



Manipulation de champs quantiques mésoscopiques

Franck Ferreyrol

► To cite this version:

Franck Ferreyrol. Manipulation de champs quantiques mésoscopiques. Autre [cond-mat.other]. Université Paris Sud - Paris XI, 2011. Français. NNT : 2011PA112029 . tel-00585534

HAL Id: tel-00585534

<https://theses.hal.science/tel-00585534>

Submitted on 28 Jul 2011

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.



Laboratoire Charles Fabry
Institut d'Optique
CNRS UMR 8501

Université Paris-Sud 11
UFR Scientifique d'Orsay

THÈSE

Spécialité :
ONDES ET MATIÈRE

présenté pour obtenir le grade de
DOCTEUR EN SCIENCES
de l'Université Paris-Sud 11

par
Franck FERREYROL

Sujet :

**MANIPULATION DE CHAMPS
QUANTIQUES MÉSCOPHIQUES**

Soutenue le 22 mars 2011 devant la commission d'examen

M.	Nicolas	CERF	Examineur
M.	Philippe	GRANGIER	Examineur
M.	Juan Ariel	LEVENSON	Président
M.	Jean-Philippe	POIZAT	Rapporteur
M.	Nicolas	TREPS	Rapporteur
Mme.	Rosa	TUALLE-BROURI	Directrice de thèse

Remerciements

Mes premiers remerciements vont à ma famille pour m'avoir soutenu et accompagné pendant toute la durée de cette thèse.

Merci ensuite à tout ceux avec qui j'ai travaillé, à commencer par Rosa Tualle-Brouri et Philippe Grangier qui m'on accueilli et encadré. Merci à Alexei Ourjountsev pour m'avoir enseigné le fonctionnement de la manip, aussi bien dans la pratique que la théorie, de nombreux points présentés dans ce manuscrit sont basés sur ce qu'il m'a appris. Je remercie aussi Florence Fuchs, outre avoir été le principal artisan du VLPC, elle a aussi travaillé sur plusieurs points de la manip et a été une très bonne conseillère. Et je ne saurais les oublier, un grand merci à Marco Barbieri et Rémi Blandino avec qui j'ai travaillé quotidiennement, ce fut un réel plaisir.

Merci aussi à tous les membres, passés et présents, du groupe d'Optique Quantique avec qui j'ai vécu des moments inoubliables. Sans leur aide, leurs conseils, les discussions que nous avons eu, ou même leur convivialité, cette thèse n'aurait certainement pas eu d'aussi bon résultats. On ne soulignera jamais assez l'importance d'une bonne ambiance de travail.

Enfin je remercie l'ensemble du personnel de l'Institut d'Optique, et notamment le personnel technique et administratif sans qui nous ne pourrions rien faire.

Table des matières

1	Introduction	11
1.1	Physique quantique et information quantique	11
1.1.1	Naissance de l'information quantique	11
1.1.1.1	La physique quantique	11
1.1.1.2	La théorie de l'information	12
1.1.1.3	L'information quantique	12
1.1.2	Variables discrètes et variables continues	14
1.2	Contexte et objectif	14
1.3	Plan de la thèse	15
I	Les outils théoriques et expérimentaux	17
2	Le champ électromagnétique quantique	19
2.1	Description du champ quantifié	20
2.1.1	Quantification du champ	20
2.1.2	Variables discrètes	22
2.1.3	Variables continues	23
2.2	Représentation des états quantiques	25
2.2.1	La matrice densité	25
2.2.2	La fonction de Wigner	26
2.2.2.1	Définition	27
2.2.2.2	Propriétés	28
2.2.2.3	Fonctions de Wigner à plusieurs modes	28
2.2.2.4	Retour à la matrice densité	29
2.2.3	Reconstruction par tomographie	29
2.2.3.1	Principe	29
2.2.3.2	Algorithmes utilisés	29
2.3	Quelques états « de base »	30
2.3.1	Les états gaussiens	30
2.3.1.1	Le vide quantique	31
2.3.1.2	Les états cohérents	31
2.3.1.3	Les états comprimés monomodes	32
2.3.1.4	Les états comprimés bimodes ou états EPR	33
2.3.1.5	Les états thermiques	34
2.3.2	Les états non gaussiens	34
2.3.2.1	Caractéristiques et intérêt	34
2.3.2.2	Les états de Fock	36

2.3.2.3	Les états chats de Schrödinger	37
2.4	Conclusion	38
3	Le dispositif expérimental	39
3.1	Présentation du dispositif	40
3.1.1	Introduction	40
3.1.2	Schéma générique	41
3.2	La source laser	42
3.2.1	Le laser impulsionnel	42
3.2.2	Diagnostic du faisceau	43
3.2.2.1	Photodiodes	43
3.2.2.2	Spectromètre	43
3.2.2.3	Auto-corrélateur	45
3.2.2.4	Profilomètre	45
3.3	Les transformations unitaires	46
3.3.1	L'optique linéaire	46
3.3.1.1	Déphasage	46
3.3.1.2	Lame semi-réfléchissante	46
3.3.1.3	Lame demi-onde et quart-d'onde	47
3.3.1.4	Cube séparateur de polarisation (PBS)	48
3.3.2	Le générateur de seconde harmonique (GSH)	48
3.3.2.1	Principe	48
3.3.2.2	Effets parasites	50
3.3.3	L'amplificateur paramétrique optique (OPA)	51
3.3.3.1	Configuration non-dégénérée	52
3.3.3.2	Configuration dégénérée	55
3.4	Détection et mesures projectives	57
3.4.1	La détection homodyne	57
3.4.1.1	Principe	57
3.4.1.2	Imperfections	58
3.4.1.3	Montage	59
3.4.2	La photodiode à avalanche	63
3.4.2.1	Caractéristiques du dispositif	63
3.4.2.2	Mesures projectives et conditionnement	63
3.5	Production de photons uniques	64
3.5.1	Production	64
3.5.2	Modélisation	64
3.5.2.1	Amplification paramétrique	66
3.5.2.2	Pertes homodynes et APD	66
3.5.2.3	Conditionnement	66
3.5.3	Utilisations	68
3.5.3.1	Taux de production et matrice densité	68
3.5.3.2	Caractérisation des imperfections	69

3.6 Conclusion	70
II Résultats expérimentaux	73
4 Génération de superpositions non locales d'états cohérents	75
4.1 Introduction	75
4.1.1 L'intrication	75
4.1.1.1 L'intrication et l'information quantique	75
4.1.1.2 Mesure de l'intrication	77
4.1.1.3 États de Bell	78
4.1.2 Principe de l'expérience	79
4.2 Réalisation expérimentale	81
4.2.1 Montage	81
4.2.2 Modélisation	82
4.2.3 Résultats	84
4.3 Comparaison avec un état de Bell discret	86
4.3.1 Montage	86
4.3.2 Modélisation	87
4.3.3 Résultats	87
4.4 Discussion	89
4.4.1 Communications à longues distances	89
4.4.2 Calcul quantique et systèmes hybrides	90
4.4.3 Conclusion	94
5 L'amplificateur sans bruit non déterministe	97
5.1 Amplification d'un signal quantique	97
5.1.1 Amplification déterministe	98
5.1.2 Principe de l'amplificateur sans bruit non déterministe	100
5.2 Réalisation expérimentale	106
5.2.1 Montage	106
5.2.2 Modélisation de l'expérience	108
5.2.2.1 Considérations générales	108
5.2.2.2 Modèle numérique	109
5.2.2.3 Modèle analytique	111
5.2.3 Résultats	115
5.3 Discussion	119
5.3.1 Cohérence avec les lois de la physique quantique	119
5.3.2 Autres expériences	120
5.3.3 Utilisations diverses	121
5.3.4 Retour sur le calcul quantique hybride	122
5.3.5 Conclusion	124
6 Mesure de la non-gaussianité d'un état	125
6.1 La non-gaussianité	125
6.1.1 États gaussiens et états non-gaussiens	125
6.1.2 Les mesures	127
6.1.3 Non gaussianité et non-classicité	130

6.2 Réalisation expérimentale	131
6.2.1 Montage	131
6.2.2 Modélisation	133
6.2.3 Résultats	135
6.3 Discussion	141
6.3.1 Non-gaussianité de quelques états courants	141
6.3.2 Non-gaussianité des processus quantiques	143
6.3.3 Non-gaussianité et information	144
6.3.4 Conclusion	145
III Amélioration du protocole expérimental	147
7 L'amplificateur femtoseconde	149
7.1 Amplification d'impulsions femtosecondes	149
7.1.1 Intérêt d'un amplificateur	149
7.1.2 Différents amplificateurs possibles	150
7.1.2.1 Amplification paramétrique	150
7.1.2.2 Amplification régénérative à base de saphir dopé au titane	151
7.1.2.3 Amplification passive à base de saphir dopé au titane	151
7.1.3 Application au Système expérimental	152
7.2 Dispositif	154
7.2.1 Vue d'ensemble	154
7.2.2 Cristaux utilisés	156
7.2.3 Refroidissement cryogénique	157
7.3 Intégration au dispositif expérimental	157
7.3.1 Faisceau amplifié	157
7.3.2 Doublage de fréquence	160
7.3.3 Amplification paramétrique	162
7.4 Travail restant	162
8 Le VLPC	165
8.1 Le comptage de photons	165
8.1.1 Introduction	165
8.1.2 Comparaison des différentes méthodes	166
8.1.3 Principe de fonctionnement du VLPC	168
8.2 Montage	170
8.2.1 Vue d'ensemble	170
8.2.2 Montage cryogénique	172
8.2.3 Acheminement du faisceau par fibre optique	173
8.3 Caractérisation	174
8.3.1 Dispositif de test	174
8.3.2 Recherche d'un point de fonctionnement	176
8.3.3 Efficacité en fonction du flux incident	177
8.4 Conclusion	179
9 Conclusion et perspectives	181

9.1 Conclusion	181
9.2 Perspectives	182
9.2.1 Améliorations du dispositif expérimental	182
9.2.2 Calcul quantique à variables continues	182
9.2.3 De nouveaux outils	183
 IV Annexes	 185
A Formulaire et considérations mathématiques	187
A.1Intégrales gaussiennes	187
A.2Pertes et amplification parasite	188
A.3Bruit électronique	188
B Le laser femtoseconde	191
B.1Critères importants	191
B.2Principe d'un laser femtoseconde	192
B.2.1 Amplification	192
B.2.2 Verrouillage de mode	192
B.2.3 Dispersion de la vitesse de groupe	193
B.2.4 Extraction des impulsions (cavity dumper)	193
C Calculs pour la comparaison des protocoles de communication	197
C.1Envoi direct : état EPR et état de Bell	197
C.2Distillation d'états EPR	197
C.3Transfert d'intrication	201
D Influence de la cadence sur l'amplification femtoseconde	203

Chapitre 1

Introduction

Sommaire

1.1 Physique quantique et information quantique	11
1.1.1 Naissance de l'information quantique	11
1.1.1.1 La physique quantique	11
1.1.1.2 La théorie de l'information	12
1.1.1.3 L'information quantique	12
1.1.2 Variables discrètes et variables continues	14
1.2 Contexte et objectif	14
1.3 Plan de la thèse	15

1.1 Physique quantique et information quantique

1.1.1 Naissance de l'information quantique

1.1.1.1 La physique quantique

La physique quantique est partout. De nombreuses disciplines scientifiques font appel à elle : l'optique et la physique atomique bien sûr, la physique de la matière condensée, avec notamment la physique du solide qui en a énormément tiré partie, la physique nucléaire et plus généralement la physique des particules, l'astrophysique, la physique des plasmas et même la chimie et plus récemment la biologie. L'incorporation de la physique quantique aux autres domaines a permis l'avènement de nombreuses applications : l'électronique, le laser, l'énergie nucléaire, l'IRM, ...

Née au début du XX^e siècle, elle a ainsi eu une part de responsabilité importante dans la plupart avancées technologiques de ce siècle. Malgré cela et le fait que la physique quantique a apporté de nouveaux concepts et présente des comportements très différents de la physique classique, et même souvent contre-intuitifs, peu d'applications exploitent directement ses propriétés. Le travail présenté dans ce document est basé sur celles-ci, autant du point de vue de leur étude fondamentale que d'un travail en amont pouvant conduire à leur utilisation dans des applications pratiques.

1.1.1.2 La théorie de l'information

Lorsque que l'on veut expliquer le fonctionnement d'une expérience de physique quantique on en revient quasiment toujours à utiliser le concept d'information. Ce qui n'est pas étonnant puisque celui-ci est intimement lié à la mesure, qui joue un rôle central dans la physique quantique.

Les fondements de la théorie de l'information ont été formalisés par Claude Shannon dans un article de 1948 [1]. L'information y est traitée d'un point de vue probabiliste, indépendamment de toute considération sur sa signification. Ainsi l'information est considérée comme un moyen de décrire l'état d'un système qui peut en prendre plusieurs. Dans ce but plusieurs grandeurs ont été définies afin de quantifier l'information dans diverses situations.

La principale quantité définie par la théorie de l'information est l'entropie. Elle représente la quantité d'information dont nous avons besoin pour connaître précisément l'état d'un système. Mais comment quantifier cette information ? Lorsque nous voulons obtenir des informations sur quelque chose nous posons des questions. Nous pouvons alors décrire l'entropie comme étant le nombre de questions qu'il faut poser pour connaître l'état du système. Pour cela nous devons toutefois nous limiter à des questions ayant un nombre fini de réponses possibles, le plus simple consistant à utiliser des questions dont la réponse est *oui* ou *non*. Ainsi pour un système ayant 2^n états possibles, nous aurons besoin de poser n questions (en utilisant par exemple la méthode de la dichotomie). De manière générale nous pouvons définir l'entropie d'un système pouvant prendre N états différents comme étant $S = \log_2(N)$.

Cette définition n'est cependant valable que si les différents états sont équiprobables. Il semble raisonnable que si ce n'est pas le cas, c'est-à-dire que certains états ont une plus forte probabilité de se réaliser, alors l'état du système est plus facile à trouver, et donc son entropie est plus faible. Considérons par exemple un système ayant trois états possibles : le premier se réalise avec une probabilité $p_1 = 1/2$ et les deux autres avec une probabilité $p_2 = p_3 = 1/4$. Pour déterminer l'état, le plus efficace consiste à demander d'abord s'il est dans l'état 1 puis si ce n'est pas le cas s'il est dans l'état 2 (ou 3). Nous pouvons alors déterminer l'état 1 avec une seule question et les deux autres avec deux questions, et il faut donc 1,5 question en moyenne pour déterminer l'état du système. En généralisant, il faut poser $-\log_2(p_i)$ questions pour déterminer un état ayant une probabilité p_i de se réaliser, soit $S = -\sum_{i=1}^N p_i \log_2(p_i)$ questions en moyenne pour déterminer l'état du système, ce qui définit l'entropie de ce dernier.

1.1.1.3 L'information quantique

Au delà des considérations précédentes, le lien entre physique et information a d'abord été effectué en rapprochant l'entropie de Shannon et l'entropie thermodynamique, puis par l'assertion de Landauer disant que « l'information est physique »¹. Il a fallu attendre les années 80 pour qu'apparaisse un lien fort entre information et physique quantique, donnant ainsi naissance à l'information quantique. Il s'agit d'une des applications directes des propriétés de la physique quantique. Les études actuelles sur les applications de l'information quantique sont essentiellement divisées en deux axes : l'*ordinateur quantique* et la *communication quantique*.

L'ordinateur quantique Il a été initié par Feynmann en partant du principe que certains systèmes quantiques ne peuvent pas être simulés par un ordinateur classique en un temps raison-

1. Landauer a démontré que l'information perdue dans un circuit électronique était dissipée sous forme de chaleur.

nable, ceci à cause de la taille de l'espace de Hilbert qui varie exponentiellement avec le nombre de degrés de liberté. Par contre ils pourraient être simulés par un système utilisant les règles de la physique quantique.

Plus généralement un ordinateur quantique permettrait d'effectuer bien plus rapidement un certain nombre de calculs difficiles. Il existe des problèmes que l'on ne sait pas résoudre autrement qu'en essayant toutes les solutions possibles, et le temps nécessaire varie alors exponentiellement avec la quantité d'information nécessaire pour décrire la solution, ce qui rend ce temps assez rapidement bien trop long pour être raisonnable. C'est d'ailleurs sur ce principe que reposent les méthodes de cryptage asymétrique utilisées dans de nombreuses communications sécurisées. Le principe de l'ordinateur quantique consiste à profiter du parallélisme offert par le principe de superposition : il est ainsi possible d'effectuer simultanément un calcul sur de très nombreuses valeurs et, par exemple, de tester toutes les solutions possibles en même temps. La mesure complique cependant les choses : on ne peut lire qu'un seul résultat et il faut s'arranger pour qu'au final on ait uniquement celui qui nous intéresse.

Les ordinateurs quantiques nécessitent donc des algorithmes spécifiques permettant de tirer parti de leurs avantages. Parmi les algorithmes connus, deux sont particulièrement intéressants : il s'agit de l'algorithme de Shor [2], qui décompose un nombre en facteurs premiers avec une accélération exponentielle par rapport au cas classique, et de l'algorithme de Grover [3], un algorithme de recherche (d'un élément dans une base de donnée ou plus généralement d'une solution à un problème) dont l'accélération est plus faible (seulement quadratique).

De manière similaire au bit en informatique classique, l'élément de base d'un ordinateur quantique est le *qubit*. Contrairement à son homologue classique, il peut prendre à la fois les valeurs 0 et 1. Pour manipuler ces qubits, les ordinateurs ont besoin de deux pièces essentielles : les portes quantiques, qui correspondent aux opérations de bases effectuées par l'ordinateur sur les qubits, et les codes correcteurs, qui permettent de corriger les erreurs introduites par les imperfections du système physique.

La communication quantique Concernant les apports de la physique quantique aux communications, la plus importante est la *cryptographie quantique*. Nous avons évoqué le fait que la sécurité de la plupart des méthodes de cryptage utilisées actuellement reposent sur des problèmes difficiles à résoudre. Non seulement un ordinateur quantique pourrait casser ces codes en résolvant facilement ces problèmes (par exemple le cryptage RSA, qui est le plus courant, est basé sur la factorisation de grands nombres), mais en plus nous n'avons aucune certitude concernant la non-existence d'algorithmes classiques permettant de les résoudre facilement. Il existe bien une méthode de cryptographie intégralement sûre, appelée code de Vernam (ou *one-time pad* en anglais), mais elle fait une très grande consommation en clés de cryptage (la clé doit être aussi longue que le message et utilisée une seule fois). L'application de l'information quantique la plus avancée est justement la principale catégorie de méthode de cryptographie quantique : la *distribution quantique de clé*, qui vise à régler le problème du partage sécurisé des clés de cryptage. Le principe consiste à utiliser la mesure quantique et l'incertitude qui en résulte afin de détecter la présence d'un espion et ainsi de réduire autant qu'on veut l'information qu'il pourrait avoir sur la clé transmise.

