonnante aux nombres de photons en effectuant plusieurs mesures successives.

b) Mesure résonnante à plusieurs atomes

On a vu précédemment qu'un atome interagissant à résonance avec les deux cavités est capable de sonder les cohérences entre deux états $|n_1, n_2\rangle$ et $|n_1 - 1, n_2 + 1\rangle$ tels qu'on peut passer de l'un à l'autre en absorbant un photon dans la première cavité et en l'émettant dans la seconde. En revanche, pour passer de $|n_1, n_2\rangle$ à $|n_1 - 2, n_2 + 2\rangle$, un atome n'est pas suffisant puisqu'il ne peut transporter qu'un quantum d'énergie. Ceci indique que pour sonder d'autres cohérences de la matrice densité, une méthode est d'augmenter le nombre d'atomes.

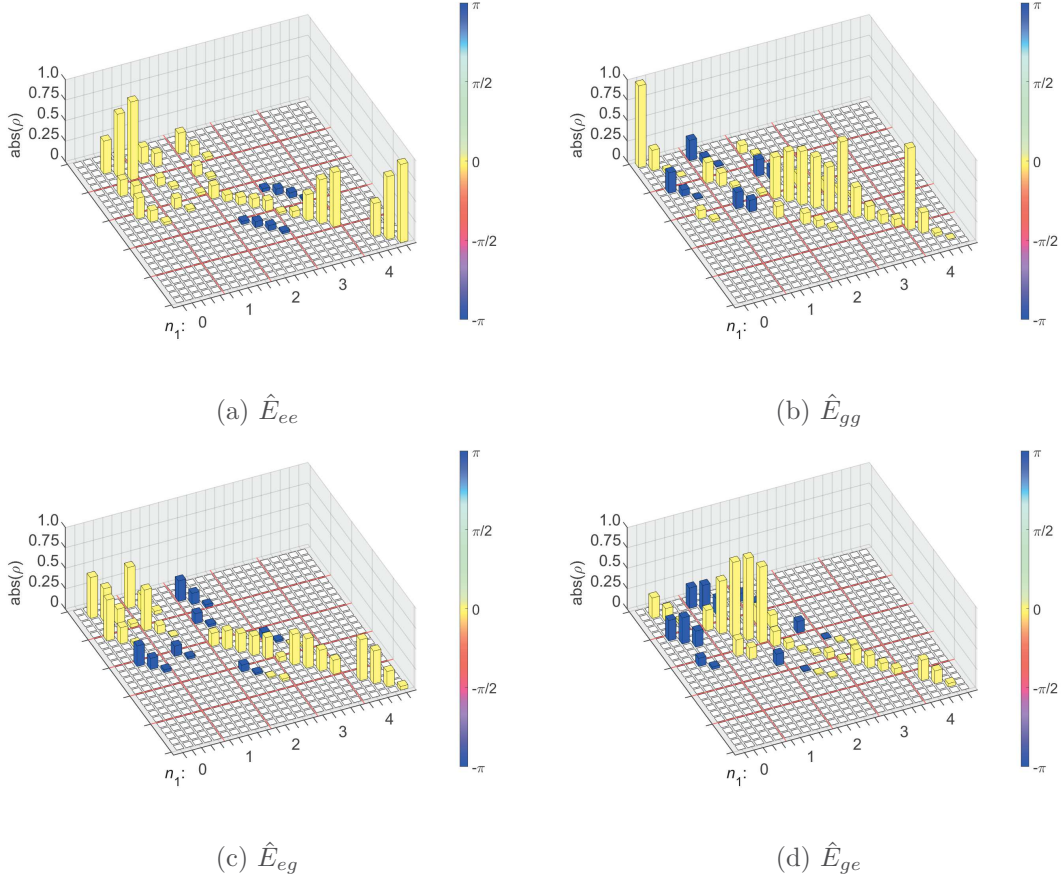


Fig. IV.2 Matrices d'effet associées aux différents résultats de mesure possibles pour deux atomes. Les espaces de Hilbert de chaque cavité sont de dimension 5. Voir dans le texte pour plus de détails.

La figure IV.2 montre les matrices d'effets :

$$\hat{E}_{ee}^{[res]} = M_e^{[res]\dagger} M_e^{[res]\dagger} M_e^{[res]} M_e^{[res]}, \quad (IV.1.8a)$$

$$\hat{E}_{gg}^{[res]} = M_g^{[res]\dagger} M_g^{[res]\dagger} M_g^{[res]} M_g^{[res]}, \quad (IV.1.8b)$$

$$\hat{E}_{eg}^{[res]} = M_e^{[res]\dagger} M_g^{[res]\dagger} M_g^{[res]} M_e^{[res]}, \quad (IV.1.8c)$$

$$\hat{E}_{ge}^{[res]} = M_g^{[res]\dagger} M_e^{[res]\dagger} M_e^{[res]} M_g^{[res]}, \quad (IV.1.8d)$$

où $\hat{E}_{ij}^{[res]}$ décrit la détection d'un atome résonnant dans $|i\rangle$, puis d'un autre dans $|j\rangle$, les temps d'interaction étant les mêmes que ceux utilisés dans la mesure à un atome. Ces matrices présentent des cohérences du même type que précédemment, mais aussi des cohérences « à deux photons » $|n_1 - 2, n_2 + 2\rangle \langle n_1, n_2|$, puisque les deux atomes sont capables d'absorber deux photons dans la première cavité pour les injecter dans la seconde.

Comme précédemment, on peut obtenir des informations intéressantes en diagonalisant ces matrices. Tout d'abord, on remarque que l'état $|0, 0\rangle$ des cavités donne systématiquement lieu au résultat de détection $\{g, g\}$ puisqu'il n'y a pas de photon à absorber.

Lorsque le champ est dans l'état $|10\rangle + |01\rangle$, le résultat de détection est toujours $\{e, g\}$, ce qui est consistant avec ce qu'on a vu précédemment. Le premier atome absorbe toujours le photon et le second trouve le champ dans l'état vide et reste donc dans son état fondamental.

De manière amusante, la matrice $\hat{E}_{ge}$ a aussi un vecteur propre associé à la valeur propre 1. Il s'agit de l'état $|13\rangle$. Dans cet état, le premier atome absorbe avec certitude le photon dans C_1 , puis il l'émet ensuite avec certitude dans C_2 . En effet, lorsqu'il interagit avec la seconde cavité, il y a 4 excitations, initialement réparties en 3 photons et une excitation atomique. Pour une durée d'interaction calibrée pour faire une impulsion $\pi/2$ avec une excitation, le système effectue une impulsion π avec 4 excitations puisque la fréquence de Rabi est multipliée par deux. Après le passage du premier atome, l'état du champ est donc $|04\rangle$. Dans cet état, l'atome n'oscille pas dans la première cavité et effectue une oscillation π dans le champ de 4 photons de la deuxième, pour la même raison que précédemment. Il sort donc toujours excité. Ces exemples, bien que ne présentant pas d'intérêt particulier pour nos mesures, illustrent la grande variété de comportements de la mesure en fonction des paramètres.

L'étude exhaustive des vecteurs propres de ces matrices est fastidieuse. Si on se restreint aux faibles nombres d'excitations, on voit par exemple que la matrice $\hat{E}_{ee}^{[res]}$ a pour vecteur propre l'état $0.49|20\rangle + 0.39|11\rangle + 0.78|02\rangle$, associé à la valeur propre 0.68. On voit que cette mesure permet en principe de mettre en évidence des cohérences non locales par l'échange de deux photons entre les cavités. Elle pourrait donc servir à étudier des états du type $|20\rangle + e^{i\phi}|02\rangle$. La population dans $|11\rangle$ dans l'état précédent est liée au fait que la durée des impulsions a été choisie pour l'état $|10\rangle + e^{i\phi}|01\rangle$ et n'est pas optimale pour des états à deux photons.

Envoyer plusieurs atomes permet, en plus de rendre la mesure sensible à plus de cohérences, d'affiner l'information extraite sur les populations. En effet, si on détecte N atomes dans l'état excité, on sait avec certitude, en l'absence d'imperfections expérimentales, que l'état du champ contenait au moins N photons. De même, si on détecte beaucoup d'atomes dans l'état fondamental la vraisemblance des états de grand nombre de photons a tendance à diminuer¹. La figure IV.3 montre la matrice d'effet pour un atome détecté dans $|e\rangle$ suivi de 9 atomes détectés dans $|g\rangle$. La comparaison avec la figure IV.1 est sans appel : la discrimination sur les nombres de photons est bien meilleure pour de faibles nombres de photons. Pour des nombres de photons plus élevés, il reste cependant quelques artefacts. Par exemple, l'état $|41\rangle$ a une vraisemblance de 0.5 pour la mesure à dix atomes comme pour celle à un atome.

On voit donc qu'il est intéressant d'ajouter des atomes supplémentaires pour fournir un complément d'information à la mesure. Dans le cas précédent, les atomes détectés dans l'état fondamental permettent d'augmenter la confiance qu'on a dans le nombre de photons, sans rien changer à l'information que le premier atome fournit sur la phase de la superposition atomique. Nous verrons dans le paragraphe IV.2.4 que cette information permet de réduire largement les ressources nécessaires pour reconstruire l'état.

En conclusion, il est donc plus intéressant d'envoyer plusieurs atomes successifs. Tout

1. Il y a quelques exceptions à cette règle. Par exemple, si on détecte tous les atomes dans $|g\rangle$, la vraisemblance de l'état $|40\rangle$ est 1, puisque les atomes n'absorbent pas dans cet état à cause de la durée des impulsions. Cependant, la vraisemblance d'un tel état fortement excité sera réduite lorsque la décohérence sera prise en compte.

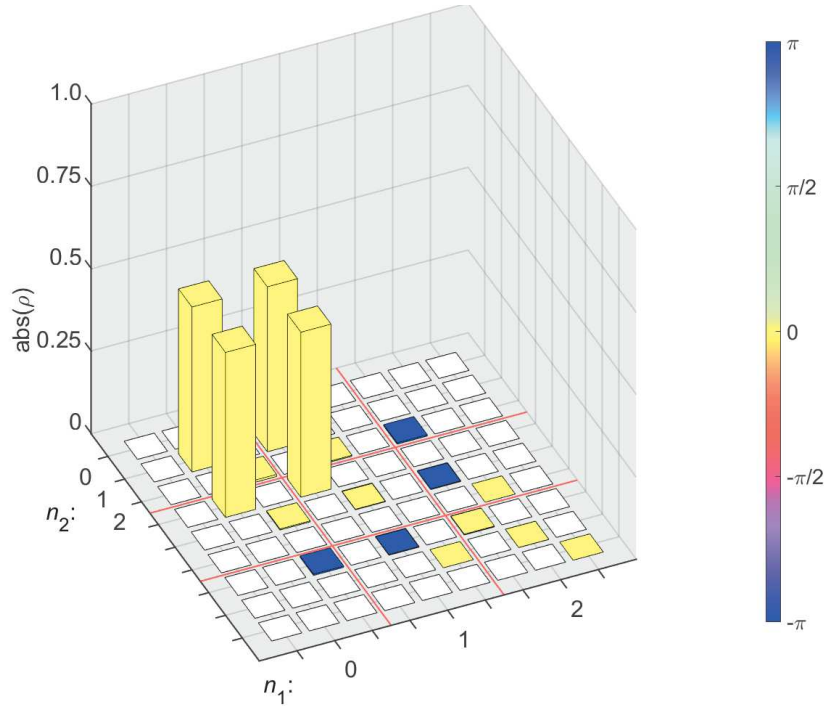


Fig. IV.3 Matrice d'effet associée à la mesure d'un atome dans $|e\rangle$ suivi de neuf atomes dans $|g\rangle$. Le principal état propre, associé à une vraisemblance de 1, est l'état $|10\rangle + |01\rangle$. La comparaison avec la matrice d'effet associée à la mesure d'un seul atome dans $|e\rangle$, à gauche dans la figure IV.1, montre que les faibles nombres de photons sont bien mieux discriminés à l'aide d'une mesure à plusieurs atomes.

d'abord, la mesure devient sensible à des cohérences à plusieurs photons. On sonde donc une plus grande partie de la matrice densité. C'est ce qu'on fera par exemple dans le protocole à deux photons de la partie IV.2.5. De plus, la mesure est meilleure pour discriminer les nombres de photons et donne donc lieu à des matrices d'effet de plus grande pureté, ce qui accélère la reconstruction. Un autre avantage de cette méthode qu'elle permet d'extraire de l'information beaucoup plus vite, ce qui sera étudié dans le paragraphe IV.2.4.

IV.1.2 La mesure dispersive

a) Mesure de la parité

On a vu précédemment que la mesure résonnante est appropriée pour mesurer une partie des cohérences non locales de l'état, mais qu'elle ne permet pas toujours de bien discriminer les populations. La mesure non destructive (QND) du nombre de photons, une technique standard dans nos expériences d'électrodynamique quantique en cavité, permet en revanche d'estimer les populations de manière efficace. Dans la mesure où les atomes ne modifient pas le nombre de photons, il est possible d'envoyer, lors d'une même réalisation, de nombreux échantillons atomiques. L'accumulation d'information fournie

par les différents atomes fait que la mesure tend, pour de grands nombres d'échantillons, vers une mesure projective du nombre de photons [52].

On a vu dans la section I.5.4 que la probabilité de mesurer $|e\rangle$ pour un atome dispersif oscille sinusoïdalement avec le nombre total de photons. Pour un déphasage de π par photon dans chaque cavité et en choisissant convenablement la phase des impulsions de Ramsey pour se placer au maximum des franges, la mesure idéale est donc décrite par :

$$M_e^{[dis]} = \cos(\hat{N}\pi/2), \quad (\text{IV.1.9a})$$

$$M_g^{[dis]} = \sin(\hat{N}\pi/2), \quad (\text{IV.1.9b})$$

ce qu'on peut récrire sous forme matricielle :

$$M_e^{[dis]} = \begin{pmatrix} \begin{array}{cccc|ccc|ccc} |0,0\rangle & |0,1\rangle & |0,2\rangle & |0,3\rangle & & & & & & & \\ 1 & & & & & & & & & & \\ & 0 & & & & & & & & & \\ & & -1 & & & & & & & & \\ & & & 0 & & & & & & & \\ & & & & \ddots & & & & & & \\ \hline & & & & & 0 & & & & & \\ & & & & & & -1 & & & & \\ & & & & & & & 0 & & & \\ & & & & & & & & \ddots & & \\ \hline & & & & & & & & & -1 & \\ & & & & & & & & & & 0 \\ & & & & & & & & & & & 1 \\ & & & & & & & & & & & & \ddots \\ & & & & & & & & & & & & & \ddots \end{array} \end{pmatrix}, \quad (\text{IV.1.10})$$

$$M_g^{[dis]} = \left(\begin{array}{cccc|ccc|ccc} |0,0\rangle & |0,1\rangle & |0,2\rangle & |0,3\rangle & & |1,0\rangle & |1,1\rangle & |1,2\rangle & & |2,0\rangle & |2,1\rangle & |2,2\rangle \\ 0 & & & & & & & & & & & \\ & 1 & & & & & & & & & & \\ & & 0 & & & & & & & & & \\ & & & -1 & & & & & & & & \\ & & & & \ddots & & & & & & & \\ \hline & & & & & 1 & & & & & & \\ & & & & & & 0 & & & & & \\ & & & & & & & -1 & & & & \\ & & & & & & & & \ddots & & & \\ \hline & & & & & & & & & 0 & & \\ & & & & & & & & & & -1 & \\ & & & & & & & & & & & 0 \\ & & & & & & & & & & & \ddots \\ & & & & & & & & & & & \ddots \end{array} \right). \quad (\text{IV.1.11})$$

Les matrices d'effets associées à ces deux mesures, qui forment un POVM s'écrivent :

$$\hat{E}_e^{dis} = \left(\begin{array}{c|c|c} \begin{array}{cccc} |0,0\rangle & |0,1\rangle & |0,2\rangle & |0,3\rangle \\ 1 & & & \\ & 0 & & \\ & & 1 & \\ & & & 0 \\ & & & \ddots \end{array} & \begin{array}{ccc} |1,0\rangle & |1,1\rangle & |1,2\rangle \\ & 0 & \\ & & 1 \\ & & & 0 \\ & & & \ddots \end{array} & \begin{array}{ccc} |2,0\rangle & |2,1\rangle & |2,2\rangle \\ & & \\ & & \\ & & \\ & & \\ & & \ddots \end{array} \\ \hline & & \\ \hline & & \end{array} \right), \quad (IV.1.12)$$

$$\hat{E}_g^{dis} = \left(\begin{array}{c|c|c} \begin{array}{cccc} |0,0\rangle & |0,1\rangle & |0,2\rangle & |0,3\rangle \\ 0 & & & \\ & 1 & & \\ & & 0 & \\ & & & 1 \\ & & & \ddots \end{array} & \begin{array}{ccc} |1,0\rangle & |1,1\rangle & |1,2\rangle \\ & 1 & \\ & & 0 \\ & & & 1 \\ & & & \ddots \end{array} & \begin{array}{ccc} |2,0\rangle & |2,1\rangle & |2,2\rangle \\ & & \\ & & \\ & & \\ & & \\ & & \ddots \end{array} \\ \hline & & \\ \hline & & \end{array} \right). \quad (IV.1.13)$$

Il apparaît très clairement sur ces deux matrices que cette mesure nous donne accès à la parité globale du champ contenu dans les deux cavités. Par exemple, dans le cas de l'état $|10\rangle + |01\rangle$, il y a exactement un photon dans les deux cavités et l'atome dispersif est toujours mesuré dans $|g\rangle$.

b) Mesure de plusieurs atomes successifs

La mesure de plusieurs atomes successifs ne présente pas d'intérêt dans le cas de la mesure idéale, puisqu'elle est projective. En revanche, en présence d'imperfections donnant lieu à un contraste de Ramsey réduit, elle permet de se rapprocher d'une mesure projective. Dans ce cas, en notant c le contraste de la mesure, les opérateurs de Kraus correspondant respectivement à une détection effective dans $|e\rangle$ ou $|g\rangle$ sont donnés par :

$$\tilde{M}_e^{[dis]} = \sqrt{\frac{1+c}{2}} M_e^{[dis]} + \sqrt{\frac{1-c}{2}} M_g^{[dis]}, \quad (\text{IV.1.14a})$$

$$\tilde{M}_g^{[dis]} = \sqrt{\frac{1-c}{2}} M_e^{[dis]} + \sqrt{\frac{1+c}{2}} M_g^{[dis]}, \quad (\text{IV.1.14b})$$

de sorte que les matrices d'effets correspondantes s'écrivent :

$$\hat{E}_e = \frac{1+c}{2} \hat{E}_e^{dis} + \frac{1-c}{2} \hat{E}_g^{dis}, \quad (\text{IV.1.15a})$$

$$\hat{E}_g = \frac{1-c}{2} \hat{E}_e^{dis} + \frac{1+c}{2} \hat{E}_g^{dis}. \quad (\text{IV.1.15b})$$

La pureté des matrices d'effets correspondante est réduite par rapport à celle du cas idéal. Examinons ce qui se passe lorsqu'on combine deux mesures de ce type. Les matrices d'effet correspondant aux différentes possibilités sont :

$$\hat{E}_{ee} = \frac{(1+c)^2}{4} \hat{E}_e^{dis} + \frac{(1-c)^2}{4} \hat{E}_g^{dis}, \quad (\text{IV.1.16a})$$

$$\hat{E}_{eg} = \frac{1-c^2}{4} \mathbb{1}, \quad (\text{IV.1.16b})$$

$$\hat{E}_{ge} = \frac{1-c^2}{4} \mathbb{1}, \quad (\text{IV.1.16c})$$

$$\hat{E}_{gg} = \frac{(1-c)^2}{4} \hat{E}_e^{dis} + \frac{(1+c)^2}{4} \hat{E}_g^{dis}. \quad (\text{IV.1.16d})$$

Une manière de quantifier l'apport des mesures successives est de regarder la pureté des matrices d'effet obtenues. Si on normalise l'opérateur $\hat{E}$ en écrivant $\tilde{E} \equiv \hat{E}/\text{Tr}(\hat{E})$, on obtient un opérateur hermitien, positif de trace 1, à savoir une matrice densité, dont on peut calculer la pureté $\text{Tr}(\tilde{E}^2)$. Plus celle-ci est élevée, plus la mesure s'approche d'une mesure projective extrayant le maximum d'information sur l'état. Lorsqu'on envoie de nombreux atomes, le nombre de résultats de mesure augmente très vite, ce qui réduit la probabilité de chacune d'entre elles. En revanche, la pureté moyenne des mesures a tendance à augmenter.

Les matrices d'effet correspondant à deux mesures dans le même état donnent des matrices réduites dont la pureté est plus grande que celle obtenue pour une seule mesure : en présence d'imperfections, mesurer deux atomes dans $|e\rangle$ ou $|g\rangle$ fournit une information plus fiable qu'en mesurer un seul. En revanche, dans le cas où on mesure un atome dans $|e\rangle$ et l'autre dans $|g\rangle$, la matrice d'effet est proportionnelle à l'identité : on n'obtient aucune information. Pour des valeurs du contraste proches de 1, la probabilité de ce second scénario est faible. Si on veut quantifier l'avantage obtenu en répétant la mesure,

il faut donc tenir compte, pour tous les résultats de mesure possibles, du gain ou de la perte de pureté ainsi que de la probabilité du résultat.

La probabilité de chaque résultat dépend en principe de l'état considéré. Afin de ne pas faire d'hypothèse sur cet état, une possibilité est de considérer la trace de la matrice d'effet. Plus celle-ci est élevée, plus la matrice joue un rôle prépondérant dans la décomposition de l'unité associée à un ensemble de résultats de mesure possibles. Dans la suite, nous considérons la moyenne des puretés des matrices d'effet réduites associées à la mesure de N_{disp} atomes dispersifs, pondérées par leur trace. La figure IV.4 montre l'évolution de cette pureté moyenne en fonction de N_{disp} , pour un contraste $c = 0.5$, c'est-à-dire pour une probabilité de mesurer la mauvaise parité de 0.25. La pureté moyenne est calculée dans un espace de Hilbert tronqué à 5 photons dans chaque cavité. La pureté de la mesure de parité idéale est alors égale à $1/18 \simeq 0.056$. La pureté moyenne s'approche de cette valeur pour un nombre d'atomes supérieur à la dizaine.

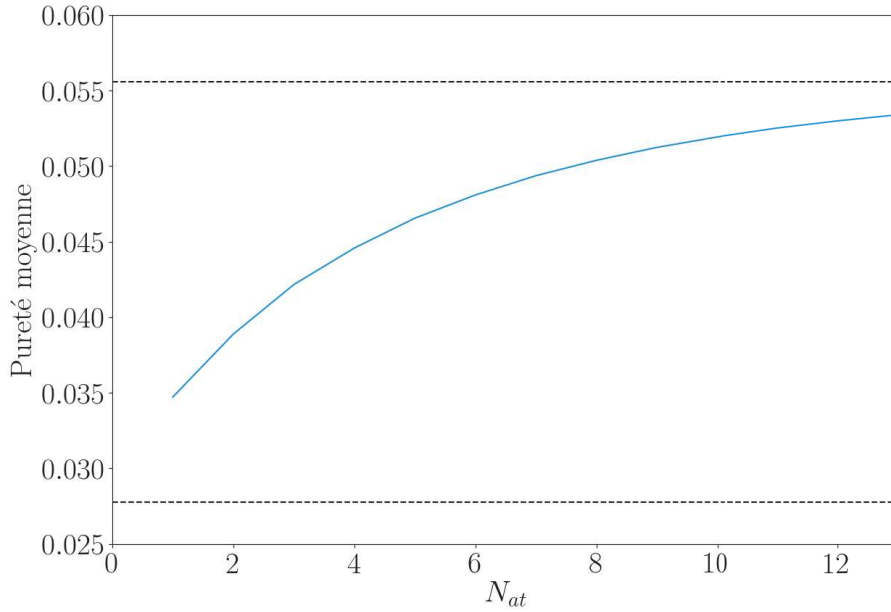


Fig. IV.4 Évolution de la pureté moyenne des matrices d'effet en fonction du nombre d'atomes détectés dans le cas d'une mesure imparfaite de contraste $c = 0.5$. La moyenne est prise sur l'ensemble des matrices d'effet correspondant aux résultats de mesure possibles, pondérées par leur trace avant renormalisation. Les traits pointillés indiquent les valeurs extrémales accessibles pour la pureté. La pureté croît avec le nombre d'atomes détectés et tend vers la valeur maximale.

La mesure précédente permet de mesurer la parité totale du champ. Cependant elle ne permet pas à elle seule de reconstruire la distribution de Fock. Dans ce but, il faudrait combiner plusieurs mesures, par exemple une mesure de parité avec des mesures modulo 4 et 8. Il faudrait aussi mesurer chaque cavité séparément. Dans la suite, on va étudier une autre méthode de mesure qui se fonde sur la mesure de parité mais qui permet de sonder toute la matrice densité : la mesure de la fonction de Wigner.

c) Mesure de la fonction de Wigner

Matrices d'effet

Dans le chapitre I.5, on a vu que la mesure de la fonction de Wigner consiste à effectuer un déplacement de l'état, puis à mesurer la parité. En pratique, on prépare l'état puis on effectue un déplacement $\hat{D}(-\alpha_1, -\alpha_2) \equiv \hat{D}_1(-\alpha_1) \otimes \hat{D}_2(-\alpha_2)$ d'amplitudes complexes $-\alpha_1$ dans C_1 et $-\alpha_2$ dans C_2 . On profite ensuite du caractère non destructif de la mesure dispersive pour le nombre de photons pour envoyer, après une durée T_0 , une série de N_{disp} atomes qui mesurent la parité, séparés par un intervalle de temps T_d , fournissant une trajectoire de mesure Y . Une telle séquence, schématisée dans la figure IV.5, fournit un ensemble de résultats de mesure $(y_i)_{i \in \llbracket 1, N_{disp} \rrbracket}$ auquel est associée la matrice d'effet :

$$\hat{E}_{wig}^Y = \frac{1}{2} \left(\mathbb{1} + \hat{W}^Y \right), \quad (\text{IV.1.17})$$

où

$$\hat{W}^Y = \hat{D}(\alpha_1, \alpha_2) \hat{U}_{T_0}^\dagger M_{y_1}^{[dis]\dagger} \hat{U}_{T_d}^\dagger \dots \hat{U}_{T_d}^\dagger M_{y_{N_{disp}}}^{[dis]\dagger} M_{y_{N_{disp}}}^{[dis]} \hat{U}_{T_d} \dots \hat{U}_{T_d} M_{y_1}^{[dis]} \hat{U}_{T_0} \hat{D}(-\alpha_1, -\alpha_2). \quad (\text{IV.1.18})$$

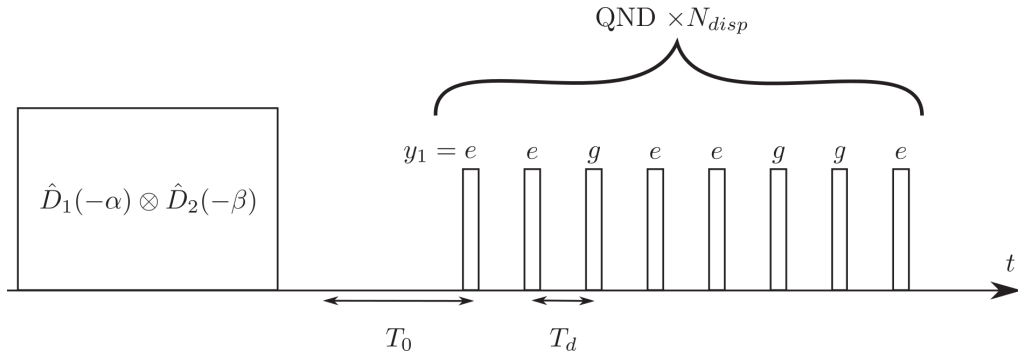


Fig. IV.5 Séquence utilisée pour mesurer la fonction de Wigner du champ. Après préparation de l'état à reconstruire, on effectue un déplacement des deux cavités déterminant le point (α, β) dans l'espace des phases auquel on va mesurer la fonction de Wigner. On envoie ensuite, après une durée T_0 une série de N_{disp} atomes séparés de T_d . L'atome k donne le résultat de mesure $y_k \in \{e, g\}$. Cette série d'atomes permet de mieux moyenner les résultats de mesure, donc de rendre les matrices d'effet plus pures.

Les évolutions libres $\hat{U}_t$ ne modifient pas les populations et commutent avec l'opérateur $\hat{N}$, tout comme les mesures dispersives qui s'écrivent comme fonctions de cet opérateur. Dans l'expression précédente, on peut donc faire commuter tous ces opérateurs. Il en résulte que les termes d'évolution se neutralisent avec leurs conjugués hermitiques. En conséquence, le temps d'attente entre les atomes dispersifs ne joue aucun rôle, contrairement à ce qu'on avait observé dans le cas des atomes résonnants². La mesure peut donc

2. Ceci ne sera plus le cas lorsqu'on prendra en compte la décohérence, puisque la cavité émet et absorbe des photons sous l'effet de sauts quantiques.

se mettre sous la forme simplifiée :

$$\hat{W}^Y = \hat{D}(\alpha_1, \alpha_2) M_{y_1}^{[dis]\dagger} \dots M_{y_{N_{disp}}}^{[dis]\dagger} M_{y_{N_{disp}}}^{[dis]} \dots M_{y_1}^{[dis]} \hat{D}(-\alpha_1, -\alpha_2). \quad (\text{IV.1.19})$$

Cette mesure est simplement constituée d'un déplacement de l'état, suivi de N_{disp} mesures dispersives successives. Comme l'opérateur déplacement est unitaire, la pureté d'une telle mesure est la même que celle obtenue sans le déplacement. L'analyse de la pureté de la mesure dispersive développée dans le paragraphe est donc toujours valide. Par ailleurs, on en déduit que les états propres de la mesure, associés au résultat $+1$ (respectivement -1) sont des états de Fock pairs (respectivement impairs) déplacés.