Les apports de la physique quantique aux communications peuvent aussi se voir à travers les quantités de la théorie de l'information, comme l'entropie conditionnelle. Partant d'un système bipartite, l'une des parties correspondant à l'envoi du message et l'autre à la réception, l'entropie conditionnelle représente la quantité d'information nécessaire pour décrire l'une des parties lorsque l'on connaît parfaitement l'autre. Pour un système composé des parties A et B elle vaut $S(B|A) = S(A, B) - S(A)$. D'un point de vue classique, elle est toujours positive : si elle est nulle

alors les parties sont parfaitement corrélées, si elle est égale à l'entropie de la partie considérée ($S(B|A) = S(B)$) alors elles sont complètement indépendantes. Revenons au cas quantique, et prenons donc un état bimode intriqué : étant bien défini son entropie est nulle. Mais si l'on ne mesure que l'un des deux modes alors le résultat est aléatoire, l'entropie de chaque mode n'est donc pas nulle. On en déduit que l'entropie conditionnelle peut-être négative ! Cette négativité correspond à une information potentielle qu'il est possible d'utiliser pour une communication ultérieure. Ainsi si Alice et Bob partagent un état intriqué, il suffit à Alice d'envoyer la partie qu'elle possède, après y avoir effectué des opérations locales, pour transmettre 2 bit d'information par qubit transmis. Cette technique s'appelle le codage super-dense.

1.1.2 Variables discrètes et variables continues

La notion de variables est un point capital aussi bien pour la physique que pour l'information. En physique elles permettent d'identifier les différents états que peut prendre un système, tandis qu'en théorie de l'information les différentes valeurs qu'elles peuvent prendre servent à coder l'information. Ces deux aspects sont d'ailleurs liés puisqu'en pratique chaque « lettre » de l'alphabet utilisé pour coder l'information correspond à un état physique ; en d'autres termes les variables informatiques sont réalisées à l'aide de variables physiques.

Les variables peuvent être discrètes ou continues, suivant le système regardé et/ou, grâce à la dualité onde-corpuscule, suivant la façon dont on le regarde.

Variables discrètes Nous sommes familiers avec les variables discrètes depuis que nous savons compter, puisqu'elle correspondent à des phénomènes dénombrables. Elles sont aussi très courantes en physique quantique puisque liée à la quantification. Ainsi nous les retrouvons notamment dans les nombres quantiques, ou le spin. Leur grand succès vient notamment des systèmes à deux niveaux, qui, de par leur simplicité, furent les premiers utilisés pour bon nombre d'applications de l'information quantique. Elles sont aussi très utilisées en information classique puisque liées au codage numérique, et à la notion de *bit*, que l'on retrouve dans tous nos appareils informatiques.

Variables continues Les variables continues correspondent quant à elles à ce qui est indénombrable. Elles sont omniprésentes en physique classique puisque la plupart des grandeurs y sont continues (position, vitesse, température, pression, ...) et restent aussi présentes dans les communications puisque liées au codage analogique encore très utilisés, notamment, dans les communications radiophoniques. On les retrouve aussi en physique quantique, et en particulier en optique quantique où elles possèdent le grand avantage que les systèmes les utilisant sont plus faciles à produire et à mesurer, de par le fait qu'ils ne nécessitent généralement que des outils communs avec l'optique classique.

1.2 Contexte et objectif

Pour réaliser les protocoles d'information quantique l'un des supports les plus utilisés est l'optique quantique : la lumière est en effet relativement facile à produire, manipuler et détecter. Ceci est encore plus vrai pour les communications quantiques grâce à la simplicité de transmission d'un signal lumineux.

Au moment de débiter cette thèse de nombreux travaux, tant théoriques qu'expérimentaux, traitaient de la réalisation de protocoles d'information quantique au moyen de l'optique quantique : téléportation, clonage, cryptographie, codage dense, portes pour le calcul quantique, code

correcteurs, ... La question était tout de même, et est toujours, plus avancée pour l'utilisation de variables discrètes que pour l'usage des variables continues, surtout d'un point de vue expérimental. De nombreuses avancées ont néanmoins été réalisées dans ce domaine durant les années précédant cette thèse, certaines de ces avancées ayant été effectuées au sein même de notre groupe.

Les variables continues ont pu profiter de l'apport de techniques auparavant réservées aux variables discrètes, car correspondant à une approche discrète. Le mélange des deux a permis de réaliser un certain nombre d'opérations impossibles avec des outils cantonnés aux variables continues. L'objectif de cette thèse s'inscrit dans cette approche consistant à mélanger les techniques propres aux variables discrètes et aux variables continues, et aborde plusieurs aspects de celle-ci. Ainsi, outre l'apport à l'information quantique à variables continues, nous verrons qu'elle permet aussi d'outrepasser les limites habituelles de la physique quantique. Un autre sujet abordé sera l'étude des mesures de non-gaussianité, qui est le fer de lance de cette approche.

Ce travail est avant tout expérimental. Néanmoins la théorie est indispensable à la bonne compréhension d'un phénomène physique. Les modèles théoriques des expériences y ont donc aussi une part importante. Ainsi tout comme cette thèse marie descriptions continue et discrète de la lumière, elle associe expérience et théorie, afin d'obtenir un aperçu le plus complet possible.

Enfin, d'un point de vue prosaïque, cette thèse et le stage qui l'a précédé, et pendant lequel ont débuté ces travaux, ont encadré le déménagement du laboratoire dans les nouveaux bâtiments de l'Institut d'Optique. Ce fut donc l'occasion idéale pour se pencher sur des améliorations importantes pour le dispositif expérimental. Ce travail de thèse a donc aussi fait la part belle à l'étude et la mise en place de ces améliorations.

1.3 Plan de la thèse

Ce manuscrit est composé de trois parties.

La première partie présente les outils, aussi bien théoriques qu'expérimentaux, indispensables à ce travail de thèse. Le chapitre 2 se concentre sur les outils théoriques et permet d'introduire les moyens utilisés pour décrire les états que nous manipulerons, ainsi que les états de départ de nos expériences et leur spécificités. Dans le chapitre 3 nous présenterons les éléments expérimentaux dont nous disposons afin de produire, manipuler et détecter nos états. Nous en profiterons pour compléter les outils théoriques en y ajoutant ceux qui décrivent l'action de ces différents éléments. Nous finirons ce chapitre en décrivant une méthode pour caractériser le fonctionnement de notre dispositif mettant en jeu l'intégralité des outils présentés.

La deuxième partie traite des expériences réalisées lors de cette thèse. La première, présentée au chapitre 4, montre un protocole permettant d'intriquer à distances deux états initialement indépendant, et ce y compris à travers un canal de communication ayant de fortes pertes. Au chapitre 5 nous parlerons de la principale expérience de cette thèse, à savoir la démonstration et la caractérisation d'un amplificateur sans bruit. Celui-ci permet d'amplifier un signal optique sans en amplifier le bruit quantique, augmentant ainsi le rapport signal sur bruit. Enfin la dernière expérience, qui est l'objet du chapitre 6, illustre et compare expérimentalement différentes mesures du caractère non-gaussien d'un état quantique.

La troisième et dernière partie est consacrée à l'amélioration du système expérimental. Pour cela il présente deux dispositifs conçus au sein du groupe et qui ont pu être testés et/ou intégrés au dispositif lors de cette thèse. La première, dont nous parlerons au chapitre 7, est un amplificateur d'impulsions femtosecondes qui nous permettra de repousser les limites imposées par la puissance de notre laser. Ayant déjà fait l'objet d'un développement et de quelques tests auparavant

nous avons pu profiter du déménagement pour l'intégrer au dispositif et effectuer des tests plus poussés. La seconde amélioration, traitée au chapitre 8, est un appareil capable de compter le nombre de photons incidents. N'étant pas à l'origine adapté pour notre longueur d'onde nous avons dû tester son fonctionnement afin de savoir s'il apporterait un réel gain.

Première partie

**Les outils théoriques et
expérimentaux**

Chapitre 2

Le champ électromagnétique quantique

Sommaire

2.1 Description du champ quantifié	20
2.1.1 Quantification du champ	20
2.1.2 Variables discrètes	22
2.1.3 Variables continues	23
2.2 Représentation des états quantiques	25
2.2.1 La matrice densité	25
2.2.2 La fonction de Wigner	26
2.2.2.1 Définition	27
2.2.2.2 Propriétés	28
2.2.2.3 Fonctions de Wigner à plusieurs modes	28
2.2.2.4 Retour à la matrice densité	29
2.2.3 Reconstruction par tomographie	29
2.2.3.1 Principe	29
2.2.3.2 Algorithmes utilisés	29
2.3 Quelques états « de base »	30
2.3.1 Les états gaussiens	30
2.3.1.1 Le vide quantique	31
2.3.1.2 Les états cohérents	31
2.3.1.3 Les états comprimés monomodes	32
2.3.1.4 Les états comprimés bimodes ou états EPR	33
2.3.1.5 Les états thermiques	34
2.3.2 Les états non gaussiens	34
2.3.2.1 Caractéristiques et intérêt	34
2.3.2.2 Les états de Fock	36
2.3.2.3 Les états chats de Schrödinger	37
2.4 Conclusion	38

L'objectif de ce chapitre est tout d'abord de présenter le cadre dans lequel s'inscrit ce travail de thèse, à savoir l'optique quantique et plus particulièrement l'information quantique à variables

continues. Nous en profiterons pour poser les diverses notations et introduire les outils qui seront utilisés tout au long de ce manuscrit. Le but de ce chapitre n'est toutefois pas de reconstruire rigoureusement l'optique quantique et ses outils, aussi de nombreux résultats ne seront pas redémontrés. Le lecteur curieux pourra se référer aux nombreux ouvrages traitant du sujet parmi lesquels [10, 41, 6, 7, 8] ainsi qu'aux diverses autres références présentées au cours de ce chapitre.

2.1 Description du champ quantifié

2.1.1 Quantification du champ

Cette section se contente d'effectuer un bref rappel sur la quantification du champ et par la même occasion de poser un certain nombre de notations. Pour une description plus détaillée on pourra se référer par exemple à [9, 10]

La description quantique du champ électromagnétique est basée sur l'application du principe de correspondance à son homologue classique, qui vérifie dans le vide l'équation de Helmholtz :

$$\Delta \vec{E} - \frac{1}{c^2} \partial_t^2 \vec{E} = 0 \quad (2.1)$$

Le champ peut, de plus, se décomposer sur une base d'ondes planes permettant alors de simplifier cette équation. En se limitant à un espace parallélépipédique de volume V , avec conditions aux limites périodiques, on pourra ainsi écrire

$$\vec{E}(\vec{r}, t) = \sum_{l=1}^2 \sum_{\vec{k}} \mathcal{E}_{l,\vec{k}} \alpha_{l,\vec{k}}(t) \vec{e}_{l,\vec{k}} e^{i\vec{k} \cdot \vec{r}} + \text{C.C.} \quad (2.2)$$

où $\alpha_{l,\vec{k}}(t)$ est l'amplitude complexe du mode défini par la polarisation l et le vecteur d'onde $\vec{k}$, les vecteurs polarisation $\vec{e}_{1,\vec{k}}$ et $\vec{e}_{2,\vec{k}}$ étant orthogonaux à $\vec{k}$ et entre eux. Le coefficient $\mathcal{E}_{l,\vec{k}}$ est défini de telle sorte qu'un champ d'amplitude unité confiné dans un volume V ait une énergie $\hbar\omega_k$, ce qui donne

$$\mathcal{E}_{l,\vec{k}} = \sqrt{\frac{\hbar\omega_k}{2\epsilon_0 V}} \quad (2.3)$$

En appliquant l'équation d'Helmholtz à l'équation 2.2 on obtient ainsi l'équation vérifiée par les amplitudes complexes de chacun des modes de la décomposition :

$$\ddot{\alpha}_{l,\vec{k}}(t) + k^2 c^2 \alpha_{l,\vec{k}}(t) = 0 \quad (2.4)$$

$$\ddot{\alpha}_{l,\vec{k}}(t) + \omega_k^2 \alpha_{l,\vec{k}}(t) = 0 \quad (2.5)$$

Dont nous retiendrons la solution à fréquence positive :

$$\alpha_{l,\vec{k}}(t) = e^{-i\omega_k t} \alpha_{l,\vec{k}} \quad (2.6)$$

Nous pouvons aussi réécrire l'équation 2.5 en introduisant les *quadratures*, qui correspondent aux parties réelles et imaginaires de l'amplitude complexe :

$$X_{l,\vec{k}}(t) = \frac{\alpha_{l,\vec{k}}(t) + \alpha_{l,\vec{k}}^*(t)}{\sqrt{2}} \quad (2.7)$$

$$P_{l,\vec{k}}(t) = \frac{\alpha_{l,\vec{k}}(t) - \alpha_{l,\vec{k}}^*(t)}{i\sqrt{2}} \quad (2.8)$$

Connaissant l'expression de l'amplitude complexe (équation 2.6), il apparaît que ces quadratures vérifient les équations

$$\dot{X}_{l,\vec{k}}(t) = \omega_k P_{l,\vec{k}}(t) \quad (2.9)$$

$$\dot{P}_{l,\vec{k}}(t) = -\omega_k X_{l,\vec{k}}(t) \quad (2.10)$$

Ces équations sont similaires à celles d'un oscillateur harmonique, dont la position et l'impulsion vérifient

$$\dot{x}(t) = \frac{p(t)}{m} \quad (2.11)$$

$$\dot{p}(t) = -\frac{m}{\omega^2} x(t) \quad (2.12)$$

Nous pouvons ainsi relier le champ électromagnétique au problème de l'oscillateur harmonique à travers la correspondance :

$$X_{l,\vec{k}}(t) \leftrightarrow \sqrt{\frac{m\omega}{\hbar}} x(t) \quad P_{l,\vec{k}}(t) \leftrightarrow \frac{p(t)}{\sqrt{m\hbar\omega}} \quad (2.13)$$

Il est alors possible de quantifier le champ électromagnétique en s'appuyant sur cette correspondance : on associe les observables $\hat{X}_{l,\vec{k}}$ et $\hat{P}_{l,\vec{k}}$ aux quadratures $X_{l,\vec{k}}(t)$ et $P_{l,\vec{k}}(t)$, avec la relation de commutation :

$$[\hat{X}_{l,\vec{k}}, \hat{P}_{l',\vec{k}'}] = i\delta_{l,l'}\delta_{\vec{k},\vec{k}'} \quad (2.14)$$

On associera ainsi aux amplitudes complexe $\alpha_{l,\vec{k}}(t)$ et $\alpha_{l,\vec{k}}^*(t)$ les opérateurs annihilation et création, définis en inversant les relations 2.7 et 2.8 :

$$\hat{a}_{l,\vec{k}} = \frac{\hat{X}_{l,\vec{k}} + i\hat{P}_{l,\vec{k}}}{\sqrt{2}} \quad \hat{a}_{l,\vec{k}}^\dagger = \frac{\hat{X}_{l,\vec{k}} - i\hat{P}_{l,\vec{k}}}{\sqrt{2}} \quad (2.15)$$

qui vérifient les relations de commutation

$$[\hat{a}_{l,\vec{k}}, \hat{a}_{l',\vec{k}'}^\dagger] = \delta_{l,l'}\delta_{\vec{k},\vec{k}'} \quad (2.16)$$

Notons que l'on peut associer à ces opérateurs l'hamiltonien :

$$\hat{H} = \hbar \sum_{l,\vec{k}} \omega_k \left(\hat{a}_{l,\vec{k}}^\dagger \hat{a}_{l,\vec{k}} + \frac{1}{2} \right) \quad (2.17)$$

À partir de là on peut en déduire l'expression de l'opérateur champ électrique, à un facteur de phase près :

$$\begin{aligned} \vec{E}(\vec{r}, t) &= \sum_{l,\vec{k}} \mathcal{E}_{l,\vec{k}} \vec{\epsilon}_{l,\vec{k}} \hat{a}_{l,\vec{k}} e^{i(\vec{k} \cdot \vec{r} - \omega_k t)} + \text{C.C.} \\ &= \sum_{l,\vec{k}} \mathcal{E}_{l,\vec{k}} \vec{\epsilon}_{l,\vec{k}} \left(\hat{X}_{l,\vec{k}} \cos(\vec{k} \cdot \vec{r} - \omega_k t) + \hat{P}_{l,\vec{k}} \sin(\vec{k} \cdot \vec{r} - \omega_k t) \right) \end{aligned} \quad (2.18)$$

Dans la pratique, nous ne travaillerons ni avec l'opérateur « champ électrique », ni avec des ondes planes, mais avec des paquets d'ondes, qui forment des modes définis par leur enveloppes $\phi_l(\vec{k})$, normalisées à l'unité. On peut alors définir l'opérateur $\hat{a}$ pour le mode qui nous intéresse :

$$\hat{a} = \sum_{l,\vec{k}} \phi_l^*(\vec{k}) \hat{a}_{l,\vec{k}} \quad (2.19)$$

et ensuite en déduire les opérateurs $\hat{a}^\dagger$, $\hat{X}$ et $\hat{P}$ pour ce même mode. Un rapide calcul nous donne les nouveaux commutateurs

$$\begin{aligned} [\hat{a}, \hat{a}^\dagger] &= \sum_{l, \vec{k}, l', \vec{k}'} \phi_l^*(\vec{k}) \phi_{l'}(\vec{k}') [\hat{a}_{l, \vec{k}}, \hat{a}_{l', \vec{k}'}^\dagger] \\ &= \sum_{l, \vec{k}} |\phi_l(\vec{k})|^2 = 1 \end{aligned} \quad (2.20)$$

$$[\hat{X}, \hat{P}] = \left[\frac{\hat{a} + \hat{a}^\dagger}{\sqrt{2}}, \frac{\hat{a} - \hat{a}^\dagger}{i\sqrt{2}} \right] = i \quad (2.21)$$

commutateurs qui sont identiques à ceux de l'oscillateur harmonique.

2.1.2 Variables discrètes

Dans la section précédente il reste un aspect de l'oscillateur harmonique dont nous n'avons pas parlé, et qui concerne l'opérateur nombre $\hat{n} = \hat{a}^\dagger \hat{a}$. Cet opérateur est responsable de la quantification au sens propre du terme : il quantifie l'énergie du champ en « paquets », les fameux *quanta d'énergie* de Planck et Einstein [11, 12], ou photons. Son spectre est évidemment l'ensemble des entiers naturels et ses états propres, $|n\rangle$, sont appelés *états de Fock*. Leur énergie vaut $E = \hbar\omega (n + \frac{1}{2})$ et ils sont non-dégénérés : pour un nombre de photons donné il existe un seul état quantique. Ces états rendent compte de l'aspect corpusculaire de la lumière.

Les opérateurs d'annihilation et de création quant à eux diminuent ou augmentent le nombre de photons suivant les relations

$$\hat{a}|n\rangle = \sqrt{n}|n-1\rangle \quad (2.22)$$

$$\hat{a}^\dagger|n\rangle = \sqrt{n+1}|n+1\rangle \quad (2.23)$$

Les variables discrètes ont de nombreuses utilisations en information quantique. Elles constituent l'équivalent du codage numérique en information classique, ce qui permet une transposition facile au cas quantique tout en conservant les avantages de ce codage. Pour cela on remplace le bit classique, qui peut prendre les valeurs 0 ou 1, par un bit quantique, ou qubit, dont les valeurs correspondent à deux états orthogonaux $|\hat{0}\rangle$ et $|\hat{1}\rangle$. Puisque nous avons un ensemble discret d'états possibles pour un mode donné, nous pouvons utiliser deux d'entre eux, à savoir les états de Fock $|0\rangle$ et $|1\rangle$, c'est-à-dire l'absence ou la présence d'un photon. Cette pratique souffre néanmoins d'un gros handicap : on ne peut pas distinguer l'état $|0\rangle$ d'un état $|1\rangle$ dont le photon aurait été perdu à cause des imperfections. Une idée plus astucieuse consiste alors à coder l'information sur le mode du photon : on peut alors distinguer les pertes des mesures concluantes. On peut utiliser plusieurs type de modes différents (polarisation, espace, temps, fréquence, ...) et il est généralement facile de passer de l'un à l'autre.

De tels systèmes présentent l'avantage de pouvoir réaliser très facilement les portes à un qubit comme la porte NOT, la porte de Phase ou la porte de Hadamard, caractérisées par les transformations unitaires suivantes (dans la base $|\hat{0}\rangle, |\hat{1}\rangle$)

$$U_{\text{NOT}} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad U_\phi = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\phi} \end{pmatrix} \quad H = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \quad (2.24)$$

Les portes à deux qubits, comme les portes contrôlées, où l'action sur le deuxième qubit (qui correspond à celle d'une porte à un qubit) n'est réalisée que si le premier qubit est dans l'état $|\hat{1}\rangle$,

sont par contre bien plus difficiles à réaliser. En effet les photons n'interagissent pas directement entre eux. La solution la plus simple pour faire interagir des faisceaux lumineux est d'utiliser des cristaux non-linéaires. Seulement ces non-linéarités sont faibles, du coup il faut avoir une forte puissance pour que les effets ne soient pas négligeables, ce qui n'est pas le cas pour les systèmes à variables discrètes qui utilisent des photons uniques. D'autres approches existent, comme par exemple l'utilisation de photons uniques comme intermédiaires entre des portes qui utiliseraient un autre système (atomes, ions, ...) [13, 14], ou encore l'utilisation des fortes non-linéarités produites en couplant les photons à des atomes uniques dans des cavités optiques résonnantes [15, 16], mais ces méthodes ne permettent pas encore de réaliser à la fois les portes et la transmission de l'information entre celles-ci. En 2001 est apparue une nouvelle approche [17, 18, 19] qui repose uniquement sur l'optique linéaire et les détecteurs de photons. Le résultat, basé sur l'interférence entre les photons, est probabiliste, mais le taux de succès peut être augmenté à l'aide de la téléportation quantique, l'idée étant de téléporter sur les qubits une porte ayant déjà fonctionné.

Un tel encodage sur des variables discrètes est aussi utilisé dans de nombreux systèmes de communication quantique dont le premier, et le plus célèbre, est le protocole BB84 [20]. Le codage y est fait aléatoirement entre les bases orthogonales $(|\hat{1}\rangle, |\hat{0}\rangle)$ et $((|\hat{0}\rangle + |\hat{1}\rangle)/\sqrt{2}, (|\hat{0}\rangle - |\hat{1}\rangle)/\sqrt{2})$. La sécurité repose sur le fait qu'un espion qui essaye de lire le message utiliserait pour cela la mauvaise base en moyenne une fois sur deux, introduisant ainsi des erreurs dans la clé transmise. Sa présence peut ainsi être détectée en comparant une partie de l'information transmise.

Les variables discrètes permettent aussi de coder l'information sur plus de deux états : trois pour les *qutrits* [21, 22, 23], un nombre quelconque (mais fixe) pour les *qudits* [24, 25].

Le plus gros défaut actuel des variables discrètes réside dans les technologies utilisées pour produire et détecter les photons uniques. Concernant les sources, il en existe deux types : les sources déterministes et les sources non déterministes. Les premières reposent sur des émetteurs uniques (atomes piégés [26, 27, 28], centre N-V du diamant [29, 30], boîtes quantiques [31, 32], ...) qui vont émettre un photon après qu'on les a excités. Si leur déclenchement peut être parfaitement contrôlé, il n'en est pas de même de la récupération du photon : ces sources ont en effet de faibles efficacités de collection. Le deuxième type de source repose sur des systèmes, comme la fluorescence paramétrique, dont l'émission est aléatoire mais peut être détectée, par exemple avec un détecteur de photon pour un système émettant les photons par paires. Malheureusement ces sources ne sont pas adaptées à des applications nécessitant un grand nombre de photons : la probabilité d'avoir une émission simultanée de toutes les sources décroît en effet de manière exponentielle avec leur nombre.

Les détecteurs de photons les plus courants peuvent uniquement détecter la présence de photons, et non les compter ; de plus, leur efficacité est aux alentours de 50 %. Il existe des détecteurs avec une meilleure efficacité, et qui pour certains peuvent compter les photons ; mais ils sont encore pour la plupart en cours de développement et/ou demandent d'importants systèmes cryogéniques.

2.1.3 Variables continues

L'aspect le plus connu de la description quantique de la lumière est la dualité onde-corpuscule : même si la quantification lui donne un caractère corpusculaire, on peut toujours décrire la lumière de manière ondulatoire. Elle est alors caractérisée, comme n'importe quelle onde, par son amplitude et sa phase qui, contrairement à la section précédente, gardent des valeurs continues.

Si classiquement ce sont l'amplitude et la phase qui prévalent pour décrire une onde, on utilisera plutôt en optique quantique le concept de quadratures, qui sont en fait les quantités

que l'on est amené à mesurer expérimentalement et qui correspondent à des observables bien définies. Rappelons que leur expression s'obtient à partir des opérateurs d'annihilation et de création :

$$\hat{X} = \frac{\hat{a} + \hat{a}^\dagger}{\sqrt{2}} \quad \hat{P} = \frac{\hat{a} - \hat{a}^\dagger}{i\sqrt{2}} \quad (2.25)$$

Si on représente graphiquement l'amplitude complexe de l'onde dans un diagramme de Fresnel, cf. figure 2.1, l'amplitude et la phase représentent les coordonnées polaires et les quadratures les coordonnées cartésiennes (à un facteur $\sqrt{2}$ près).

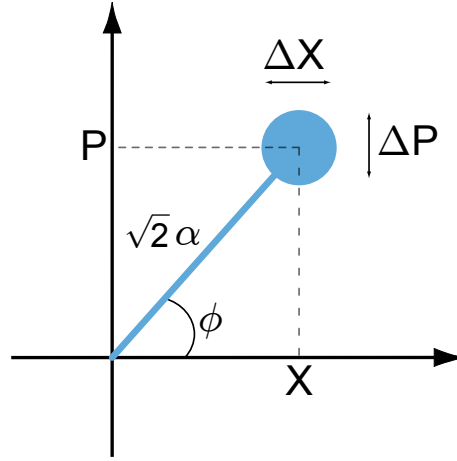


FIGURE 2.1: Représentation de Fresnel de l'amplitude complexe d'une onde

L'amplitude d'une onde est une grandeur absolue, mais sa phase est une grandeur relative : elle ne peut être définie que par rapport à une référence. Changer la phase de référence revient à tourner notre repère dans l'espace des phases, ce qui transforme les quadratures. De cette manière il est possible de définir une infinité de quadratures en fonction de la phase θ servant de référence pour les mesures :

$$\hat{X}_\theta = \hat{X} \cos(\theta) + \hat{P} \sin(\theta) \quad (2.26)$$

$$\hat{P}_\theta = -\hat{X} \sin(\theta) + \hat{P} \cos(\theta) \quad (2.27)$$

Jusqu'ici nous n'avons pas vu de « spécificité quantique » pour le comportement des quadratures. Comme souvent elle réside dans la non-commutation de ces dernières (c.f. équation 2.14). Elles respectent donc une inégalité de Heisenberg qui empêche de les mesurer simultanément [134]

$$\Delta X \Delta P \geq \frac{1}{2} \quad (2.28)$$

Toutes les mesures seront ainsi entachées d'un bruit quantique, représenté par un disque sur le diagramme de Fresnel. Dans le cas discret, ce bruit correspond au « bruit de grenaille » produit par les temps d'arrivée aléatoires des photons d'un faisceau cohérent sur le détecteur. Seulement ces fluctuations sont toujours présentes, y compris dans le vide, ce que l'on peut relier à l'énergie non nulle du vide [34].

Pour ce qui est de l'information et de la communication quantique, les avantages et inconvénients sont relativement opposés à ceux des variables discrètes : on peut facilement produire certains états et les détecter, mais le codage est par contre bien plus compliqué. Côté source,

la plupart des états de base (que l'on verra section 2.3.1) peuvent être produits de manière déterministe et avec de faibles pertes. Quant aux détecteurs, ils sont constitués de simples photodiodes ce qui, outre la simplicité, offre de bien meilleures efficacités quantiques. À contrario coder l'information sur des états non orthogonaux et dont la mesure est intrinsèquement bruitée pose évidemment de nombreux problèmes. C'est d'ailleurs pour des raisons semblables que dans les communications classiques on passe de l'analogique (codage continu) au numérique (codage discret).

Ainsi les protocoles de distribution quantique de clés à variables continues n'ont besoin que d'un simple laser [35, 36] ; par contre leurs performances sont limitées par les algorithmes de traitement de l'information utilisés pour produire la clé à partir des données mesurées.

Si le codage à variables continues est bien connu classiquement pour les communications, c'est moins le cas pour le calcul informatique. Transposer le calcul quantique à un codage continu s'avère donc bien plus compliqué, et les propositions n'ont d'ailleurs encore guère dépassé le stade théorique [37, 38]. Une solution plus simple consiste à utiliser à nouveau des qubits en les codant sur des états dont le recouvrement est suffisamment faible pour que l'on puisse les distinguer correctement. Les états utilisés sont généralement deux états cohérents de phases opposées $|\alpha\rangle$ et $|- \alpha\rangle$ [39]. Si la porte NOT reste simple à réaliser (simple déphasage), ce n'est pas le cas des autres portes, y compris à un qubit.

2.2 Représentation des états quantiques

Les vecteurs d'états $|\psi\rangle$ ne peuvent représenter que des états parfaitement préparés, ou *états purs*. Seulement dans la pratique les imperfections expérimentales conduisent à des mélanges statistiques de différents états purs, et il nous faut donc des outils capables de représenter de tels états, appelés aussi *états mixtes*.

2.2.1 La matrice densité

Si on veut représenter notre état quantique en terme de variables discrètes, l'outil le plus adapté est la *matrice densité*, ou *opérateur densité*, $\hat{\rho}$. Pour un mélange statistique de N états purs $|\psi_k\rangle$ ayant chacun une probabilité p_k (celles-ci étant évidemment normées), on définit la matrice densité par :

$$\hat{\rho} = \sum_{k=1}^N p_k |\psi_k\rangle \langle \psi_k| \quad (2.29)$$

Il s'agit d'un opérateur hermitien, $\hat{\rho}^\dagger = \hat{\rho}$, dont la trace vaut 1. Son expression matricielle dépend évidemment de la base utilisée, la plus commune étant la base des états de Fock. Dans une base donnée, ses éléments diagonaux, les *populations*, donnent la probabilité de mesurer les différents états de la base et les éléments non-diagonaux, appelés *cohérences*, donnent une information sur la cohérence quantique, c'est-à-dire le fait que l'on ait une superposition quantique et non un mélange statistique.

Dans les cas pratiques on utilise généralement des bases finies de dimension peu élevée. Il est alors pratique de représenter graphiquement les coefficients de la matrice densité, la lecture se révélant assez aisée comme le montre la figure 2.2 représentant l'état $\frac{|0\rangle - |1\rangle}{\sqrt{2}}$

Enfin de nombreuses propriétés de l'état sont accessibles simplement en prenant la trace d'un produit de deux opérateurs. Dans le cas de deux états purs nous retrouvons alors le recouvrement entre ceux-ci, et cette opération est par extension appelée formule de recouvrement. Ainsi, pour

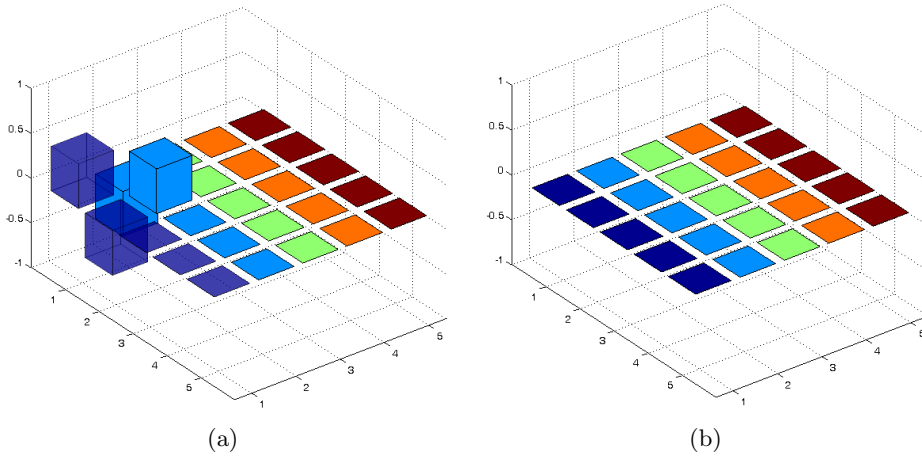


FIGURE 2.2: Matrice densité de l'état $\frac{|0\rangle - |1\rangle}{\sqrt{2}}$ (a) Partie réelle (b) Partie imaginaire

un état défini par la matrice densité $\hat{\rho}$, la valeur moyenne d'un opérateur vaut

$$\langle \hat{O} \rangle = \text{Tr}(\hat{\rho} \hat{O}) \quad (2.30)$$

La fidélité avec un état pur $|\phi\rangle$, qui définit sa « ressemblance » avec ce dernier, vaut ¹

$$\mathcal{F} = \langle \phi | \hat{\rho} | \phi \rangle = \text{Tr}(\hat{\rho} |\phi\rangle \langle \phi|) \quad (2.31)$$

Et nous avons enfin pour la pureté de l'état :

$$\mathcal{P} = \text{Tr}(\hat{\rho}^2) \quad (2.32)$$

Celle-ci permet de quantifier le niveau de pureté de l'état : on peut facilement se convaincre (en diagonalisant la matrice densité) que $\mathcal{P} \leq 1$ et qu'il n'y a égalité que dans le cas d'un état pur. On notera d'ailleurs que la matrice densité d'un état pur est un projecteur ($\hat{\rho}^2 = \hat{\rho}$).

Enfin la matrice densité permet aussi de calculer l'entropie de l'état :

$$S = -\text{Tr}(\hat{\rho} \log_2(\hat{\rho})) \quad (2.33)$$

L'entropie est liée au degré de mélange statistique. Dans le cadre de la physique statistique quantique ce mélange est d'origine thermodynamique, elle est donc liée au désordre, tout comme l'entropie thermodynamique classique. Dans le cadre de l'information quantique il provient des différents états utilisés pour encoder les messages, l'entropie donne donc la quantité d'information qui peut être transmise.

2.2.2 La fonction de Wigner

La matrice densité est par contre très peu adaptée aux variables continues, ce que l'on peut comprendre aisément : les matrices sont des outils discrets, or on ne peut représenter une quantité

1. Cette formule n'est valable que pour la fidélité avec un état pur, dans le cas général elle vaut $\mathcal{F} = \left(\text{Tr} \left(\sqrt{\sqrt{\hat{\rho}_1} \hat{\rho}_2 \sqrt{\hat{\rho}_1}} \right) \right)^2$. Le calcul nécessite alors de diagonaliser les matrices densité.

continue par un objet discret sans tronquer cette quantité (d'où par exemple l'importance du choix du taux d'échantillonnage lors de la numérisation d'un signal analogique).

La physique quantique étant intrinsèquement probabiliste il est tentant de chercher une densité de probabilité dans l'espace des phases, d'une manière analogue à la physique statistique. Malheureusement, qui dit densité de probabilité dans l'espace des phases dit densité de probabilité d'avoir une certaine valeur de x et une certaine valeur de p , et donc de pouvoir mesurer les deux simultanément, ce qui violerait l'inégalité d'Heisenberg. Il est néanmoins possible de définir ce que l'on appelle des *densités de quasi-probabilité* qui ont un comportement proche d'une vraie distribution de probabilité tout en permettant de décrire les états quantiques. Cette définition étant floue on peut ainsi trouver de nombreux candidats, le plus connu, et certainement le plus utilisé, dont nous parlerons dans cette section est la *fonction de Wigner* introduite par E. Wigner en 1932 [40]. Il existe d'autres distributions connues, la *fonction P* et la *fonction Q*, et pour plus d'informations sur celles-ci ou de manière générale sur l'espace des phases en physique quantique on pourra se référer par exemple à [8, 41].

2.2.2.1 Définition

Tout comme classiquement la densité de probabilité dans l'espace des phases est la transformée de Fourier de la fonction caractéristique $\langle e^{-i\mu x - i\nu p} \rangle$, la fonction de Wigner est la transformée de Fourier d'un équivalent quantique de cette fonction caractéristique $2\langle e^{-i\mu \hat{X} - i\nu \hat{P}} \rangle$ [42, 43]. Ceci donne

$$\widetilde{W}(\mu, \nu) = \text{Tr} \left(\hat{\rho} e^{-i\mu \hat{X} - i\nu \hat{P}} \right) \quad (2.34a)$$

$$= e^{-i\mu\nu/2} \int_{-\infty}^{\infty} \langle x | \hat{\rho} e^{-i\nu \hat{P}} e^{-i\mu \hat{X}} | x \rangle dx \quad (2.34b)$$

$$= e^{-i\mu\nu/2} \int_{-\infty}^{\infty} e^{-i\mu x} \langle x | \hat{\rho} | x + \nu \rangle dx \quad (2.34c)$$

$$= \int_{-\infty}^{\infty} e^{-i\mu x} \langle x - \nu/2 | \hat{\rho} | x + \nu/2 \rangle dx \quad (2.34d)$$

où on a utilisé la formule de Baker-Hausdorff pour effectuer la transformation $e^{-i\mu \hat{X} - i\nu \hat{P}} = e^{-i\frac{\mu\nu}{2}} e^{-i\nu \hat{P}} e^{-i\mu \hat{X}}$

En effectuant une transformée de Fourier inverse on obtient la définition de la fonction de Wigner à partir de la matrice densité [8] :

$$W(x, p) = \frac{1}{2\pi} \int e^{i\nu p} \langle x - \nu/2 | \hat{\rho} | x + \nu/2 \rangle d\nu \quad (2.35)$$

De par sa construction il parait normal qu'elle ait un certain nombre des propriétés d'une distribution de probabilité. Ainsi elle est réelle et normalisée : $\iint W(x, p) dx dp = 1$. De plus les distributions de probabilités des quadratures s'obtiennent simplement en projetant la fonction de Wigner sur la quadrature considérée. Plus généralement pour une quadrature quelconque on a :

$$P_{\theta}(x_{\theta}) = \int W(x_{\theta} \cos(\theta) - p_{\theta} \sin(\theta), x_{\theta} \sin(\theta) + p_{\theta} \cos(\theta)) dp_{\theta} \quad (2.36)$$

2. Les opérateurs $\hat{X}$ et $\hat{P}$ ne commutent pas on peut définir plusieurs fonctions caractéristiques quantiques, chacune donnant une nouvelle distribution de quasi-probabilité.

2.2.2.2 Propriétés

Fonction de Wigner d'opérateur L'équation 2.35 peut être étendue à n'importe quel opérateur. Dans le cas où celui-ci peut s'exprimer sous la forme $\hat{A} = f(\hat{X}) + g(\hat{P})$ une décomposition en série entière montre que la fonction de Wigner de l'opérateur vaut

$$W_{\hat{A}}(x, p) = \frac{1}{2\pi} (f(x) + g(p)) \quad (2.37)$$

De plus en effectuant le changement de variable $\nu \mapsto -\nu$ dans l'équation 2.35 on peut vérifier que $W_{\hat{A}^\dagger}(x, p) = W_{\hat{A}}(x, p)^*$. La fonction de Wigner d'un opérateur hermitien est donc bien réelle.