Universalité de la mesure

La connaissance de la fonction de Wigner est identique à celle de la matrice densité. En mesurant la fonction de Wigner dans tout l'espace, on peut donc en principe déterminer la matrice densité. Ceci n'est bien entendu pas possible en pratique : on ne mesure la fonction de Wigner que sur un ensemble de points bien choisis. En revanche, si on mesure la fonction de Wigner en des points bien choisis pour obtenir suffisamment de matrices d'effets indépendantes pour générer l'ensemble des matrices densité, on peut procéder à une reconstruction. Le choix des points de mesure est une question complexe. Remarquons qu'il est nécessaire de faire des injections simultanées dans les deux cavités pour être sensible aux cohérences non locales. Il faut par ailleurs utiliser au moins deux phases différentes d'injection pour être sensible aux parties réelle et imaginaire des cohérences. Il suffit donc en général de prendre des points réalisant un quadrillage dans l'espace des phases à quatre dimensions pour obtenir un ensemble de mesures complet. Lorsqu'on dispose d'informations supplémentaires sur l'état, par exemple la distribution du nombre de photons, ou la présence de motifs particulièrement contrastés dans sa fonction de Wigner, on peut choisir les points de mesure pour limiter le nombre de mesures nécessaires et améliorer la convergence de la reconstruction.

Exemples de matrices d'effet

La figure IV.6 montre plusieurs exemples de matrices $\hat{W}(\alpha_1, \alpha_2)$. Pour $\alpha_1 = \alpha_2 = 0$, il n'y a pas de déplacement et la mesure revient à une mesure de parité. Pour $\alpha_1 = 0$, la mesure est sensible aux populations et aux cohérences dans la cavité 2, mais seulement à la parité dans la cavité 1. Elle n'est par ailleurs pas sensible aux cohérences non locales. La situation est symétrique lorsqu'on effectue $\alpha_2 = 0$. Lorsqu'on effectue un déplacement dans chaque cavité on est sensible à tous les termes de la matrice densité. On voit enfin, en comparant les figures IV.6d et IV.6e, que la phase des coefficients de la matrice peut être contrôlée avec la phase des impulsions de la mesure.

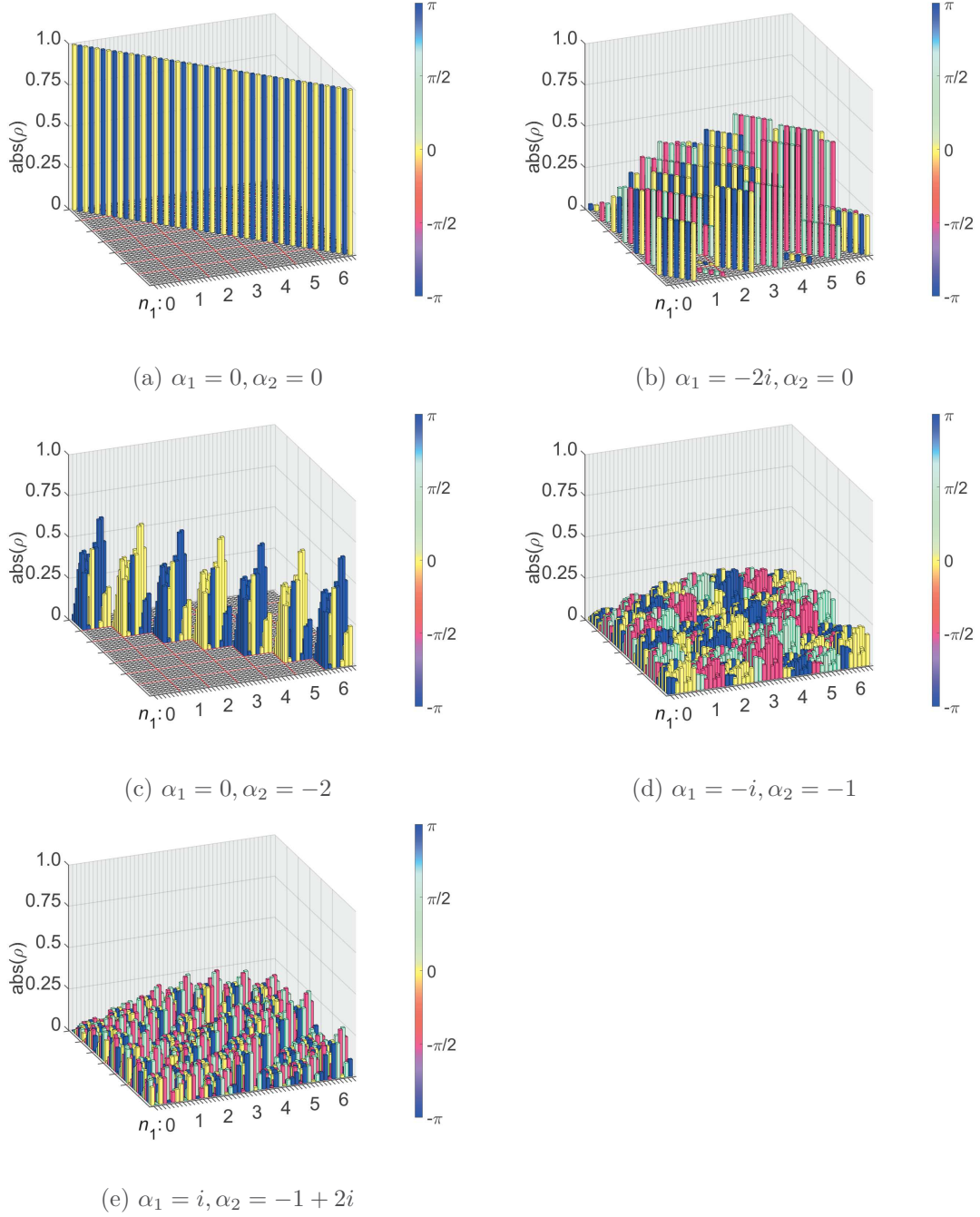


Fig. IV.6 Matrices $\hat{W}(\alpha_1, \alpha_2)$ associées à la mesure de la fonction de Wigner pour plusieurs déplacements (α, β) .

d) Effet de la relaxation sur les matrices d'effet

L'évolution du champ contenu dans une cavité en présence de relaxation est décrite par l'équation I.1.36. Le hamiltonien, proportionnel au nombre de photons, ne commute pas avec les sauts quantiques qui font varier ce dernier. On peut cependant rendre compte de l'évolution par Trotterisation : pour un temps court $\tau \ll T_{ci}, T_H$, où T_H est un temps caractéristique de l'évolution hamiltonienne, on peut négliger le commutateur entre le hamiltonien et les opérateurs de Kraus décrivant les sauts quantiques. On peut donc décrire l'évolution du champ dans la cavité i entre les instants 0 et t en concaténant des intervalles de durée τ durant lesquels on peut séparer l'évolution hamiltonienne de celle due aux opérateurs de Kraus :

$$\hat{\rho}(t) = [\mathbb{T}_\tau \circ \mathbb{R}_\tau]^{t/\tau} (\hat{\rho}(0)), \quad (\text{IV.1.20})$$

où les super-opérateurs $\mathbb{R}_\tau$ et $\mathbb{T}_\tau$ décrivent respectivement l'évolution hamiltonienne et la relaxation du champ et ont pour expressions :

$$\mathbb{R}_\tau(\hat{\rho}) = \hat{R}\hat{\rho}\hat{R}^\dagger, \quad (\text{IV.1.21a})$$

$$\mathbb{T}_\tau(\hat{\rho}) = \mathbb{T}_{1,\tau} \otimes \mathbb{T}_{2,\tau}(\hat{\rho}), \quad (\text{IV.1.21b})$$

$$\mathbb{T}_{i,\tau}(\hat{\rho}_i) = \hat{\rho}_i + \tau \sum_{\mu} \left(\hat{L}_{i,\mu} \hat{\rho}_i \hat{L}_{i,\mu}^\dagger - \frac{1}{2} \hat{\rho}_i \hat{L}_{i,\mu}^\dagger \hat{L}_{i,\mu} - \frac{1}{2} \hat{L}_{i,\mu}^\dagger \hat{L}_{i,\mu} \hat{\rho}_i \right), \quad (\text{IV.1.21c})$$

où $\hat{R}$ est l'opérateur de rotation lié à l'évolution hamiltonienne IV.1.6 et où les $\hat{L}_{i,\mu}$ sont les opérateurs de Kraus associés à perte d'un photon et à l'absorption d'un photon thermique, définis par les équations I.1.37. Les super-opérateurs adjoints associés à ces cartes quantiques ont pour expression :

$$\mathbb{R}_\tau^*(\hat{E}) = \hat{R}^\dagger \hat{E} \hat{R}, \quad (\text{IV.1.22a})$$

$$\mathbb{T}_\tau^*(\hat{E}) = \mathbb{T}_{1,\tau}^* \otimes \mathbb{T}_{2,\tau}^*(\hat{E}), \quad (\text{IV.1.22b})$$

$$\mathbb{T}_{i,\tau}^*(\hat{E}_i) = \hat{E}_i + \tau \sum_{\mu} \left(\hat{L}_{i,\mu}^\dagger \hat{E}_i \hat{L}_{i,\mu} - \frac{1}{2} \hat{E}_i \hat{L}_{i,\mu}^\dagger \hat{L}_{i,\mu} - \frac{1}{2} \hat{L}_{i,\mu}^\dagger \hat{L}_{i,\mu} \hat{E}_i \right). \quad (\text{IV.1.22c})$$

Le super-opérateur adjoint pour la rotation est simplement une rotation en sens inverse. En revanche, la décohérence agit de manière fort différente sur la matrice d'effet et la matrice densité. En effet, le super-opérateur $\mathbb{T}_{i,\tau}^*$ ne remplit pas les critères III.3.8 pour être une carte quantique : il ne préserve pas la trace. Il en résulte que la trace de la matrice d'effet n'est pas toujours égale à 1. Dans la suite, on va donner quelques idées sur la rétro-évolution suivie par la matrice d'effet.

Tout d'abord, on peut remarquer que l'identité est un vecteur propre de $\mathbb{T}_{i,\tau}^*$. Il en résulte que si la matrice d'information ne contient aucune information sur le champ à l'instant t , elle n'en contient pas plus à l'instant $t - \tau$. Si on ne fait pas de mesure, la matrice d'effet reste donc proportionnelle à l'identité. Si on fait des mesures, la décohérence entre ces mesures a aussi tendance à ramener la matrice d'effet vers l'identité. On ne donnera pas de preuve de ce résultat, qui nécessite certaines conditions sur les opérateurs de saut. On l'illustrera sur un exemple. La figure IV.7 montre l'évolution du nombre de photons moyen, de la trace et de la pureté de la matrice d'effet dans le cas où on a effectué une mesure projective sur l'état $|0\rangle$ à l'instant $t = T_c$. Dans ce cas, la matrice d'effet à la

fin de l'évolution est la matrice $|0\rangle\langle 0|$. La rétro-évolution fait tendre cette matrice vers une matrice proportionnelle à l'identité de l'espace tronqué. Pour que les résultats soient corrects, il faut donc choisir le nombre maximal de photons de sorte que les populations des états tronqués restent négligeables. Pour les temps considérés ici, se limiter à un nombre maximal de 9 photons permet d'avoir moins de 1% des populations dans les états tronqués. On voit alors que la trace et le nombre de photons croissent de manière accélérée au fur et à mesure que l'on peuple tous les états excités. La pureté en revanche décroît lors de la rétro-évolution, pour atteindre une valeur d'environ 22% après T_c . Il apparaît donc qu'une mesure effectuée plusieurs T_c après la préparation de l'état apporte peu d'information sur l'état initial du champ.

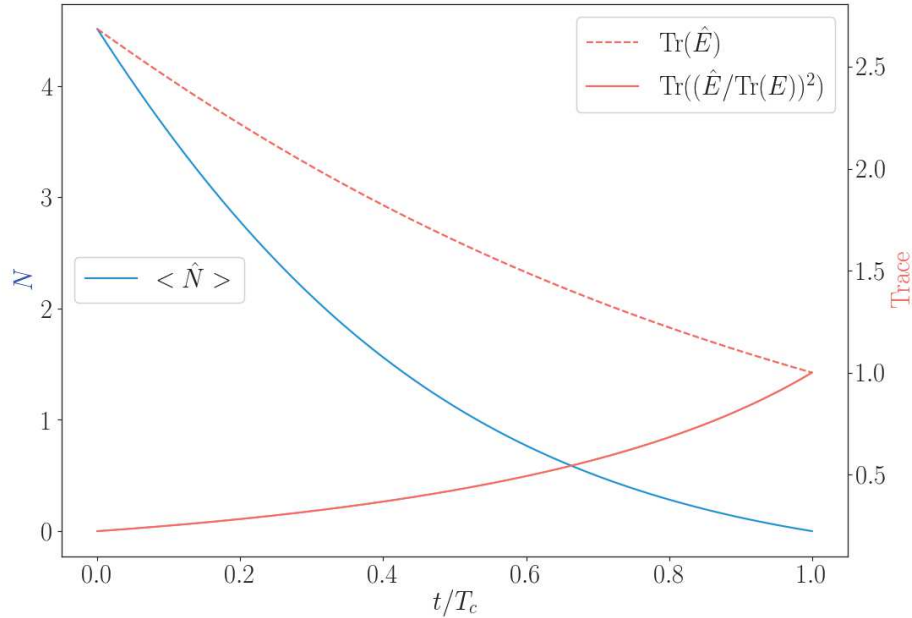


Fig. IV.7 Effet de la décohérence sur la matrice d'effet. La rétro-évolution a lieu de droite à gauche. Le champ est mesuré à l'instant $t/T_c = 1$ dans l'état $|0\rangle$. La matrice d'effet relaxe à partir de cet instant vers un état proportionnel à l'identité. L'espace de Hilbert est tronqué à 9 photons, de sorte que les populations de la matrice dans les états de Fock tronqués représentent moins de 1% de la somme des populations. Le nombre de photons thermiques est pris égal à 0,05. La courbe bleue montre la valeur moyenne du nombre de photons qui croît lorsqu'on remonte le temps. La courbe rouge pointillée représente la trace, qui n'est pas préservée par l'évolution. La courbe rouge pleine montre la pureté de la matrice d'effet, qui tend vers la pureté associée à l'identité normalisée pour l'espace tronqué.

IV.1.3 Prise en compte des imperfections de la détection

Les erreurs de détection, décrites dans le paragraphe I.4.2.c, doivent être prises en compte dans les matrices d'effet. Rappelons leur effet sur la mesure. Un échantillon atomique contient en moyenne μ atomes, suivant une loi de Poisson. La probabilité d'avoir n atomes dans un échantillon est donc donnée par :

$$P_a(n) = e^{-\mu} \frac{\mu^n}{n!}. \quad (\text{IV.1.23})$$

En pratique, on a utilisé au maximum 0.3 atomes par échantillon, ce qui donne une probabilité de 3.7% d'avoir deux atomes ou plus et une probabilité de 0,4% d'avoir 3 atomes ou plus. On ne considérera donc que la possibilité d'avoir deux atomes. Lorsque les trajectoires de mesures contiennent des échantillons atomiques contenant plus de deux atomes, on ne les prendra pas en compte. Pour une interaction $\nu \in \{res, dis\}$, et un résultat de mesure μ , l'évolution du champ est donnée par un super-opérateur $\mathbb{L}_\mu^\nu$ qui vaut :

$$\mathbb{L}_\emptyset^{[\nu]} = P_a(0) \mathbb{I}, \quad (\text{IV.1.24a})$$

$$\mathbb{L}_g^{[\nu]} = P_a(1) \mathbb{M}_g^{[\nu]}, \quad (\text{IV.1.24b})$$

$$\mathbb{L}_e^{[\nu]} = P_a(1) \mathbb{M}_e^{[\nu]}, \quad (\text{IV.1.24c})$$

$$\mathbb{L}_{gg}^{[\nu]} = P_a(2) \mathbb{M}_{gg}^{[\nu]}, \quad (\text{IV.1.24d})$$

$$\mathbb{L}_{eg}^{[\nu]} = 2P_a(2) \mathbb{M}_{eg}^{[\nu]}, \quad (\text{IV.1.24e})$$

$$\mathbb{L}_{ee}^{[\nu]} = P_a(2) \mathbb{M}_{ee}^{[\nu]}. \quad (\text{IV.1.24f})$$

On voit qu'on n'obtient aucune information si on ne mesure pas d'atome, tandis que si on mesure plusieurs atomes dans une certaine configuration, on obtient l'information correspondante. Les probabilités devant les super-opérateurs correspondant aux mesures montrent que le hasard classique donnant la distribution de Poisson a pour effet de rendre chaque résultat moins probable mais ne change pas l'information acquise lors d'une mesure donnée.

Par ailleurs, l'efficacité de détection, c'est à dire la probabilité de détecter un atome, vaut $\epsilon = 0,5$. À cause de cette imperfection, le nombre d'atomes détectés ne correspond pas nécessairement au nombre d'atomes réel. Le détecteur est aussi susceptible de déclencher des cascades électroniques en l'absence d'atomes, notamment à cause de phénomènes de photoémission. Ces phénomènes se produisent environ une fois par minute. Pour être pris en compte, ils doivent par ailleurs avoir lieu dans une fenêtre de détection atomique, large de quelques microsecondes. On peut donc les négliger. Le nombre d'atomes détectés est donc toujours inférieur ou égal au nombre d'atomes réel. Enfin, un atome dans $|e\rangle$ (resp. $|g\rangle$) a une probabilité η_e^d (η_g^d) d'être détecté dans $|g\rangle$ ($|e\rangle$).

Ces imperfections font que pour un résultat de mesure détecté μ' , il y a plusieurs résultats réels correspondants μ . La carte quantique correspondante est alors un mélange des cartes quantiques correspondant aux différents résultats de mesure possibles :

$$\mathbb{S}_{\mu'(\hat{\rho})}^\nu = \sum_\mu P(\mu'|\mu) \mathbb{L}_\mu^\nu(\hat{\rho}), \quad (\text{IV.1.25})$$

où $P(\mu'|\mu)$ est la probabilité d'obtenir μ' sachant que le vrai résultat de la mesure est μ . La matrice stochastique $P(\mu'|\mu)$ est donnée table IV.1.

$\mu' \setminus \mu$	$\emptyset$	g	e	gg	ee	ge
$\emptyset$	1	$1-\epsilon$	$1-\epsilon$	$(1-\epsilon)^2$	$(1-\epsilon)^2$	$(1-\epsilon)^2$
g	0	$\epsilon(1-\eta_g^d)$	$\epsilon\eta_e^d$	$2\epsilon(1-\epsilon)(1-\eta_g^d)$	$2\epsilon(1-\epsilon)\eta_e^d$	$\epsilon(1-\epsilon)(1-\eta_g^d+\eta_e^d)$
e	0	$\epsilon\eta_g^d$	$\epsilon(1-\eta_e^d)$	$2\epsilon(1-\epsilon)\eta_g^d$	$2\epsilon(1-\epsilon)(1-\eta_e^d)$	$\epsilon(1-\epsilon)(1-\eta_e^d+\eta_g^d)$
gg	0	0	0	$\epsilon^2(1-\eta_g^d)^2$	$\epsilon^2(\eta_e^d)^2$	$\epsilon^2\eta_e^d(1-\eta_g^d)$
ge	0	0	0	$2\epsilon^2\eta_g^d(1-\eta_g^d)$	$2\epsilon^2\eta_e^d(1-\eta_e^d)$	$\epsilon^2((1-\eta_g^d)(1-\eta_e^d)+\eta_g^d\eta_e^d)$
ee	0	0	0	$\epsilon^2(\eta_g^d)^2$	$\epsilon^2(1-\eta_e^d)^2$	$\epsilon^2\eta_g^d(1-\eta_e^d)$

TABLE IV.1 – Matrice stochastique donnant la probabilité d'obtenir le résultat μ' sachant μ .

IV.2 Reconstruction de l'état $|10\rangle + |01\rangle$

Dans cette section, on applique le protocole de reconstruction décrit dans le chapitre III et les méthodes de calcul de matrices d'effet de la section précédente pour effectuer la tomographie de l'état $|10\rangle + |01\rangle$. On s'attache à montrer les propriétés des matrices d'effet obtenues, ainsi que celles de la fonction de vraisemblance pour chaque reconstruction, ainsi que les propriétés de l'état reconstruit et les incertitudes sur ce dernier. On étudie d'abord le protocole de mesure avec un atome dispersif, puis la mesure de la fonction de Wigner. On montre ensuite que la combinaison de ces deux méthodes permet d'effectuer une tomographie. Enfin, on montre les résultats obtenus pour la mesure de plusieurs échantillons d'atomes résonnants.

IV.2.1 Mesure résonnante seule

Dans un premier temps, on applique la reconstruction à la mesure d'un seul atome résonnant décrite dans la section II. Pour cette mesure, on prépare l'état avec un atome injecteur, puis on envoie, après un délai variable T_0 , un seul échantillon d'atomes sondes. Comme le battement des deux cavités a une fréquence de 17 kHz, on effectue la mesure pour 17 valeurs de T_0 allant de 500 μ s à 660 μ s, ce qui permet d'échantillonner une période d'oscillation pour l'état $|10\rangle + |01\rangle$. Les matrices d'effet sont calculées avec la méthode présentée dans la section IV.1.1.a, en prenant en compte la décohérence entre la préparation et la mesure. On garde les trajectoires expérimentales pour lesquelles on a détecté un atome injecteur dans l'état $|g\rangle$. Par ailleurs, on ne garde que les trajectoires pour lesquelles on a détecté un seul atome sonde et on exclut celles où on a détecté un atome dans l'intervalle de détection entre celui de $|e\rangle$ et celui de $|g\rangle$. Les échantillons atomiques contiennent en moyenne $\mu = 0.28$ atomes. Les mesures ont été effectuées à la température de 0.8 K, pour laquelle les temps de vie valent $T_{c1} = 20$ ms et $T_{c2} = 50$ ms. Afin d'éviter les problèmes liés à la taille finie de l'espace de Hilbert, la reconstruction est effectuée dans un espace de dimension 8 pour chaque cavité, qu'on tronque ensuite à la dimension 5. On dispose de 361200 séquences exécutées, ce qui donne, après les filtrages décrits précédemment, 3913 répétitions utilisables pour lesquelles on calcule les matrices d'effet.

Puisqu'on a effectué les mesures pour 17 délais différents, on obtient 34 matrices d'effet différentes qui correspondent chacune à l'un des délais et à la mesure de l'atome dans $|e\rangle$ ou $|g\rangle$. Si un coefficient est nul pour toutes les matrices densité, alors le terme correspondant

de la matrice densité ne joue aucune rôle dans la log-vraisemblance, comme il apparaît sur l'expression III.3.44c. La figure IV.8 présente la matrice de sensibilité correspondant à cet ensemble de mesures. Les éléments de la matrice coloriés en noir sont ceux pour lesquels au moins l'une des matrices d'effet a un coefficient non nul. Les autres termes sont ceux pour lesquels toutes les matrices d'effet valent 0. La vraisemblance est donc indépendante de ces termes. Les matrices d'effet ont des termes non nuls pour les populations et pour les cohérences $|m+1, n-1\rangle\langle m, n|$. Cependant, on n'a que 34 matrices indépendantes alors qu'il y a 57 coefficients sondés. La mesure ne devrait donc pas permettre de déterminer avec certitude tous ces coefficients.

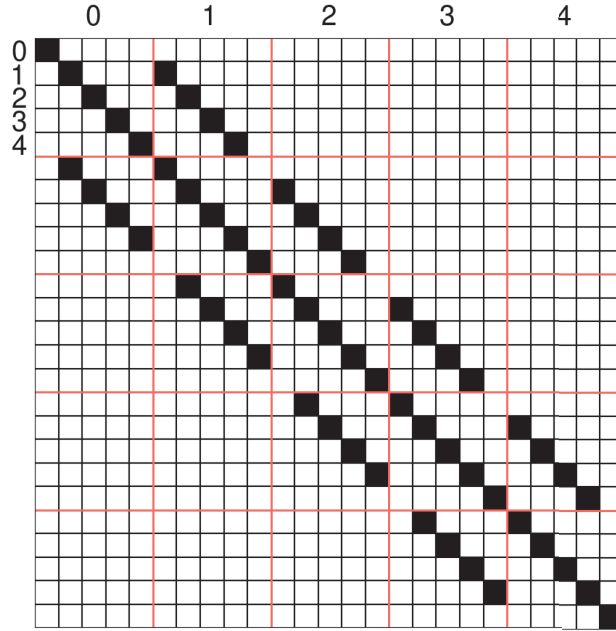


Fig. IV.8 Matrice de sensibilité de la mesure résonnante. Les emplacements de la matrice en noir sont ceux pour lesquels au moins une des matrices d'effet a un coefficient non nul, c'est-à-dire ceux sur lesquels la mesure apporte de l'information. Les autres emplacements correspondent à des coefficients de la matrice densité qui ne jouent aucun rôle dans la vraisemblance des résultats mesurés. Les mesures sont donc insensibles à ces coefficients. La seule limite apportée à leur valeur est donnée par les contraintes sur la matrice densité. Il apparaît que la mesure est en principe sensible aux populations et aux cohérences correspondant à l'échange d'un photon entre les deux cavités.

Résultats de la reconstruction

La matrice $\hat{\rho}_{\text{ML}}$ obtenue lors de la reconstruction est présentée figure IV.9. On peut d'emblée remarquer que cette matrice présente de nombreux termes non nuls, ainsi que des barres d'erreurs très importantes, à cause de l'indétermination qui vient d'être évoquée. Son rang est égal à 15 pour cette reconstruction. On peut donc l'écrire sous la forme d'une décomposition :

$$\hat{\rho}_{\text{ML}} = \sum_{i=1}^{15} p_i |\psi_i\rangle \langle \psi_i|, \quad (\text{IV.2.1})$$

où les $|\psi_i\rangle$ sont des états orthogonaux deux à deux. On note $|\psi_1\rangle$, $|\psi_2\rangle$ et $|\psi_3\rangle$ les trois états de plus grand poids dans cette décomposition. Ces états sont représentés figure IV.9. L'état $|\psi_4\rangle \equiv |0\rangle$ est aussi élément de la décomposition. Les poids associés respectivement à $|\psi_1\rangle$, $|\psi_2\rangle$, $|\psi_3\rangle$ et $|\psi_4\rangle$ sont 0.300, 0.230, 0.115 et 0.031. Ces populations n'ont pas de relation simple avec les log-vraisemblances associées à ces états qui valent respectivement -12590 , -12440 , -12410 et -14419 . L'état $|\psi_1\rangle$, proche de $|10\rangle + |01\rangle$ est celui qui a le poids le plus important, bien que sa vraisemblance soit plus faible que celle des deux états suivants dans la décomposition. La figure IV.10 permet de le comprendre. Elle présente le signal d'oscillation mesuré, ainsi que les signaux attendus pour les états $|\psi_1\rangle$, $|\psi_2\rangle$ et $|\psi_3\rangle$, ainsi que pour $\hat{\rho}_{\text{ML}}$. L'état $|\psi_1\rangle$ donne des oscillations dont le contraste est trop élevé pour rendre correctement compte du signal observé. Il n'est donc pas étonnant que sa vraisemblance soit plus faible que celle des états $|\psi_2\rangle$ et $|\psi_3\rangle$, qui approchent mieux les données. L'état $\hat{\rho}_{\text{ML}}$, qui combine tous ces états, est celui qui rend le mieux compte du signal observé.