Linéarité L'équation 2.35 est linéaire en $\hat{\rho}$, on peut donc en conclure que la fonction de Wigner d'un mélange statistique se décompose exactement de la même manière que pour une matrice densité :

$$W_{p_1\hat{\rho}_1+\dots+p_N\hat{\rho}_N} = p_1W_{\hat{\rho}_1} + \dots + p_NW_{\hat{\rho}_N} \quad (2.38)$$

Ce comportement est d'ailleurs identique à celui d'une distribution de probabilité.

De la même manière on peut aussi déterminer la fonction de Wigner d'une superposition linéaire de deux états purs $\alpha|\psi\rangle + \beta|\varphi\rangle$ avec la normalisation qui impose $|\alpha|^2 + |\beta|^2 + \text{Re}(2\alpha^*\beta\langle\psi|\varphi\rangle) = 1$. Celle-ci vaut

$$W = |\alpha|^2 W_\psi + |\beta|^2 W_\varphi + 2\text{Re}(\alpha\beta^*W_{|\psi\rangle\langle\varphi|}) \quad (2.39)$$

On peut noter que dans le cas de deux états orthogonaux la norme du dernier terme est nulle.

Recouvrement Une des propriétés les plus remarquables de la fonction de Wigner est la simplicité de l'expression du recouvrement entre deux opérateurs [41, 8]

$$\text{Tr}(\hat{A}\hat{B}) = 2\pi \iint W_{\hat{A}}(x, p)W_{\hat{B}}(x, p)dx dp \quad (2.40)$$

Outre que cela permet de prouver la normalisation de la fonction de Wigner d'un état, elle nous donne aussi, comme pour la matrice densité, l'expression de la valeur moyenne d'un opérateur la fidélité et la pureté

$$\langle\hat{O}\rangle = 2\pi \iint W(x, p)W_{\hat{O}}(x, p)dx dp \quad (2.41)$$

$$\mathcal{F} = 2\pi \iint W(x, p)W_\phi(x, p)dx dp \quad (2.42)$$

$$\mathcal{P} = 2\pi \iint W(x, p)^2 dx dp \quad (2.43)$$

2.2.2.3 Fonctions de Wigner à plusieurs modes

Il est possible de généraliser la définition de la fonction de Wigner à des états multimodes

$$\begin{aligned} W(x_1, p_1, \dots, x_n, p_n) &= \frac{1}{(2\pi)^n} \int \dots \int e^{i\nu_1 p_1 + \dots + i\nu_n p_n} \\ &\times \left\langle x_1 - \frac{\nu_1}{2}, \dots, x_n - \frac{\nu_n}{2} \left| \hat{\rho} \right| x_1 + \frac{\nu_1}{2}, \dots, x_n + \frac{\nu_n}{2} \right\rangle d\nu_1 \dots d\nu_n \end{aligned} \quad (2.44)$$

On remarque aisément que le produit tensoriel de plusieurs opérateurs revient à multiplier les fonctions de Wigner. La trace partielle sur un ou plusieurs modes consiste quant à elle à intégrer les quadratures des modes considérés.

Enfin si on effectue un changement de variables défini par $\hat{\vec{V}}' = M\hat{\vec{V}}$, où $\vec{V}$ et $\vec{V}'$ sont les vecteurs constitués par les quadratures des modes considérés avant et après le changement de variables, alors la nouvelle fonction de Wigner vaut

$$W'(\vec{V}') = W(M^{-1}\vec{V}') \quad (2.45)$$

2.2.2.4 Retour à la matrice densité

Dans certains cas la fonction de Wigner ne suffit pas et il faut repasser par la matrice densité. Nous aurons donc besoin d'une formule pour donner les coefficients de la matrice densité à partir de la fonction de Wigner. Pour cela on peut à nouveau utiliser la formule du recouvrement, en effet

$$\langle m | \hat{\rho} | n \rangle = \text{Tr}(\hat{\rho} | n \rangle \langle m |) = \iint W(x, p) W_{|n\rangle\langle m|}(x, p) dx dp \quad (2.46)$$

Comme nous pourrions le démontrer à l'aide de l'expression de la fonction d'onde, en terme de quadratures, des états de Fock donnée un peu plus loin (section 2.3.2.2)

$$\begin{aligned} W_{|n\rangle\langle m|}(x, p) &= \frac{1}{\sqrt{\pi 2^{n+m} n! m!}} \int e^{i\nu p} \left[\frac{\partial^n}{\partial t^n} e^{-t^2 + 2(x - \frac{\nu}{2})t} \right]_{t=0} e^{-\frac{(x - \frac{\nu}{2})^2}{2}} \\ &\times \left[\frac{\partial^m}{\partial z^m} e^{-z^2 + 2(x + \frac{\nu}{2})z} \right]_{z=0} e^{-\frac{(x + \frac{\nu}{2})^2}{2}} d\nu \end{aligned} \quad (2.47)$$

En réunissant ensemble les formules on obtient

$$\langle m | \hat{\rho} | n \rangle = \left[\frac{\partial_t^m \partial_z^n}{\sqrt{\pi 2^{m+n} m! n!}} \iiint e^{-t^2 - z^2 - x^2 - \frac{\nu^2}{4} + 2x(t+z) + \nu(t-z+ip)} W(x, p) dx d\nu dp \right]_{t=z=0} \quad (2.48)$$

Et donc au final

$$\langle m | \hat{\rho} | n \rangle = \frac{2}{\sqrt{2^{m+n} m! n!}} \left[\frac{\partial^{n+m}}{\partial t^m \partial z^n} \iint W(x, p) e^{-x^2 - p^2 + 2x(t+z) - 2ip(t-z) - 2tz} dx dp \right]_{t=z=0} \quad (2.49)$$

2.2.3 Reconstruction par tomographie

2.2.3.1 Principe

Les mesures ne nous donnent évidemment pas accès directement à la fonction de Wigner d'un état mais uniquement aux distributions de probabilité des quadratures. Comme le dit Ulf Leonhardt dans son ouvrage [41] nous ne pouvons voir que les ombres des états quantiques tout comme les « prisonniers éclairés » de la caverne de Platon. Néanmoins si nous n'avons accès qu'aux ombres il est possible de reconstruire la description de l'état à partir de plusieurs ombres. C'est exactement ce qui est fait lors d'une tomographie médicale : à partir des mesures d'absorption de rayons X suivant différentes orientations il est possible de reproduire la cartographie des organes. Cette méthode peut être transposée dans le domaine de l'optique quantique afin de reconstruire la fonction de Wigner à partir des distributions de probabilité des quadratures [44].

2.2.3.2 Algorithmes utilisés

Au cours de ce travail de thèse deux algorithmes ont été utilisés pour reconstruire la fonction de Wigner ; il en existe cependant de nombreux autres dont nous ne parlerons pas ici.

Transformée de Radon inverse Il s'agit de la méthode la plus communément employée ; elle a été formulée en 1917 par le mathématicien Johann Radon qui est à l'origine des processus de tomographie. Elle s'obtient en inversant l'équation 2.36 ce qui donne l'opérateur intégral

$$W(x, p) = \frac{1}{2\pi^2} \int_0^\pi \int_{-\infty}^\infty P_\theta(y) K(x \cos(\theta) + p \sin(\theta) - y) dy d\theta \quad (2.50)$$

où K est le noyau défini par

$$K(y) = \frac{1}{2} \int_{-\infty}^\infty |\epsilon| e^{i\epsilon y} d\epsilon \quad (2.51)$$

Il faut néanmoins prendre certaines précautions lors de la mise en œuvre numérique à cause de l'échantillonnage et du bruit lié à l'expérience, sans compter le fait que le noyau est une distribution et non une fonction bien définie. On introduira donc une fonction de lissage soigneusement choisie afin d'effectuer le calcul numérique.

Maximum de vraisemblance L'algorithme de maximum de vraisemblance, « Maximal Likelihood » ou « MaxLik » en anglais, cherche de manière itérative à obtenir la matrice densité qui donne les distributions de probabilités des quadratures les plus proches de celles mesurées [45, 46]. Cette proximité est donnée par la fonction de vraisemblance

$$\mathcal{L} = \prod_{j=1}^N \left(p_j^{f_j} \right) \quad (2.52)$$

où N est le nombre de mesures, p_j la probabilité d'obtenir chaque résultat pour la matrice densité considérée et f_j le nombre d'occurrences expérimentales de chaque résultat. En pratique les données obtenues sont divisées en histogrammes pour chaque quadratures, les mesures correspondent donc à l'ensemble des *bin* des histogrammes.

Afin de trouver la matrice densité qui maximise cette vraisemblance on utilisera un opérateur [47]

$$\hat{R}(\hat{\rho}) = \sum_{j=1}^N \left(\frac{f_j}{p_j} \hat{\Pi}_j \right) \quad (2.53)$$

où $\hat{\Pi}_j$ est le projecteur sur le résultat j . Pour la matrice densité qui maximise la vraisemblance cet opérateur est proportionnel à l'identité puisque f_j est alors proportionnel à p_j , elle est donc stable sous l'action de celui-ci. L'algorithme consiste simplement à partir de la matrice identité et à appliquer l'opérateur $\hat{R}$ (en renormalisant) jusqu'à atteindre le point fixe.

$$\hat{\rho}^{(n+1)} = \frac{\hat{R}(\hat{\rho}^{(n)}) \hat{\rho}^{(n)} \hat{R}(\hat{\rho}^{(n)})}{\text{Tr} \left(\hat{R}(\hat{\rho}^{(n)}) \hat{\rho}^{(n)} \hat{R}(\hat{\rho}^{(n)}) \right)} \quad (2.54)$$

Le grand avantage de cette méthode est qu'il est possible d'incorporer les pertes dues au système de mesure dans le calcul des probabilités p_j ; de cette manière on peut reconstruire la fonction de Wigner corrigée de ces pertes, c'est-à-dire celle de l'état réellement produit.

2.3 Quelques états « de base »

2.3.1 Les états gaussiens

Chaque point de vue, discret et continu, a des états qui sont plus simples à décrire. Dans le cas discret, comme nous l'avons déjà vu, se sont les états de Fock. Pour les variables continues se sont les états gaussiens, c'est-à-dire ceux dont la fonction de Wigner est une gaussienne.

Les états gaussiens sont caractérisés par les valeurs moyennes, $\langle x \rangle$ et $\langle p \rangle$, et les écart-types, Δx et Δp , des quadratures. Leur fonction de Wigner vaut donc

$$W(x, p) = \frac{1}{\pi} e^{-\frac{(x-\langle x \rangle)^2}{2\Delta x^2} - \frac{(p-\langle p \rangle)^2}{2\Delta p^2}} \quad (2.55)$$

Les états gaussiens purs sont les états d'incertitude minimale, ou *états minimaux*. Une méthode simple pour s'en convaincre consiste à calculer la pureté d'un état gaussien en utilisant l'équation 2.43. On trouve ainsi

$$\mathcal{P} = \frac{1}{2\Delta x \Delta p} \quad (2.56)$$

On en déduit donc que

$$\mathcal{P} = 1 \Leftrightarrow \Delta x \Delta p = \frac{1}{2} \quad (2.57)$$

On peut alors introduire un *facteur de compression* s tel que $\Delta x^2 = \frac{s}{2}$ et $\Delta p^2 = \frac{1}{2s}$.

Cette section décrit les principaux états gaussiens qui sont les états de bases de la majeure partie des expériences d'optique quantique avec des variables continues.

2.3.1.1 Le vide quantique

L'état vide $|0\rangle$, décrivant une absence de photon, est l'état d'énergie minimale, c'est donc l'état fondamental du champ électromagnétique. Ses fluctuations sont symétriques par rotation de la phase : $\Delta x^2 \Delta p^2 = \frac{1}{2}$, elles sont appelés bruit quantique, « shot noise limit » ou SNL.

La fonction d'onde du vide vaut

$$\psi_0(x) = \langle x | \psi_0 \rangle = \frac{1}{4\pi} e^{-\frac{x^2}{2}} \quad (2.58)$$

et sa fonction de Wigner

$$W_0(x, p) = \frac{1}{\pi} e^{-x^2 - p^2} \quad (2.59)$$

2.3.1.2 Les états cohérents

Les états cohérents sont ce qui se rapproche le plus des ondes planes progressives monochromatiques utilisées en optique classique. Ils doivent donc avoir une amplitude et une phase aussi bien définie que possible dans les limites permises par la relation de Heisenberg. L'expression du champ électromagnétique, donné par l'équation 2.18, conduit à définir ces états comme étant les états propres de l'opérateur d'annihilation

$$\hat{a}|\alpha\rangle = \alpha|\alpha\rangle \quad (2.60)$$

Ce sont les états produits par un laser très au dessus du seuil. Leur fonction de Wigner vaut

$$W_\alpha(x, p) = \frac{1}{\pi} e^{-(x-\sqrt{2}\text{Re}(\alpha))^2 - (p-\sqrt{2}\text{Im}(\alpha))^2} \quad (2.61)$$

Ils correspondent en fait à la translation du vide dans l'espace des phases, ce qui se réalise mathématiquement en appliquant l'opérateur de déplacement

$$\hat{D}_\alpha = e^{i\sqrt{2}(\text{Re}(\alpha)\hat{P} - \text{Im}(\alpha)\hat{X})} \quad (2.62)$$

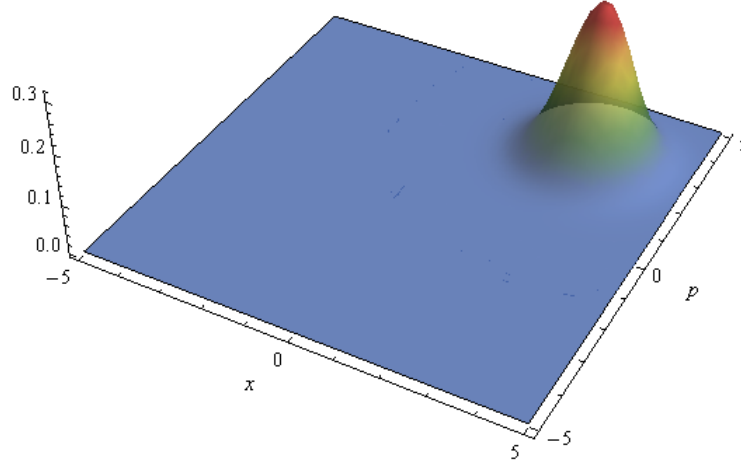


FIGURE 2.3: Fonction de Wigner d'un état cohérent

À partir de l'équation 2.60 on peut aussi trouver la décomposition d'un état cohérent sur la base de Fock :

$$|\alpha\rangle = e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle \quad (2.63)$$

Celle-ci nous donne la probabilité de mesurer n photons :

$$P_{\alpha}(n) = |\langle n|\alpha\rangle|^2 = e^{-|\alpha|^2} \frac{|\alpha|^{2n}}{n!} \quad (2.64)$$

On retrouve alors la distribution poissonnienne, qui s'interprète en terme de photons indépendants arrivant aléatoirement sur le détecteur.

Pour finir, les états cohérents ne sont pas orthogonaux entre-eux, le recouvrement entre deux d'entre eux valant

$$\langle \beta|\alpha\rangle = e^{-\frac{|\beta-\alpha|^2}{2}} \quad (2.65)$$

Ils forment néanmoins une base de l'espace de Hilbert, mais en raison de cette non orthogonalité la décomposition n'est alors pas unique.

2.3.1.3 Les états comprimés monomodes

Jusqu'ici nous avons uniquement vu des états dont le bruit est symétrique par rapport à la phase, mais il est aussi possible de comprimer le bruit sur l'une des quadratures et par conséquent de l'amplifier sur l'autre afin de préserver l'inégalité de Heisenberg. De tels états sont appelés états comprimés et leur fonction de Wigner est donnée par

$$W_s(x, p) = \frac{1}{\pi} e^{-\frac{(x - \sqrt{2} \operatorname{Re}(\alpha))^2}{s} - \frac{(p - \sqrt{2} \operatorname{Im}(\alpha))^2}{1/s}} \quad (2.66)$$

où s est le facteur de compression précédemment introduit. Mathématiquement ils s'obtiennent en appliquant l'opérateur de compression $\hat{S} = e^{\frac{r}{2}(\hat{a}^2 - (\hat{a}^\dagger)^2)}$ aux états cohérents, r étant le *paramètre de compression* défini par $s = e^{-2r}$. Le plus utilisé de ces états est le *vide comprimé*,

correspondant à la compression du vide. Contrairement à ce que son nom pourrait laisser croire,

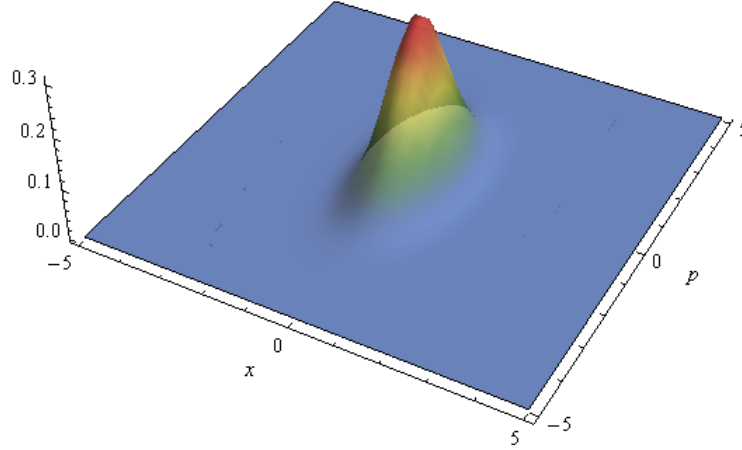


FIGURE 2.4: Fonction de Wigner d'un vide comprimé

bien que les moyennes des champs restent nuls, ce n'est plus un vide de photons. Le fait de comprimer le vide lui ajoute des photons ; de manière générale le vide est le seul état qui n'en a pas, donc dès qu'on le transforme en un autre état on en fait forcément apparaître.

Une décomposition sur la base de Fock nous permet de voir que le vide comprimé est constitué uniquement de nombres pairs de photons

$$|\psi_s\rangle = \frac{1}{\sqrt{\cosh(r)}} \sum_{n=0}^{\infty} \sqrt{C_n^{2n}} \left(-\frac{\tanh(r)}{2} \right)^n |2n\rangle \quad (2.67)$$

2.3.1.4 Les états comprimés bimodes ou états EPR

Une des particularités de la physique quantique est la notion d'intrication. En terme de variables continues le plus simple consiste à effectuer la compression/amplification du bruit sur les corrélations et anti-corrélations des quadratures. Plus exactement, les corrélations entre les quadratures x des deux modes sont comprimées et leurs anti-corrélations amplifiées, tandis que pour les quadratures p c'est l'inverse qui se produit. Autrement dit, en mesurant une quadrature d'un mode on a une bonne idée, d'autant meilleure que la compression est importante, de la valeur de cette même quadrature dans l'autre mode. Dans la limite d'une compression infinie on retrouve l'état utilisé par Einstein, Podolsky et Rosen dans leur article de 1935 [48] pour conclure à un caractère incomplet de la physique quantique.

De tels états peuvent être produits soit en mélangeant deux vides comprimés dans des directions orthogonales sur une lame semi-réfléchissante 50/50 [49, 50, 51, 52] (cf. section 3.3.1.2), soit directement à partir d'un amplificateur paramétrique non dégénéré [122, 67, 55, 56, 57] (cf. section 3.3.3.1). La première méthode nous permet de calculer la distribution de Wigner :

$$W_{\text{EPR}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2} e^{-\frac{(x_1-x_2)^2}{2s} - \frac{(x_1+x_2)^2}{2/s} - \frac{(p_1-p_2)^2}{2/s} - \frac{(p_1+p_2)^2}{2s}} \quad (2.68)$$

La décomposition sur la base de Fock nous montre un autre aspect intéressant des états EPR : chacun des deux modes contient le même nombre de photons. Cet état permet donc de produire des états de Fock de manière probabiliste en comptant les photons dans l'un des deux modes.

$$|\psi_{\text{EPR}}\rangle = \frac{1}{\cosh(r)} \sum_{n=0}^{\infty} \tanh^n(r) |n, n\rangle \quad (2.69)$$

2.3.1.5 Les états thermiques

Il s'agit de l'état obtenu, entre autres, en ne regardant qu'un mode d'un état EPR, l'autre étant « perdu dans l'environnement ». Sa fonction de Wigner s'obtient donc en intégrant sur un des modes

$$W_{\text{th}}(x, p) = \frac{1}{\pi \cosh(2r)} e^{-\frac{x^2 + p^2}{\cosh(2r)}} \quad (2.70)$$

On remarque que cet état a une variance plus importante que la limite quantique, ce n'est donc pas un état pur. On peut obtenir de la même manière sa matrice densité en effectuant une trace partielle sur celle de l'état EPR :

$$\hat{\rho}_{\text{th}} = \frac{1}{\cosh^2(r)} \sum_{n=0}^{\infty} (\tanh^{2n}(r) |n\rangle \langle n|) \quad (2.71)$$

Nous pouvons voir que les cohérences sont nulles dans la base de Fock : il s'agit uniquement d'un mélange statistique d'états nombres.

Enfin son nom vient du fait qu'il décrit un champ électromagnétique en équilibre avec un réservoir thermodynamique à température T ; dans ce cas nous avons

$$\tanh^2(r) = e^{-\frac{\hbar\omega}{k_B T}} \quad (2.72)$$

2.3.2 Les états non gaussiens

2.3.2.1 Caractéristiques et intérêt

Jusqu'à présent, les fonctions de Wigner que nous avons introduites se comportent comme des distributions de probabilité : les fonctions de Wigner des états gaussiens sont en effet des fonctions positives qui pourraient parfaitement être des distributions de probabilité. Et, bien qu'ils aient des caractéristiques quantiques comme le bruit imposé par l'inégalité de Heisenberg et l'intrication des quadratures, on se doute bien qu'il est possible d'aller encore plus loin, de « pousser » la fonction de Wigner au delà des limites en dehors desquelles elle cesse de ressembler à une densité de probabilité et ainsi faire apparaître son caractère fondamentalement quantique.

Anticipons un peu sur la section suivante et regardons la distribution de probabilité des quadratures d'un état à un photon

$$P_1(x) = \frac{2}{\sqrt{\pi}} x^2 e^{-x^2} \quad (2.73)$$

Celle-ci est valable pour toutes les quadratures et elle est nulle à l'origine. Quel objet pourrait projeter la même ombre dans toutes les directions sans jamais rien projeter au centre ? Si toutes les projections sont identiques quel que soit l'angle c'est forcément que l'objet est invariant par rotation. Or si rien n'est projeté c'est, dans un monde classique, qu'il n'y a rien à projeter, et ce sur tout l'axe de la projection. Du coup à cause de la symétrie par rotation il ne devrait rien

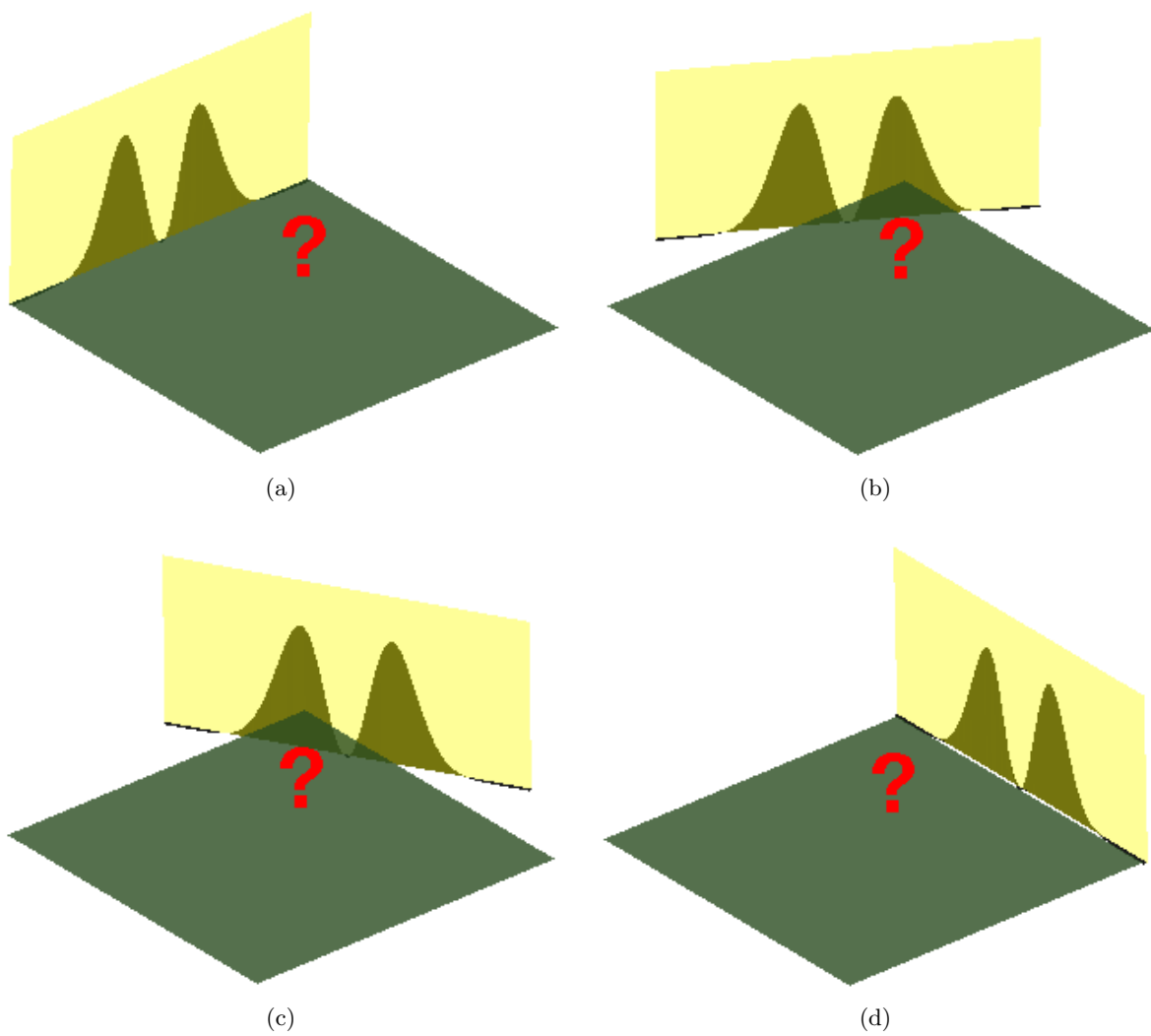


FIGURE 2.5: Tomographie d'un état à un photon (a) $\theta = 0^\circ$ (b) $\theta = 30^\circ$ (c) $\theta = 60^\circ$ (d) $\theta = 90^\circ$

y avoir. Mais il y a bien quelque chose de projeté sur les bords. Il ne reste donc qu'une seule possibilité : ce qui est au centre de l'objet annule ce qui est à ses bord. On a bien le comportement que l'on cherchait : à l'extrémité de la projection on ne passe que par les bords de l'objet donc on a quelque chose, tandis qu'au milieu de la projection on passe également par le centre qui vient compenser et on a rien. En termes concrets la fonction de Wigner d'un état à un photon est négative en son centre.

Théorème de Hudson-Piquet On sait maintenant que certains états ont une fonction de Wigner négative, mais on peut se demander s'ils sont nombreux. Les états gaussiens n'en font clairement pas partie mais qu'en est-il des autres ? C'est là qu'intervient le théorème de Hudson-Piquet : tout état pur non gaussien a une fonction de Wigner négative [58]. C'est donc loin d'être une propriété marginale, et en pratique elle est d'ailleurs souvent utilisée pour caractériser la non-classicité d'un état. Il n'y a évidemment pas de telle règle pour les mélanges statistiques qui peuvent être divers et variés.

Intérêt des états non gaussien Cette « signature quantique » ne relève pas uniquement de la curiosité scientifique : les états non gaussiens possèdent de nombreux atouts que n'ont pas leur homologues gaussiens. À tel point que certaines opérations comme par exemple la *distillation d'intrication* [59, 60, 61, 62], le calcul quantique [63, 64] ou les codes correcteurs d'erreur [65] sont impossibles sans eux.

2.3.2.2 Les états de Fock

Comme on pouvait s'y attendre, les états de Fock étant plus liés à une description discrète, leur expression en terme de variables continues est plus compliquée. Ainsi on peut les construire facilement d'un point de vue discret en appliquant plusieurs fois l'opérateur de création

$$|n\rangle = \frac{(\hat{a}^\dagger)^n}{\sqrt{n!}} |0\rangle \quad (2.74)$$

Par contre leur description continue fait intervenir des polynômes de degré proportionnel au nombre de photons. En effet la fonction d'onde d'un état de Fock est

$$\phi_n(x) = \frac{1}{\pi^{1/4} \sqrt{2^n n!}} H_n(x) e^{-\frac{x^2}{2}} \quad (2.75)$$

où H_n sont les polynômes de Hermite

$$H_n(x) = \left[\frac{\partial^n}{\partial t^n} e^{-t^2 + 2xt} \right]_{t=0} = (-1)^n e^{x^2} \frac{d^n}{dt^n} e^{-x^2} \quad (2.76)$$

Et la fonction de Wigner de l'état $|n\rangle$ vaut

$$W_n = \frac{(-1)^n}{\pi} e^{-x^2 - p^2} L_n(2(x^2 + p^2)) \quad (2.77)$$

où L_n sont les polynômes de Laguerre

$$L_n(x) = \sum_{k=0}^n \left(C_n^k \frac{(-x)^k}{k!} \right) \quad (2.78)$$

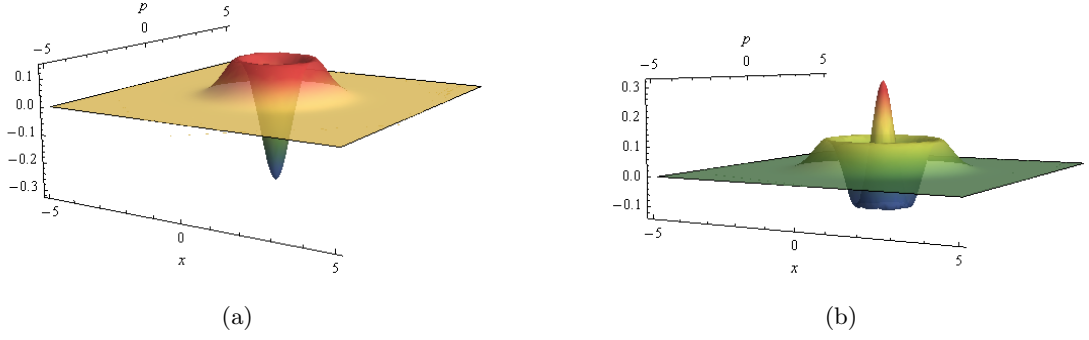


FIGURE 2.6: Fonction de Wigner des états de Fock (a) État à un photon (b) État à deux photons

Les états de Fock sont invariants par rotation dans l'espace des phase : ils sont liés à l'énergie du champ et nous retrouvons donc la propriété, déjà présente classiquement, que l'intensité ne dépend pas de la phase. De plus on peut remarquer que pour un état nombre impair la fonction de Wigner est négative à l'origine et la probabilité de mesurer une quadrature nulle (et donc un champ nul) est nulle.

2.3.2.3 Les états chats de Schrödinger

Le dernier état que nous mentionnerons est la transposition de l'expérience de pensée qui est sans nul doute la plus connue du grand public dans le domaine de la physique quantique : le chat de Schrödinger [66]. Suite à celle-ci le nom de « chat de Schrödinger » a été donné aux états correspondant à la superposition d'états considérés comme classiques et discernables classiquement. Dans le cas de l'optique quantique il s'agit d'une superposition d'états cohérents de phases opposées

$$|\psi_{\text{chat}}\rangle = \frac{|\alpha\rangle + e^{i\theta} |-\alpha\rangle}{\sqrt{2(1 + \cos(\theta) e^{-2|\alpha|^2})}} \quad (2.79)$$

Nous pouvons alors trouver la fonction de Wigner d'un tel état grâce à l'équation 2.39 :

$$W_{\text{chat}}(x, p) = \frac{e^{-x^2 - p^2}}{\pi(1 + \cos(\theta) e^{-2|\alpha|^2})} \left(e^{-2|\alpha|^2} \cosh(2\sqrt{2}\alpha x) + \cos(2\sqrt{2}\alpha p - \theta) \right) \quad (2.80)$$

Parmi les différents chats possibles deux cas sont particulièrement intéressants : il s'agit des « chats pairs », correspondant à $\theta = 0$, et des « chats impairs », correspondant à $\theta = \pi$. Leur nom vient du fait que leur décomposition sur la base de Fock fait uniquement intervenir des nombres de photons respectivement pairs et impairs :

$$|C_+\rangle = \frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2(1 + e^{-2|\alpha|^2})}} = \frac{\sqrt{2}e^{-|\alpha|^2/2}}{\sqrt{1 + e^{-2|\alpha|^2}}} \sum_{k=0}^{\infty} \frac{\alpha^{2k}}{\sqrt{(2k)!}} |2k\rangle \quad (2.81)$$

$$|C_-\rangle = \frac{|\alpha\rangle - |-\alpha\rangle}{\sqrt{2(1 - e^{-2|\alpha|^2})}} = \frac{\sqrt{2}e^{-|\alpha|^2/2}}{\sqrt{1 - e^{-2|\alpha|^2}}} \sum_{k=0}^{\infty} \frac{\alpha^{2k+1}}{\sqrt{(2k+1)!}} |2k+1\rangle \quad (2.82)$$

Les fonctions de Wigner des chats pairs et impairs sont illustrés figure 2.7. Nous pouvons y

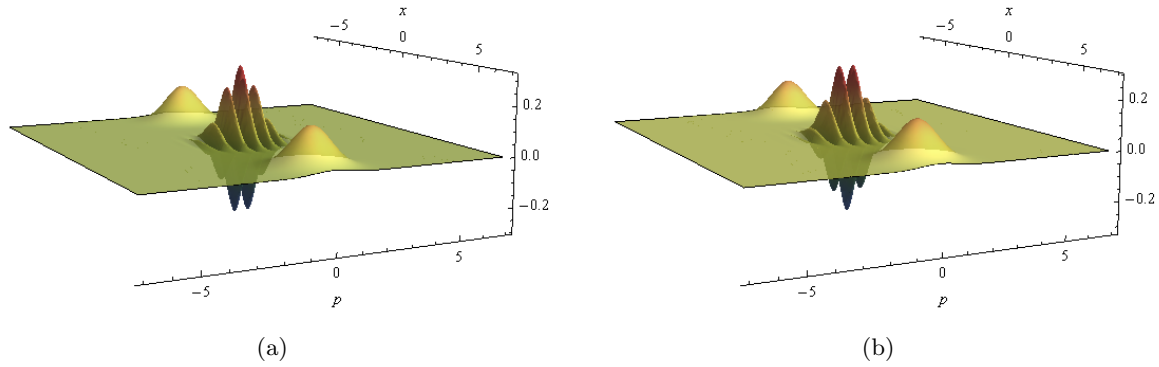


FIGURE 2.7: Fonctions de Wigner des états chats de Schrödinger pour $\alpha^2 = 15$ (a) Chat pair (b) Chat impair

voir deux « bosses » parfaitement distinctes correspondant aux deux états « classiques » (chat vivant et chat mort) ainsi que des franges d'interférences qui sont la preuve de la superposition quantique. Un mélange statistique donnerait les deux mêmes bosses mais sans ces franges.

Ces états sont particulièrement utiles pour le calcul quantique avec des qubits à états cohérents puisque les chats pairs et impairs correspondent respectivement aux superpositions $(|\hat{0}\rangle + |\hat{1}\rangle)/\sqrt{2}$ et $(|\hat{0}\rangle - |\hat{1}\rangle)/\sqrt{2}$ lors d'un codage sur la phase d'un état cohérent. Enfin, il est aussi possible de coder le qubit sur la parité du chat ; nous retrouvons alors un codage sur des états orthogonaux, mais en contre-partie la mesure de l'état est plus compliquée.

2.4 Conclusion

Nous disposons maintenant des outils théoriques nécessaires afin de décrire les états quantiques de la lumière autant d'un point de vue discret que d'un point de vue quantique. Nous avons aussi mentionné l'intérêt qu'il y a à mélanger les deux approches, et nous exploiterons ce mélange au cours des différentes expériences présentées dans ce manuscrit.

Chapitre 3

Le dispositif expérimental

Sommaire

3.1 Présentation du dispositif	40
3.1.1 Introduction	40
3.1.2 Schéma générique	41
3.2 La source laser	42
3.2.1 Le laser impulsionnel	42
3.2.2 Diagnostic du faisceau	43
3.2.2.1 Photodiodes	43
3.2.2.2 Spectromètre	43
3.2.2.3 Auto-corrélateur	45
3.2.2.4 Profilomètre	45
3.3 Les transformations unitaires	46
3.3.1 L'optique linéaire	46
3.3.1.1 Déphasage	46
3.3.1.2 lame semi-réfléchissante	46
3.3.1.3 lame demi-onde et quart-d'onde	47
3.3.1.4 Cube séparateur de polarisation (PBS)	48
3.3.2 Le générateur de seconde harmonique (GSH)	48
3.3.2.1 Principe	48
3.3.2.2 Effets parasites	50
3.3.3 L'amplificateur paramétrique optique (OPA)	51
3.3.3.1 Configuration non-dégénérée	52
3.3.3.2 Configuration dégénérée	55
3.4 Détection et mesures projectives	57
3.4.1 La détection homodyne	57
3.4.1.1 Principe	57
3.4.1.2 Imperfections	58
3.4.1.3 Montage	59
3.4.2 La photodiode à avalanche	63
3.4.2.1 Caractéristiques du dispositif	63
3.4.2.2 Mesures projectives et conditionnement	63
3.5 Production de photons uniques	64
3.5.1 Production	64

3.5.2	Modélisation	64
3.5.2.1	Amplification paramétrique	66
3.5.2.2	Pertes homodynes et APD	66
3.5.2.3	Conditionnement	66
3.5.3	Utilisations	68
3.5.3.1	Taux de production et matrice densité	68
3.5.3.2	Caractérisation des imperfections	69
3.6	Conclusion	70

3.1 Présentation du dispositif

3.1.1 Introduction

Dans le chapitre précédent nous avons vu les outils théoriques utilisés pour ce travail de thèse, passons maintenant aux outils expérimentaux. Nous avons appris qu'un certain nombre d'opérations essentielles à l'information quantique nécessitent des états et/ou des opérations non gaussiens. Nous devons donc disposer d'outils qui ont un caractère non gaussien.