Sensibilité de la vraisemblance à un déplacement

Il apparaît clairement sur la figure IV.10 que plusieurs combinaisons des différents états intervenant dans $\hat{\rho}_{\text{ML}}$ sont susceptibles de donner lieu aux mêmes oscillations. Les mesures sous-déterminent l'état, ce qui explique les barres d'erreur élevées sur les composantes de $\hat{\rho}_{\text{ML}}$. Une façon de le visualiser est de regarder les variations de la vraisemblance lorsqu'on modifie les poids associés à $|\psi_1\rangle$, $|\psi_2\rangle$ et $|\psi_3\rangle$. Dans ce but, on part de la matrice $\hat{\rho}_{\text{ML}}$, à laquelle on fait subir un déplacement, dans l'espace des opérateurs hermitiens. Afin de s'assurer que la matrice obtenue est toujours une matrice densité, ce déplacement est pris proportionnel à une matrice Λ de trace nulle et on projette l'état obtenu sur l'ensemble des matrices densité. On obtient alors l'état $\hat{\rho}_h^\Lambda$:

$$\hat{\rho}_h^\Lambda = \mathcal{P}_{\mathcal{D}(\mathcal{H})} (\hat{\rho}_{\text{ML}} + h\Lambda). \quad (\text{IV.2.2})$$

À titre d'exemple, on a considéré la façon dont la vraisemblance varie lorsqu'on fait des déplacements qui modifient la population p_{01} . La figure IV.11 montre les variations de la vraisemblance selon trois directions Λ_1 , Λ_2 et Λ_3 définies comme suit :

$$\Lambda_1 = |\psi_1\rangle \langle \psi_1| - |\psi_2\rangle \langle \psi_2|, \quad (\text{IV.2.3a})$$

$$\Lambda_2 = |\psi_1\rangle \langle \psi_1| - 1.46 |\psi_2\rangle \langle \psi_2| + 0.46 |\psi_3\rangle \langle \psi_3|, \quad (\text{IV.2.3b})$$

$$\Lambda_3 = |\psi_1\rangle \langle \psi_1| - |\psi_3\rangle \langle \psi_3|. \quad (\text{IV.2.3c})$$

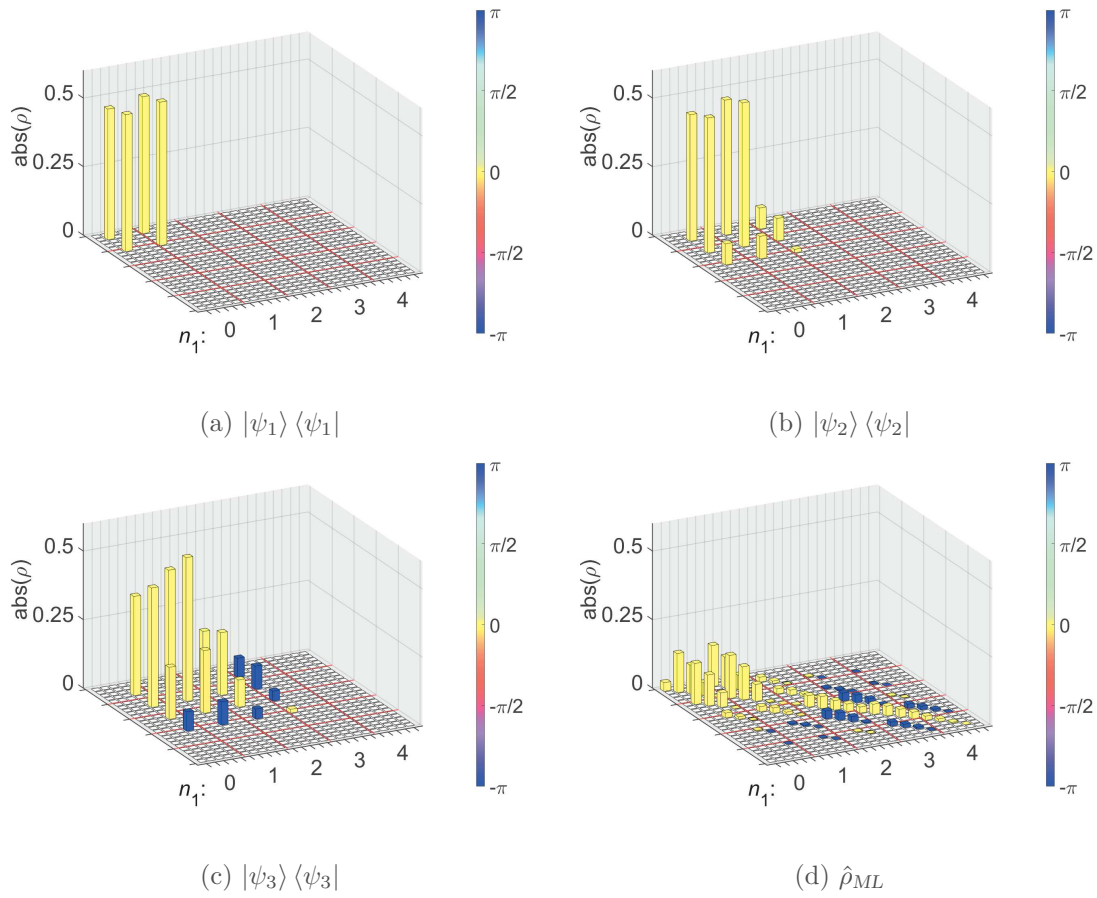


Fig. IV.9 Matrices densité de $\hat{\rho}_{ML}$ et des principales matrices entrant dans sa décomposition en états purs orthogonaux : $\hat{\rho}_{ML} = 0.300 |\psi_1\rangle\langle\psi_1| + 0.230 |\psi_2\rangle\langle\psi_2| + 0.115 |\psi_3\rangle\langle\psi_3| + \dots$

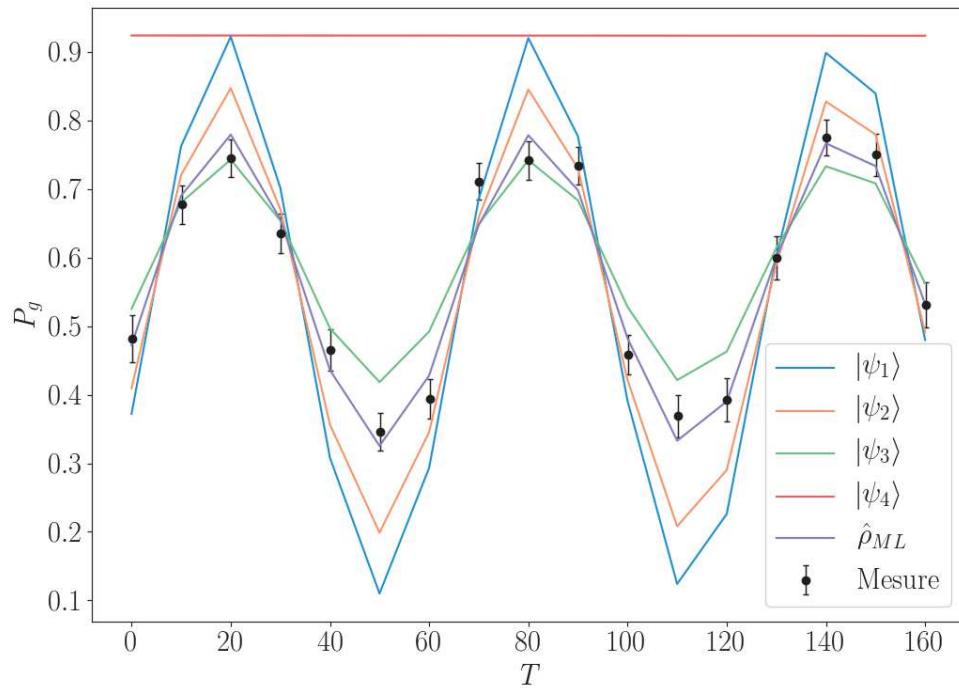


Fig. IV.10 Signal mesuré pour la mesure résonnante avec un seul atome et signaux attendus pour les états $|\psi_i\rangle \langle \psi_i|$ introduits précédemment et pour $\hat{\rho}_{ML}$. L'état $\hat{\rho}_{ML}$, qui est une somme pondérée des autres états, est celui qui ajuste le mieux les résultats mesurés.

Les poids sur les différents états apparaissant dans Λ_2 ont été choisis de sorte que la vraisemblance soit la plus plate possible dans la direction obtenue. Le fait que les variations soient plus faibles dans la direction donnée par Λ_2 s'explique de la manière suivante : lorsqu'on réduit le poids de $|\psi_1\rangle$ dans l'état et qu'on augmente celui de $|\psi_2\rangle$, le contraste des oscillations diminue. On contrebalance cet effet en diminuant aussi le poids de $|\psi_3\rangle$, qui donne des oscillations moins contrastées.

Pour obtenir une information quantitative sur la sensibilité dans la direction Λ , on s'intéresse à la courbure de la vraisemblance selon cette direction. On effectue un ajustement quadratique de la courbe obtenue en faisant varier h et on prend la racine carrée de l'inverse du coefficient du second ordre. On appelle S_Λ la valeur ainsi obtenue :

$$S_\Lambda = \sqrt{\left. \frac{d^2}{dh^2} \mathcal{L}(\hat{\rho}_{\text{ML}} + h\Lambda) \right|_{h=0}}. \quad (\text{IV.2.4})$$

Les régressions effectuées pour les trois directions précédentes sont présentées figure IV.12 et donnent les valeurs :

$$S_{\Lambda_1} = 0.183, \quad (\text{IV.2.5a})$$

$$S_{\Lambda_2} = 0.482, \quad (\text{IV.2.5b})$$

$$S_{\Lambda_3} = 0.153. \quad (\text{IV.2.5c})$$

Ces résultats ont le même ordre de grandeur que la déviation $\sqrt{\sigma^2(p_{01})}$ calculée en utilisant la formule III.4.7, qui vaut 0,704. Remarquons que cette dernière fait intervenir toutes les directions de l'ensemble des matrices densité, tandis que les valeurs estimées par les régressions précédentes ne concernent chacune qu'une direction.

Les régressions effectuées permettent aussi de donner un ordre de grandeur de l'incertitude liée aux conditions d'arrêt de notre algorithme. Le fait que les b_i sont non nuls indique qu'il existe des matrices densité de vraisemblance supérieure à celle de $\hat{\rho}_{\text{ML}}$ déterminée par notre algorithme. Dans la direction Λ_2 , le maximum est atteint pour $h = -0.012$. Le gain de vraisemblance n'est cependant que de $8 \cdot 10^{-3}$.

La courbe correspondant à Λ_1 présente un second maximum local $\hat{\rho}_s$ pour $h = 0.39$, de vraisemblance $\mathcal{L}(\hat{\rho}_s) = \mathcal{L}(\hat{\rho}_{\text{ML}}) - 0.13$, très proche de celle de $\hat{\rho}_{\text{ML}}$. Cet état a une population nulle dans $|10\rangle$ et $|01\rangle$ et une population de 0.61 dans $|\psi_2\rangle$. Le fait que la courbe de vraisemblance suivant Λ_1 n'est pas concave est dû à la projection qui fait que la courbe suivie en faisant varier h n'est pas une ligne droite. Si on regarde la vraisemblance suivant la ligne droite entre $\hat{\rho}_{\text{ML}}$ et $\hat{\rho}_s$, on retrouve la concavité.

Sensibilité vis-à-vis des cohérences

Les cohérences entre $|10\rangle$ et $|01\rangle$ sont les termes qui nous intéressent le plus dans la matrice densité. Dans le paragraphe précédent, on les faisait varier simultanément avec les populations p_{10} et p_{01} . Cependant, elles peuvent aussi varier indépendamment. On va donc effectuer la même étude que précédemment en suivant désormais la direction $\Lambda_c = |0\rangle\langle 1| + |1\rangle\langle 0|$. Comme précédemment, on déplace l'état puis on le reprojette sur $\mathcal{D}(\mathcal{H})$. La figure IV.13 présente la vraisemblance, de même que le rang et les termes pertinents de la matrice densité. Le rang de la matrice densité permet de visualiser les

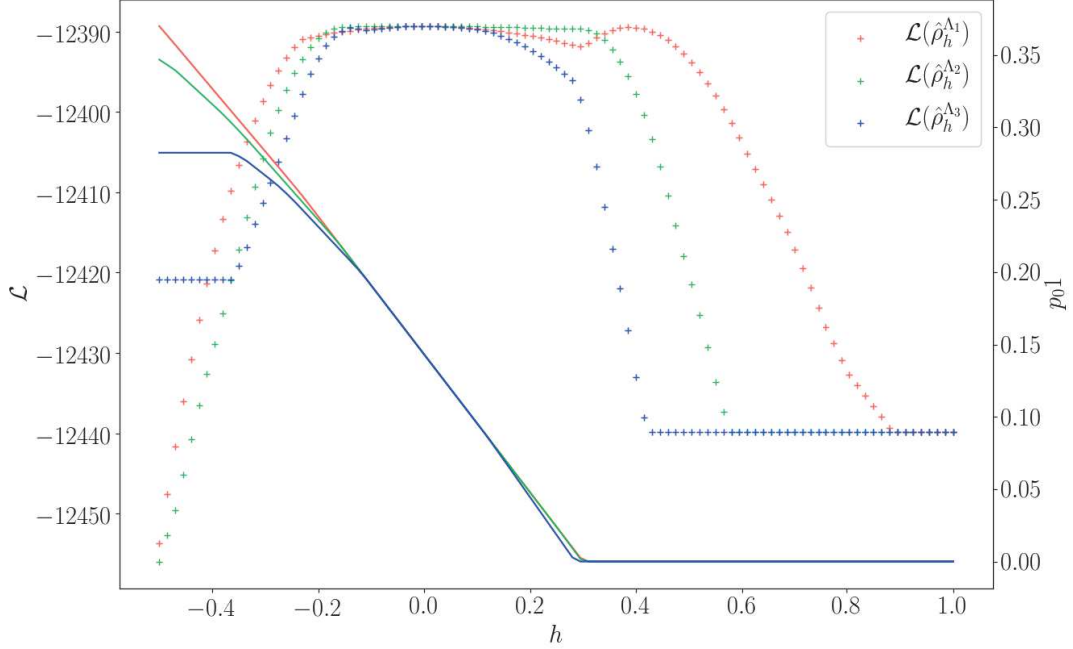


Fig. IV.11 Vraisemblance en fonction de l'amplitude h du déplacement pour les trois directions Λ_1 , Λ_2 et Λ_3 de l'ensemble des matrices hermitiennes de trace 1. Après déplacement, les matrices obtenues sont reprojctées sur $\mathcal{D}(\mathcal{H})$. Les déplacements Λ_1 , Λ_2 et Λ_3 correspondent respectivement aux courbes rouges, vertes et bleues. Les points donnent la vraisemblance, tandis que les courbes pleines donnent la valeur de la population p_{01} après projection. La projection fait que les populations n'évoluent pas linéairement avec le déplacement. Par exemple, pour h suffisamment grand, p_{01} est constante égale à 0. Le plateau de vraisemblance autour de la valeur $h = 0$, en particulier dans la direction Λ_2 , est responsable de l'indétermination de la population.

points où le déplacement fait passer d'une couche de $\mathcal{D}(\mathcal{H})$ de rang fixé à une autre couche de rang différent. En ces points, la vraisemblance et les coefficients de $\hat{\rho}_h^{\Lambda_c}$ ne sont pas réguliers, à cause des contraintes induites sur les coefficients qui rendent la courbe paramétrique de $(\hat{\rho}_h^{\Lambda_c})_h$ irrégulière. De manière intéressante, $\hat{\rho}_{ML}$ se situe exactement sur un de ces points, ce qui traduit le fait que les cohérences à un photon sont maximales. Il suffit donc de les réduire, en choisissant h négatif, arbitrairement petit, pour que le rang de $\hat{\rho}_h^{\Lambda_c}$ passe à 16. Pour $h \in]-0.28, 0[$, seules les cohérences évoluent dans la matrice densité et la pente de la courbe $\rho_{10,01}$ est exactement 1, en conséquence de l'absence de contraintes imposées par le bord de $\mathcal{D}(\mathcal{H})$. Lorsque h atteint -0.28 ou 0, il est nécessaire d'augmenter p_{01} et p_{10} pour que les cohérences puissent continuer à croître, ce qui donne lieu une rupture de pente pour $\rho_{10,01}$ et une diminution de p_{00} . On peut remarquer que toutes les courbes décrivant les coefficients de $\hat{\rho}_h^{\Lambda_c}$ sont symétriques par rapport au point

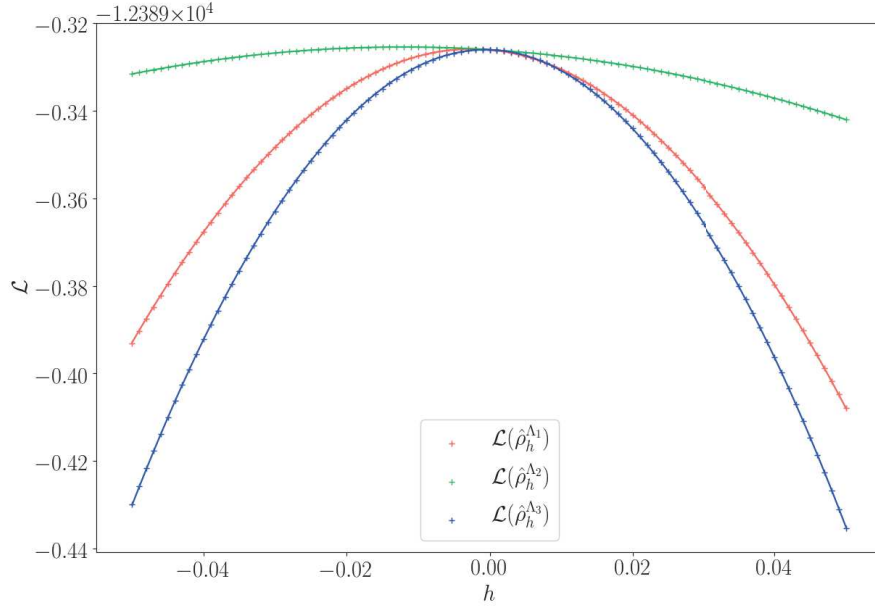


Fig. IV.12 Agrandissement de la figure IV.11 et ajustements quadratiques des vraisemblances. Les ajustements ont pour équation $y(h) = a_i h^2 + b_i h + \mathcal{L}\hat{\rho}_{ML}$ et on obtient $a_1 = -29.7$, $a_2 = -4.3$ et $a_3 = -42.6$ ainsi que $b_1 = -0.150$, $b_2 = -0.104$ et $b_3 = -0.051$.

pour lequel les cohérences s'annulent ($h=-0.15$). Ceci est dû au fait que l'ensemble des matrices densité est symétrique par rapport au changement de signe des cohérences. En revanche, la vraisemblance est plus grande pour $h > 0$ puisque le signal mesuré correspond à un état dont la phase est positive.

Contrairement à ce qu'on a fait précédemment, on ne peut pas effectuer d'interpolation quadratique de la vraisemblance autour de son maximum, à cause de l'absence de régularité causée par le changement de rang de $(\hat{\rho}_h^{\Lambda_c})_h$. On peut en revanche regarder pour quelles valeurs de la cohérence la log-vraisemblance est diminuée d'une unité par rapport à celle de $\hat{\rho}_{ML}$. On obtient alors un intervalle $[0.13, 0.17]$. Cet intervalle est bien plus étroit que ceux donnés pour les populations dans le paragraphe précédent. Ceci est dû au fait qu'on fait varier les cohérences $|10\rangle\langle 01|$, ce qui modifie le signal d'interférence attendu, sans modifier les autres cohérences. Comme on ne modifie pas les autres cohérences pour contrebalancer cet effet, comme on le faisait précédemment, la vraisemblance chute plus abruptement.

Sensibilité par rapport à la phase des cohérences.

La propriété la plus importante de l'état $|10\rangle + e^{i\varphi}|01\rangle$, dans le cadre de l'interférométrie, est la phase φ de la superposition. La courbe verte de la figure IV.14 montre l'évolution de la vraisemblance lorsqu'on transforme l'état $\hat{\rho}_{ML}$ en l'état $\hat{\rho}_\theta^{10} \equiv \hat{R}_{10}(\theta) \hat{\rho}_{ML} \hat{R}_{10}^\dagger(\theta)$,

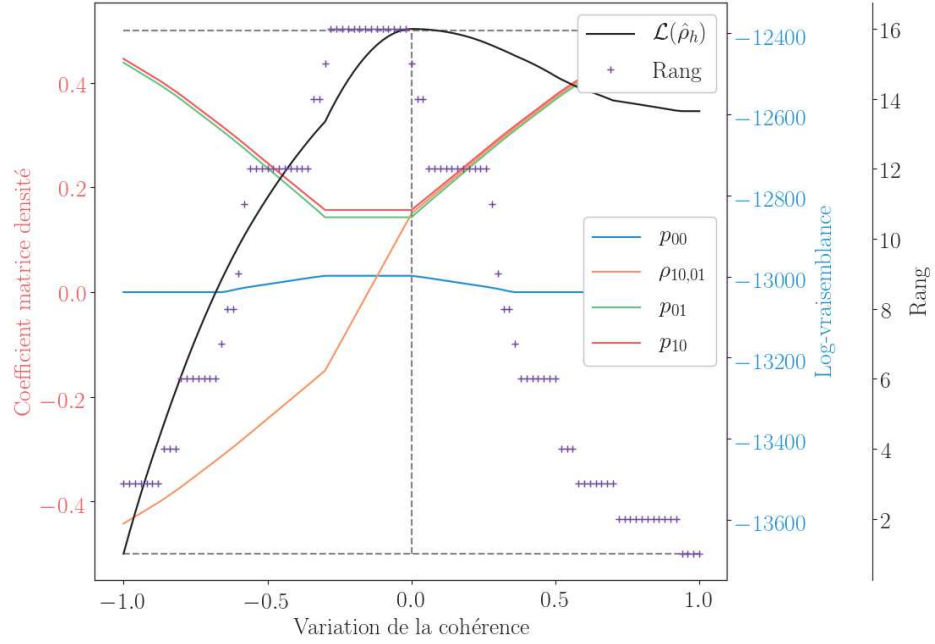


Fig. IV.13  volution du rang et des coefficients de la matrice densit , et de la vraisemblance lorsqu'on d place l' tat suivant la direction Λ_c , puis qu'on projette l' tat sur $\mathcal{D}(\mathcal{H})$. Le trait vertical en pointill  indique la valeur de h correspondant   $\hat{\rho}_{ML}$ obtenu pour les mesures r sonnantes. Les traits en pointill  horizontaux indiquent les valeurs extr males des coh rences.

o 

$$\hat{R}_{10}(\theta) \equiv \sum_{\substack{(n_1, n_2) \neq (1,0) \\ (n_1, n_2) \neq (0,1)}} |n_1, n_2\rangle \langle n_1, n_2| + e^{i\theta} |0, 1\rangle \langle 0, 1| + e^{-i\theta} |1, 0\rangle \langle 1, 0| \quad (\text{IV.2.6})$$

L'application de $\hat{R}_{10}(\theta)$ permet de changer la phase des coh rences entre $|10\rangle$ et $|01\rangle$ uniquement. La diff rence de log-vraisemblance entre l' tat $\hat{\rho}_{ML}$ et l' tat obtenu en mettant la phase oppos e sur les coh rences $\rho_{10,01}$ vaut 229. En effectuant une r gression parabolique autour du maximum, on trouve un coefficient du second ordre $a^{(10)}/2 = -63$, ce qui correspond   un  cart-type sur la phase de 0.09 rad. On peut cependant remarquer que l' tat $\hat{\rho}_{ML}$ contient aussi de l'information sur la phase de l'oscillation dans ses autres coh rences. La courbe rouge de la figure IV.14 montre la vraisemblance associ e aux  tats $\hat{\rho}_\theta \equiv \hat{R}(\theta)\hat{\rho}_{ML}\hat{R}^\dagger(\theta)$, obtenus en effectuant une rotation globale des coh rences. Cette rotation correspond   l' volution hamiltonienne de l' tat. C'est donc elle qui nous donne la sensibilit  du signal   la phase du signal d'oscillation des atomes r sonnantes. La diff rence de log-vraisemblance entre l' tat $\hat{\rho}_{ML}$ et l' tat de phase oppos e vaut 810. En effectuant une r gression parabolique autour du maximum, on trouve un coefficient au second ordre

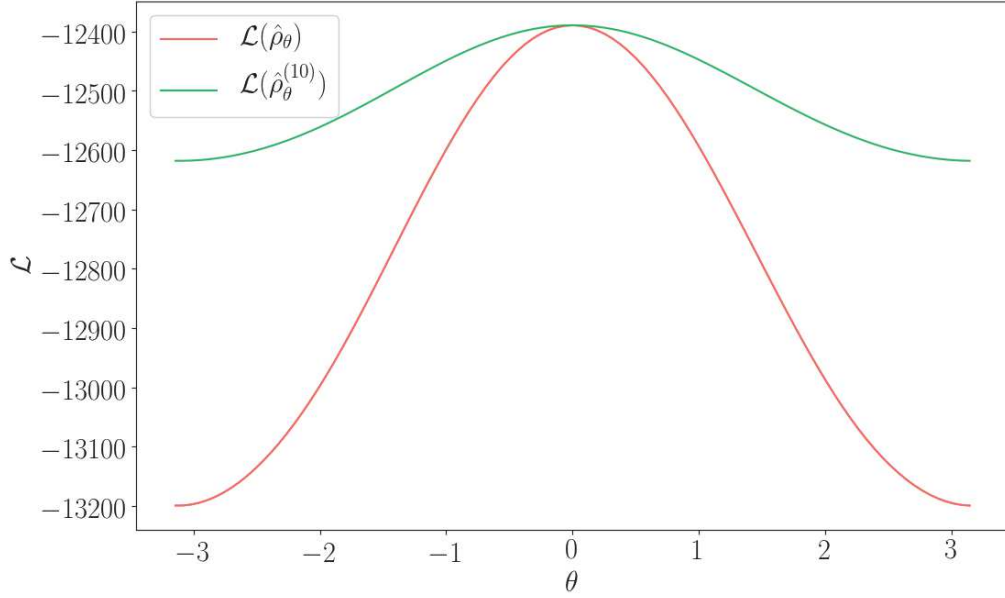


Fig. IV.14 Variations de la vraisemblance de l'état lorsqu'on modifie la phase de ses cohérences. Pour la courbe verte, on modifie la phase des cohérences $\rho_{10,01}$ uniquement. Pour la courbe rouge, on effectue une rotation globale de toutes les cohérences de l'état. La régression parabolique ($y = a/2\theta^2 + b\theta + \mathcal{L}(\hat{\rho}_{ML})$) a pour coefficients $a^{(10)}/2 = -63$ et $b^{(10)} = 0.27$ pour la courbe verte et $a/2 = -222$ et $b = 0.90$ pour la courbe rouge.

$a/2 = -222$, ce qui correspond à un écart-type sur la phase de 0.05 rad. On voit donc que la sensibilité de la mesure à la phase de l'oscillation de P_g est plus importante que ce qui est encodé dans les cohérences $\rho_{10,01}$ seulement.

IV.2.2 Mesures locales de la fonction de Wigner

Pour les mesures locales de la fonction de Wigner, on prépare initialement l'état à l'aide d'un atome injecteur. On déplace ensuite le champ dans l'une des deux cavités. Les mesures dispersives sont enfin effectuées avec 40 échantillons atomiques pour lesquels l'atome subit deux impulsions $\pi/2$ dans les zones de Ramsey R_1 et R_3 . Les phases par photon valent 3.23 ± 0.1 rad dans la cavité 1 et 3.32 ± 0.1 rad dans la cavité 2, afin de réaliser une mesure de parité. La différence de phase entre ces impulsions est choisie de sorte que le contraste soit maximum entre la valeur de P_g correspondant à 0 photon et celle correspondant à 1 photon. L'échantillon préparant l'état contient en moyenne 0.21 atome réel pour la mesure dans la cavité 1 et 0.25 pour celle dans la cavité 2. Les échantillons d'atomes utilisés pour la mesure contiennent en moyenne 0.11 atome réel pour la mesure dans C_1 et 0.14 pour celle dans C_2 . Les mesures ont été effectuées à la température de 0.8 K, pour laquelle les temps de vie valent $T_{c1} = 20$ ms et $T_{c2} = 50$ ms. Afin d'éviter les problèmes liés à la taille finie de l'espace de Hilbert, la reconstruction est effectuée dans

un espace de dimension 8 pour chaque cavité, qu'on tronque ensuite à la dimension 5. La figure IV.15 montre les amplitudes des déplacements utilisés. On a exclu les déplacements d'amplitude complexe supérieure à 2 pour lesquels la troncature de l'espace de Hilbert donne des matrices d'effet incorrectes. On dispose de 124000 réalisations avec injection dans C_1 et de 120200 réalisations avec injection dans C_2 . À l'exception des cas où on a moins de deux atomes sur l'ensemble des échantillons, les mesures sont toutes différentes, à cause du très grand nombre de combinaisons de résultats possibles. On ne peut donc plus regrouper les mesures et il faut prendre chaque mesure en compte séparément. On ne garde que les réalisations pour lesquelles on a détecté un unique atome injecteur dans $|g\rangle$. On dispose alors de 5634 mesures avec injection dans C_1 et de 6535 mesures avec injection dans C_2 . En combinant ces deux ensembles de mesures, on dispose donc de 12169 matrices d'effet. La figure IV.16a montre la matrice de sensibilité de la mesure de Wigner locale. Celle-ci est sensible aux populations et aux cohérences locales.

Par rapport au calcul des matrices d'effet présenté dans la section IV.1, la réalisation expérimentale pour l'état préparé présente une difficulté. En effet, la préparation de l'état ne faisant intervenir aucune injection micro-onde, il n'y a pas de référence de phase pour les injections de mesure. La phase de chaque mesure est donc aléatoire et on n'a aucun moyen d'y accéder. Ceci n'est en principe pas gênant pour déterminer notre état, qui est invariant par rotation de phase dans une cavité, lorsque l'amplitude d'injection dans l'autre cavité est nulle.