Le premier point à regarder est la production des états de base. En regardant de plus près la « zoologie » présentée en section 2.3 on peut s'apercevoir que beaucoup d'états, aussi bien gaussiens que non gaussiens, peuvent être produits à partir d'une paire EPR :

- en recombinaison les deux faisceaux sur une lame semi-réfléchissante on obtient deux vides comprimés sur deux quadratures orthogonales.
- en comptant le nombre de photons présents dans un des faisceaux on produit un état de Fock dans l'autre.
- en « jetant » un faisceau dans l'environnement on prépare un état thermique.

Mais pour cela nous avons besoin d'un moyen de produire cet état EPR. Le plus simple consiste à utiliser l'*amplification paramétrique* [122, 67] qui est un effet non linéaire du second ordre.

Les non-linéarités, y compris celles du second ordre, nécessitent de grandes puissances optiques afin d'observer des effets significatifs. Celles-ci peuvent être obtenues de deux manières :

- Soit en plaçant le cristal non-linéaire dans une cavité optique résonnante afin que le faisceau laser passe un très grand nombre de fois à travers celui-ci, multipliant ainsi l'effet non-linéaire. Cette technique donne des faisceaux de très bonne qualité et permet d'obtenir de fortes compressions [68, 69] et de très bonnes efficacités homodynes [126]. Par contre elle demande de maintenir verrouillées de nombreuses cavités (pour les cristaux non-linéaires mais aussi pour le filtrage spectral) ainsi que d'appliquer un filtrage temporel aux données fournies par les détecteurs afin de définir le mode laser étudié [71].
- Soit en utilisant des impulsions ultra-brèves afin d'avoir de fortes puissances crêtes. Outre l'absence de cavité, cette méthode permet de mieux définir les modes considérés, que l'on peut directement assimiler aux impulsions laser. Cette caractéristique s'avère extrêmement pratique que se soit pour effectuer les expériences ou les calculs pour les modéliser. Elle simplifie par ailleurs le conditionnement multiple.

La solution que nous avons retenue est la deuxième, qui offre de plus un échantillonnage naturel pour les applications de communication quantique.

Pour ce qui est des opérations non gaussiennes, la première méthode à avoir été proposée consiste à utiliser des effets non-linéaires du troisième ordre comme l'*effet Kerr* [72]. Cette

méthode a l'avantage d'être déterministe, malheureusement il n'existe pas actuellement de matériaux capables de produire ces effets de manière suffisamment intense. Une autre technique, non déterministe cette fois, consiste à utiliser des mesures projectives pour effectuer des opérations non gaussiennes [73, 74]. C'est cette méthode que nous utiliserons ; les mesures les plus simples à mettre en œuvre étant destructives (absorption du faisceau par la photodiode), nous aurons besoin d'effectuer celles-ci sur des modes auxiliaires intriqués avec le reste du système.

3.1.2 Schéma générique

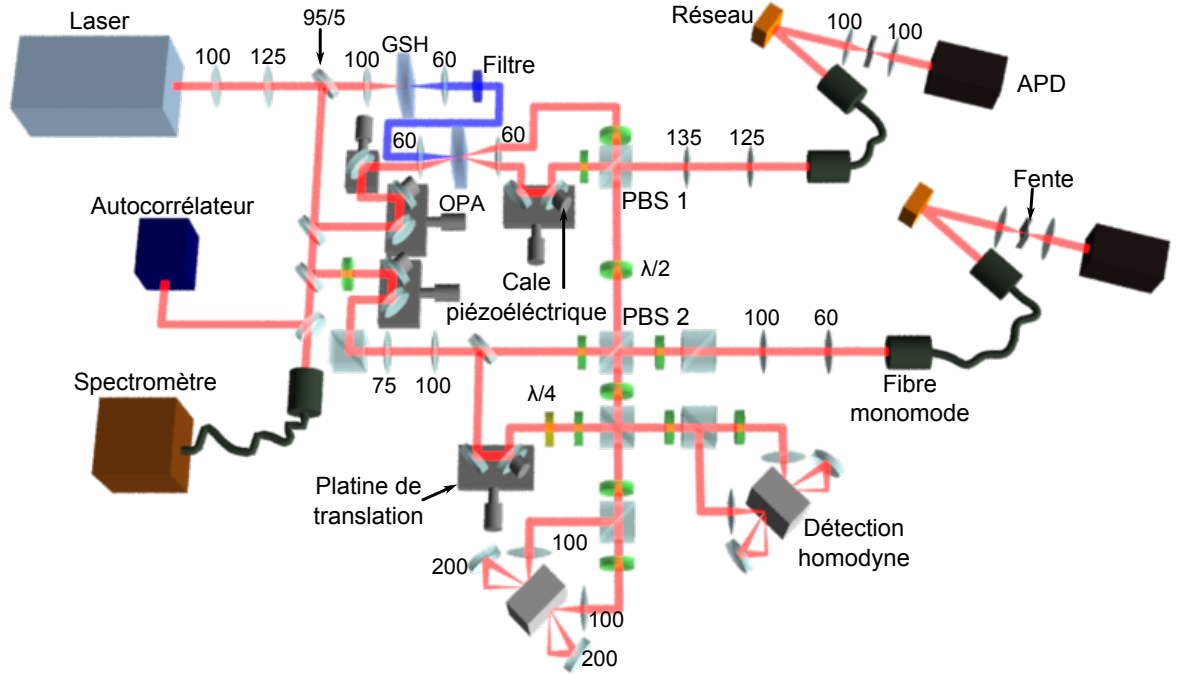


FIGURE 3.1: Schéma du dispositif expérimental

La figure 3.1 illustre le schéma expérimental générique utilisé pour les différentes expériences présentées dans ce manuscrit. Le dispositif peut se séparer en trois parties :

1. Le laser ainsi que les appareils permettant d'en effectuer le diagnostic.
2. Les différents éléments d'optique permettant de manipuler la lumière : lames semi-réfléchissantes, lames d'ondes, cubes séparateurs de polarisation et cristaux non-linéaires.
3. Les systèmes de détections composés de deux APD et deux détections homodynes.

Les impulsions laser sont ainsi d'abord créées à l'intérieur d'un laser titane-saphir fonctionnant en cavité fermée et extraites de celui-ci par un *cavity dumper* (cellule de Bragg). La majeure partie de ces impulsions (95%) est doublée par un premier cristal non-linéaire, afin de pomper ensuite l'amplificateur paramétrique qui produit nos états quantiques de base. L'autre partie sert au diagnostic du faisceau, aux réglages, à former l'oscillateur local servant de référence pour les détections homodynes et à générer les états cohérents utilisés dans les expériences. Tous ces états sont alors manipulés par des éléments d'optique linéaire et des opérations non gaussiennes réalisées par des mesures projectives. Enfin les états ainsi transformés sont caractérisés par une ou deux détections homodynes.

L'utilisation de lames demi-onde couplées à des cubes séparateurs de polarisation pour jouer le rôle des lames séparatrices présente deux avantages majeurs :

- Elle permet de régler à volonté la transmission et la réflexion, donnant toute sa versatilité au dispositif : il suffit de tourner quelques lames d’ondes, bloquer quelques faisceaux (ou les atténuer pour passer d’un faisceau de réglage à un petit état cohérent) et éventuellement retirer ou replacer le cube *PBS 2* pour passer d’une expérience à une autre, ou des réglages à l’expérience.
- Elle permet aussi d’utiliser un encodage des modes en polarisation et donc de les superposer spatialement. Cette technique permet ainsi d’éviter de nombreux asservissements en phase : deux modes superposés spatialement suivent en effet le même chemin optique et donc subissent le même déphasage.

Voyons maintenant plus en détails les différentes parties du dispositif ainsi que leurs effets sur les états quantiques.

3.2 La source laser

3.2.1 Le laser impulsif

Le laser utilisé pour ce travail de thèse est un *Tiger-CD* produit par la société *Time-Bandwidth Product*. Bien qu’étant originellement un modèle commercial il a été poussé à son maximum, au delà des spécifications garanties par le constructeur, afin d’avoir le plus de puissance possible. Il est pompé à l’aide d’un laser continu à 532 nm (Nd :YAG doublé) *Verdi V5* de la société *Coherent*, venant remplacer le laser de pompe d’origine, un *Melles-Griot 58 GLS 309* dont le profil spatial n’était pas assez bon ($M^2 \approx 1,1$ pour le V5 au lieu de $M^2 \approx 4$ pour le Melles-Griot). La puissance de pompe est d’environ 3,1 W ; elle peut-être légèrement modifiée afin de trouver un point de fonctionnement plus stable. Le verrouillage de mode s’effectue à l’aide d’un absorbant saturable semi-conducteur appelé SESAM (pour *SEmiconductor Saturable Absorber Mirror*) [75]. Enfin la cellule de Bragg (*Neos N13389*) est pilotée par un boîtier électronique (*Neos N64389-SYN*) permettant de régler la fréquence d’extraction ainsi que la puissance, la phase et le moment d’émission de l’onde RF.

Longueur d’onde	varie typiquement entre 845 et 852 nm
Largeur spectrale	varie autour de 4,2 nm
Durée des impulsions	entre 170 et 230 fs
Produit temps-fréquence	$\Delta\tau_{FWHM} \cdot \Delta\nu_{FWHM} \approx 0,35$ (limite de Fourier à 0,315 pour un profil en sécante hyperbolique)
Puissance moyenne	varie entre 32 et 42 mW
Cadence de répétition	800 kHz (peut être changée)
Énergie par impulsion	entre 40 et 50 nJ
Puissance crête	entre 190 et 300 kW

TABLE 3.1: *Caractéristiques du Tiger-CD*

Les caractéristiques du laser sont résumés dans la table 3.1. Les intervalles donnés dans cette table ne correspondent pas à des paramètres pleinement réglables, mais plutôt aux valeurs qui peuvent être prises lorsque le laser est bien réglé. Nous ne pouvons en effet pas pousser le laser en dehors de ses spécifications en espérant qu’il se règle aisément.

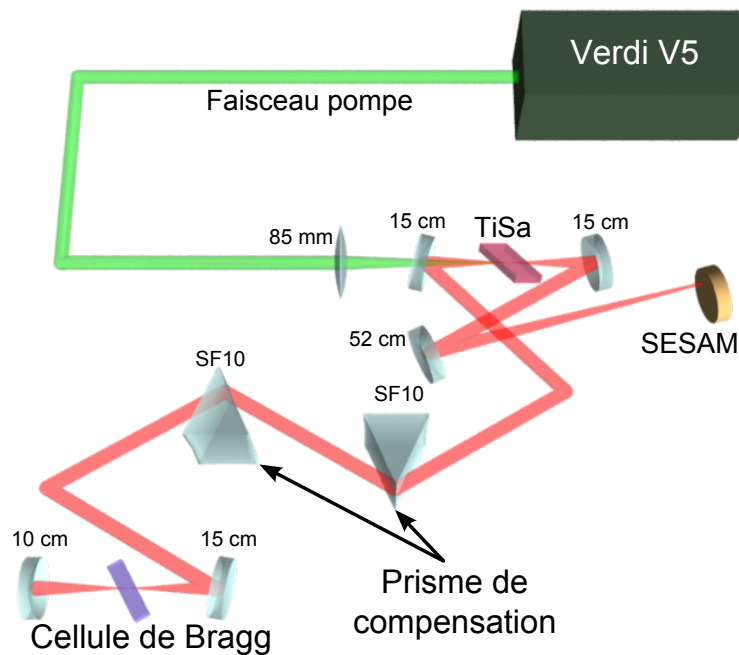


FIGURE 3.2: Cavit  du Tiger-CD

3.2.2 Diagnostic du faisceau

3.2.2.1 Photodiodes

Des photodiodes permettent de surveiller la puissance moyenne du laser, qui est le principal crit re de r glage du *Tiger-CD*. Le signal doit  tre r gulier et bien « droit ». Un signal lentement variable est g n ralement signe que la phase de l'onde RF est mal r gl e, tandis que des sauts tr s courts indiquent que sa puissance est trop forte. Des variations plus rapides ou un signal compl tement bruit  indiquent le plus souvent soit un mauvais r glage du moment d'extraction soit un probl me dans l'alignement de la cavit . Le laser a aussi tendance   privil gier deux modes : celui d sir  et un autre bien moins intense ; la puissance moyenne permet donc de distinguer les sauts entre ces deux modes.

Enfin l'utilisation d'une photodiode rapide coupl e   un oscilloscope   forte bande passante (2,5 GHz) permet de d tecter la pr sence de doubles impulsions, signe d'un mauvais r glage de la cavit .

3.2.2.2 Spectrom tre

Le spectre est surveill    l'aide d'un spectrom tre optique *Anritsu MS900 1B1* ayant une r solution de 0,1 nm. Il permet de v rifier que le mode spectral du laser est bon (cf. figure 3.3) : le pic principal doit  tre bien plus grand que les pics secondaires. On se contente g n ralement de l'amener dans la zone de bon fonctionnement (cf. caract ristiques). Le spectrom tre est donc  galement utile pour conna tre la longueur d'onde afin de r gler la temp rature des cristaux non lin aires (cf. section 3.3.2).

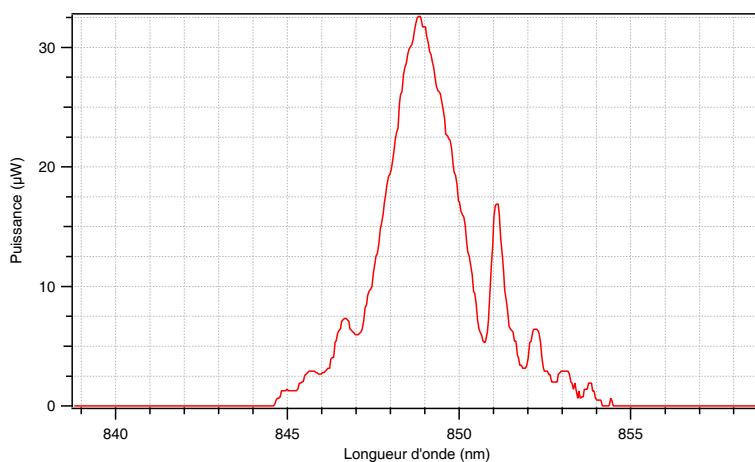


FIGURE 3.3: Spectres du laser femtoseconde

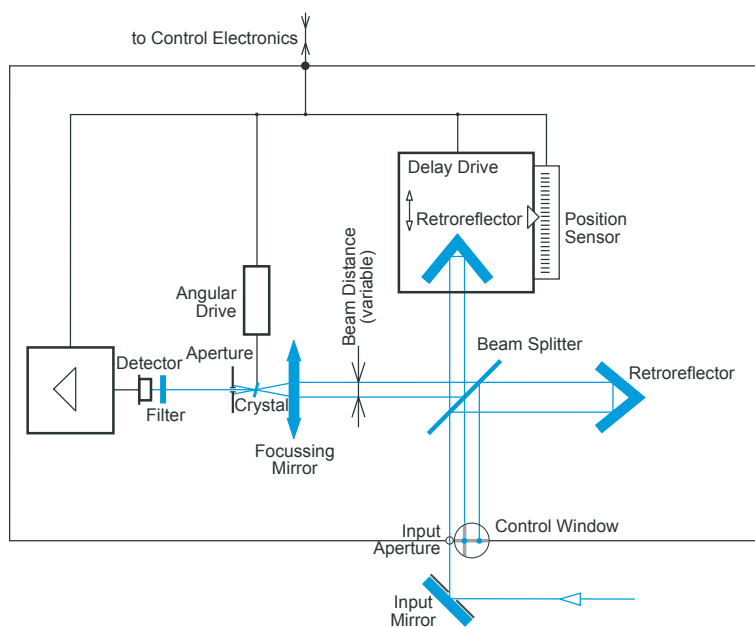


FIGURE 3.4: Fonctionnement de l'auto-corrélateur

3.2.2.3 Auto-corrélateur

La durée des impulsions est mesurée à l'aide d'un auto-corrélateur *PulseCheck 15* de la société *APE GmbH*. Le principe de fonctionnement (cf. figure 3.4) consiste à effectuer, dans un cristal non-linéaire du second ordre, la superposition non-colinéaire de deux parties du faisceau séparées par une lame semi-réfléchissante et décalées spatialement et temporellement. On obtient alors un signal à fréquence double $I_{2\omega}(\tau) = \int I_{\omega}(t)I_{\omega}(t + \tau)dt$ qui, mesurée par un photo-multiplicateur, peut être reconstruit en balayant le décalage temporel τ . En effectuant une hypothèse sur la forme du profil temporel on peut alors déduire la largeur de l'impulsion. L'hypothèse effectuée ici est celle d'un profil en sécante hyperbolique¹ ; dans ce cas la largeur à mi-hauteur de l'impulsion vaut : $\Delta\tau_{\text{FWHM}} = \Delta\tau_{\text{FWHM,autocorr}}/1,543$. Un second jeu de cristal et de photo-multiplicateur permet d'effectuer une corrélation croisée entre le faisceau à 850 nm et celui obtenu par doublage de fréquence afin de mesurer la durée de ce dernier.

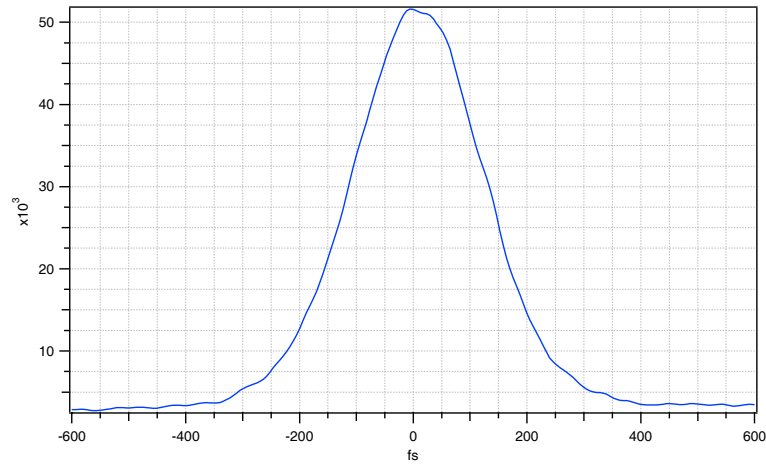


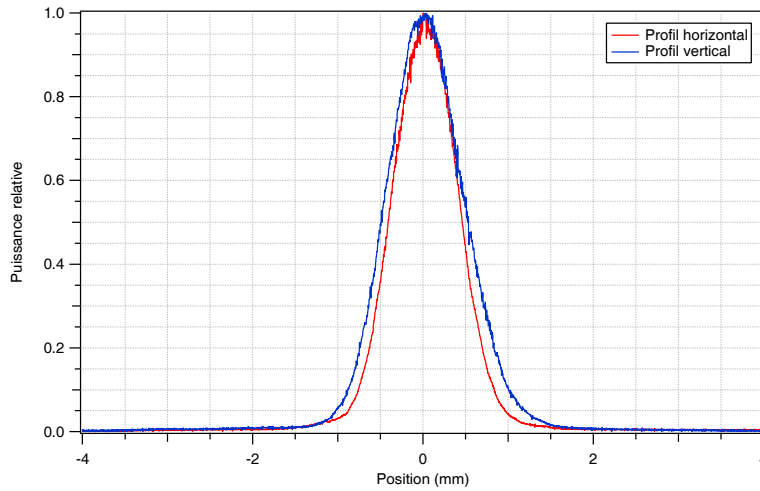
FIGURE 3.5: Signal d'auto-corrélation

L'auto-corrélateur permet en plus de détecter les doubles impulsions trop rapprochées pour être vues par la photodiode rapide, le délai entre les deux bras de l'auto-corrélateur pouvant être balayé manuellement sur une plage d'environ 100 ps.