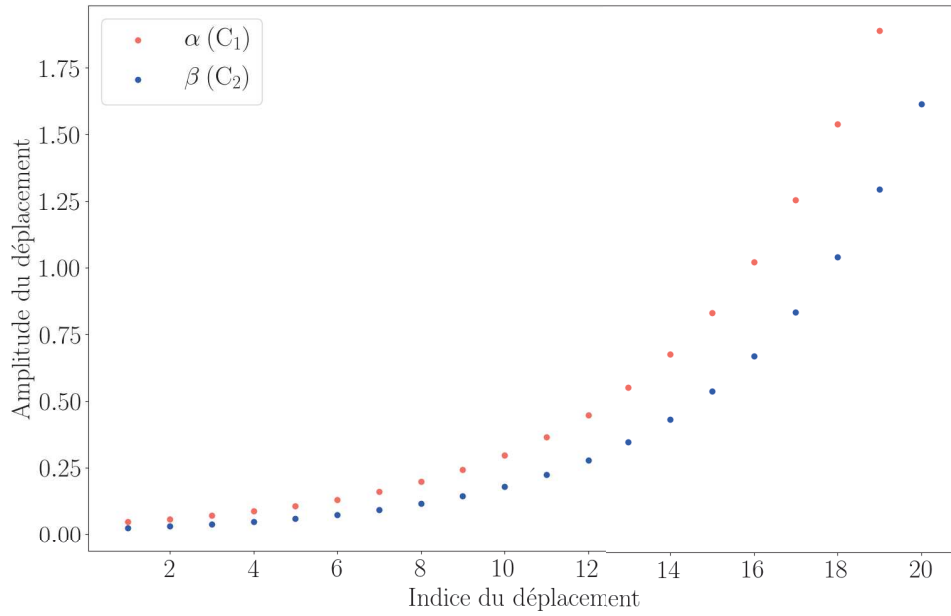
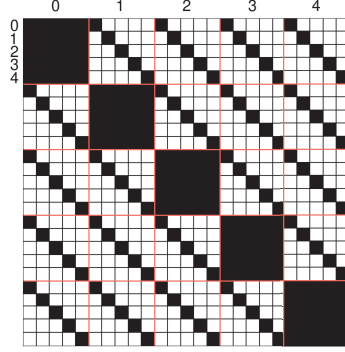
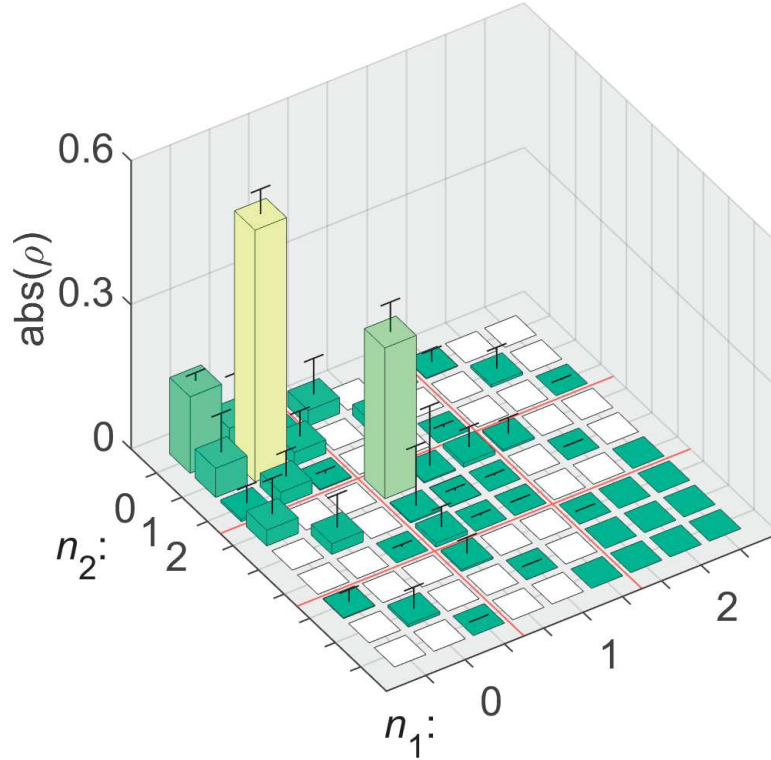


Fig. IV.15 Amplitudes des déplacements effectués dans la première cavité, en rouge et dans la deuxième cavité, en bleu.



(a) Matrice de sensibilité de la mesure locale de Wigner. Il apparaît que la mesure est en principe sensible aux populations et aux cohérences locales.



(b) Matrice $\hat{\rho}_{ML}$ obtenue en utilisant les mesures locales de la fonction de Wigner. Les principaux termes sont $p_{00} = 0.21 \pm 0.02$, $p_{01} = 0.48 \pm 0.05$ et $p_{10} = 0.32 \pm 0.06$. Les cohérences obtenues sont dues au bruit sur la mesure et dépendent des phases aléatoires choisies pour l'injection lors de la reconstruction. Si on répète la reconstruction pour un grand nombre de tirage de ces phases, elles se moyennent à zéro. Les barres d'erreur ont été calculées avec la formule III.4.15. On ne montre pas les barres d'erreur sur les éléments en blanc, auxquels la mesure est insensible.

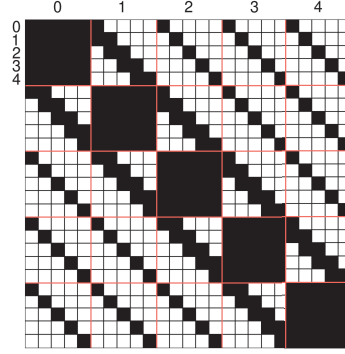
Fig. IV.16 Résultats de la reconstruction avec les mesures de Wigner locales.

Résultats de la reconstruction

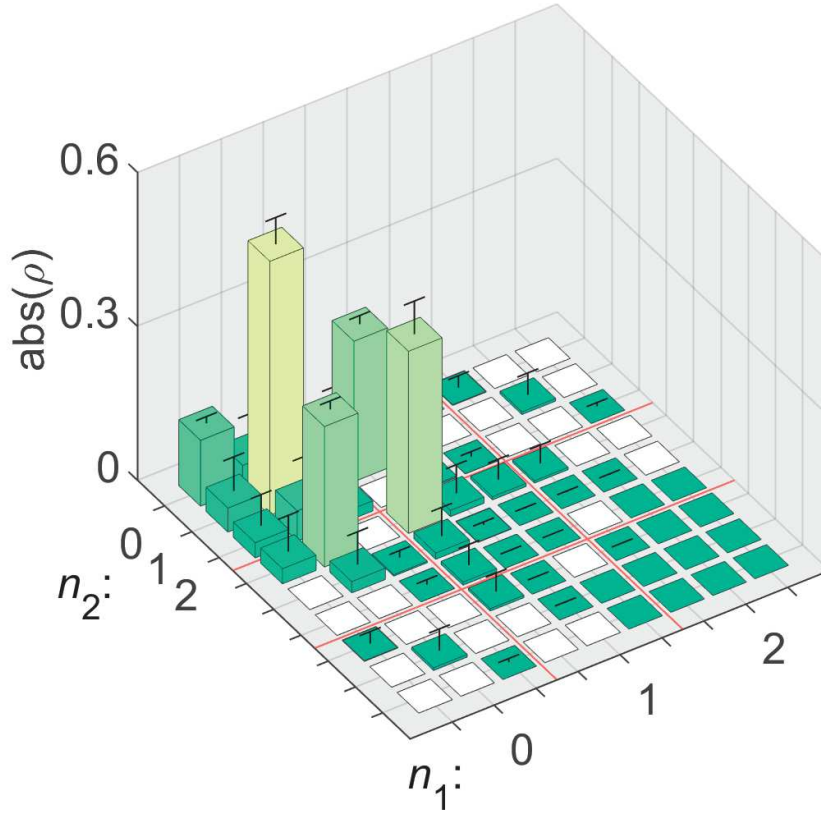
La matrice densité $\hat{\rho}_{\text{ML}}$ obtenue pour l'ensemble des mesures locales de la fonction de Wigner est montrée figure IV.16b.

IV.2.3 Mesure complète

La reconstruction par maximum de vraisemblance permet d'agréger des données résultant de mesures différentes. On peut donc combiner les résultats des mesures résonnantes et dispersives. On a alors 16082 matrices d'effet. La figure IV.17a montre la matrice de sensibilité de cette reconstruction. La matrice $\hat{\rho}_{\text{ML}}$ obtenue est présentée figure IV.17b. Elle est de rang 3. La figure IV.18 présente trois matrices $\hat{\rho}_1$, $\hat{\rho}_2$ et $\hat{\rho}_3$ permettant de décomposer $\hat{\rho}_{\text{ML}}$ sous la forme $\hat{\rho}_{\text{ML}} = 0.714 \hat{\rho}_1 + 0.186 \hat{\rho}_2 + 0.099 \hat{\rho}_3$. $\hat{\rho}_1$ a une fidélité de 0.986 avec $|10\rangle + |01\rangle$. La fidélité entre $\hat{\rho}_{\text{ML}}$ et $|10\rangle + |01\rangle$ vaut quant à elle 0.836. La phase de l'état a été arbitrairement ramenée à 0.



(a) Matrice de sensibilité lorsqu'on combine mesure résonnante et mesures de Wigner locales.



(b) Matrice $\hat{\rho}_{ML}$ obtenue en agrégeant les données des mesures résonnantes et dispersives. La matrice est de rang 3. Les principaux termes sont $p_{00} = 0.13 \pm .013$, $p_{01} = 0.50 \pm 0.05$, $p_{10} = 0.35 \pm 0.06$ et $\rho_{10,01} = 0.27 \pm 0.02 \pm 0.02i$. Les barres d'erreur ont été calculées avec la formule [III.4.15](#).

Fig. IV.17 Résultats de la reconstruction avec les atomes résonnants et les mesures de Wigner locales.

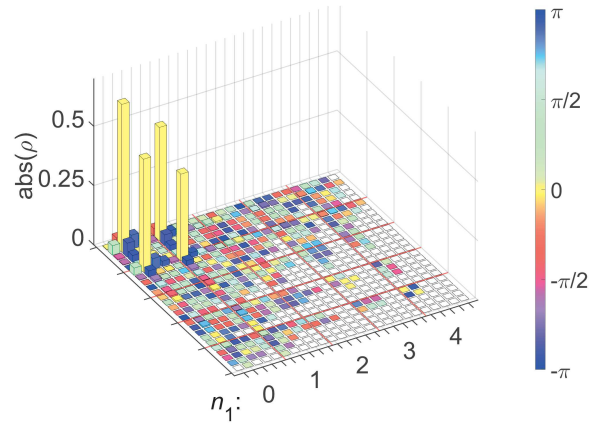
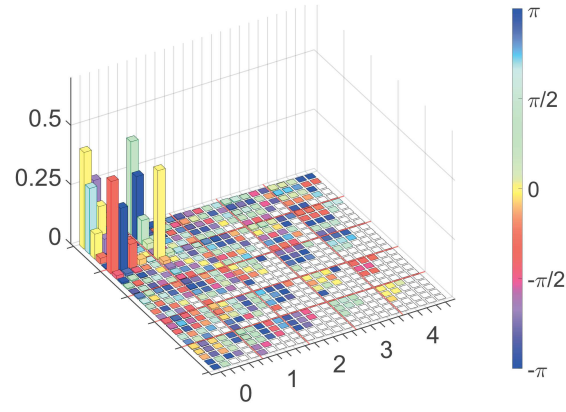
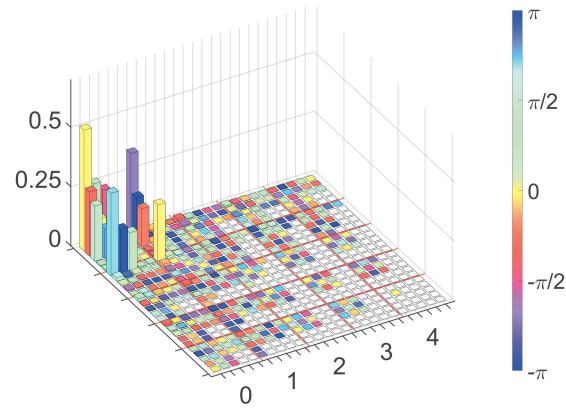
(a) $|\psi_1\rangle\langle\psi_1|$ (b) $|\psi_2\rangle\langle\psi_2|$ (c) $|\psi_3\rangle\langle\psi_3|$

Fig. IV.18 Principales matrices entrant dans la décomposition de $\hat{\rho}_{ML}$ en états purs orthogonaux.

La figure IV.19 est l'équivalent de la figure IV.13 dans le cas où on prend en compte l'ensemble des mesures. Elle montre l'évolution de l'état et de la fidélité lorsqu'on fait varier la cohérence $|10\rangle\langle 01|$, puis qu'on reprojette la matrice obtenue sur l'ensemble des matrices densité. Comme précédemment, la matrice est située au bord de l'ensemble des matrices densité et son rang augmente de 3 à 4 lorsqu'on réduit la cohérence $|10\rangle\langle 01|$. L'intervalle de cohérences pour lequel la vraisemblance diminue de moins d'une unité par rapport à son maximum est $[0.26, 0.29]$.

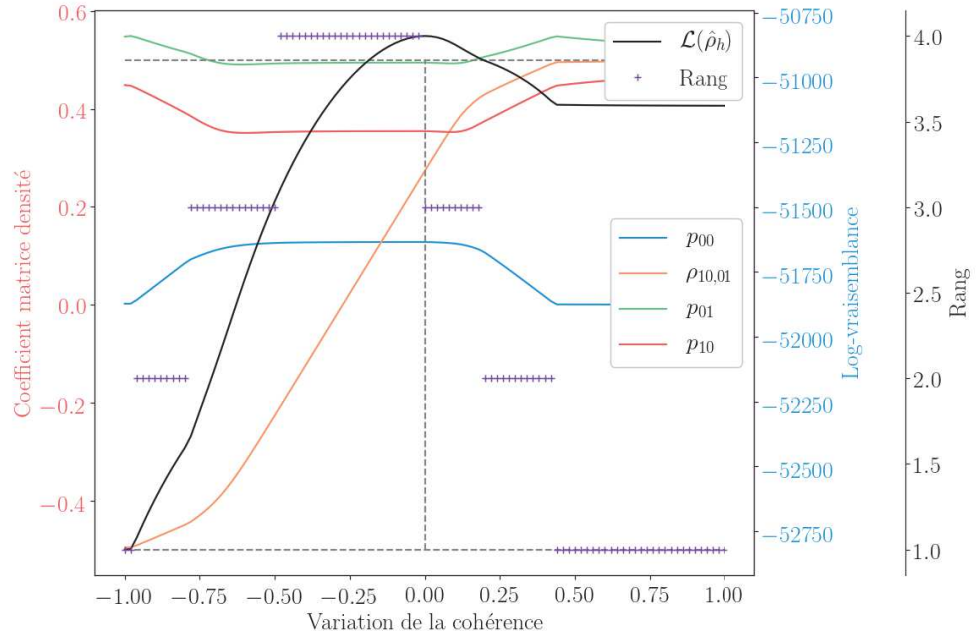


Fig. IV.19 Évolution du rang et des coefficients de la matrice densité, et de la vraisemblance lorsqu'on déplace l'état suivant la direction Λ_c , puis qu'on projette l'état sur $\mathcal{D}(\mathcal{H})$. Le trait vertical en pointillé indique la valeur de h correspondant à $\hat{\rho}_{ML}$ obtenu pour l'ensemble des mesures résonnantes et dispersives. Les traits en pointillé horizontaux indiquent les valeurs extrémales des cohérences.

Sensibilité par rapport à la phase des cohérences.

La figure IV.20 présente la vraisemblance obtenue lorsqu'on fait varier la phase des cohérences de l'état reconstruit, de la même manière que ce qui a été fait pour la figure IV.14.

La courbe verte correspond à une modification de la phase de la cohérence $\rho_{10,01}$ uniquement. L'écart entre la vraisemblance de l'état reconstruit et celui obtenu pour $\theta = \pi$ vaut 770. En effectuant une régression parabolique autour du maximum, on trouve

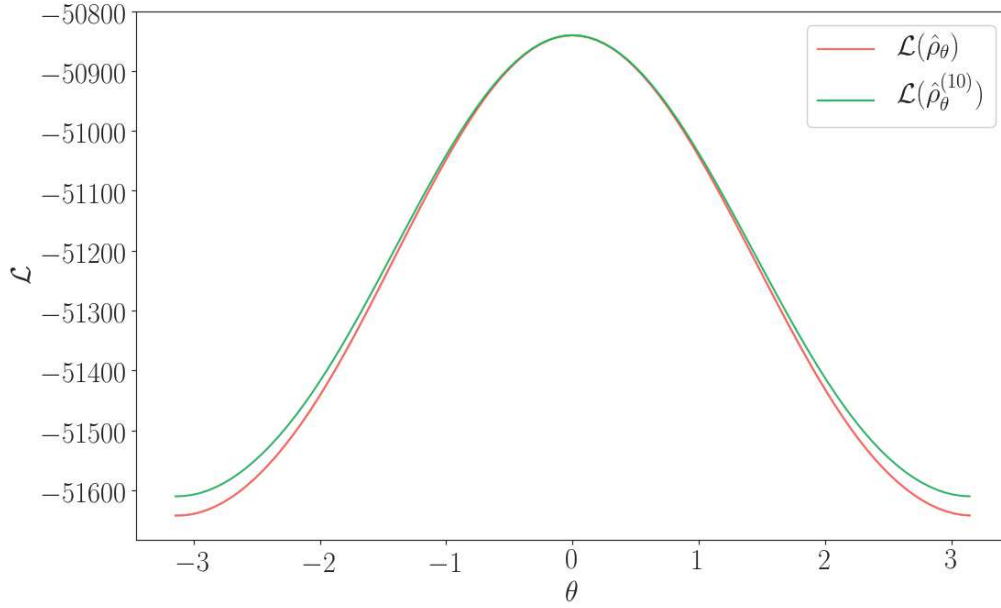


Fig. IV.20 Variations de la vraisemblance de l'état lorsqu'on modifie la phase de ses cohérences. Pour la courbe verte, on modifie la phase des cohérences $\rho_{10,01}$ uniquement. Pour la courbe rouge, on effectue une rotation globale de l'état. La régression parabolique ($y = a/2\theta^2 + b\theta + \mathcal{L}(\hat{\rho}_{ML})$) a pour coefficients $a^{(10)}/2 = -212$ et $b^{(10)} = 0.37$ pour la courbe verte et $a/2 = -218$ et $b = 0.88$ pour la courbe rouge.

un coefficient du second ordre $a^{(10)}/2 = -212$, ce qui correspond à un écart-type sur la phase de 0.05 rad. On peut cependant remarquer que l'état $\hat{\rho}_{ML}$ contient aussi de l'information sur la phase de l'oscillation dans ses autres cohérences. La courbe rouge montre la vraisemblance associée aux états $\hat{\rho}_\theta \equiv \hat{R}(\theta)\hat{\rho}_{ML}\hat{R}^\dagger(\theta)$, obtenus en effectuant une rotation globale d'angle θ de l'état. La différence de log-vraisemblance entre l'état $\hat{\rho}_{ML}$ et l'état de phase opposée vaut 802. En effectuant une régression parabolique autour du maximum, on trouve un coefficient au second ordre $a/2 = -218$, ce qui correspond à un écart-type sur la phase de 0.05 rad. Les courbes verte et rouge sont très proches : la prise en compte des atomes QND permet de ramener toute l'information sur la phase de l'oscillation dans la cohérence $\rho_{10,01}$. Par comparaison avec ce qu'on obtenait pour les atomes résonnants seuls, on voit que la courbure de la vraisemblance lorsqu'on effectue une rotation globale des cohérence est similaire : les atomes QND n'ont pas apporté d'information supplémentaire sur la phase de l'oscillation.

IV.2.4 Reconstruction avec plusieurs atomes résonnants

On a vu que la mesure résonnante ne suffit pas à déterminer combien de photons la cavité contient, ceci à cause des différentes possibilités donnant lieu au même signal et à cause du fait que l'atome peut emporter au plus un photon. En envoyant plusieurs

atomes sondes, on peut contourner ces difficultés. En effet, si les atomes sont préparés dans $|g\rangle$ et si les cavités contiennent au plus un photon, on ne peut trouver au plus qu'un atome excité. Plus généralement, le nombre d'atomes trouvés dans l'état excité dans une même trajectoire nous renseigne sur la distribution de photons, comme on l'a vu dans la section IV.1.1.b. Ceci nous apporte un éclairage intéressant sur la reconstruction. Même si la mesure résonnante avec un atome sonde est en principe la bonne manière d'étudier notre état, il est toujours avantageux d'ajouter, après celle-ci, d'autres mesures pour extraire davantage d'information. Ces mesures n'ont pas besoin d'être des mesures non destructives, tant qu'on est capable de décrire l'effet de la mesure avec des opérateurs de Kraus.

Dans cette section, on utilise cette stratégie pour reconstruire l'état. L'état est préparé comme précédemment par un atome injecteur issu d'un échantillon contenant en moyenne 0.24 atomes. On ne garde que les trajectoires pour lesquelles l'injecteur est détecté dans $|g\rangle$. On envoie ensuite, après environ 700 μs , 40 échantillons d'atomes qui effectuent les mêmes impulsions que celles utilisées précédemment pour la mesure résonnante. Ces échantillons sont séparés de 200 μs ³. Ils contiennent en moyenne 0.29 atomes détectés dans $|e\rangle$ ou $|g\rangle$ ⁴. On effectue la reconstruction en prenant en compte l'efficacité $\epsilon = 0.5$ du détecteur, les erreurs de mesure $\eta_e^d = 0.09$ et $\eta_g^d = 0.075$, la décohérence entre la préparation et la détection ainsi qu'entre les échantillons atomiques, à la température de 1.6 K.

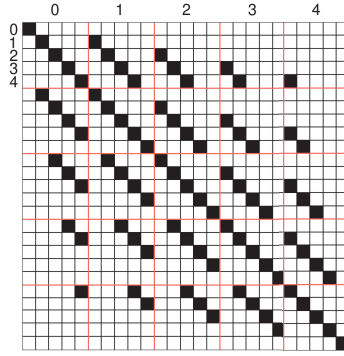
Le temps T_0 d'attente entre la préparation et le premier atome résonnant prend 21 valeurs régulièrement espacées entre 700 μs et 800 μs . On dispose de 237000 trajectoires expérimentales. On ne garde que les trajectoires pour lesquelles un seul atome injecteur a été détecté dans $|g\rangle$. Après cette sélection, il reste 18460. Avec le grand nombre d'échantillons atomiques, la probabilité de trouver plusieurs atomes dans le même échantillon est élevé. Dans notre cas, il y a 8134 trajectoires qui contiennent au moins un échantillon dans lequel on a détecté au moins deux atomes. On doit donc tenir compte des interactions à plusieurs atomes. Les opérateurs de Kraus permettant de décrire l'interaction résonnante de deux atomes avec la cavité sont décrits dans l'annexe D. Lorsqu'un échantillon contient trois atomes ou plus, on a choisi, pour simplifier le code, de le traiter comme un échantillon dans lequel on a détecté un atome dans $|e\rangle$ et un atome dans $|g\rangle$, après s'être assurés que les résultats n'étaient pas altérés, par rapport à une situation où l'échantillon n'est pas pris en compte. Dans notre jeu de données, 703 trajectoires sont concernées par cette situation.

La figure IV.21a montre la matrice de sensibilité pour la mesure avec des multiples échantillons résonnants. Utiliser plusieurs atomes permet d'être sensible à des cohérences pour lesquelles plusieurs photons sont échangés entre les cavités. Nous utiliserons cette propriétés dans la section suivante pour tenter de reconstruire l'état $|20\rangle + |02\rangle$.

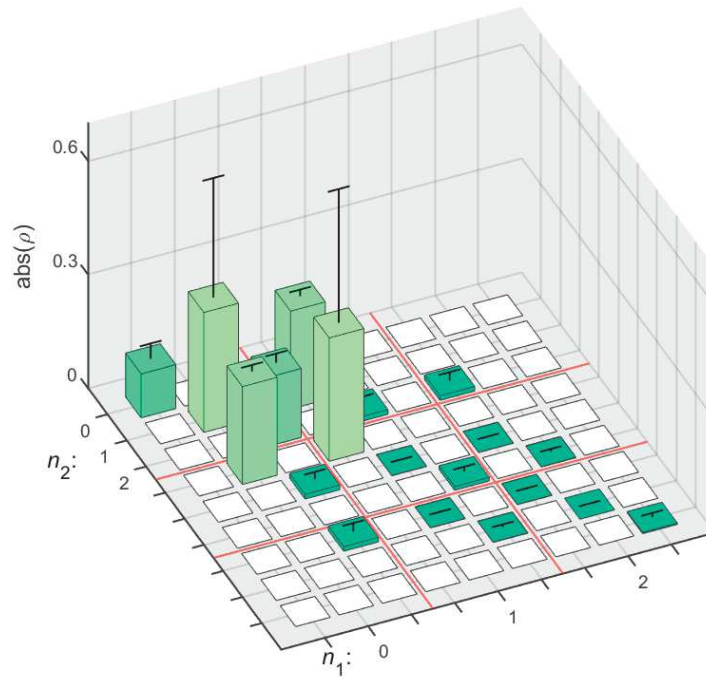
La matrice densité obtenue après reconstruction est montrée figure IV.21b. Les principaux termes sont $p_{00} = 0.09$, $p_{01} = 0.36$, $p_{02} = 0.09$, $p_{10} = 0.40$, $|\rho_{10,01}| = 0.26$. Elle est

3. On est obligé de laisser plus de temps entre les échantillons d'atomes résonnants qu'entre des échantillons d'atomes dispersifs. En effet si les atomes résonnants sont trop proches, le potentiel appliqué pour contrôler la fréquence de transition d'un atome peut perturber l'atome adjacent.

4. Pour simplifier, on ne prend pas en compte les atomes détectés dans l'intervalle situé entre les intervalles de détection de $|e\rangle$ et $|g\rangle$, ce qui arrive dans 2 % des cas.



(a) Matrice de sensibilité de la mesures avec de multiples atomes résonnants. Utiliser plusieurs atomes permet d'être sensible à plus de cohérences.



(b) Matrice densité $\hat{\rho}_{ML}$ obtenue par reconstruction avec plusieurs échantillons d'atomes résonnants. Les principaux termes valent $p_{00} = 0.12 \pm 0.04$, $p_{01} = 0.32 \pm 0.32$, $p_{10} = 0.33 \pm 0.35$, $p_{02} = 0.18 \pm 0.02$ et $\rho_{10,01} = 0.16 \pm 0.02 + (0.20 \pm 0.02)i$. Les barres d'erreur ont été calculées avec la formule III.4.15.

Fig. IV.21 Résultats de la reconstruction avec de multiples échantillons d'atomes résonnants.

de rang 13. On peut la développer sous la forme :

$$\hat{\rho}_{\text{ML}} = \sum_{i=1}^{13} p_i |\psi_i\rangle \langle \psi_i|. \quad (\text{IV.2.7})$$

Les matrices densité des principaux états intervenant dans cette décomposition sont données sur la figure IV.21b. Les états $|\psi_1\rangle$ et $|\psi_2\rangle$ sont assez proches respectivement de $|10\rangle + e^{i\phi}|01\rangle$ et $|10\rangle + e^{i\phi+\pi}|01\rangle$ avec $\phi = -0.90$ rad. L'état $|\psi_3\rangle$ est l'état vide $|00\rangle$. L'état $|\psi_4\rangle$ est proche de l'état $|0, 2\rangle$. Les poids qui leur sont associés sont $p_1 = 0.65$, $p_2 = 0.12$, $p_3 = 0.09$, $p_4 = 0.06$. On peut aussi remarquer la présence de $|\psi_6\rangle$, très proche de $|4, 2\rangle$, avec un poids $p_6 = 0.03$. Il s'agit d'un artefact lié au fait que la mesure ne discrimine pas totalement les nombres de photons. Le résultat de cette reconstruction est assez proche de celui obtenu précédemment avec les mesures dispersives et des mesures résonnantes à un atome. On peut noter en particulier l'excellent accord sur la valeur absolue des cohérences. On observe cependant des artefacts dans les populations des états de Fock.

Sensibilité par rapport à la phase des cohérences.

La figure IV.22 présente la vraisemblance obtenue lorsqu'on fait varier la phase des cohérences de l'état reconstruit, de la même manière que ce qui a été fait pour la figure IV.14.

La courbe verte correspond à une modification de la phase de la cohérence $\rho_{10,01}$ uniquement. L'écart entre la vraisemblance de l'état reconstruit et celui obtenu pour $\theta = \pi$ vaut 1086. En effectuant une régression parabolique autour du maximum, on trouve un coefficient du second ordre $a^{(10)}/2 = -251$, ce qui correspond à un écart-type sur la phase de 0.04 rad. On peut cependant remarquer que l'état $\hat{\rho}_{\text{ML}}$ contient aussi de l'information sur la phase de l'oscillation dans ses autres cohérences. La courbe rouge montre la vraisemblance associée aux états $\hat{\rho}_\theta \equiv \hat{R}(\theta)\hat{\rho}_{\text{ML}}\hat{R}^\dagger(\theta)$, obtenus en effectuant une rotation globale de l'état. La différence de log-vraisemblance entre l'état $\hat{\rho}_{\text{ML}}$ et l'état de phase opposée vaut 1101. En effectuant une régression parabolique autour du maximum, on trouve un coefficient au second ordre $a/2 = -254$, ce qui correspond à un écart-type sur la phase de 0.04 rad. Les courbes verte et rouge sont très proches, comme dans le cas de la reconstruction avec les atomes résonnants uniques et les QND. Les courbes correspondant aux deux façons de modifier la phase sont presque confondues : toute l'information sur l'oscillation est contenue dans la cohérence $\rho_{10,01}$. Envoyer plusieurs échantillons d'atomes résonnants permet donc de pallier le manque d'information de la mesure résonnante sur le nombre de photons. Par ailleurs, la sensibilité sur la phase est légèrement meilleure que pour les mesures avec un seul échantillon d'atomes résonnant. Cependant, cette sensibilité dépend du nombre de répétitions de la mesure et il faut donc comparer le nombre de mesures effectuées, ce qui est l'objet de la première remarque du paragraphe suivant.