3.2.2.4 Profilomètre

Le mode spatial peut être caractérisé en sortie du laser à l'aide d'un profilomètre *BP109-UV* de *Thorlabs*, qui, comme son nom l'indique, mesure le profil spatial du faisceau suivant deux axes orthogonaux. Sa dégradation provient quasiment systématiquement d'un mauvais réglage de la cellule de Bragg. Les différents faisceaux utilisés pour les réglages sont trop peu puissants pour le profilomètre, leurs modes spatiaux doivent donc être vérifiés visuellement (à l'aide d'une caméra infrarouge), les dégradations sont essentiellement dues à un mauvais centrage sur les optiques.

1. La sécante hyperbolique est l'inverse du cosinus hyperbolique, à noter que lorsque l'on parle d'un profil en sécante hyperbolique il s'agit en réalité du carré d'une sécante hyperbolique.

FIGURE 3.6: *Profil spatial*

3.3 Les transformations unitaires

En physique quantique l'évolution d'un système est décrite par des transformations unitaires et linéaires. D'un point de vue pratique ces transformations sont effectuées à partir des divers éléments optiques usuels ; l'unitarité induit la réversibilité de ces transformations, qui se retrouve dans le principe de retour inverse de la lumière. Attention cependant à ne pas confondre linéarité par rapport au champ électrique et linéarité de l'évolution quantique. Ainsi l'optique non-linéaire correspond tout autant à des transformations quantiques linéaires que l'optique linéaire plus traditionnelle.

Les transformations unitaires que nous utiliserons présentent en outre la particularité de s'exprimer comme une transformation linéaire des quadratures, qui peut être représentée sous forme matricielle et ainsi utilisée pour calculer la fonction de Wigner de l'état issu de cette transformation (cf. section 2.2.2.3).

3.3.1 L'optique linéaire

3.3.1.1 Déphasage

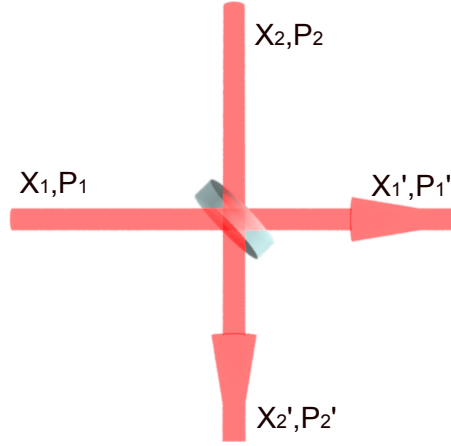
La plus simple de ces transformations est certainement le déphasage, puisqu'elle ne demande aucun élément optique. Pour effectuer un déphasage de θ il suffit de laisser le faisceau se propager sur un chemin optique de $L = \frac{\lambda}{2\pi}\theta$. Les quadratures se transforment alors selon :

$$\begin{pmatrix} x' \\ p' \end{pmatrix} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{pmatrix} \begin{pmatrix} x \\ p \end{pmatrix} \quad (3.1)$$

Il est possible de créer des déphasages modifiables à l'aide de cales piézoélectriques.

3.3.1.2 lame semi-réfléchissante

Il s'agit d'une lame de verre traitée de manière à réfléchir une fraction bien déterminée du faisceau. Ces lames sont caractérisées par leur transmittivité t et leur réflectivité r , ou par leur

FIGURE 3.7: *Lame semi-réfléchissante*

réflexion $R = r^2$ et leur transmission $T = t^2$ (qui vérifient $R + T = 1$ en l'absence de pertes).

$$\begin{pmatrix} x_1' \\ p_1' \\ x_2' \\ p_2' \end{pmatrix} = \begin{pmatrix} t & 0 & r & 0 \\ 0 & t & 0 & r \\ -r & 0 & t & 0 \\ 0 & -r & 0 & t \end{pmatrix} \begin{pmatrix} x_1 \\ p_1 \\ x_2 \\ p_2 \end{pmatrix} \quad (3.2)$$

La réflexion de l'une des voies inverse la phase, ce qui est indispensable pour assurer la conservation de l'énergie.

Elles ont peu de pertes ($< 0,25 \%$), et celles que nous utilisons dans notre dispositif sont traitées afin de minimiser la dispersion des impulsions dans les couches diélectriques assurant la réflexion (c'est aussi le cas pour les miroirs utilisés).

3.3.1.3 Lame demi-onde et quart-d'onde

Ce sont des lames biréfringentes utilisées pour modifier la polarisation en fonction de leur orientation, donnée par l'angle θ entre l'axe lent et la verticale (ou de manière identique entre l'axe rapide et l'horizontale). La lame demi-onde ajoute une phase π entre l'axe rapide et l'axe lent, ce qui donne la transformation suivante

$$\begin{pmatrix} x_H' \\ p_H' \\ x_V' \\ p_V' \end{pmatrix} = \begin{pmatrix} \cos(2\theta) & 0 & \sin(2\theta) & 0 \\ 0 & \cos(2\theta) & 0 & \sin(2\theta) \\ -\sin(2\theta) & 0 & \cos(2\theta) & 0 \\ 0 & -\sin(2\theta) & 0 & \cos(2\theta) \end{pmatrix} \begin{pmatrix} x_H \\ p_H \\ x_V \\ p_V \end{pmatrix} \quad (3.3)$$

On remarque qu'elle a la même action qu'une lame semi-réfléchissante mais sur des modes séparés en polarisation.

Les lames quart-d'onde ajoutent une phase $\pi/2$ entre l'axe rapide et l'axe lent, transformant les polarisations linéaires en polarisations elliptiques et vice-versa

$$\begin{pmatrix} x_H' \\ p_H' \\ x_V' \\ p_V' \end{pmatrix} = \begin{pmatrix} 1 & \cos(2\theta) & 0 & \sin(2\theta) \\ -\cos(2\theta) & 1 & -\sin(2\theta) & 0 \\ 0 & \sin(2\theta) & 1 & -\cos(2\theta) \\ -\sin(2\theta) & 0 & \cos(2\theta) & 1 \end{pmatrix} \begin{pmatrix} x_H \\ p_H \\ x_V \\ p_V \end{pmatrix} \quad (3.4)$$

3.3.1.4 Cube séparateur de polarisation (PBS)

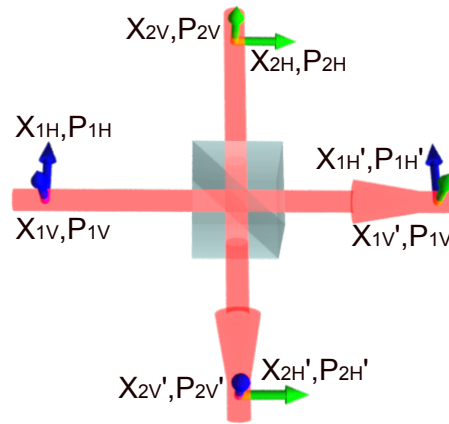


FIGURE 3.8: Cube séparateur de polarisation

C'est un cube composé de deux prismes séparés par un traitement diélectrique qui permet de réfléchir la polarisation verticale tandis que la polarisation horizontale, qui arrive sous incidence de Brewster, est transmise. Il est malheureusement impossible d'avoir à la fois une bonne transmission (nécessite d'avoir peu de couches diélectriques) et un bon taux de réjection (nécessite d'en avoir beaucoup). Ils ont typiquement une transmission de 95–98 % (pour les meilleurs) et des taux de réjection allant de 1/100 à 1/1000.

Ils servent notamment à séparer les polarisations, ce qui est indispensables en pratique si l'on veut utiliser des lames demi-onde comme lames semi-réfléchissantes. De plus, en envoyant un second faisceau dans l'autre entrée, il est possible de superposer les modes ainsi séparés avec deux nouveaux modes afin d'effectuer de nouveaux mélanges pour la suite de l'expérience.

3.3.2 Le générateur de seconde harmonique (GSH)

La génération d'états à la longueur d'onde du laser par un amplificateur paramétrique demande que ce dernier soit pompé par un faisceau à fréquence double, et il nous faut donc générer celui-ci.

3.3.2.1 Principe

La génération de secondes harmoniques est un processus non-linéaire du second ordre qui transforme deux photons en un photon de fréquence double ; dans notre cas nous passons donc de 850 à 425 nm. À cette relation correspondant à la conservation de l'énergie s'ajoute une condition d'accord de phase imposée par la conservation de l'impulsion. On a ainsi $\vec{k}_{\text{bleu}} = 2\vec{k}_{\text{IR}}$ pour un accord de phase colinéaire de type I. Cette condition impose une égalité des indices de réfraction qui peut être satisfaite en utilisant des matériaux biréfringents ; l'accord de phase s'obtient alors en tournant le cristal et /ou en changeant la température.

Afin de conserver l'aspect monomode des impulsions, i.e. pour que le mode soit aussi peu dégradé que possible, il convient de respecter certains critères :

- Fort coefficient non linéaire $\chi^{(2)}$ afin de réduire au minimum l'épaisseur du cristal, et donc ses effets sur le mode du faisceau, tout en ayant des effets non-linéaires significatifs.
- Faible décalage entre les directions de propagation des faisceaux (*walk-off*) de manière à limiter les problèmes de recouvrement spatial des impulsions

- Faible dispersion de vitesse de groupe (GVD : *Group Velocity Dispersion*), qui correspond à l'étalement de l'impulsion dans un milieu dispersif.
- Faible désaccord de vitesse de groupe (GVM : *Group Velocity Mismatch*) qui induirait un problème de recouvrement supplémentaire : si les impulsions bleue et infrarouge se déplacent à des vitesses différentes lors de la génération, alors l'impulsion de pompe bleue n'amplifiera plus le signal dès qu'elle ne lui sera plus superposée, et générera toute une série d'impulsions de vide comprimé décalées temporellement les unes par rapport aux autres, générant ainsi une seule impulsion bien plus large temporellement.

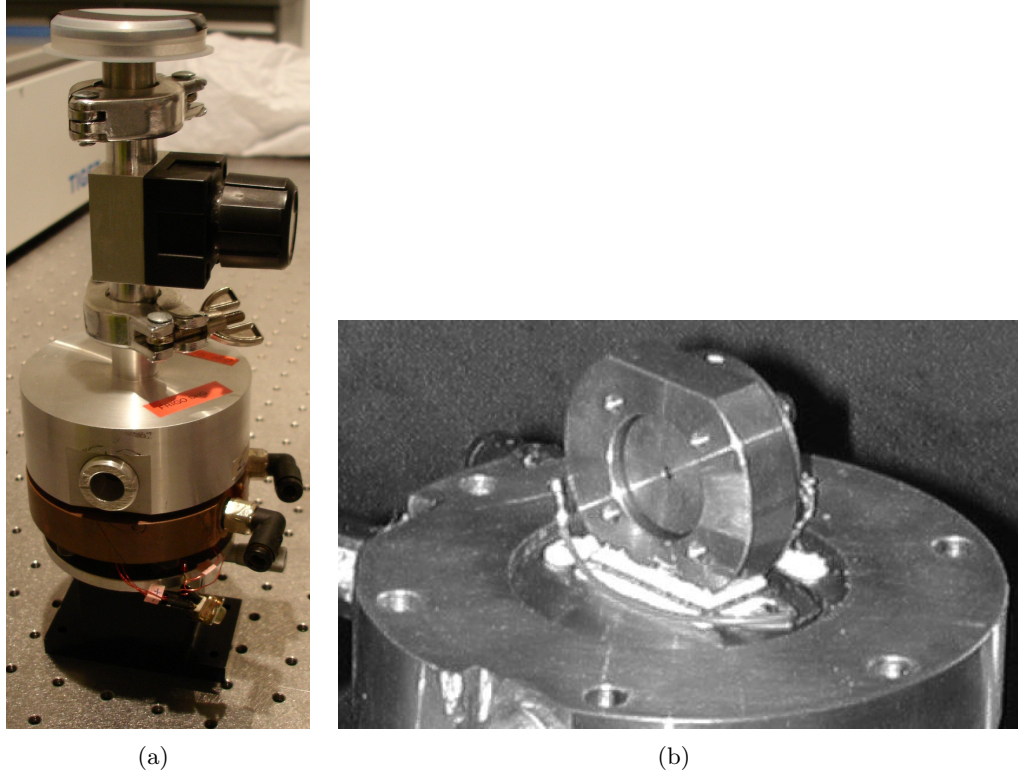


FIGURE 3.9: Enceinte à vide contenant un cristal de KNbO_3 (a) Vue extérieure (b) Support du cristal

Les cristaux utilisés, que se soit pour la génération de seconde harmonique ou pour l'amplification paramétrique présentée un peu plus loin, sont en niobate de potassium (KNbO_3) et proviennent de la société *FEE GmbH*. Ils sont traités anti-reflet aux longueurs d'ondes concernées. Ce sont des cristaux biréfringents *biaxiaux*. La coupe étant effectuée perpendiculairement à l'axe a , on peut alors se ramener au cas uniaxe en considérant l'axe b comme étant l'axe ordinaire et l'axe c comme étant l'axe extraordinaire, appelé aussi axe optique. La propagation étant orthogonale à celui-ci l'accord de phase est non critique (l'indice dépend peu de l'axe de propagation) ce qui induit une absence de *walk-off*. L'accord de phase est de type I (ou *ooe*) : le faisceau infrarouge est polarisé suivant l'axe ordinaire et le faisceau bleu suivant l'axe extraordinaire ; il est réglé en ajustant la température du cristal (figure 3.10) [76]. Pour cela ce dernier est refroidi par effet Peltier et placé dans une enceinte à vide afin d'éviter la condensation (figure 3.9). Les caractéristiques du cristal sont données table 3.2.

Le faisceau est focalisé de manière à avoir un waist $w_0 = 16 \mu\text{m}$ au niveau du cristal, ce qui donne un paramètre confocal de $2z_R = 2\pi n w_0^2 / \lambda \approx 4 \text{ mm}$. Celui-ci est très grand devant

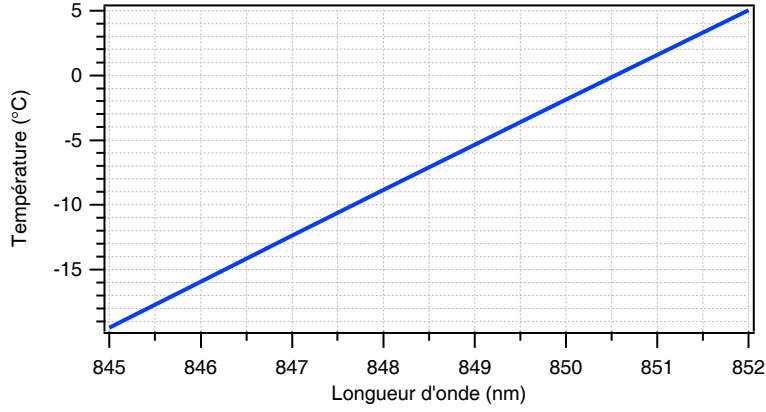


FIGURE 3.10: Température d'accord de phase en fonction de la longueur d'onde du laser

Coefficient non linéaire	$d_{\text{eff}} = \frac{\xi^{(2)}}{2} = 12 \text{ pm/V}$
Indice de réfraction	$n_{b,\omega} = n_{c,2\omega} = 2,281$
GVD	$0,38 \text{ fs}^2/\mu\text{m}$
GVM	$1,2 \text{ ps/mm}$
Absorption ²	$< 0,002 \text{ cm}^{-1}$
Seuil de dommage	2 GW/cm^2

TABLE 3.2: Caractéristiques des cristaux de KNbO_3

l'épaisseur du cristal qui est de $L = 100 \mu\text{m}$ et nous pouvons donc négliger la diffraction et considérer le front d'onde comme plan. Dans ce cas l'efficacité de conversion théorique vaut [77, 78] :

$$\eta_{\text{GSH}} = \tanh^2 \left(\frac{L}{L_{\text{NL}}} \right) \quad (3.5)$$

où la longueur d'interaction non-linéaire est donnée par

$$L_{\text{NL}} = \sqrt{\frac{w_0^2 n^3 \epsilon_0 c \lambda^2}{16 \pi d_{\text{eff}}^2 P_{\text{crte}}}} \quad (3.6)$$

3.3.2.2 Effets parasites

Idéalement l'efficacité de conversion devrait atteindre 90% [81]. Dans la pratique la plus forte efficacité observée est de 32%, et il faut généralement la baisser aux alentours de 10–15% en diminuant la focalisation afin de limiter la dégradation du mode.

Les effets parasites limitant l'efficacité sont proportionnels à l'intensité du faisceaux infrarouge. Ainsi pour une forte focalisation (par exemple $8 \mu\text{m}$) nous devons soit baisser la puissance (figure 3.11(a)) soit défocaliser le faisceau (figure 3.11(b)) afin d'obtenir le maximum d'efficacité. Ces effets parasites sont essentiellement dus à l'absorption à deux photons [81, 82]. Les performances maximales en termes de doublage sont alors obtenues avec un *waist* de $16 \mu\text{m}$. Néanmoins, cette focalisation entraîne une déformation spatiale qui devient visible au bout d'environ 1 h d'utilisation dans la même zone du cristal. Cette distorsion vient principalement des effets photoréfractifs [79], qui induisent des modifications de l'indice de réfraction. Afin de garder un faisceau à peu près monomode spatialement nous devons donc défocaliser le faisceau jusqu'à

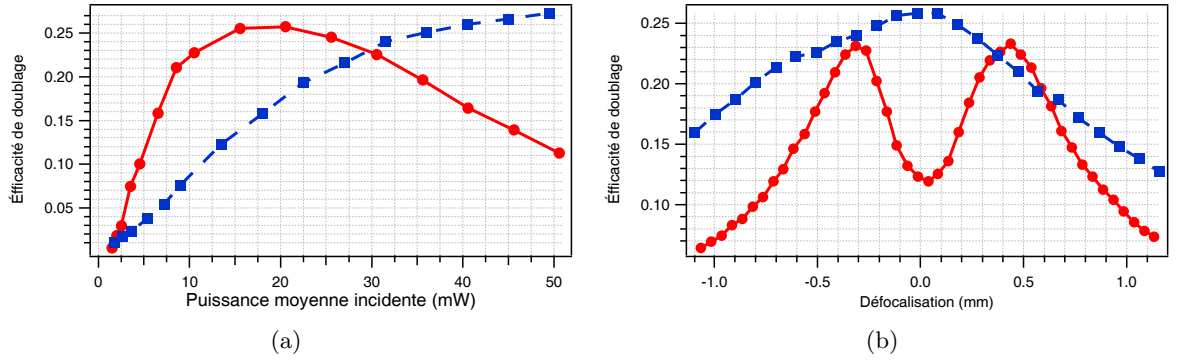


FIGURE 3.11: Efficacité de doublage pour des waist de 16 μm (carrés) et 8 μm (disques) (a) en fonction de la puissance moyenne, le waist étant au centre du cristal (b) en fonction de la défocalisation (à puissance constante)

ce qu'il ait une taille d'environ 40 μm au niveau du cristal. De plus des platines de translations permettent de translater transversalement le cristal de manière à changer la zone du cristal utilisée, qui finit par s'user avec le temps.