Deux remarques importantes sur la reconstruction avec de multiples échantillons résonnants

La reconstruction obtenue avec de multiples atomes résonnants est bien plus rapide. En effet, dans le cas où on ne mesure qu'avec un atome résonnant, on n'obtient de l'in-

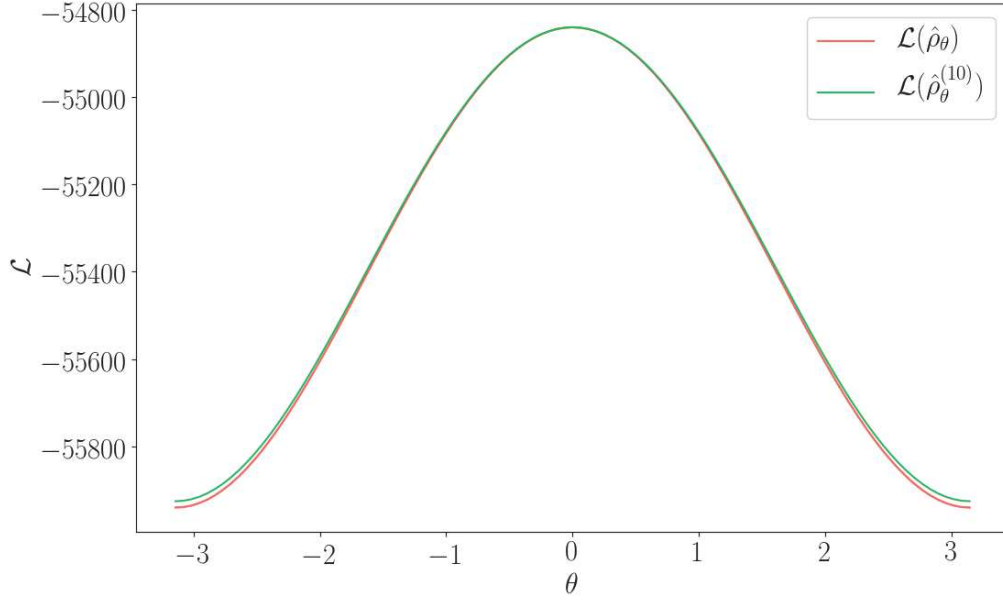


Fig. IV.22 Variations de la vraisemblance de l'état lorsqu'on modifie la phase de ses cohérences. Pour la courbe verte, on modifie la phase des cohérences $\rho_{10,01}$ uniquement. Pour la courbe rouge, on effectue une rotation globale de toutes les cohérences de l'état. La régression parabolique ($y = a/2\theta^2 + b\theta + \mathcal{L}(\hat{\rho}_{ML})$) a pour coefficients $a^{(10)}/2 = -251$ et $b^{(10)} = 0.16$ pour la courbe verte et $a/2 = -255$ et $b = 0.20$ pour la courbe rouge.

formation que lorsqu'on détecte un atome dans l'échantillon de l'atome sonde. Pour les nombres moyens d'atomes utilisés, l'échantillon est vide la plupart du temps et on n'obtient aucune information. Lorsqu'on envoie de nombreux échantillons, il est en revanche très probable qu'au moins l'un d'entre eux contienne un atome. Les atomes détectés dans le deuxième échantillon, lorsque le premier échantillon est vide donnent aussi un signal oscillant. Celui-ci a cependant une phase différente de celle du signal obtenu pour les atomes du premier échantillon, puisque l'état a évolué pendant les 200 μs entre les deux échantillons. Par ailleurs, le contraste de l'oscillation est légèrement plus faible, à cause de la décohérence, mais aussi de la probabilité que le premier échantillon contienne un atome non détecté.

La figure IV.23a montre la probabilité de détection des atomes sonde dans $|g\rangle$, en fonction du délai T_0 entre la préparation de l'état et la mesure et du numéro de l'échantillon atomique. On observe des oscillations pour les atomes des 5 premiers échantillons environ. Après un nombre d'échantillons plus important, les atomes n'apportent plus d'information sur la phase de l'état à cause du brouillage effectué par les atomes antérieurs non détectés. Ils continuent cependant d'apporter de l'information sur le nombre de photons. Si un atome est détecté, les atomes suivants n'apportent plus d'information sur les cohérences à un photon, mais permettent d'obtenir de l'information sur le nombre de photons. La figure IV.23b montre les oscillations associées aux 4 premiers échantillons atomiques.

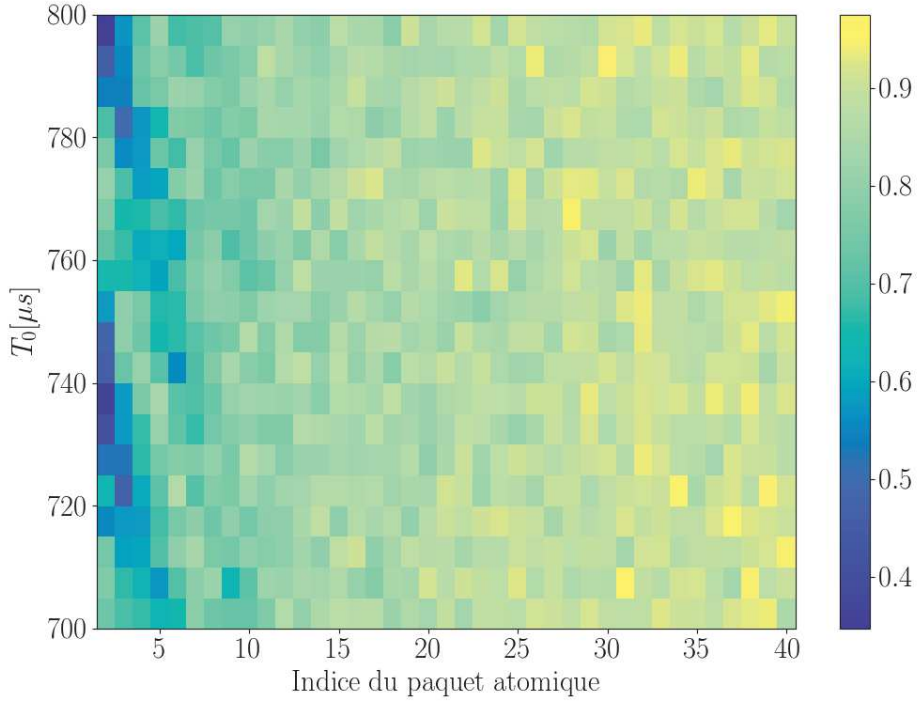
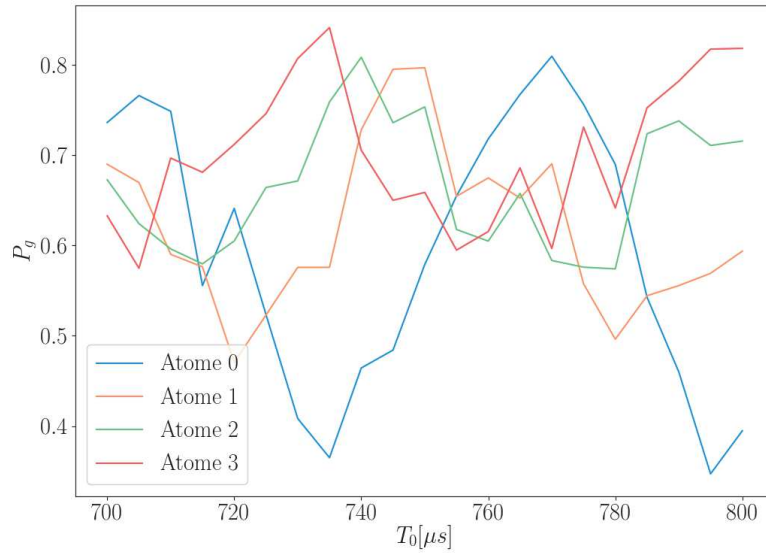
Chacun des points de ces courbes est moyenné sur environ 130 détections atomiques, ce qui explique la variance des signaux. Au total, on dispose de 108323 atomes détectés qui contribuent effectivement à la reconstruction, pour 237000 répétitions de l'expérience. À titre de comparaison, rappelons qu'on disposait de 3913 atomes détectés pour la mesure résonnante avec un seul atome, alors qu'on avait réalisé 361200 répétitions. L'utilisation de plusieurs atomes successifs permet donc de démultiplier l'information obtenue sur l'état et de réduire le temps d'acquisition des mesures pour obtenir des résultats comparables.

L'autre remarque concerne le caractère partiel de la mesure. Dans le cas de la mesure avec des résonnants multiples, il apparaît que la mesure n'est pas complète. Il n'y a pas assez d'information pour reconstruire correctement la distribution du nombre de photons. En revanche, la mesure contient beaucoup d'informations sur la phase de l'état quantique. Elle est donc adaptée lorsqu'on veut accéder aux cohérences, par exemple pour réaliser une interférence d'interférométrie, mais qu'on n'a pas besoin d'une reconstruction fiable des populations. On pourrait aussi accéder à l'information sur la phase en procédant à un ajustement de la phase pour chaque échantillon atomique et en prenant en compte l'évolution de l'état entre les mesures. Cette méthode nécessite ensuite d'effectuer une moyenne sur les échantillons atomiques en mettant un poids. Cependant, une fois que la méthode de reconstruction a été calibrée, elle est plus facile à utiliser, puisqu'elle prend en compte toute l'information dont on dispose sur l'évolution du système et qu'elle agrège tous les résultats de mesure. En particulier, elle leur fournit un poids dans la détermination de la phase, lié à la pureté des mesures. Afin de vérifier que la phase extraite de cette reconstruction partielle correspond bien à la phase de l'état, nous en avons effectué une étude systématique, qui est détaillée dans la section suivante.

IV.2.5 Reconstruction de l'état $|20\rangle + |02\rangle$

Lorsque nous avons reconstruit l'état $|10\rangle + |01\rangle$ avec des séries d'atomes résonnants, nous avons aussi tenté de préparer l'état $|20\rangle + |02\rangle$ en utilisant la méthode décrite dans la section II.1.4. La séquence expérimentale commence comme pour la préparation de $|10\rangle + |01\rangle$. Le premier injecteur est préparé dans un échantillon atomique contenant en moyenne $\mu = 0,25$ atome. Il effectue une impulsion $\pi/2$ dans la première cavité et une impulsion π dans la deuxième. Le deuxième injecteur est préparé 200 μs après le premier, avec $\mu = 0,23$. Il effectue une impulsion 2π dans la première cavité, puis une impulsion $\pi/\sqrt{2}$ dans la deuxième. L'impulsion dans la deuxième cavité a été choisie de sorte que l'atome émette avec certitude lorsqu'elle contient un photon. Lorsqu'elle est vide, on s'attend à ce que la première cavité contienne un photon et que l'atome émette dedans. Utiliser une impulsion plus courte permet de mieux injecter puisque le contraste diminue avec le temps d'interaction entre l'atome et la cavité.

La mesure est effectuée par 40 échantillons d'atomes résonnants, avec $\mu = 0,35$, identiques à ceux utilisés pour la mesure de $|10\rangle + |01\rangle$, qui effectuent des impulsions π dans la première cavité et $\pi/2$ dans la deuxième. On a effectué au total 285700 préparations de l'état. Après sélection des répétitions pour lesquelles on détecte exactement un atome dans $|g\rangle$ pour chaque échantillon injecteur, il reste 1001 trajectoires. Le résultat de la reconstruction est donné sur la figure IV.24. Ces résultats sont préliminaires et nous ont servi à vérifier que la reconstruction avec plusieurs échantillons d'atomes résonnants permet de donner des résultats pour d'autres états. On n'a pas cherché à optimiser la

(a) P_g 

(b) Coupes de la figure précédente, pour les 4 premiers échantillons atomiques.

Fig. IV.23 Probabilité de mesurer les atomes sonde dans $|g\rangle$ en fonction du numéro de l'échantillon atomique, en abscisse, et du temps T_0 entre la préparation et le premier atome sonde, en ordonnée. P_g présente des oscillations en fonction de T_0 pour les premiers échantillons atomiques. Lorsque le numéro de l'échantillon augmente, la probabilité d'avoir déjà détecté un atome sonde augmente aussi. Le contraste des oscillations décroît donc. Passé environ 5 échantillons, la mesure n'apporte donc plus d'information sur les cohérences.

préparation de l'état. Par ailleurs, la mesure effectuée est plus adaptée à la mesure de $|10\rangle + |01\rangle$ qu'à celle de $|20\rangle + |02\rangle$. L'état présente de nombreuses imperfections, dont on ne donnera pas d'étude quantitative. Parmi celles-ci, on remarque la présence résiduelle de $|10\rangle + |01\rangle$ et $|00\rangle$. Enfin, l'état comporte une forte proportion de $|11\rangle$, cohérente avec les autres états à deux photons. Une explication de cette présence pourrait être que, dans le cas où l'atome échoue à émettre dans la première cavité contenant un photon, il effectue une impulsion $\pi/\sqrt{2}$ dans la deuxième qui le fait émettre avec une probabilité 0,8. Malgré toutes ces imperfections, la présence de cohérences entre $|20\rangle$ et $|02\rangle$ indique que l'ensemble des mesures contient des traces d'une oscillation à 2δ due aux cohérences entre $|20\rangle$ et $|02\rangle$. Cette oscillation est très difficile à mettre en évidence par une mesure directe, puisqu'elle nécessite deux atomes pour la préparation et deux atomes pour la détection. C'est uniquement parce qu'on a plusieurs échantillons d'atomes sondes que le jeu de données permet de la mettre en évidence indirectement. Ces résultats préliminaires indiquent que la tomographie par trajectoires pourrait être un outil pour extraire de l'information de jeux de données sur lesquels on ne peut pas effectuer de moyennage.

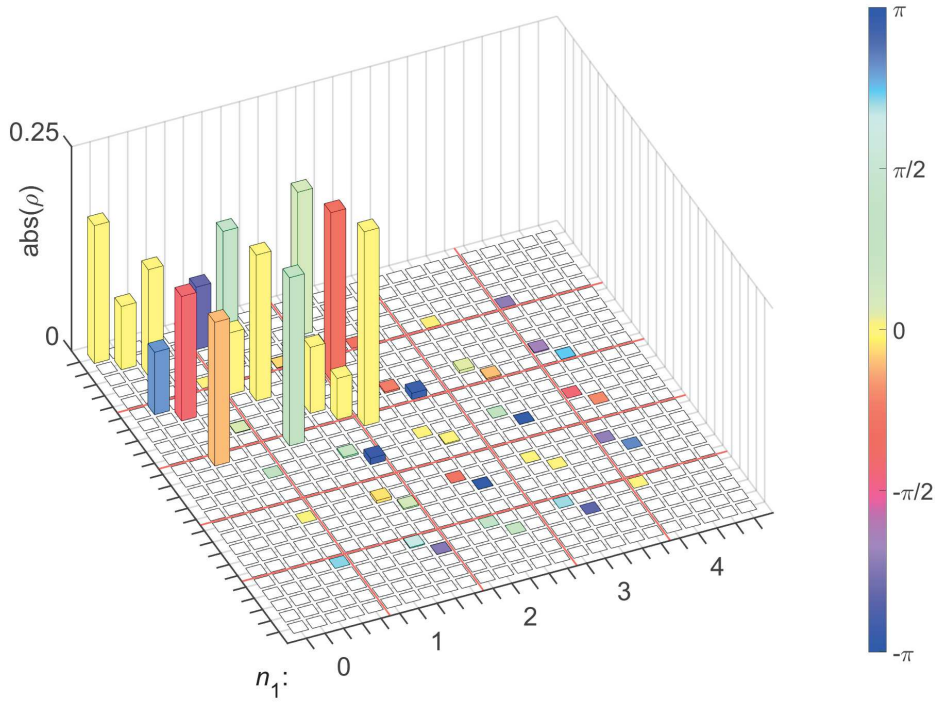


Fig. IV.24 État reconstruit pour la préparation de $|20\rangle + |02\rangle$.

IV.3 Reconstruction partielle

IV.3.1 Motivation

Dans la section IV.2.4, on a présenté une méthode qui permet d'obtenir une estimation correcte de l'état, tant qu'on s'intéresse aux cohérences entre $|10\rangle$ et $|01\rangle$ et pas au détail de la distribution de Fock. On a vu que cette méthode présente l'intérêt d'extraire plus rapidement l'information sur la phase de ces cohérences, en envoyant plus d'atomes qui y sont sensibles. Par ailleurs, cette méthode n'utilise que des atomes résonnants et ne nécessite donc pas l'utilisation d'injections micro-onde, ni d'atomes dispersifs. Elle est donc plus facile d'emploi que la reconstruction qui utilise des mesures locales de la fonction de Wigner en plus des atomes résonnants. Cependant, on sait que les matrices d'effet obtenues ne sont pas suffisantes pour espérer effectuer une tomographie complète de l'état. Par ailleurs, la présence d'artefacts sur les populations nous indique qu'elle n'est pas en mesure de reconstruire totalement la matrice densité. Dans ce chapitre, nous tentons de justifier qu'une telle tomographie, partielle, peut cependant être fiable pour étudier la cohérence $\rho_{10,01}$, qui est une des propriétés les plus intéressantes de l'état préparé, en particulier dans une perspective interférométrique. Cette approche peut aisément se généraliser à d'autres situations typiques en information quantique. En effet, la tomographie d'états quantiques est souvent un préalable à leur utilisation en vue de réaliser des protocoles d'information quantique. Dans ces protocoles, une information (le résultat d'un calcul, la valeur d'un paramètre physique qu'on souhaite mesurer) est souvent encodée dans une phase quantique ou, ce qui revient au même à un changement de base de mesure près, dans une population.

Afin d'assurer que la valeur obtenue pour la phase de $\rho_{10,01}$ par la reconstruction de la section IV.2.4 est fiable, on a préparé plusieurs états superposés avec des phases différentes, calibrées indépendamment. On a ensuite utilisé notre méthode de reconstruction et on a comparé la phase calibrée à celle obtenue par la reconstruction. Dans le but d'étudier le comportement statistique de la reconstruction et de le comparer aux barres d'erreur données par la reconstruction, on a aussi effectué une étude statistique systématique de la variance de la phase reconstruite.

IV.3.2 Préparation d'états de phases différentes et résultats obtenus.

a) Protocole et calibration de la phase

Comme on l'a vu dans la section II.1, la phase ϕ de l'état $|10\rangle + e^{i\phi}|01\rangle$ dépend de la phase de la superposition atome-cavité $|e, 0, 0\rangle + |g, 1, 0\rangle$ juste avant le transfert d'intrication de l'atome vers la deuxième cavité. On peut donc faire varier ϕ en modifiant la phase quantique relative entre les deux états atomiques $|e\rangle$ et $|g\rangle$ lors du trajet entre les deux cavités. Dans ce but, on applique un potentiel ajustable V_r sur une électrode de la zone de Ramsey située entre les deux cavités. Le déplacement des niveaux par effet Stark quadratique dépendant de ce potentiel, on peut alors modifier la phase de la superposition atomique. Afin de calibrer la variation de la phase de la superposition atomique avec V_r , on mesure la variation de la phase des franges de Ramsey entre R_1 et R_3 lorsque V_r varie.

La figure IV.25 montre la variation de la phase atomique avec le potentiel appliqué. La phase varie de manière quadratique à cause de l'effet Stark. La courbe bleue est un

ajustement prenant en compte cette dépendance. On utilise cet ajustement pour calibrer la phase de l'état préparé.

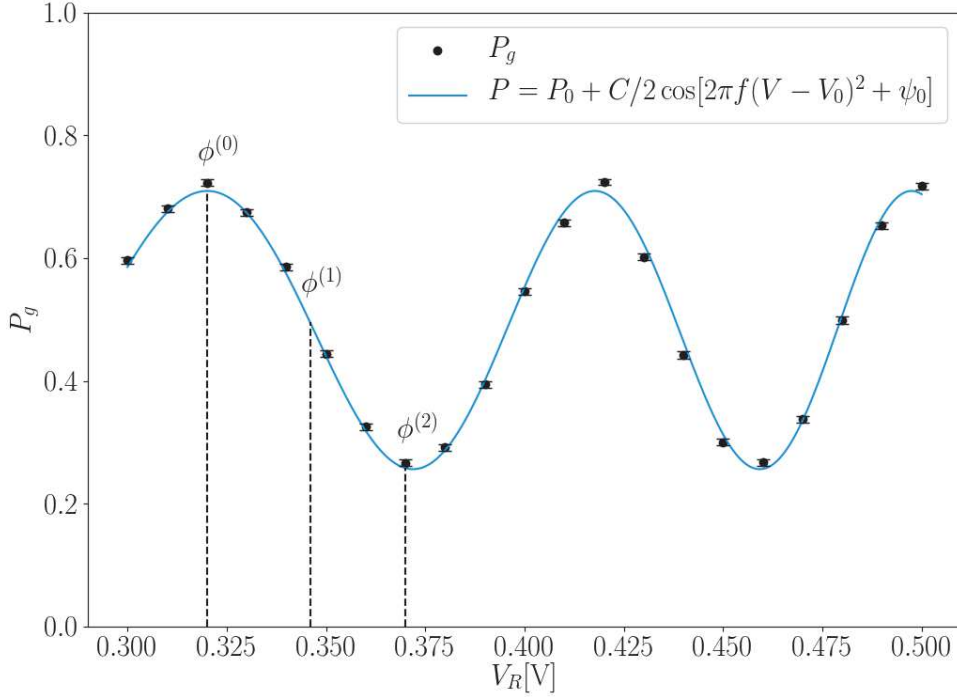


Fig. IV.25 Probabilité de détecter les atomes QND dans $|g\rangle$ en fonction du potentiel V_r appliqué à l'électrode R_2 pendant leur traversée de l'expérience. La courbe en bleu correspond à un ajustement par la fonction $V_r \rightarrow P_0 + C/2 \cos[2\pi f(V_r - V_0)^2 + \psi_0]$ avec les paramètres $P_0 = 0.483 \pm 0.001$, $C = 0.454 \pm 0.003$, $f = 12.9 \pm 0.4 \text{ V}^{-2}$, $V_0 = -0.03 \pm 0.01 \text{ V}$ et $\psi_0 = 3.2 \pm 0.5 \text{ rad}$. Les barres d'erreurs importantes sur f et ψ_0 sont dues au fait que ces deux paramètres sont liés et qu'on peut obtenir une courbe très proche de notre ajustement en les faisant varier conjointement. Si on fixe toutes les valeurs des paramètres qui jouent un rôle dans la phase sauf ψ_0 et qu'on fait l'ajustement, on trouve des barres d'erreurs de l'ordre de $\pm 0.01 \text{ rad}$ sur la phase. Les points $\phi^{(0)}$, $\phi^{(1)}$ et $\phi^{(2)}$ ont été utilisés pour effectuer la reconstruction de trois états quantiques $\hat{\rho}^{(0)}$, $\hat{\rho}^{(1)}$ et $\hat{\rho}^{(2)}$, différant uniquement par la phase des cohérences.

On choisit pour référence de phase la phase de l'état reconstruit dans le dernier paragraphe de la session précédente. Cette phase est obtenue pour $V_r = 0.320 \pm 0.001 \text{ V}$. Notons $\hat{\rho}_0$ l'état correspondant. On a également préparé deux états $\hat{\rho}^{(1)}$ et $\hat{\rho}^{(2)}$ avec les potentiels respectifs $V_r^{(1)} = 0.346 \pm 0.001 \text{ V}$ et $V_r^{(2)} = 0.370 \pm 0.001 \text{ V}$, correspondant aux phases respectives $\phi^{(1)} = 1.52 \pm 0.06 \text{ rad}$ et $\phi^{(2)} = 3.02 \pm 0.06 \text{ rad}$. Ces potentiels sont appliqués uniquement pendant le passage de l'atome injecteur. Les atomes sondes voient le

même potentiel pour les trois reconstructions. On dispose de 23888 matrices d'effet pour la reconstruction de $\hat{\rho}^{(1)}$. Pour $\hat{\rho}^{(2)}$, on en a 14467.

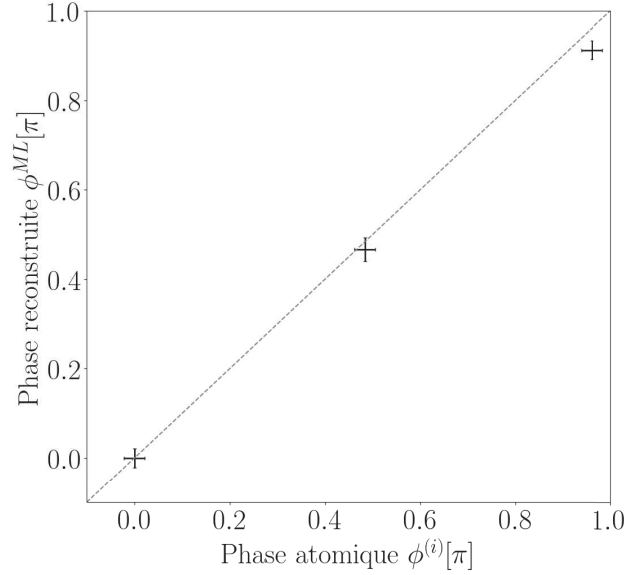


Fig. IV.26 Phases reconstruites par la tomographie par trajectoires en fonction des phases atomiques calibrées par les franges de la figure IV.25, en unités de π . Dans les deux cas, on a pris pour origine des phases la valeur obtenue pour la première des trois mesures, ce qui explique pourquoi il n'y a que deux points. Les barres d'erreur verticales sont celles obtenues par la reconstruction avec la formule IV.3.2. Les barres d'erreur horizontales correspondent à l'incertitude sur la calibration due à l'incertitude sur V_r , qui vaut 1 mV.

La figure IV.27 montre les matrices densité reconstruites. Les matrices sont très proches les unes des autres, à l'exception de leur phase. Si on prend la valeur absolue pour rendre leurs phases quantiques identiques, on trouve les fidélités :

$$\mathcal{F}(\hat{\rho}_{\text{ML}}^{(0)}, \hat{\rho}_{\text{ML}}^{(1)}) = 0.96, \quad (\text{IV.3.1a})$$

$$\mathcal{F}(\hat{\rho}_{\text{ML}}^{(0)}, \hat{\rho}_{\text{ML}}^{(2)}) = 0.95, \quad (\text{IV.3.1b})$$

$$\mathcal{F}(\hat{\rho}_{\text{ML}}^{(1)}, \hat{\rho}_{\text{ML}}^{(2)}) = 0.99. \quad (\text{IV.3.1c})$$

La figure IV.26 montre les phases mesurées $\phi_0^{ML} = 0.91 \pm 0.07$ rad, $\phi_1^{ML} = 2.37 \pm 0.08$ rad et $\phi_2^{ML} = 3.77 \pm 0.06$ rad correspondant respectivement à $\phi^{(0)}$, $\phi^{(1)}$ et $\phi^{(2)}$. Les phases extraites de la reconstruction sont en accord avec les phases calibrées. Notre méthode permet donc d'extraire correctement l'information sur la phase de $\rho_{10,01}$.

b) Incertitudes sur la phase

Les barres d'erreur fournies par la reconstruction sont censées nous donner une estimation de la façon dont une grandeur mesurée varie, lorsqu'on répète plusieurs fois la

procédure entière de reconstruction. Afin de tester si elles remplissent correctement ce rôle, nous les avons comparées aux barres d'erreur statistiques correspondant à plusieurs reconstructions du même état. Dans ce but, on a utilisé les données utilisées pour la reconstruction de $\hat{\rho}^{(1)}$. On en a extrait 18000 trajectoires. Ces trajectoires ont ensuite été séparées en groupes de R trajectoires, pour différentes valeurs de R entre 100 et 9000. Pour chaque valeur de R , on tire aléatoirement les groupes de trajectoires. On effectue ensuite une reconstruction pour chaque groupe, dont on extrait une valeur de la phase ϕ . On calcule alors l'écart-type $\tilde{\sigma}_{\phi,R}$. Par ailleurs, notre méthode nous permet de calculer des barres d'erreur sur les éléments de la matrice densité, à l'aide de la formule III.4.13. On peut donc en tirer une barre d'erreur $\sigma(\phi)$ sur la phase, sous la forme :

$$\sigma(\phi) = \frac{\sqrt{(y_{10,01}\sigma(x_{10,01}))^2 + (x_{10,01}\sigma(y_{10,01}))^2}}{r_{10,01}^2}, \quad (\text{IV.3.2})$$

où $x_{10,01} = \Re(\rho_{10,01})$, $y_{10,01} = \Im(\rho_{10,01})$ et $r_{10,01}^2 = x_{10,01}^2 + y_{10,01}^2$.

Afin d'obtenir des résultats plus significatifs, on applique une méthode de *bootstrapping* en répétant la procédure précédente 4 fois avec des tirages indépendants des groupes de trajectoires. On note alors $\langle \sigma(\phi) \rangle_R$ la moyenne de $\sigma(\phi)$ pour tous les groupes de R trajectoires des 4 tirages. On compare cette valeur à la moyenne $\overline{\tilde{\sigma}_{\phi,R}}$ sur les 4 réalisations de l'écart-type.

La figure IV.28 montre les résultats de cette étude. On peut tout d'abord remarquer que les valeurs de $\langle \sigma(\phi) \rangle_R$ sont proches des incertitudes statistiques lorsqu'on répète la reconstruction. Elles fournissent donc une estimation correcte de celle-ci. On peut cependant remarquer qu'elles sont légèrement surestimées. Ceci est dû au fait qu'elles prennent uniquement en compte les variations locales de la vraisemblance au voisinage de son maximum. Elles ne permettent donc pas de prendre entièrement en compte les contraintes sur la matrice densité, qui limitent les états accessibles à la reconstruction et donc les plages de variation accessibles à leurs paramètres. On peut aussi remarquer que $\overline{\tilde{\sigma}_{\phi,R}}$ suit la décroissance attendue en $1/\sqrt{R}$, comme l'atteste l'ajustement en bleu. L'insert de la figure IV.28 montre un histogramme des résultats obtenus pour la phase, pour 800 reconstructions effectuées avec $R = 800$. La courbe bleue est une gaussienne de largeur $\overline{\tilde{\sigma}_{\phi,R}}$. On voit que l'intervalle d'incertitude de la valeur calibrée, en vert est en accord avec celui donné par la reconstruction. On peut de plus constater que la barre d'erreur $\langle \sigma(\phi) \rangle_R$, figurée par la gaussienne rouge, est proche de l'incertitude statistique, mais légèrement plus élevée.

En conclusion, la phase quantique $\phi_{ML}^{(1)}$ extraite de la reconstruction correspond bien à la phase calibrée $\phi^{(1)}$. Son incertitude a de plus le comportement statistique attendu et on peut en donner une borne supérieure à l'aide de $\sigma(\phi)$. La reconstruction, même si elle ne permet pas en théorie de reconstruire parfaitement l'état, fournit donc des résultats valides pour la phase en question.

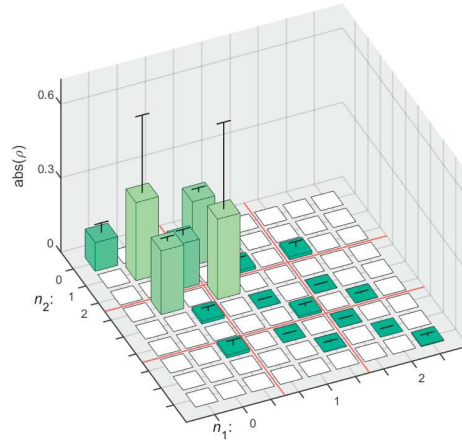
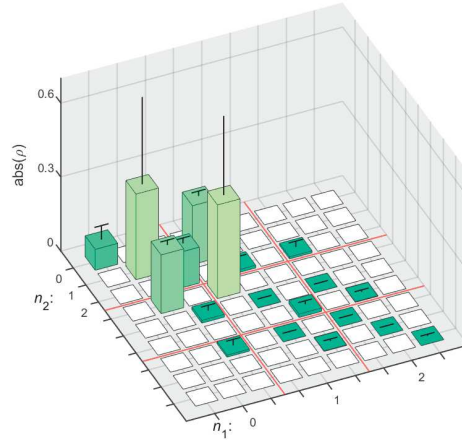
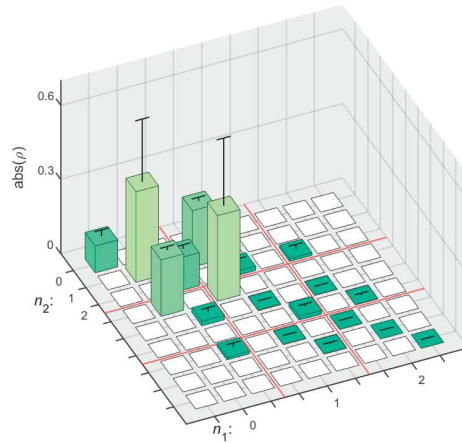
(a) $\hat{\rho}_{ML}(\phi^{(0)})$ (b) $\hat{\rho}_{ML}(\phi^{(1)})$ (c) $\hat{\rho}_{ML}(\phi^{(2)})$

Fig. IV.27 Matrices densité reconstruites pour les phases quantiques $\phi^{(0)} = 0$ rad, $\phi^{(1)} = 1.52$ rad et $\phi^{(2)} = 3.02$ rad. Pour la phase $\phi^{(1)}$, les barres d'erreur sur les populations sont tronquées. Elles valent $\sigma(\rho_{01,01}) = 0.57$ et $\sigma(\rho_{10,10}) = 0.62$.

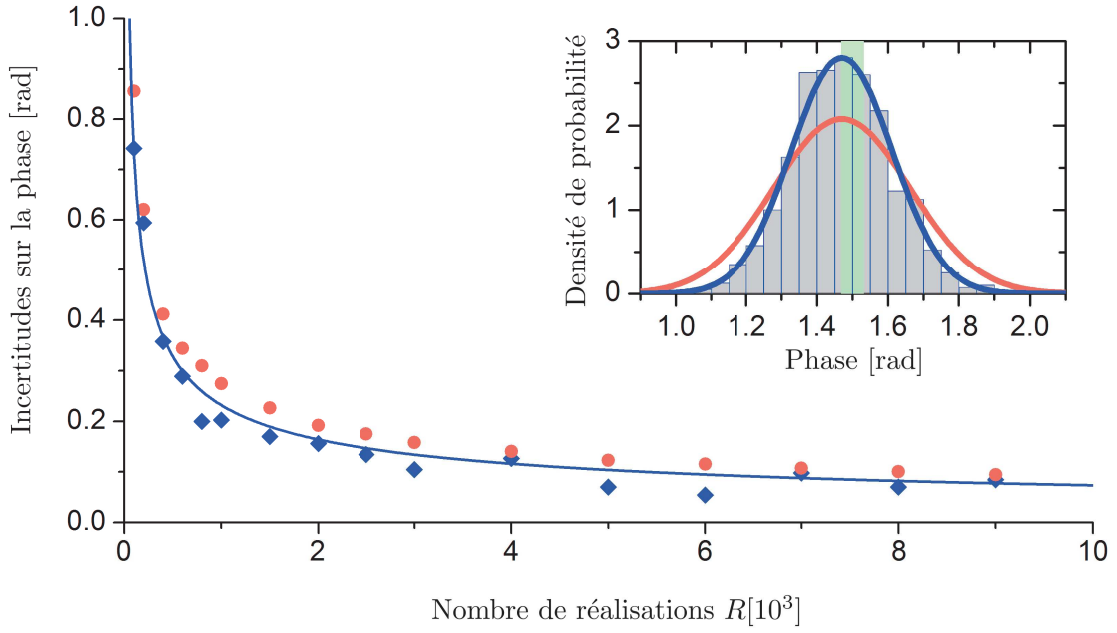


Fig. IV.28 Comparaison de la moyenne $\langle \sigma(\phi) \rangle_R$ des barres d'erreur estimées avec notre méthode (en rouge), avec la moyenne de l'écart-type $\tilde{\sigma}_{\phi,R}$ des phases extraites de la reconstruction (en bleu), pour différentes valeurs de la taille des groupes de trajectoires utilisés pour la reconstruction. La ligne bleue est un ajustement des points bleus par la fonction $x \rightarrow A/\sqrt{x}$, qui correspond à la décroissance attendue pour les incertitudes. On peut remarquer que l'écart-type est proche des barres d'erreur que nous calculons, tout en étant systématiquement plus faible. Insert : Distribution de ϕ pour 800 reconstructions effectuées avec $R = 800$. La ligne rouge est une gaussienne de largeur $\langle \sigma(\phi) \rangle_R$ et la ligne bleue une gaussienne de largeur $\tilde{\sigma}_{\phi,R}$. La bande verte est l'intervalle de confiance mesuré indépendamment pour ϕ .

Conclusion

Dans les travaux présentés dans ce manuscrit, nous avons étudié en détails la tomographie par trajectoires. Nous avons illustré l'intérêt de cette méthode qui permet de combiner des mesures différentes, qui apportent des informations distinctes sur la matrice densité du système, en prenant en compte les imperfections de mesure. Nous avons également mis en évidence l'intérêt d'effectuer des séries de mesures sur chaque réalisation de l'état afin de suivre sa trajectoire quantique et de maximiser l'information extraite. Les matrices d'effet utilisées pour effectuer la reconstruction par maximum de vraisemblance constituent par ailleurs un outil précieux pour visualiser l'information obtenue sur une trajectoire, ce qui est utile pour imaginer des protocoles de mesure plus efficaces.

En prenant pour exemple un état intriqué de deux cavités micro-onde, nous avons montré que cette technique permet d'estimer correctement l'état quantique d'un système de grande dimension, en intégrant toute la connaissance du dispositif expérimental et de ses imperfections. Nous avons mis en évidence l'intérêt d'utiliser des mesures adaptées à l'état, en l'occurrence des mesures sensibles aux seules cohérences de la matrice densité susceptibles d'être non nulles. Celles-ci permettent de reconstruire l'état sans mesurer toutes les composantes de la matrice densité. On peut par ailleurs utiliser une approche adaptative en déterminant dans un premier temps les populations non nulles, avant d'effectuer les mesures sur les cohérences entre ces populations. Nous avons pu tirer de la reconstruction des barres d'erreur sur les composantes de la matrice densité, afin d'estimer les incertitudes sur chacune d'entre elles. Il est par ailleurs possible, avec les outils actuels de calculer des barres d'erreur sur n'importe quelle fonction linéaire de la matrice densité.

Nous avons aussi étudié la possibilité de réaliser des tomographies partielles, pour obtenir de l'information sur certaines parties de la matrice densité qui présentent un intérêt particulier. Cette méthode s'est avérée particulièrement avantageuse pour extraire la phase quantique d'un état $|10\rangle + e^{i\phi}|01\rangle$. En comparant la phase reconstruite à celle déterminée *a priori*, nous avons établi l'exactitude de la mesure. Nous avons par ailleurs effectué une analyse statistique des reconstructions afin de montrer la pertinence des barres d'erreur obtenues sur la phase de la reconstruction.

Perspectives

La méthode de tomographie par trajectoires que nous avons utilisée présente encore des questions ouvertes du point de vue théorique. Des travaux sont actuellement en cours afin de généraliser le calcul des barres d'erreur à des observables qui ne sont pas des fonctions linéaires de la matrice densité. Ceci permettra en particulier de déterminer des

barres d'erreur sur des estimations de l'entropie d'états quantiques.

Par ailleurs, la méthode a été développée initialement pour des ensembles de mesures qui engendrent tout l'ensemble des matrices densité. Nous avons mis en évidence dans ces travaux qu'elle peut cependant fournir des résultats corrects sans que cette condition soit remplie, lorsque l'état à reconstruire est proche du bord de l'ensemble des matrices densité. En effet, si certaines populations sont proches de zéro, il n'est pas nécessaire de mesurer les cohérences associées pour reconstruire précisément l'état. Un travail d'extension du cadre de la tomographie par trajectoire à ces situations serait particulièrement intéressant, dans la mesure où les dimensions des espaces de Hilbert des systèmes utilisés pour le traitement d'information quantique sont amenées à croître, ce qui augmente grandement le nombre de mesures à effectuer.

Le dispositif expérimental sert actuellement à réaliser des expériences de thermodynamique quantique dans lesquelles on tente d'appliquer à des systèmes quantiques les concepts développés pour étudier les machines thermiques classiques. Le grand degré de contrôle dont nous disposons sur l'interaction entre les atomes et les cavités permet d'imaginer de nombreux protocoles expérimentaux. Notre capacité de réaliser de nombreuses mesures et d'effectuer des tomographies du champ rend possible des études détaillées d'identités thermodynamiques, dans lesquelles on peut mesurer tous les termes.

Dans le cadre de la thermodynamique quantique, les deux cavités peuvent jouer le rôle de réservoirs. Nous sommes capables de préparer facilement des états thermiques des cavités en les couplant à une source micro-onde dont la phase est modulée aléatoirement. Il est alors possible de définir la température T des cavités. On peut aussi définir une température pour un ensemble d'atomes à deux niveaux. Nous étudions l'interaction des atomes avec les cavités pour mesurer des identités thermodynamiques lorsque ces différents systèmes interagissent entre eux.

La première expérience que nous envisageons fait intervenir un démon de Maxwell quantique. Ce dernier est l'analogue du démon de l'expérience de pensée de Maxwell qui est capable, à l'aide d'un trou qu'il peut obturer, de trier les particules d'un gaz dans deux compartiments différents selon qu'elles ont une énergie cinétique élevée ou faible. Cette expérience de pensée semble à première vue violer la seconde loi de la thermodynamique si on ne prend pas en compte l'information acquise par le démon lors de ses observations de la vitesse des particules et la production d'entropie qui l'accompagne.

La figure A présente le protocole de l'expérience. Les deux cavités C_1 et C_2 sont initialement préparées dans des états thermiques avec des températures respectives T_1 et T_2 , de sorte que la première cavité soit plus froide que la deuxième. Elles interagissent par l'intermédiaire d'un atome, initialement préparé dans l'état $|g\rangle$. L'atome effectue un passage adiabatique dans la première cavité de sorte qu'il reste dans son état fondamental si elle est vide et qu'il absorbe un photon sinon. Dans la deuxième cavité, il a la même évolution si il est dans son état fondamental. Si il est excité, en revanche il émet dans tous les cas.

Le démon de Maxwell est joué dans notre expérience par une impulsion π appliquée dans la cavité R_1 . Cette impulsion transfère l'état $|g\rangle$ vers l'état de Rydberg inférieur $|f\rangle$. Une fois dans cet état, l'atome n'a plus de transition à résonance avec les cavités et n'interagit plus avec elles. L'atome $|e\rangle$ n'est quant à lui pas affecté par cette impulsion. On peut donc la voir comme une mesure virtuelle de $|g\rangle$. L'atome n'interagit donc avec C_2 que si il est dans $|e\rangle$.

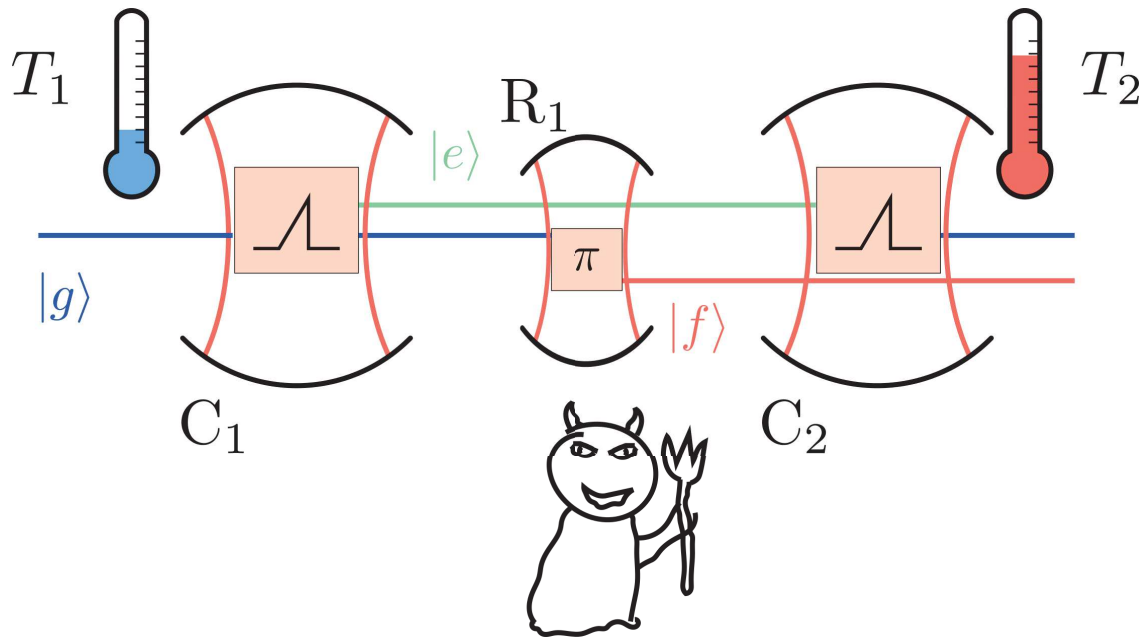


Fig. A Principe d'une expérience à deux cavités faisant intervenir un démon de Maxwell autonome. Description dans le texte.

Grâce au démon on fait donc en sorte de ne garder que les interactions dans lesquelles l'atome injecte de l'énergie dans la deuxième cavité. De cette manière, l'énergie circule toujours de la source froide vers la source chaude. Comme on est capable de réaliser des tomographies de l'état des deux cavités, on peut reconstruire les corrélations entre elles et vérifier des identités thermodynamiques. Le démon de Maxwell que nous proposons ici est un démon autonome puisqu'il fonctionne sans nécessiter de faire une mesure sur le système et d'effectuer une rétroaction. Dans une expérience préliminaire, nous avons déjà réalisé un démon de Maxwell autonome sur l'interaction entre un atome et une cavité. Ces expériences devraient nous permettre dans un premier temps de tester la pertinence des concepts de thermodynamique dans le cadre de notre dispositif.

Pour la suite, deux voies semblent prometteuses. La possibilité de réaliser de nombreuses mesures sur une trajectoire permet d'envisager des expériences appliquant au domaine quantique les principes de la thermodynamique stochastique, dans laquelle on étudie des trajectoires individuelles d'un système subissant une transformation thermodynamique. Le cadre de la tomographie par trajectoires trouverait alors peut-être une extension naturelle. D'un autre côté, notre capacité à préparer des états intriqués des deux cavités devrait permettre de réaliser des expériences étudiant le rôle des cohérences non locales dans l'échange d'énergie entre deux systèmes.

Annexe A

Quelques ingrédients théoriques sur l'information de Fisher

Dans cette annexe, on donne quelques éléments mathématiques sur la manipulation de l'information de Fisher. On commence par démontrer la borne de Cramer-Rao [III.1.6](#). On donne ensuite la preuve de l'équivalence [A.7](#) entre deux formulations de l'information de Fisher. On donnera aussi quelques remarques sur le calcul analytique et numérique de l'information de Fisher.

Borne de Cramer-Rao et expressions équivalentes de l'information de Fisher

Commençons par considérer le cas où on souhaite estimer un paramètre θ à l'aide d'une mesure décrite par une variable aléatoire X . On dispose d'un estimateur Θ_{est} dont la moyenne s'écrit :

$$\mathbb{E}[\Theta_{\text{est}}(X)] = \int dx p(x|\theta) \theta_{\text{est}}(x). \quad (\text{A.1})$$

Si l'estimateur est non biaisé, on a :

$$\mathbb{E}[\Theta_{\text{est}}(X)] = \theta, \quad (\text{A.2a})$$

$$\int dx p(x|\theta) (\theta_{\text{est}}(x) - \theta) = 0. \quad (\text{A.2b})$$

En différentiant par rapport à θ , il vient :

$$\int dx \left[\frac{\partial}{\partial \theta} p(x|\theta) (\theta_{\text{est}}(x) - \theta) - p(x|\theta) \right] = 0, \quad (\text{A.3a})$$

$$\int dx \frac{\partial}{\partial \theta} p(x|\theta) (\theta_{\text{est}}(x) - \theta) = 1, \quad (\text{A.3b})$$

$$\int dx p(x|\theta) \underbrace{\frac{\partial}{\partial \theta} \log p(x|\theta)}_A \underbrace{(\theta_{\text{est}}(x) - \theta)}_B = 1, \quad (\text{A.3c})$$

où on reconnaît dans la dernière équation le produit scalaire $\langle A, B \rangle = \mathbb{E}[AB]$ entre deux vecteurs de probabilité A et B . En appliquant l'inégalité de Cauchy-Schwartz, il vient

$$\mathbb{E} \left[\left(\frac{\partial}{\partial \theta} \log p(x|\theta) \right)^2 \right] \underbrace{\mathbb{E} [(\theta_{\text{est}}(x) - \theta)^2]}_{\Delta^2 \Theta_{\text{est}}} \geq 1, \quad (\text{A.4a})$$

$$\Delta^2 \Theta_{\text{est}} \geq \frac{1}{F(\theta)}. \quad (\text{A.4b})$$

Cette inégalité est appelée borne de Cramer-Rao. Elle donne une limite inférieure à la variance de tout estimateur non biaisé de θ . La fonction F apparaissant au dénominateur du membre de droite est appelée *information de Fisher* et a plusieurs expressions équivalentes :

$$F(\theta) \equiv \mathbb{E} \left[\left(\frac{\partial}{\partial \theta} \log p(x|\theta) \right)^2 \right] \quad (\text{A.5a})$$

$$= \mathbb{E} [V(X, \theta)^2], \quad (\text{A.5b})$$

où $V(X, \theta) \equiv \partial \log p(x|\theta) / \partial \theta$ est la fonction de score. Calculons son espérance :

$$\mathbb{E} [V(X, \theta)] = \int dx p(x|\theta) \frac{\partial \log p(x|\theta)}{\partial \theta} = \int dx \frac{\partial p(x|\theta)}{\partial \theta} \quad (\text{A.6a})$$

$$= \frac{\partial}{\partial \theta} \int dx p(x|\theta) = 0. \quad (\text{A.6b})$$

La fonction de score a une espérance nulle. On en déduit que la fonction de Fisher est la variance de la fonction de score. On peut établir une autre expression utile de l'information de Fisher :

$$\begin{aligned} F(\theta) &= \int dx p(x|\theta) \left[\frac{\partial}{\partial \theta} \log p(x|\theta) \right]^2, \\ &= \int dx p(x|\theta) \frac{1}{p(x|\theta)^2} \left[\frac{\partial}{\partial \theta} p(x|\theta) \right]^2, \\ &= - \int dx \left\{ -p(x|\theta) \frac{1}{p(x|\theta)^2} \left[\frac{\partial}{\partial \theta} p(x|\theta) \right]^2 \right\} - \underbrace{\int dx \frac{\partial^2}{\partial \theta^2} (p(x|\theta)) p(x|\theta) \frac{1}{p(x|\theta)}}_0, \\ &= - \int dx p(x|\theta) \frac{\partial}{\partial \theta} \left(\frac{\partial}{\partial \theta} p(x|\theta) \frac{1}{p(x|\theta)} \right), \\ &= - \int dx p(x|\theta) \frac{\partial^2}{\partial \theta^2} \log p(x|\theta), \\ F(\theta) &= -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \log p(x|\theta) \right]. \end{aligned} \quad (\text{A.7})$$

Cette expression donne une interprétation géométrique très simple à l'information de Fisher. Il s'agit de l'espérance de la courbure de la fonction de log-vraisemblance.

Calcul numérique de l'information de Fisher

Des exemples de calculs analytiques de l'information de Fisher sont donnés dans les sections II.2.2 et III.1.2.a. Dans la première de ces deux sections, on montre l'expression de l'information de Fisher lorsqu'on souhaite déterminer la phase ϕ d'un signal sinusoïdal de contraste C (formule II.11) :

$$F(\phi) = P_e \left(\frac{d \log P_e}{d\phi} \right)^2 + P_g \left(\frac{d \log P_g}{d\phi} \right)^2, \quad (\text{A.8a})$$

$$= \frac{C^2 \sin^2 \phi}{1 - C^2 \cos^2 \phi}. \quad (\text{A.8b})$$

Lorsque le contraste est maximal, l'information de Fisher vaut 1, quelle que soit la phase de l'interféromètre. Ce résultat a de quoi surprendre. En effet, lorsqu'on se place au sommet d'une frange, la dérivée du signal s'annule et on est insensible à la phase, au premier ordre. Cet effet est contrebalancé, dans la limite d'une infinité de mesures, par le fait qu'on obtient exactement toujours le même résultat, puisque la probabilité vaut 1, ce qui garantit qu'on est au sommet de la frange. Cette situation n'est cependant pas très réaliste. Dès lors que le contraste est différent de 1, l'information de Fisher s'annule pour les extrema des franges. Pour la calculer numériquement, il est donc utile de mettre un contraste fini, ce qui évite par ailleurs les formes indéterminées. Il est aussi important de bien choisir le pas pour calculer les dérivées numériques puisque l'information de Fisher peut varier très abruptement au voisinage d'un extremum du signal, si celui-ci est très contrasté.

Annexe B

Construction de l'intervalle de confiance de Clopper-Pearson

Supposons qu'on souhaite estimer la probabilité θ d'une loi de Bernoulli. On effectue une série de n mesures décrites par les variables aléatoires $(X_i)_{i \in \llbracket 1, n \rrbracket}$. On note x le nombre de résultats positifs. La valeur estimée pour θ est

$$\theta_{\text{est}} = \frac{x}{n}. \quad (\text{B.1})$$

On souhaite donner la borne supérieure θ_2 de l'intervalle de confiance $[\theta_1, \theta_2]$, de seuil d'erreur α . On choisit d'avoir une probabilité égale de se tromper par valeurs positives ou par valeurs négatives. La probabilité d'obtenir un résultat x tel que $\theta \geq \theta_2(x)$ admet donc la majoration suivante

$$p(\theta \geq \theta_2(X)) \leq \frac{\alpha}{2}. \quad (\text{B.2})$$

La figure B.1 montre le principe de la construction des barres de Clopper-Pearson. n est fixé et on a mis en abscisse les résultats x possibles. L'axe des ordonnées représente les valeurs possibles du paramètre θ . Les croix représentent l'estimateur choisi lorsqu'on obtient le résultat x . Il s'agit ici de l'estimateur Maxlike

$$\theta_{\text{est}} = \frac{x}{n}. \quad (\text{B.3})$$

On se place à θ fixé et on veut tracer la condition B.2. Dans ce but, on cherche $x_{\leq}$ tel que :

$$p(X \leq x_{\leq}) \leq \frac{\alpha}{2} \leq p(X \leq x_{\leq} + 1). \quad (\text{B.4})$$

$x_{\leq}$ est donc le plus petit x tel que

$$\sum_{x=0}^{x_{\leq}} \binom{n}{x} \theta^x (1-\theta)^{n-x} \leq \frac{\alpha}{2}. \quad (\text{B.5})$$

Comme $p(X \leq x_{\leq})$ est une fonction décroissante de θ , qui vaut 1 en 0 et 0 en 1, la condition B.4 est vérifiée par le même $x_{\leq}$ pour un intervalle de valeurs de θ . Pour la même raison

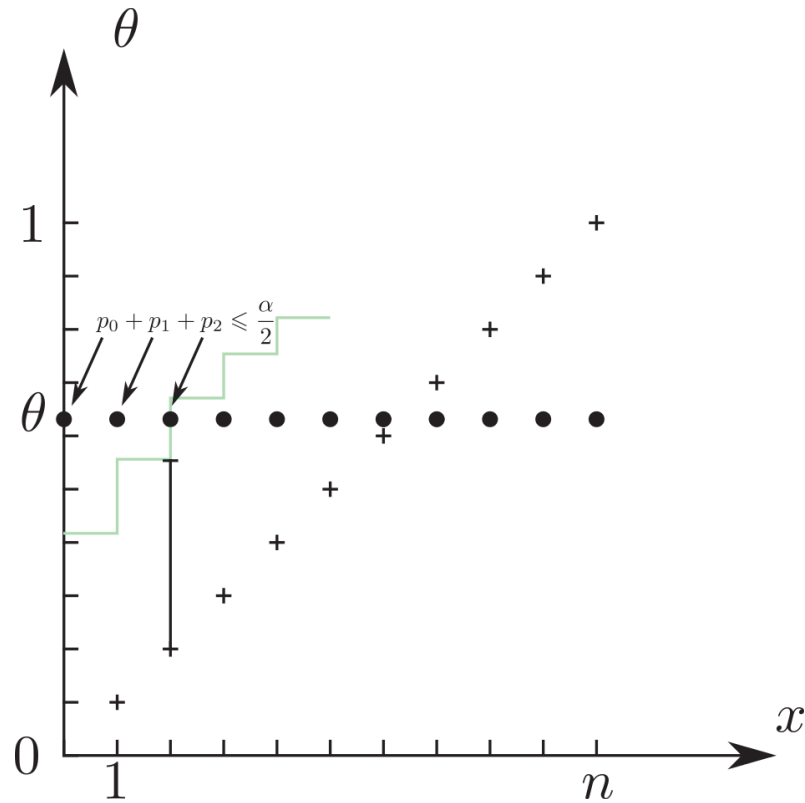


Fig. B.1 Principe de la construction des bornes de Clopper-Pearson. Le diagramme montre les résultats déterminés en fonction du résultat du tirage (croix). La ligne verte représente la condition B.2 sur la borne supérieure de l'intervalle. À θ donné, les résultats possibles sont représentés par les ronds noirs. L'estimation est surestimée lorsqu'on tire un des points situés au dessus de la ligne verte.

$x_{\leq}$ est une fonction croissante de θ . Ceci a pour conséquence la forme d'escalier de $x_{\leq}$, en vert sur la figure B.1.

La construction de l'intervalle d'erreur se déduit de l'ensemble des $x_{\leq}$ qu'on vient de trouver. Puisque tous les $x_{\leq}$ vérifient la condition B.2, on peut prendre comme valeur supérieure θ_2 de l'intervalle de confiance pour x le bas de la marche telle que $x_{\leq} = x$, comme illustré par la demi-barre d'erreur tracée figure B.1. Pour θ donné, si on tire une valeur x telle que θ est situé au dessus de la barre d'erreur pour x , alors $x \leq x_{\leq}$ par construction. On en déduit que la probabilité de tirer un tel événement est inférieure à $\alpha/2$.

On peut construire de la même manière la borne inférieure θ_1 de l'intervalle de confiance.

Annexe C

Purification d'états intriqués

Considérons le cas de l'interférométrie avec l'état $|10\rangle + |01\rangle$. La phase acquise par l'état vaut δT_0 . Pour être le plus sensible possible à δ , on souhaite donc étendre la durée entre la préparation et la détection du photon. On est cependant limité par la décohérence. Le photon est perdu après un temps caractéristique T_c . Supposons que la durée de mesure soit fixée à T_0 pour des raisons pratiques. L'état du système est alors :

$$\hat{\rho} = \begin{pmatrix} & |0,0\rangle & |0,1\rangle & |1,0\rangle & |1,0\rangle \\ \begin{pmatrix} 1 - e^{-T_0/T_c} & 0 & 0 & 0 \\ 0 & e^{-T_0/T_c} & e^{-T_0/T_c - i\phi} & 0 \\ 0 & e^{-T_0/T_c + i\phi} & e^{-T_0/T_c} & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \end{pmatrix}. \quad (\text{C.1})$$

Avec une probabilité $1 - e^{-T_0/T_c}$, le photon a été perdu et il n'y a plus rien entre les miroirs. Dans ce cas, il n'y a aucun intérêt à envoyer un atome sonde puisqu'il va être détecté dans $|g\rangle$ et réduire le contraste des franges. Avant d'effectuer la mesure, on voudrait donc filtrer les événements pour ne garder que ceux où il reste un photon entre les miroirs. Une manière de procéder consiste à envoyer un atome QND mesurant la parité du nombre de photons global. La mesure d'un atome dans $|e\rangle$ est alors décrite par l'opérateur de Kraus :

$$M_e^{[dis]} = \begin{pmatrix} & |0,0\rangle & |0,1\rangle & |1,0\rangle & |1,0\rangle \\ \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix} \end{pmatrix} \quad (\text{C.2})$$

et celle d'un atome dans $|e\rangle$ par

$$M_g^{[dis]} = \begin{pmatrix} & |0,0\rangle & |0,1\rangle & |1,0\rangle & |1,0\rangle \\ \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \end{pmatrix}. \quad (\text{C.3})$$

Si les cavités contiennent encore un photon, l'atome est donc détecté dans $|g\rangle$, tandis qu'il est détecté dans $|e\rangle$ si elles sont vides. On ne souhaite donc garder que les cas où l'atome

est détecté dans $|g\rangle$. Dans ces cas, l'état du champ après détection est :

$$\frac{M_g^{[dis]} \hat{\rho} M_g^{[dis]}}{\text{Tr}(M_g^{[dis]} \hat{\rho})} = \begin{pmatrix} & |0,0\rangle & |0,1\rangle & |1,0\rangle & |1,0\rangle \\ \begin{matrix} 0 \\ 0 \\ 0 \\ 0 \end{matrix} & \begin{matrix} 0 \\ 1 \\ e^{i\phi} \\ 0 \end{matrix} & \begin{matrix} 0 \\ e^{-i\phi} \\ 1 \\ 0 \end{matrix} & \begin{matrix} 0 \\ 0 \\ 0 \\ 0 \end{matrix} \end{pmatrix}. \quad (\text{C.4})$$

L'état est purifié et ses cohérences sont de nouveau maximales. Notons que cette technique permet aussi de filtrer les erreurs de préparation dans lesquelles l'atome injecteur est initialement dans $|g\rangle$ et laisse les cavités dans l'état vide. La mesure de parité permet de d'éliminer les populations dans $|00\rangle$ et $|11\rangle$, tout en laissant intact les vecteurs de Vect($|10\rangle + |01\rangle, |10\rangle - |01\rangle$), qu'on appellera l'espace métrologique, par analogie avec l'espace de calcul des codes correcteurs d'erreur quantiques.

Le cas précédent est très simple et on peut se demander si on peut stabiliser des espaces métrologiques plus compliqués. On va donner un exemple à partir de la tentative de préparer l'état $|20\rangle + |02\rangle$ de la section IV.2.5. On a vu que la principale pollution de l'état était alors l'état $|11\rangle$. On ne peut pas appliquer la même stratégie que précédemment puisque l'état a ici le même nombre de photons que les états de l'espace métrologique. On peut cependant utiliser le fait que, pour les états de l'espace métrologique, l'atome interagit en même temps avec les deux photons, tandis qu'il interagit successivement avec chacun d'entre eux pour l'état $|11\rangle$.

Le principe de la purification des états 2002 est expliqué sur les figures C.1 et C.2. On utilise des atomes mesurant le nombre de photons modulo 4. L'interaction d'un atome avec une cavité donne alors lieu à une phase $\pi/2$ par photon. L'atome est préparé dans l'état $|+_x\rangle$. Il interagit d'abord avec la première cavité. Si le champ est dans l'état $|02\rangle$, il ne se passe rien (Pour simplifier, on ne tient pas compte du déplacement de Lamb et de la phase indépendante du nombre de photons acquise par l'atome). Si il est dans l'état $|20\rangle$, la superposition atomique acquiert une phase π et se trouve alors dans l'état $|-_x\rangle$. Si le champ est dans l'état $|11\rangle$, la phase de la superposition atomique devient $\pi/2$ et correspond à l'état $|+_y\rangle$.

Dans la zone de Ramsey R_2 , on applique à l'atome une impulsion π , dont la phase est choisie de sorte que la rotation de la sphère de Bloch s'effectue autour de l'état $|+_x\rangle$. Pour les états de l'espace métrologique, l'atome est dans $|\pm_x\rangle$ qui n'est pas affecté par cette rotation. En revanche, lorsque le champ est dans $|11\rangle$, l'atome est dans l'état $|+_y\rangle$ qui est transformé en $|-_y\rangle$. L'atome interagit enfin avec C_2 . Si le champ est dans l'état $|20\rangle$, l'atome voit un champ vide et reste dans $|-_x\rangle$. Si il est dans $|02\rangle$, il interagit avec les deux photons et la superposition acquiert une phase π . L'atome termine donc aussi sa course dans l'état $|-_x\rangle$. En revanche, lorsque le champ est dans l'état $|11\rangle$, il acquiert une phase $\pi/2$ et son état final est $|+_x\rangle$.

On termine en mesurant l'atome dans la base $\{|+_x\rangle, |-_x\rangle\}$. Lorsque le champ est dans $|11\rangle$, on obtient le résultat positif, tandis qu'on obtient le résultat négatif lorsque il est dans un état de l'espace métrologique. En utilisant le fait que l'atome interagit successivement avec les deux photons dans le premier cas, tandis qu'il interagit simultanément dans le second, on a réussi à introduire une phase π entre les deux états finaux de l'atome. Contrairement au cas de la purification à un photon présentée précédemment, l'état du champ ne reste pas inchangé dans ce protocole. En effet, les composantes $|20\rangle$ et $|02\rangle$

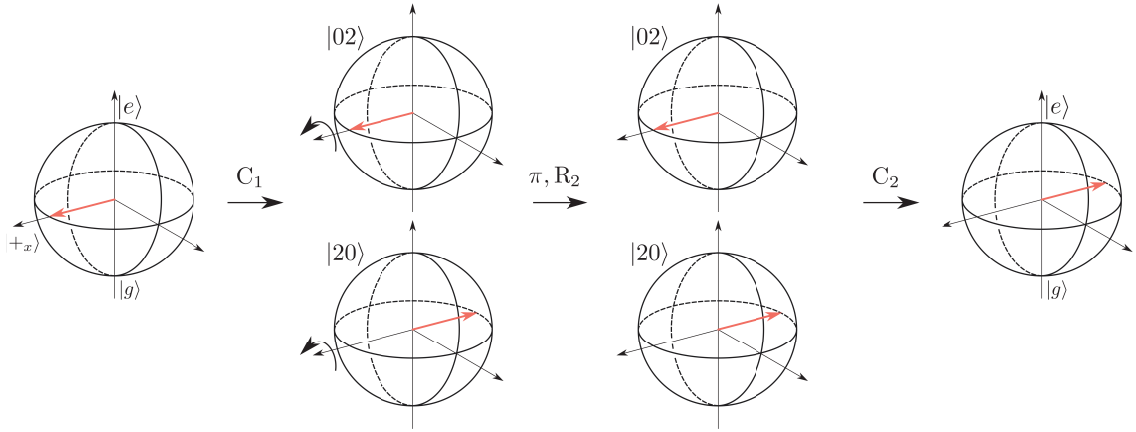


Fig. C.1 Schéma de fonctionnement de la purification lorsque le champ est dans un état 2002. Quelle que soit la cavité dans laquelle le photon se trouve, l'atome finit sa course dans l'état $|-x\rangle$.

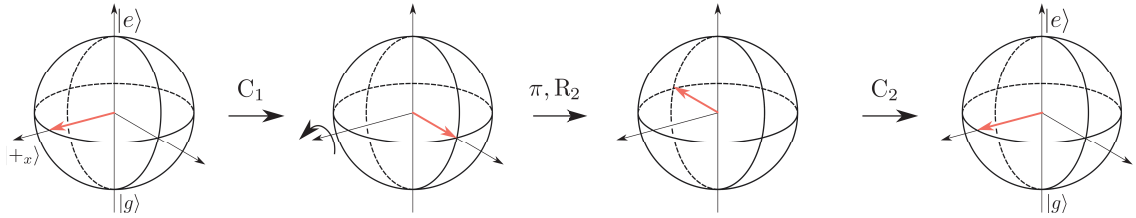


Fig. C.2 Schéma de fonctionnement de la purification lorsque le champ est dans l'état 11. À cause de l'impulsion micro-onde, l'atome acquiert une phase π supplémentaire par rapport aux cas où les deux photons se trouvent dans la même cavité. Il a donc pour état final $|+x\rangle$.

acquièrent une phase quantique relative π lors de l'impulsion π . Cette phase est cependant indépendante de la phase de l'état 2002 et ajoute simplement un décalage fixe sur les franges de l'interféromètre.

Cette même méthode permet par ailleurs d'éliminer les populations dans $|00\rangle$, puisque l'atome reste dans ce cas dans l'état $|+x\rangle$. Si le champ contient par ailleurs des composantes à 1 photon, on peut les filtrer en effectuant une mesure globale de la parité du champ. On peut donc en principe éliminer toutes les impuretés de l'état en utilisant uniquement deux atomes QND.

Annexe D

Interaction de deux atomes résonnants avec un mode du champ

Dans cette annexe, on calcule les opérateurs de Kraus décrivant l'interaction d'un échantillon atomique contenant deux atomes avec la cavité C_j . On va supposer que les deux atomes ont la même fréquence de transition et interagissent de la même manière avec la cavité. En supposant que la pulsation de Rabi est imaginaire pure, ce qu'on fera uniquement dans cette annexe, le hamiltonien de Jaynes-Cummings pour deux atomes s'écrit :

$$\begin{aligned} \hat{H}_{JC}^{(2)} = & \frac{\hbar\omega_{at}}{2}\hat{\sigma}_z^{(1)} + \frac{\hbar\omega_{at}}{2}\hat{\sigma}_z^{(2)} + \hbar\omega_{cj} \left(\hat{a}_j^\dagger \hat{a}_j + \frac{1}{2} \right) + \frac{\hbar}{2} \left(\Omega_{0,j} \hat{\sigma}_+^{(1)} \hat{a}_j + \Omega_{0,j} \hat{\sigma}_-^{(1)} \hat{a}_j^\dagger \right) \\ & + \frac{\hbar}{2} \left(\Omega_{0,j} \hat{\sigma}_+^{(2)} \hat{a}_j + \Omega_{0,j} \hat{\sigma}_-^{(2)} \hat{a}_j^\dagger \right), \end{aligned} \quad (D.1)$$

où les opérateurs $\hat{\sigma}_z^{(k)}$, $\hat{\sigma}_+^{(k)}$ et $\hat{\sigma}_-^{(k)}$ correspondent respectivement à $\hat{\sigma}_z$, $\hat{\sigma}_+$ et $\hat{\sigma}_-$ pour l'atome k .

Le principe du calcul est le suivant :

- On diagonalise le hamiltonien $\hat{H}_{JC}^{(2)}$ par bloc, en traitant séparément les cas à 0 ou 1 excitation ;
- On en tire la matrice d'évolution dans la base des états propres de $\hat{H}_{JC}^{(2)}$;
- On passe de cette base à la base de Fock.

Plaçons-nous à résonance. Lorsqu'il n'y a aucune excitation, le hamiltonien est simplement nul :

$$\hat{H}_{JC}^{(2)} |g, g, 0\rangle = 0. \quad (D.2)$$

Cas d'une excitation

Dans le cas où le système contient une excitation, il y a trois vecteurs de base $|e, g, 0\rangle$, $|g, e, 0\rangle$ et $|g, g, 1\rangle$. Le hamiltonien s'écrit, dans cette base :

$$\hat{H}_{JC}^{(2)} = \frac{\hbar}{2} \begin{pmatrix} |e,g,0\rangle & |g,e,0\rangle & |g,g,1\rangle \\ 0 & 0 & \Omega_{0,j} \\ 0 & 0 & \Omega_{0,j} \\ \Omega_{0,j} & \Omega_{0,j} & 0 \end{pmatrix}. \quad (D.3)$$

Cette matrice se diagonalise de façon évidente. Ses vecteurs propres sont :

$$|\phi_0\rangle = \frac{1}{\sqrt{2}} (|eg0\rangle - |ge0\rangle), \quad (\text{D.4a})$$

$$|\psi_{\pm}\rangle = \frac{1}{\sqrt{2}} \left(\frac{|eg0\rangle + |ge0\rangle}{\sqrt{2}} \pm |gg1\rangle \right), \quad (\text{D.4b})$$

associés aux énergies :

$$E_0 = 0, \quad (\text{D.5a})$$

$$E_{\pm} = \pm \frac{\Omega_{0,j}\sqrt{2}}{2}. \quad (\text{D.5b})$$

Le premier état est un *état sombre*, qui n'échange pas d'énergie avec le champ. L'opérateur d'évolution s'écrit donc dans cette base

$$\hat{U}^{(2)}(t) = \text{diag} \begin{pmatrix} |\phi_0\rangle & |\psi_+\rangle & |\psi_-\rangle \\ 1, & e^{-i\frac{\Omega_{0,j}\sqrt{2}}{2}t}, & e^{i\frac{\Omega_{0,j}\sqrt{2}}{2}t} \end{pmatrix}. \quad (\text{D.6})$$

En revenant dans la base initiale, on a, pour 0 ou 1 excitation :

$$\hat{U}^{(2)}(t) = \frac{1}{2} \begin{pmatrix} |g,g,0\rangle & |e,g,0\rangle & |g,e,0\rangle & |g,g,1\rangle \\ 2 & 0 & 0 & 0 \\ 0 & \cos\theta_0 + 1 & \cos\theta_0 - 1 & -i\sqrt{2}\sin\theta_0 \\ 0 & \cos\theta_0 - 1 & \cos\theta_0 + 1 & -i\sqrt{2}\sin\theta_0 \\ 0 & -i\sqrt{2}\sin\theta_0 & -i\sqrt{2}\sin\theta_0 & 2\cos\theta_0 \end{pmatrix}, \quad (\text{D.7})$$

où $\theta_0 = \Omega_{0,j}/\sqrt{2}t$.

Multiplicités supérieures

Considérons le cas de $n + 1$ excitations, avec $n \geq 1$. En représentation d'interaction, le hamiltonien s'écrit, dans la base $\{|ee, n-1\rangle, |eg, n\rangle, |ge, n\rangle, |gg, n+1\rangle\}$:

$$\hat{H}_{JC}^{(2)} = \frac{\hbar\Omega_{0,j}}{2} \begin{pmatrix} |ee, n-1\rangle & |eg, n\rangle & |ge, n\rangle & |gg, n+1\rangle \\ 0 & \sqrt{n} & \sqrt{n} & 0 \\ \sqrt{n} & 0 & 0 & \sqrt{n+1} \\ \sqrt{n} & 0 & 0 & \sqrt{n+1} \\ 0 & \sqrt{n+1} & \sqrt{n+1} & 0 \end{pmatrix}. \quad (\text{D.8})$$

Comme précédemment, il y a un état sombre

$$|\phi_0\rangle = \frac{|eg, n\rangle - |ge, n\rangle}{\sqrt{2}}, \quad (\text{D.9})$$

d'énergie nulle. Un autre état propre évident du hamiltonien est

$$|\psi_0\rangle = \frac{1}{\sqrt{2n+1}} \left(\sqrt{n+1} |ee, n-1\rangle - \sqrt{n} |gg, n+1\rangle \right), \quad (\text{D.10})$$

dont l'énergie est aussi nulle. En cherchant deux vecteurs propres orthogonaux aux précédents, on trouve :

$$|\psi_{\pm}\rangle = \frac{1}{\sqrt{4n+2}} \left(\pm\sqrt{n} |ee, n-1\rangle + \sqrt{2n+1} \frac{|eg, n\rangle + |ge, n\rangle}{\sqrt{2}} \pm \sqrt{n+1} |gg, n+1\rangle \right), \quad (\text{D.11})$$

associés aux énergies propres :

$$E_{\pm} = \pm \frac{\Omega_{0,j}}{2} \sqrt{2(2n+1)}. \quad (\text{D.12})$$

L'opérateur d'évolution s'écrit donc dans cette base

$$\hat{U}^{(2)}(t) = \text{diag} \begin{pmatrix} |\psi_0\rangle|\phi_0\rangle & |\psi_+\rangle & |\psi_-\rangle \\ 1, & 1, & e^{-i\frac{\Omega_{0,j}\sqrt{2(2n+1)}}{2}t}, & e^{i\frac{\Omega_{0,j}\sqrt{2(2n+1)}}{2}t} \end{pmatrix}. \quad (\text{D.13})$$

On en tire l'opérateur d'évolution dans la base initiale

$$\hat{U}^{(2)}(t) = \frac{1}{4n+2} \times \begin{pmatrix} |ee, n-1\rangle & |eg, n\rangle & |ge, n\rangle & |gg, n+1\rangle \\ 2n+2+2n\cos\theta_n & -i\sqrt{2n(2n+1)}\sin\theta_n & -i\sqrt{2n(2n+1)}\sin\theta_n & \sqrt{2n(2n+2)}(\cos\theta_n-1) \\ -i\sqrt{2n(2n+1)}\sin\theta_n & (2n+1)(\cos\theta_n+1) & (2n+1)(\cos\theta_n-1) & -i\sqrt{(2n+1)(2n+2)}\sin\theta_n \\ -i\sqrt{2n(2n+1)}\sin\theta_n & (2n+1)(\cos\theta_n-1) & (2n+1)(\cos\theta_n+1) & -i\sqrt{(2n+1)(2n+2)}\sin\theta_n \\ \sqrt{2n(2n+2)}(\cos\theta_n-1) & -i\sqrt{(2n+1)(2n+2)}\sin\theta_n & -i\sqrt{(2n+1)(2n+2)}\sin\theta_n & 2n+(2n+2)\cos\theta_n \end{pmatrix}, \quad (\text{D.14})$$

où $\theta_n = \Omega_{0,j}\sqrt{2(2n+1)}/2t$.

On en tire les opérateurs de Kraus $\hat{R}_{i,j}$ correspondant à deux atomes dans i détectés dans j après l'interaction :

$$\hat{R}_{ee,ee} = \sum_n \frac{n+2 + (n+1) \cos \theta_{n+1}}{2n+3} |n\rangle \langle n|, \quad (\text{D.15})$$

$$\hat{R}_{ee,eg} = \frac{-i}{\sqrt{2}} \sum_n \sqrt{\frac{n+1}{2n+3}} \sin \theta_{n+1} |n+1\rangle \langle n|, \quad (\text{D.16})$$

$$\hat{R}_{ee,ge} = \hat{R}_{ee,eg}, \quad (\text{D.17})$$

$$\hat{R}_{ee,gg} = \sum_n \frac{\sqrt{(n+1)(n+2)}}{2n+3} (\cos \theta_{n+1} - 1) |n+2\rangle \langle n|, \quad (\text{D.18})$$

$$\hat{R}_{eg,ee} = \frac{-i}{\sqrt{2}} \sum_n \sqrt{\frac{n+1}{2n+3}} \sin \theta_{n+1} |n\rangle \langle n+1|, \quad (\text{D.19})$$

$$\hat{R}_{ge,ee} = \hat{R}_{eg,ee}, \quad (\text{D.20})$$

$$\hat{R}_{eg,eg} = \sum_n \frac{1}{2} (\cos \theta_n + 1) |n\rangle \langle n|, \quad (\text{D.21})$$

$$\hat{R}_{ge,ge} = \hat{R}_{eg,ge} \quad (\text{D.22})$$

$$\hat{R}_{eg,ge} = \sum_n \frac{1}{2} (\cos \theta_n - 1) |n\rangle \langle n|, \quad (\text{D.23})$$

$$\hat{R}_{ge,eg} = \hat{R}_{eg,ge}, \quad (\text{D.24})$$

$$\hat{R}_{eg,gg} = \frac{-i}{\sqrt{2}} \sum_n \sqrt{\frac{n+1}{2n+1}} \sin \theta_n |n+1\rangle \langle n|, \quad (\text{D.25})$$

$$\hat{R}_{ge,gg} = \hat{R}_{eg,gg}, \quad (\text{D.26})$$

$$\hat{R}_{gg,ee} = \sum_n \frac{\sqrt{(n+1)(n+2)}}{2n+3} (\cos \theta_{n+1} - 1) |n\rangle \langle n+2|, \quad (\text{D.27})$$

$$\hat{R}_{gg,eg} = \frac{-i}{\sqrt{2}} \sum_n \sqrt{\frac{n+1}{2n+1}} \sin \theta_n |n\rangle \langle n+1|, \quad (\text{D.28})$$

$$\hat{R}_{gg,ge} = \hat{R}_{gg,eg}, \quad (\text{D.29})$$

$$\hat{R}_{gg,gg} = |0\rangle \langle 0| + \cos \theta_1 |1\rangle \langle 1| + \sum_{n \geq 1} ((2n+2) \cos \theta_{n-1} + 2n) |n\rangle \langle n|, \quad (\text{D.30})$$

où $\theta_{-1} \equiv 0$.

Annexe E

Calcul de fonctions de Wigner à deux modes

Dans cet annexe, nous montrons comment on peut calculer des fonctions de Wigner à deux modes à partir des fonctions de Wigner monomodes. Nous utilisons ensuite ce résultat pour calculer les fonctions de Wigner des états *noon*.

Calcul de fonctions de Wigner à deux modes

La fonction de Wigner à deux modes de l'état $\hat{\rho}$ est définie par

$$W^{\hat{\rho}}(\alpha_1, \alpha_2) = \text{Tr} \left(\hat{D}_1(-\alpha_1) \hat{D}_2(-\alpha_2) \hat{\rho} \hat{D}_1(\alpha_1) \hat{D}_2(\alpha_2) \hat{P} \right), \quad (\text{E.1})$$

où $\hat{P} = \hat{P}_1 \otimes \hat{P}_2$. On peut décomposer $\hat{\rho}$ sur la base de Fock de l'espace à deux modes :

$$\hat{\rho} = \sum_{\substack{n_1, n'_1 \\ n_2, n'_2}} \rho_{n_1, n'_1, n_2, n'_2} |n_1\rangle \langle n'_1| \otimes |n_2\rangle \langle n'_2|. \quad (\text{E.2})$$

Le calcul de la trace donne alors :

$$W^{\hat{\rho}}(\alpha_1, \alpha_2) = \sum_{k_1, k_2} \sum_{\substack{n_1, n'_1 \\ n_2, n'_2}} \rho_{n_1, n'_1, n_2, n'_2} \langle k_1 | \hat{D}_1(-\alpha_1) | n_1 \rangle \langle n'_1 | \hat{D}_1(\alpha_1) \hat{P}_1 | k_1 \rangle \times \\ \langle k_2 | \hat{D}_2(-\alpha_2) | n_2 \rangle \langle n'_2 | \hat{D}_2(\alpha_2) \hat{P}_2 | k_2 \rangle. \quad (\text{E.3})$$

Les termes de droite ressemblent à ceux de fonctions de Wigner monomode, à la seule différence que $|n_i\rangle \langle n'_i|$ n'est pas nécessairement une matrice densité. On peut donc factoriser cette expression sous la forme :

$$W^{\hat{\rho}}(\alpha_1, \alpha_2) = \sum_{\substack{n_1, n'_1 \\ n_2, n'_2}} \rho_{n_1, n'_1, n_2, n'_2} W^{n_1, n'_1}(\alpha_1) W^{n_2, n'_2}(\alpha_2), \quad (\text{E.4})$$

où

$$W^{n, n'}(\alpha) \equiv \text{Tr} \left(\hat{D}^{(1)}(-\alpha) |n\rangle \langle n'| \hat{D}^{(1)}(\alpha) \hat{P}^{(1)} \right), \quad (\text{E.5})$$

est une fonction de Wigner à un mode généralisée avec $\hat{D}^{(1)}$ et $\hat{P}^{(1)}$ les opérateurs déplacement et parité monomode. On peut donc calculer la fonction de Wigner à deux modes comme un produit de fonctions de Wigner à un mode. Cette approche permet de faire des calculs analytiques de fonctions de Wigner. Elle est aussi particulièrement efficace pour faire des calculs numériques de fonctions de Wigner lorsque la matrice densité de $\hat{\rho}$ est une matrice creuse dans la base de Fock.

Fonctions de Wigner des états *noon*

La fonction de Wigner de l'état $|n0\rangle + |0n\rangle$ peut se factoriser sous la forme précédente. Elle fait alors intervenir 4 termes :

$$W^{|n0\rangle+|0n\rangle}(\alpha_1, \alpha_2) = \frac{1}{2} \left[W^{n,n}(\alpha_1)W^{0,0}(\alpha_2) + W^{0,0}(\alpha_1)W^{n,n}(\alpha_2) + W^{n,0}(\alpha_1)W^{0,n}(\alpha_2) + W^{0,n}(\alpha_1)W^{n,0}(\alpha_2) \right] \quad (\text{E.6})$$

La fonction de Wigner d'un état de Fock peut se calculer à partir de la définition de la fonction de Wigner à partir de la fonction caractéristique symétrique. Cette fonction est définie pour un état comme la valeur moyenne de l'opérateur parité dans cet état :

$$C_S^{\hat{\rho}}(\lambda) = \left\langle \hat{D}(\lambda) \right\rangle_{\hat{\rho}}. \quad (\text{E.7})$$

Pour un état de Fock $|n\rangle$, la fonction caractéristique symétrique vaut :

$$\begin{aligned} C_S^{|n\rangle\langle n|}(\lambda) &= \langle n | e^{\lambda \hat{a}^\dagger} e^{-\lambda^* \hat{a}} | n \rangle e^{-|\lambda|^2/2}, \\ &= \sum_{p,q=0}^{\infty} \langle n | (\lambda \hat{a}^\dagger)^p (-\lambda^* \hat{a})^q | n \rangle e^{-|\lambda|^2/2}, \\ &= \sum_{p,q=0}^{\infty} \frac{1}{p!q!} \delta_{p,q} \langle n | (\lambda \hat{a}^\dagger)^p (-\lambda^* \hat{a})^q | n \rangle e^{-|\lambda|^2/2}, \\ &= \sum_{p=0}^n (-1)^p \frac{n!}{p!^2(n-p)!} |\lambda|^{2p} e^{-|\lambda|^2/2}, \\ &= \mathcal{L}_n(|\lambda|^2) e^{-|\lambda|^2/2}, \end{aligned} \quad (\text{E.8})$$

où on a introduit les polynômes de Laguerre :

$$\mathcal{L}_n(|x|^2) = \sum_{p=0}^n (-1)^p \frac{n!}{p!^2(n-p)!} x^p. \quad (\text{E.9})$$

Les pseudo-fonctions caractéristiques associées aux cohérences se calculent de la même manière :

$$\begin{aligned} C_S^{|0\rangle\langle n|}(\lambda) &= \text{Tr}(|0\rangle \langle n| \hat{D}(\lambda)) \\ &= \langle n | \hat{D}(\lambda) | 0 \rangle \\ &= \langle n | |\lambda\rangle \\ &= \frac{\lambda^n}{\sqrt{(n!)}} e^{-|\lambda|^2/2}. \end{aligned} \quad (\text{E.10})$$

La fonction de Wigner est définie comme la transformée de Fourier complexe de la fonction caractéristique symétrique¹ :

$$W(\alpha) = \frac{1}{2\pi} \int d^2\lambda e^{\alpha\lambda^* - \alpha^*\lambda} C_S(\lambda). \quad (\text{E.11})$$

Afin de calculer les transformées de Fourier, on peut s'aider de la formule [?] :

$$\int d^2\lambda f(\lambda) e^{\alpha\lambda^* - |\lambda|^2/2} = 2\pi f(2\alpha). \quad (\text{E.12})$$

Il vient :

$$W^{n,n}(\alpha) = \frac{2}{\pi} \mathcal{L}_n(4|\alpha|^2) e^{-2|\alpha|^2}, \quad (\text{E.13})$$

$$W^{n,0}(\alpha) = \frac{2}{\pi} \frac{(2\alpha)^n}{\sqrt{n!}} e^{-2|\alpha|^2}. \quad (\text{E.14})$$

$$(\text{E.15})$$

On a par ailleurs $W^{n,0}(\alpha) = W^{n,0}(\alpha)^*$. La fonction de Wigner de l'état $|10\rangle + |01\rangle$ s'écrit donc :

$$W^{|10\rangle+|01\rangle}(\alpha_1, \alpha_2) = [2(\alpha_1 + \alpha_2)(\alpha_1^* + \alpha_2^*) - 1] e^{-2(|\alpha_1|^2 + |\alpha_2|^2)}. \quad (\text{E.16})$$

1. Rappelons qu'on a enlevé le facteur de normalisation $2/\pi$ de la fonction de Wigner.

Bibliographie

- [1] R. HORODECKI, P. HORODECKI, M. HORODECKI & K. HORODECKI; « Quantum Entanglement »; *Reviews of Modern Physics* **81**, p. 865–942 (2009). [viii](#)
- [2] D. P. DiVINCENZO; « The Physical Implementation of Quantum Computation »; *Fortschritte der Physik* **48**, p. 771–783 (2000). ISSN 1521-3978. [viii](#)
- [3] M. PARIS & J. REHACEK (éditeurs); *Quantum State Estimation*; Lecture Notes in Physics (Springer-Verlag, Berlin Heidelberg) (2004); ISBN 978-3-540-22329-0. [ix](#)
- [4] Y. TEO; *Introduction to Quantum-State Estimation* (2015); ISBN 978-981-4678-84-1. [ix](#)
- [5] A. SILBERFARB, P. S. JESSEN & I. H. DEUTSCH; « Quantum State Reconstruction via Continuous Measurement »; *Physical Review Letters* **95**, p. 030 402 (2005). [ix](#)
- [6] D. GROSS, Y.-K. LIU, S. T. FLAMMIA, S. BECKER & J. EISERT; « Quantum State Tomography via Compressed Sensing »; *Physical Review Letters* **105**, p. 150 401 (2010). [ix](#), 103
- [7] Y. S. TEO, H. ZHU, B.-G. ENGLERT, J. ŘEHÁČEK & Z. HRADIL; « Quantum-State Reconstruction by Maximizing Likelihood and Entropy »; *Physical Review Letters* **107**, p. 020 404 (2011). [ix](#)
- [8] B. QI, Z. HOU, L. LI, D. DONG, G. XIANG & G. GUO; « Quantum State Tomography via Linear Regression Estimation »; *Scientific Reports* **3**, p. 3496 (2013). ISSN 2045-2322. [ix](#)
- [9] A. KYRILLIDIS, A. KALEV, D. PARK, S. BHOJANAPALLI, C. CARAMANIS & S. SANGHAVI; « Provable Compressed Sensing Quantum State Tomography via Non-Convex Methods »; *npj Quantum Information* **4**, p. 1–7 (2018). ISSN 2056-6387. [ix](#)
- [10] G. TORLAI, G. MAZZOLA, J. CARRASQUILLA, M. TROYER, R. MELKO & G. CARLEO; « Neural-Network Quantum State Tomography »; *Nature Physics* **14**, p. 447–450 (2018). ISSN 1745-2481. [ix](#)
- [11] A. I. LVOVSKY, H. HANSEN, T. AICHELE, O. BENSON, J. MLYNEK & S. SCHILLER; « Quantum State Reconstruction of the Single-Photon Fock State »; *Physical Review Letters* **87**, p. 050 402 (2001). [ix](#)
- [12] H. HÄFFNER, W. HÄNSEL, C. F. ROOS, J. BENHELM, D. CHEK-AL-KAR, M. CHWALLA, T. KÖRBER, U. D. RAPOL, M. RIEBE, P. O. SCHMIDT, C. BECHER, O. GÜHNE, W. DÜR & R. BLATT; « Scalable Multiparticle Entanglement of Trapped Ions »; *Nature* **438**, p. 643–646 (2005). ISSN 1476-4687. [ix](#)

- [13] G. ZAMBRA, A. ANDREONI, M. BONDANI, M. GRAMEGNA, M. GENOVESE, G. BRIDA, A. ROSSI & M. G. A. PARIS ; « Experimental Reconstruction of Photon Statistics without Photon Counting » ; *Physical Review Letters* **95**, p. 063 602 (2005). [ix](#)
- [14] W.-B. GAO, C.-Y. LU, X.-C. YAO, P. XU, O. GÜHNE, A. GOEBEL, Y.-A. CHEN, C.-Z. PENG, Z.-B. CHEN & J.-W. PAN ; « Experimental Demonstration of a Hyper-Entangled Ten-Qubit Schrödinger Cat State » ; *Nature Physics* **6**, p. 331–335 (2010). ISSN 1745-2481. [ix](#)
- [15] D. GIOVANNINI, J. ROMERO, J. LEACH, A. DUDLEY, A. FORBES & M. J. PADGETT ; « Characterization of High-Dimensional Entangled Systems via Mutually Unbiased Measurements » ; *Physical Review Letters* **110**, p. 143 601 (2013). [ix](#)
- [16] N. BENT, H. QASSIM, A. A. TAHIR, D. SYCH, G. LEUCHS, L. L. SÁNCHEZ-SOTO, E. KARIMI & R. W. BOYD ; « Experimental Realization of Quantum Tomography of Photonic Qudits via Symmetric Informationally Complete Positive Operator-Valued Measures » ; *Physical Review X* **5**, p. 041 006 (2015). [ix](#)
- [17] C. WANG, Y. Y. GAO, P. REINHOLD, R. W. HEERES, N. OFEK, K. CHOU, C. AXLINE, M. REAGOR, J. BLUMOFF, K. M. SLIWA, L. FRUNZIO, S. M. GIRVIN, L. JIANG, M. MIRRAHIMI, M. H. DEVORET & R. J. SCHOELKOPF ; « A Schrödinger Cat Living in Two Boxes » ; *Science* **352**, p. 1087–1091 (2016). ISSN 0036-8075, 1095-9203. [ix](#)
- [18] C. A. RIOFRÍO, D. GROSS, S. T. FLAMMIA, T. MONZ, D. NIGG, R. BLATT & J. EISERT ; « Experimental Quantum Compressed Sensing for a Seven-Qubit System » ; *Nature Communications* **8**, p. 1–8 (2017). ISSN 2041-1723. [ix](#)
- [19] A. STEFFENS, C. A. RIOFRÍO, W. MCCUTCHEON, I. ROTH, B. A. BELL, A. McMILLAN, M. S. TAME, J. G. RARITY & J. EISERT ; « Experimentally Exploring Compressed Sensing Quantum Tomography » ; *Quantum Science and Technology* **2**, p. 025 005 (2017). ISSN 2058-9565. [ix](#), 103
- [20] C. GROSS & I. BLOCH ; « Quantum Simulations with Ultracold Atoms in Optical Lattices » ; *Science* **357**, p. 995–1001 (2017). ISSN 0036-8075, 1095-9203. [ix](#)
- [21] H. HÄFFNER, C. F. ROOS & R. BLATT ; « Quantum Computing with Trapped Ions » ; *Physics Reports* **469**, p. 155–203 (2008). ISSN 0370-1573. [ix](#)
- [22] J.-W. ZHOU, P.-F. WANG, F.-Z. SHI, P. HUANG, X. KONG, X.-K. XU, Q. ZHANG, Z.-X. WANG, X. RONG & J.-F. DU ; « Quantum Information Processing and Metrology with Color Centers in Diamonds » ; *Frontiers of Physics* **9**, p. 587–597 (2014). ISSN 2095-0470. [ix](#)
- [23] G. WENDIN ; « Quantum Information Processing with Superconducting Circuits : A Review » ; *Reports on Progress in Physics* **80**, p. 106 001 (2017). ISSN 0034-4885. [ix](#)
- [24] J. L. O'BRIEN, A. FURUSAWA & J. VUČKOVIĆ ; « Photonic Quantum Technologies » ; *Nature Photonics* **3**, p. 687–695 (2009). ISSN 1749-4893. [ix](#)
- [25] C. DONG, Y. WANG & H. WANG ; « Optomechanical Interfaces for Hybrid Quantum Networks » ; *National Science Review* **2**, p. 510–519 (2015). ISSN 2095-5138. [ix](#)
- [26] X. ZHANG, H.-O. LI, G. CAO, M. XIAO, G.-C. GUO & G.-P. GUO ; « Semiconductor Quantum Computation » ; *National Science Review* **6**, p. 32–54 (2019). ISSN 2095-5138. [ix](#)

- [27] S. HAROCHE & J.-M. RAIMOND ; *Exploring the Quantum : Atoms, Cavities, and Photons* (Oxford University Press) (2006) ; ISBN 978-0-19-170862-6. ix, 4, 9, 26, 109
- [28] P. SIX, P. CAMPAGNE-IBARCQ, I. DOTSENKO, A. SARLETTE, B. HUARD & P. ROUCHON ; « Quantum State Tomography with Noninstantaneous Measurements, Imperfections, and Decoherence » ; *Physical Review A* **93** (2016). ISSN 2469-9926, 2469-9934. ix, xi, xiii, 129
- [29] S. DELÉGLISE, I. DOTSENKO, C. SAYRIN, J. BERNU, M. BRUNE, J.-M. RAIMOND & S. HAROCHE ; « Reconstruction of Non-Classical Cavity Field States with Snapshots of Their Decoherence » ; *Nature* **455**, p. 510–514 (2008). ISSN 1476-4687. x, 23, 83, 102
- [30] S. DELÉGLISE ; *Reconstruction Complète d'états Non-Classiques Du Champ En Électrodynamique Quantique En Cavit * ; Thesis ; Paris 6 (2009). x, 23, 103, 104
- [31] S. GAMMELMARK, B. JULSGAARD & K. M LMER ; « Past Quantum States of a Monitored System » ; *Physical Review Letters* **111** (2013). ISSN 0031-9007, 1079-7114. x
- [32] C. GUERLIN ; *Mesure Quantique Non Destructive R p t e de La Lum re :  tats de Fock et Trajectoires Quantiques* ; Thesis ; Paris 6 (2007). xi, 23
- [33] T. RYBARCZYK, B. PEAUDECERF, M. PENASA, S. GERLICH, B. JULSGAARD, K. M LMER, S. GLEYZES, M. BRUNE, J. M. RAIMOND, S. HAROCHE & I. DOTSENKO ; « Forward-Backward Analysis of the Photon-Number Evolution in a Cavity » ; *Physical Review A* **91**, p. 062 116 (2015). xi, 113
- [34] T. RYBARCZYK ; *Application de la m thode de l' tat pass    la mesure quantique non-destructive du nombre de photons dans une cavit * ; Th se de doctorat ; ENS (2014). xii, 23, 113
- [35] A. RAUSCHENBEUTEL, P. BERTET, S. OSNAGHI, G. NOGUES, M. BRUNE, J. M. RAIMOND & S. HAROCHE ; « Controlled Entanglement of Two Field Modes in a Cavity Quantum Electrodynamics Experiment » ; *Physical Review A* **64**, p. 050 301 (2001). ISSN 1050-2947, 1094-1622 ; quant-ph/0105062. xii, xiii, 63
- [36] H. J. KIMBLE, M. DAGENAIS & L. MANDEL ; « Photon Antibunching in Resonance Fluorescence » ; *Physical Review Letters* **39**, p. 691–695 (1977). 3
- [37] C. COHEN-TANNOUDJI, J. DUPONT-ROC & G. GRYNBERG ; *An Introduction to Quantum Electrodynamics* (Wiley, New York) (1992). 3
- [38] E. T. JAYNES & F. W. CUMMINGS ; « Comparison of Quantum and Semiclassical Radiation Theories with Application to the Beam Maser » ; *Proceedings of the IEEE* **51**, p. 89–109 (1963). ISSN 0018-9219. 16
- [39] P. BERTET ; *Atomes et Cavit  : Compl mentarit  et Fonctions de Wigner* ; Thesis ; Paris 6 (2002). 19
- [40] T. MEUNIER ; *Oscillations de Rabi Induites Par Un Renversement Du Temps : Un Test de La Coh rence d'une Superposition Quantique M soscopique* ; Thesis ; Paris 6 (2004). 19
- [41] F. ASS MAT ; *Manipulation d' tats Quantiques de La Lum re Par l'interm diaire d'un Atome de Rydberg Unique* ; Thesis ; Paris 6 (2004). 19

- [42] M. BRUNE, P. NUSSENZVEIG, F. SCHMIDT-KALER, F. BERNARDOT, A. MAALI, J. M. RAIMOND & S. HAROCHE ; « From Lamb Shift to Light Shifts : Vacuum and Subphoton Cavity Fields Measured by Atomic Phase Sensitive Detection » ; *Physical Review Letters* **72**, p. 3339–3342 (1994). 19
- [43] G. GRYNBERG, A. ASPECT, C. FABRE & C. COHEN-TANNOUDJI ; *Introduction to Quantum Optics : From the Semi-Classical Approach to Quantized Light* (Cambridge University Press) (2010). 21, 69
- [44] R. G. HULET & D. KLEPPNER ; « Rydberg Atoms in Circular States » ; *Physical Review Letters* **51**, p. 1430 (1983). 23
- [45] P. BOSLAND, A. ASPART, E. JACQUES & M. RIBEAUDEAU ; « Preparation and RF Tests of L-Band Superconducting Niobium-Coated Copper Cavities » ; *IEEE Transactions on Applied Superconductivity* **9**, p. 896–899 (1999). ISSN 1051-8223. 23
- [46] S. GLEYZES ; *Vers La Préparation de Cohérences Quantiques Mésoscopiques : Réalisation d'un Montage à Deux Cavités Supraconductrices* ; Thesis ; Paris 6 (2007). 23, 26
- [47] S. GLEYZES, S. KUHR, C. GUERLIN, J. BERNU, S. DELÉGLISE, U. BUSK HOFF, M. BRUNE, J.-M. RAIMOND & S. HAROCHE ; « Quantum Jumps of Light Recording the Birth and Death of a Photon in a Cavity » ; *Nature* **446**, p. 297–300 (2007). ISSN 1476-4687. 23
- [48] J. BERNU, S. DELÉGLISE, C. SAYRIN, S. KUHR, I. DOTSENKO, M. BRUNE, J. M. RAIMOND & S. HAROCHE ; « Freezing Coherent Field Growth in a Cavity by the Quantum Zeno Effect » ; *Physical Review Letters* **101**, p. 180 402 (2008). 23
- [49] C. SAYRIN, I. DOTSENKO, X. ZHOU, B. PEAUDECERF, T. RYBARCZYK, S. GLEYZES, P. ROUCHON, M. MIRRAHIMI, H. AMINI, M. BRUNE, J.-M. RAIMOND & S. HAROCHE ; « Real-Time Quantum Feedback Prepares and Stabilizes Photon Number States » ; *Nature* **477**, p. 73–77 (2011). ISSN 1476-4687. 23
- [50] B. PEAUDECERF, T. RYBARCZYK, S. GERLICH, S. GLEYZES, J. M. RAIMOND, S. HAROCHE, I. DOTSENKO & M. BRUNE ; « Adaptive Quantum Nondemolition Measurement of a Photon Number » ; *Physical Review Letters* **112**, p. 080 401 (2014). 23
- [51] M. PENASA, S. GERLICH, T. RYBARCZYK, V. MÉTILLON, M. BRUNE, J. M. RAIMOND, S. HAROCHE, L. DAVIDOVICH & I. DOTSENKO ; « Measurement of a Microwave Field Amplitude beyond the Standard Quantum Limit » ; *Physical Review A* **94**, p. 022 313 (2016). 23
- [52] J. BERNU ; *Mesures QND En Électrodynamique Quantique En Cavité : Production et Décohérence d'états de Fock : Effet Zénon* ; Thèse de doctorat. 23, 142
- [53] C. SAYRIN ; *Préparation et Stabilisation d'un Champ Non Classique En Cavité Par Rétroaction Quantique* ; Thesis ; Paris 6 (2011). 23
- [54] X. ZHOU ; *Champ Asservi Sur Un État de Fock Par Rétroaction Quantique Utilisant Des Corrections à Photon Unique* ; Thesis ; Paris 6 (2012). 23
- [55] B. PEAUDECERF ; *Mesure Adaptative Non Destructive Du Nombre de Photons Dans Une Cavité* ; Thesis ; Paris 6 (2013). 23
- [56] M. PENASA ; *Mesure Au-Delà de La Limite Quantique Standard de l'amplitude d'un Champ Électromagnétique Dans Le Domaine Micro-Onde* ; Thesis ; Paris 6 (2016). 23

- [57] P. NUSSENZVEIG ; *Measurements of Fields at the Single Photon Level by Atom Interferometry* ; Theses ; Université Pierre et Marie Curie - Paris VI (1994). 34
- [58] D. A. STECK ; « Rubidium 85 D Line Data » ; disponible en ligne sur <http://steck.us/alkalidata>. 34
- [59] A. AUFFEVE-GARNIER ; *Oscillation de Rabi à La Frontière Classique-Quantique Génération de Chats de Schrödinger* ; Thesis ; Paris 6 (2004). 36
- [60] B. PEAUDE CERF, C. SAYRIN, X. ZHOU, T. RYBARCZYK, S. GLEYZES, I. DOTSENKO, J. M. RAIMOND, M. BRUNE & S. HAROCHE ; « Quantum Feedback Experiments Stabilizing Fock States of Light in a Cavity » ; *Physical Review A* **87**, p. 042320 (2013). 37
- [61] M. GROSS & J. LIANG ; « Is a Circular Rydberg Atom Stable in a Vanishing Electric Field ? » ; *Physical Review Letters* **57**, p. 3160–3163 (1986). 39
- [62] N. F. RAMSEY ; « A Molecular Beam Resonance Method with Separated Oscillating Fields » ; *Physical Review* **78**, p. 695–699 (1950). 50
- [63] G. NOGUES, A. RAUSCHENBEUTEL, S. OSNAGHI, M. BRUNE, J. M. RAIMOND & S. HAROCHE ; « Seeing a Single Photon without Destroying It » ; *Nature* **400**, p. 239–242 (1999). ISSN 1476-4687. 51
- [64] C. GUERLIN, J. BERNU, S. DELÉGLISE, C. SAYRIN, S. GLEYZES, S. KUHR, M. BRUNE, J.-M. RAIMOND & S. HAROCHE ; « Progressive Field-State Collapse and Quantum Non-Demolition Photon Counting » ; *Nature* **448**, p. 889–893 (2007). ISSN 1476-4687. 51
- [65] S. DELÉGLISE, I. DOTSENKO, C. SAYRIN, J. BERNU, M. BRUNE, J.-M. RAIMOND & S. HAROCHE ; « Reconstruction of Non-Classical Cavity Field States with Snapshots of Their Decoherence » ; *Nature* **455**, p. 510–514 (2008). ISSN 1476-4687. 63
- [66] P. MILMAN, A. AUFFEVE, F. YAMAGUCHI, M. BRUNE, J. M. RAIMOND & S. HAROCHE ; « A Proposal to Test Bell's Inequalities with Mesoscopic Non-Local States in Cavity QED » ; *The European Physical Journal D - Atomic, Molecular, Optical and Plasma Physics* **32**, p. 233–239 (2005). ISSN 1434-6079. 63, 70
- [67] K. BANASZEK & K. WÓDKIEWICZ ; « Testing Quantum Nonlocality in Phase Space » ; *Physical Review Letters* **82**, p. 2009–2013 (1999). 70
- [68] H. JEONG, W. SON, M. S. KIM, D. AHN & Č. BRUKNER ; « Quantum Nonlocality Test for Continuous-Variable States with Dichotomic Observables » ; *Physical Review A* **67**, p. 012106 (2003). 70
- [69] C. WANG, Y. Y. GAO, P. REINHOLD, R. W. HEERES, N. OFEK, K. CHOU, C. AXLINE, M. REAGOR, J. BLUMOFF, K. M. SLIWA, L. FRUNZIO, S. M. GIRVIN, L. JIANG, M. MIRRAHIMI, M. H. DEVORET & R. J. SCHOELKOPF ; « A Schrödinger Cat Living in Two Boxes » ; *Science* **352**, p. 1087–1091 (2016). ISSN 0036-8075, 1095-9203. 70
- [70] M. GESSNER, L. PEZZÈ & A. SMERZI ; « Efficient Entanglement Criteria for Discrete, Continuous, and Hybrid Variables » ; *Physical Review A* **94** (2016). ISSN 2469-9926, 2469-9934 ; 1608.02421. 74, 76
- [71] S. L. BRAUNSTEIN & C. M. CAVES ; « Statistical Distance and the Geometry of Quantum States » ; *Physical Review Letters* **72**, p. 3439–3443 (1994). 76

- [72] B. M. ESCHER, R. L. DE MATOS FILHO & L. DAVIDOVICH ; « Quantum Metrology for Noisy Systems » ; *Brazilian Journal of Physics* **41**, p. 229–247 (2011). ISSN 1678-4448. 76
- [73] R. A. FISHER ; « A Mathematical Examination of the Methods of Determining the Accuracy of Observation by the Mean Error, and by the Mean Square Error » ; *Monthly Notices of the Royal Astronomical Society* **80**, p. 758–770 (1920). ISSN 0035-8711. 91
- [74] FISHER R. A. & RUSSELL EDWARD JOHN ; « On the Mathematical Foundations of Theoretical Statistics » ; *Philosophical Transactions of the Royal Society of London. Series A, Containing Papers of a Mathematical or Physical Character* **222**, p. 309–368 (1922). 91
- [75] H. CRAMÉR ; *Mathematical Methods of Statistics* (Princeton University Press, Princeton) (1999) ; ISBN 978-0-691-08004-8 978-0-691-00547-8 ; oCLC : 185436716. 91
- [76] C. RADHAKRISHNA RAO ; « Information and the Accuracy Attainable in the Estimation of Statistical Parameters » ; *Bulletin of the Calcutta Mathematical Society* **37**, p. 81–91 (1945)ISSN 0008-0659. 91
- [77] M. S. BYRD & N. KHANEJA ; « Characterization of the Positivity of the Density Matrix in Terms of the Coherence Vector Representation » ; *Physical Review A* **68** (2003). ISSN 1050-2947, 1094-1622. 95
- [78] G. KIMURA ; « The Bloch Vector for N-Level Systems » ; *Physics Letters A* **314**, p. 339–349 (2003). ISSN 03759601 ; *quant-ph/0301152*. 95, 98
- [79] Z. HRADIL ; « Quantum-State Estimation » ; *Physical Review A* **55**, p. R1561–R1564 (1997). 100
- [80] S. P. BOYD & L. VANDENBERGHE ; *Convex Optimization* (Cambridge University Press, Cambridge, UK ; New York) (2004) ; ISBN 978-0-521-83378-3. 101
- [81] Z. HRADIL, J. ŘEHÁČEK, J. FIURÁŠEK & M. JEŽEK ; « 3 Maximum-Likelihood Methods in Quantum Mechanics » ; dans M. PARIS & J. ŘEHÁČEK (rédacteurs), « Quantum State Estimation », , tome 649p. 59–112 (Springer Berlin Heidelberg, Berlin, Heidelberg) (2004) ; ISBN 978-3-540-22329-0 978-3-540-44481-7. 102
- [82] A. ACHARYA, T. KYPRAIOS & M. GUŢĂ ; « A Comparative Study of Estimation Methods in Quantum Tomography » ; *Journal of Physics A: Mathematical and Theoretical* **52**, p. 234001 (2019). ISSN 1751-8121. 103
- [83] S. HAROCHE, M. BRUNE & J. RAIMOND ; « Measuring Photon Numbers in a Cavity by Atomic Interferometry : Optimizing the Convergence Procedure » ; *Journal of Physics B Atomic and Molecular Physics* **2** (1992). 103
- [84] S. STRAUPE ; « Adaptive Quantum Tomography » ; *JETP Letters* **104**, p. 510–522 (2016). ISSN 0021-3640, 1090-6487 ; *1610.02840*. 103
- [85] D. AHN, Y. S. TEO, H. JEONG, F. BOUCHARD, F. HUFNAGEL, E. KARIMI, D. KOUTNY, J. REHACEK, Z. HRADIL, G. LEUCHS & L. L. SANCHEZ-SOTO ; « Adaptive Compressive Tomography with No a Priori Information » ; *arXiv :1812.05289 [quant-ph]* (2018)*1812.05289*. 103
- [86] L. PEREIRA, L. ZAMBRANO, J. CORTÉS-VEGA, S. NIKLITSCHKEK & A. DELGADO ; « Adaptive Quantum Tomography in High Dimensions » ; *Physical Review A* **98**, p. 012339 (2018). 103

- [87] G. STRUCHALIN, E. KOVLAKOV, S. STRAUPE & S. KULIK; « Adaptive Quantum Tomography of High-Dimensional Bipartite Systems »; *Physical Review A* **98** (2018). ISSN 2469-9926, 2469-9934; [1804.05226](#). 103
- [88] G. LINDBLAD; « On the Generators of Quantum Dynamical Semigroups »; *Communications in Mathematical Physics* **48**, p. 119–130 (1976). ISSN 0010-3616, 1432-0916. 110
- [89] S. GAMMELMARK, B. JULSGAARD & K. MØLMER; « Past Quantum States of a Monitored System »; *Physical Review Letters* **111**, p. 160401 (2013). 113
- [90] P. SIX; « Estimation d'état et de paramètres pour les systèmes quantiques ouverts »; p. 139. 117, 120
- [91] Z. HRADIL, J. ŘEHÁČEK, J. FIURÁŠEK & M. JEŽEK; « 3 Maximum-Likelihood Methods in Quantum Mechanics »; dans M. PARIS & J. ŘEHÁČEK (rédacteurs), « Quantum State Estimation », , tome 649p. 59–112 (Springer Berlin Heidelberg, Berlin, Heidelberg) (2004); ISBN 978-3-540-22329-0 978-3-540-44481-7. 118
- [92] R. BLUME-KOHOUT; « Optimal, Reliable Estimation of Quantum States »; *New Journal of Physics* **12**, p. 043034 (2010). ISSN 1367-2630. 122
- [93] J. WANG, V. B. SCHOLZ & R. RENNER; « Confidence Polytopes in Quantum State Tomography »; arXiv :1808.09988 [quant-ph] (2018)[1808.09988](#). 123, 125
- [94] M. CHRISTANDL & R. RENNER; « Reliable Quantum State Tomography »; *Physical Review Letters* **109** (2012). ISSN 0031-9007, 1079-7114. 123, 125, 126
- [95] R. BLUME-KOHOUT; « Robust Error Bars for Quantum Tomography »; arXiv :1202.5270 [quant-ph] (2012)[1202.5270](#). 123, 125
- [96] C. J. CLOPPER & E. S. PEARSON; « THE USE OF CONFIDENCE OR FIDUCIAL LIMITS ILLUSTRATED IN THE CASE OF THE BINOMIAL »; *Biometrika* **26**, p. 404–413 (1934). ISSN 0006-3444. 123
- [97] D. SUESS, t. RUDNICKI, T. O. MACIEL & D. GROSS; « Error Regions in Quantum State Tomography : Computational Complexity Caused by Geometry of Quantum States »; *New Journal of Physics* **19**, p. 093013 (2017). ISSN 1367-2630. 125

RÉSUMÉ

La reconstruction d'états quantiques, ou tomographie, joue un rôle central dans les technologies quantiques, afin de caractériser les opérations effectuées et d'extraire de l'information sur les états résultats de traitements d'information quantique. Les méthodes répandues de tomographie reposent généralement sur des mesures idéales, effectuées une seule fois sur chaque préparation de l'état d'intérêt. Dans ce travail, nous utilisons une nouvelle méthode, appelée tomographie par trajectoires, qui consiste à enregistrer, pour chaque réalisation de l'état, la trajectoire quantique suivie par le système à l'aide d'une série de mesures successives du système, en présence d'imperfections expérimentales et de décohérence. On extrait alors plus d'information sur l'état à reconstruire et on est capable, à partir d'un ensemble de mesures accessibles données, de créer des mesures plus générales.

À l'aide des techniques de l'électrodynamique quantique en cavité, nous avons préparé des états intriqués de photons micro-onde délocalisés sur deux modes distants. Nous avons ensuite reconstruit ces états par tomographie par trajectoires, dans un espace de Hilbert de grande dimension. Nous montrons que cette méthode permet de reconstruire l'état, de développer des stratégies de mesure adaptées pour accélérer l'extraction d'information sur les cohérences quantiques d'intérêt et qu'elle fournit une estimation de l'incertitude sur les coefficients de la matrice densité reconstruite.

MOTS CLÉS

Électrodynamique quantique en cavité, atomes de Rydberg, tomographie d'états quantiques, trajectoires quantiques, états noon, intrication, décohérence, mesure QND

ABSTRACT

Quantum state estimation, or tomography, is a key component of quantum technologies, allowing to characterize quantum operations and to extract information on the results of quantum information processes. The usual tomography techniques rely on ideal, single-shot measurements of the unknown state. In this work, we use a new approach, called trajectory quantum tomography, where the quantum trajectory of each realization of the state is recorded through a series of measurements, including experimental imperfections and decoherence. This strategy increases the extracted amount of information and allows to build new measurements for a set of feasible measurements.

Using the tools of cavity quantum electrodynamics, we have prepared entangled states of microwave photons spread on two separated modes. We have then performed a trajectory tomography of these states, in a large Hilbert space. We have proved that this method allows to estimate the state, to develop faster strategies for extracting information on specific coherences of the state and to compute error bars on the components of the estimated density matrix.

KEYWORDS

Cavity quantum electrodynamics, Rydberg atoms, quantum state tomography, quantum trajectories, noon states, entanglement, decoherence, QND measurement