Par ailleurs, les expériences nécessitant des états avec extrêmement peu de photons, les photons résiduels à 850 nm doivent être négligeables. Il faut donc filtrer ceux-ci en sortie du cristal, ce qui rajoute des pertes et fait donc baisser à nouveau la quantité de bleu. La série de filtres utilisés conduit à une transmission de 10^{-16} pour l'infrarouge et 80% pour le bleu. Au final la puissance moyenne de bleu peut atteindre jusqu'à 9 mW, mais elle doit généralement être limitée aux alentours de 3,5 mW, soit 4,4 nJ par impulsion avec notre taux de répétition de 800 kHz.

3.3.3 L'amplificateur paramétrique optique (OPA)

L'amplification paramétrique optique est le processus inverse du précédant : un photon du faisceau pompe est transformé en deux photons dans les modes *signal* et *complémentaire* (ou *idler*) de fréquences ω_s et ω_c . Elle est décrite par l'hamiltonien

$$\hat{H} = \frac{i\hbar}{\tau_{NL}} (\hat{a}_p^\dagger \hat{a}_s \hat{a}_c + \hat{a}_p \hat{a}_s^\dagger \hat{a}_c^\dagger) \quad (3.7)$$

Le montage utilisé est exactement le même que pour le GSH. La focalisation est telle que le waist de la pompe vaut $w_p = \sqrt{2}w_s$, où w_s est le waist du faisceau sonde injecté dans le mode signal. Ce facteur, différent du facteur théoriquement optimal [83], a été trouvé empiriquement afin d'obtenir un bon compromis entre gain paramétrique et qualité du mode spatial [84]. Le facteur optimal demande en effet que la pompe soit plus petite que le signal, mais si le signal ne voit pas un front d'onde de la pompe à peu près plat l'effet non-linéaire devient plus intense au centre du faisceau signal qu'en son pourtour, ce qui déforme évidemment le mode spatial et même diminue l'efficacité d'amplification (on appelle ce phénomène *Gain Induced Diffraction* ou *GID*) [85].

Ce processus a lieu

- soit de manière similaire à l'émission stimulée (on parle d'*amplification paramétrique*) en envoyant un faisceau sonde dans le mode signal. Les fréquences sont alors bien définies ; dans notre cas nous avons $\omega_s = \omega_c = \omega_p/2$

- soit de manière similaire à l'émission spontanée (*fluorescence paramétrique*). Dans ce cas l'émission est multimode, la seule contrainte étant $\omega_s + \omega_c = \omega_p$, et il sera donc nécessaire de filtrer la sortie. Nous pouvons noter que de la même manière que l'émission spontanée correspond à une émission stimulée du vide, la fluorescence paramétrique est une amplification paramétrique du vide.

Il possède aussi deux configurations possibles :

- la configuration dégénérée, où les modes signal et complémentaires sont identiques
- la configuration non-dégénérée, où ils sont différents (dans notre cas ils sont séparés spatialement). L'accord de phase $\vec{k}_p = \vec{k}_s + \vec{k}_c$ est là aussi assez large pour permettre une émission spatialement multimode dans le cas de la fluorescence paramétrique.

3.3.3.1 Configuration non-dégénérée

Dans la pratique, cette configuration s'obtient en baissant la température d'environ 6°C par rapport à celle du GSH ; l'accord de phase a alors lieu pour des vecteur d'onde $\vec{k}_s$ et $\vec{k}_c$ ayant un angle d'environ 5° avec celui de la pompe.

La pompe peut être considérée comme étant classique (correspondant à un état cohérent $|\alpha_p\rangle$ qui satisfait l'approximation $\hat{a}_p |\alpha_p\rangle \approx \hat{a}_p^\dagger |\alpha_p\rangle$), l'hamiltonien devient alors

$$\hat{H}_{\text{nd}} = \frac{i\hbar}{\tau_{\text{NL}}} (\hat{a}_s \hat{a}_c + \hat{a}_s^\dagger \hat{a}_c^\dagger) \quad (3.8)$$

En calculant l'évolution des opérateurs en représentation de Heisenberg nous pouvons en déduire la transformation sur les quadratures :

$$\begin{pmatrix} x_s' \\ p_s' \\ x_c' \\ p_c' \end{pmatrix} = \begin{pmatrix} \cosh(r) & 0 & -\sinh(r) & 0 \\ 0 & \cosh(r) & 0 & \sinh(r) \\ -\sinh(r) & 0 & \cosh(r) & 0 \\ 0 & \sinh(r) & 0 & \cosh(r) \end{pmatrix} \begin{pmatrix} x_s \\ p_s \\ x_c \\ p_c \end{pmatrix} \quad (3.9)$$

où $r = \frac{\tau}{\tau_{\text{NL}}}$ avec τ qui est le temps d'interaction donné par $\tau = L/v_g$, où v_g est la vitesse de groupe.

Dans le cas où l'on envoie un faisceau sonde dans le mode signal, celui-ci est amplifié d'un gain en intensité $g = \cosh^2(r)$, le faisceau complémentaire en sortie du cristal est lui aussi proportionnel au faisceau sonde, mais le facteur de proportionnalité vaut $g - 1$. En l'absence d'un tel faisceau, l'application de la transformation au vide nous donne l'état EPR évoqué au début de ce chapitre. La propriété de cet état en terme de nombre de photons se voit aisément dans le mode de production : les photons sont émis par paires, un photon dans le mode signal et un autre dans le complémentaire ; il y a donc bien autant de photons dans chacun des deux modes. L'intrication entre les modes signal et complémentaire est d'autant plus grande que le paramètre de compression r est grand. Cette dernière peut être rapprochée du cas classique où la relation $\varphi_p = \varphi_s + \varphi_c$ apparaissant entre les phases implique $X_c = X_s$ et $P_c = -P_s$ si on prend $\varphi_p = 0$ comme référence des phases.

L'intrication peut être observée expérimentalement en mesurant les corrélations et anticorrélations des deux faisceaux. Ceci est en fait identique à la production de deux vides comprimés par recombinaison de la paire EPR sur une lame semi-réfléchissante : on peut en effet voir le calcul des corrélations et anti-corrélations comme un moyen de mélanger numériquement les faisceaux pour produire des vides comprimés (ou la lame séparatrice comme un moyen optique d'obtenir ces corrélations et anti-corrélations). D'ailleurs, bien que bimode, nous n'avons qu'une seule phase à faire varier. En effet seule la moyenne des phases des quadratures modifie les

données mesurées. Si on modifie l'écart entre les phases des deux faisceaux tout en gardant constante cette moyenne, c'est-à-dire qu'on déphase le signal de $\delta\theta/2$ et le complémentaire de $-\delta\theta/2$, l'application de la transformation de déphasage (équation 3.1) laisse inchangée la fonction de Wigner de l'état donnée à l'équation 3.21. En fonction de cette phase moyenne θ les variances des corrélations $x_- = (x_1 - x_2)/2$ et anticorrélations $x_+ = (x_1 + x_2)/2$ sont donc soit comprimées soit amplifiées. L'équation 3.21 nous donne leurs valeurs :

$$\Delta^2 x_-(\theta) = \cos^2(\theta) \frac{s}{2} + \sin^2(\theta) \frac{1}{2s} \quad (3.10)$$

$$\Delta^2 x_+(\theta) = \sin^2(\theta) \frac{s}{2} + \cos^2(\theta) \frac{1}{2s} \quad (3.11)$$

La figure 3.12 montre les variances de ces corrélations et anticorrélations, comparées au bruit du vide et de l'état thermique obtenu en ne regardant que l'un des modes de l'état EPR. Nous pouvons ainsi vérifier que la variance de ce dernier est égale à la moyenne des variances des corrélations et anticorrélations :

$$\Delta^2 x_{\text{th}} = \frac{\Delta^2 x_-(\theta) + \Delta^2 x_+(\theta)}{2} = \frac{s + 1/s}{4} \quad (3.12)$$

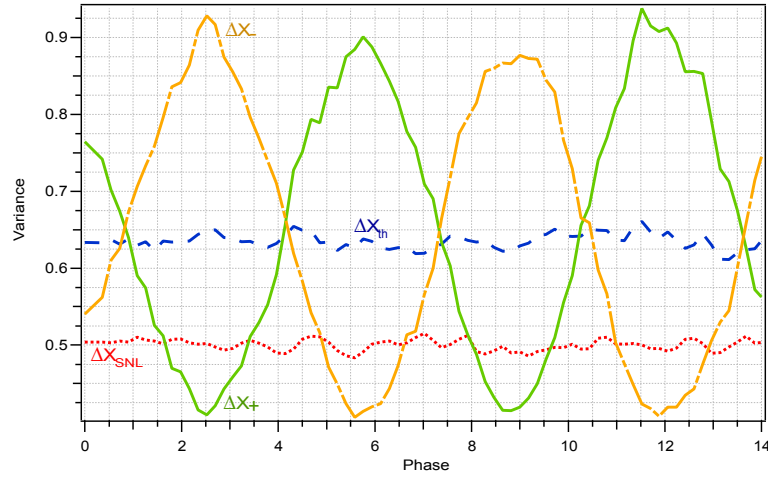


FIGURE 3.12: Observation de l'intrication des états EPR

La transformation 3.9 correspond cependant au cas idéal. Dans la réalité, la distribution par paire entre les deux modes n'est pas parfaite. Une explication, peut-être un peu simpliste, consiste à considérer que certains photons des paires paramétriques sont perdus dans d'autres modes à cause d'un mauvais recouvrement entre la pompe et le signal. Du coup on peut distinguer deux types d'amplification dans le cristal l'amplification paramétrique par paire de photon et une amplification « simple » (i.e. sans paire de photons), et non désirée, ou seul un des modes est amplifié. Ce mauvais recouvrement peut être dû à un caractère multimode du faisceau de pompe à l'entrée de l'OPA ou alors à des effets parasites présents dans le cristal, de manière similaire au GSH. De plus ces effets parasites peuvent aussi contribuer à la perte d'un des photons d'une paire d'une autre manière : le BLIIRA (*Blue Induced InfraRed Absorption*) [80] entraîne l'absorption d'environ 5% des faisceaux infrarouges. Nous pourrions donc modéliser un OPA réel par un OPA idéal de gain g suivi de deux amplificateurs, un dans chacun des modes de sortie, correspondant à des OPA dont le mode complémentaire serait perdu dans l'environnement (figure 3.13) ; comme le processus est symétrique par rapport au signal et au complémentaire nous pouvons considérer

que les deux amplificateurs parasites ont le même gain en intensité $h = \cosh^2(\gamma r)$, γ étant le rapport entre les efficacités des deux processus. D'un point de vu plus formel, et plus général, ce modèle possède l'avantage de pouvoir décrire tout état gaussien bimode [86].

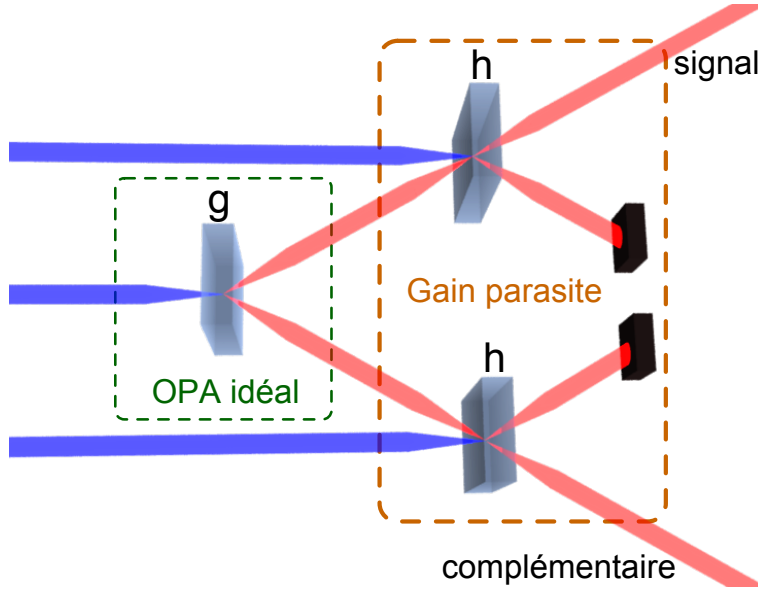


FIGURE 3.13: Modélisation de l'OPA réel

Au final le gain paramétrique vaut $G = gh = \cosh^2(r) \cosh^2(\gamma r)$, et varie de manière approximativement linéaire avec la puissance de pompe : $G \approx 1 + r^2(1 + \gamma^2) = 1 + \alpha^2 P_{\text{pompe}}$ avec $\alpha = 0,21 \text{ mW}^{-1/2}$ comme le montre la figure 3.14. Notons enfin qu'en réduisant la puissance de

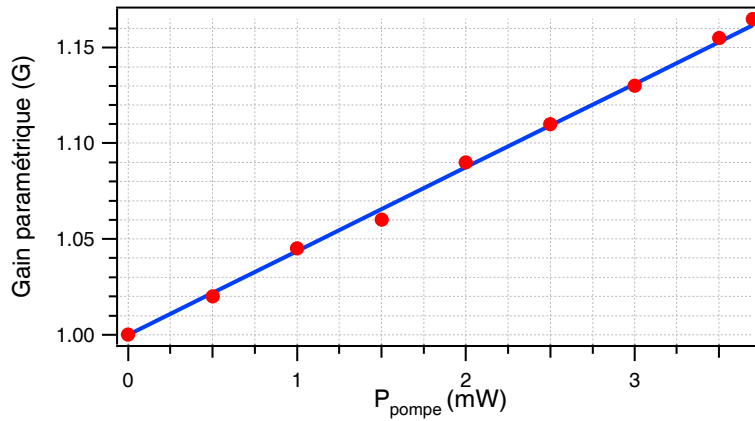


FIGURE 3.14: Évolution du gain paramétrique en fonction de la puissance de pompe pour des focalisations réduites dans les cristaux.

pompe afin d'obtenir un gain paramétrique de $G = 1,11$, le BLIIRA, qui peut facilement être estimé en tournant de 90° la polarisation de la pompe afin de détruire l'accord de phase, est réduit jusqu'à 2,5%.

3.3.3.2 Configuration dégénérée

Dans cette configuration les faisceaux signal et complémentaire sont dans le même mode, et vont donc interférer. Si on revient au cas classique, on voit que la phase du complémentaire est égale à la différence de phase entre la pompe et le signal ; le fait que l'interférence soit constructive ou destructive va donc dépendre de cette différence. Nous avons donc une amplification dépendante de la phase.

Repassons au cas quantique : l'hamiltonien se simplifie maintenant en

$$\hat{H}_d = \frac{i\hbar}{\tau_{NL}} (\hat{a}_s^2 + (\hat{a}_s^\dagger)^2) \quad (3.13)$$

ce qui nous donne la transformation suivante

$$\begin{pmatrix} x_s' \\ p_s' \end{pmatrix} = \begin{pmatrix} e^{-r} & 0 \\ 0 & e^r \end{pmatrix} \begin{pmatrix} x_s \\ p_s \end{pmatrix} \quad (3.14)$$

Cette transformation est strictement équivalente à l'opérateur compression pour un facteur de compression $s = e^{-2r}$: la quadrature x_s est comprimée d'un facteur s et la quadrature p_s amplifiée d'un facteur $1/s$. Nous avons donc un autre moyen de produire un état comprimé, et en particulier un vide comprimé, que celui présenté précédemment. Comme pour l'amplification non-dégénérée, sa décomposition dans la base de Fock s'explique elle aussi bien par l'émission des photons par paires : il y a forcément un nombre pair de photons dans le mode de sortie.

En étendant l'équation 3.14 à deux modes avec des compressions sur des quadratures pour ces deux modes, nous pouvons remarquer que, mathématiquement, il est possible d'obtenir cette relation et l'équation 3.9 l'une à partir de l'autre, en les combinant avec la transformation de la lame semi-réfléchissante (équation 3.2). Nous retrouvons donc bien les possibilités, déjà évoquées, de créer un état EPR en mélangeant deux vides comprimés sur des quadratures orthogonales et à l'inverse de produire ces derniers en recombinaison les deux modes d'un état EPR.

Nous pouvons définir les gains minimum et maximum en intensité, qui dans le cas idéal valent $g_{\min} = e^{-2r} = s$ et $g_{\max} = e^{2r} = \frac{1}{s}$, leur produit vérifiant $g_{\min}g_{\max} = 1$. Dans le cas réel nous pouvons reprendre la modélisation des imperfections utilisée pour l'amplification non-dégénérée, qui revient cette fois à placer un unique amplificateur parasite en sortie de l'amplificateur idéal. Les gains observés deviennent donc

$$\begin{aligned} g_{\min} &= e^{-2r} \cosh^2(\gamma r) = hs \\ g_{\max} &= e^{2r} \cosh^2(\gamma r) = \frac{h}{s} \end{aligned} \quad (3.15)$$

Ils peuvent être calculés expérimentalement à partir de la variance d'un état comprimé, en faisant varier la phase de la quadrature observée (figure 3.15). La variance vaut alors :

$$\Delta^2 x(\theta) = \cos^2(\theta) g_{\min} + \sin^2(\theta) g_{\max} \quad (3.16)$$

La mesure de ces gains nous permet de déduire les paramètres caractérisant l'amplification :

$$r = \frac{1}{4} \ln \left(\frac{g_{\max}}{g_{\min}} \right) \quad (3.17)$$

$$\gamma = \frac{\operatorname{arccosh}((g_{\min}g_{\max})^{1/4})}{r} \quad (3.18)$$

La 3.16(a) montre ces gains en fonction de la puissance de pompe ; ils donnent un paramètre γ à peu près constant. Le paramètre de compression montre la présence d'effets parasites,

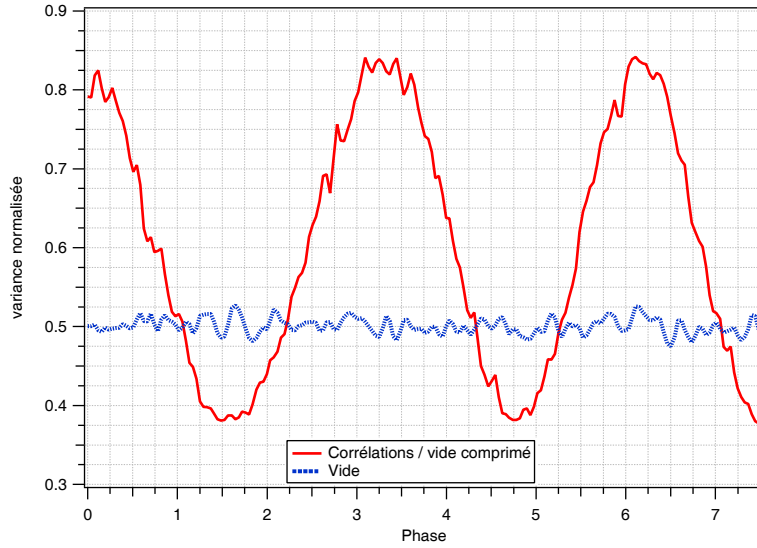


FIGURE 3.15: Variance d'un état comprimé, comparée à celle du vide.

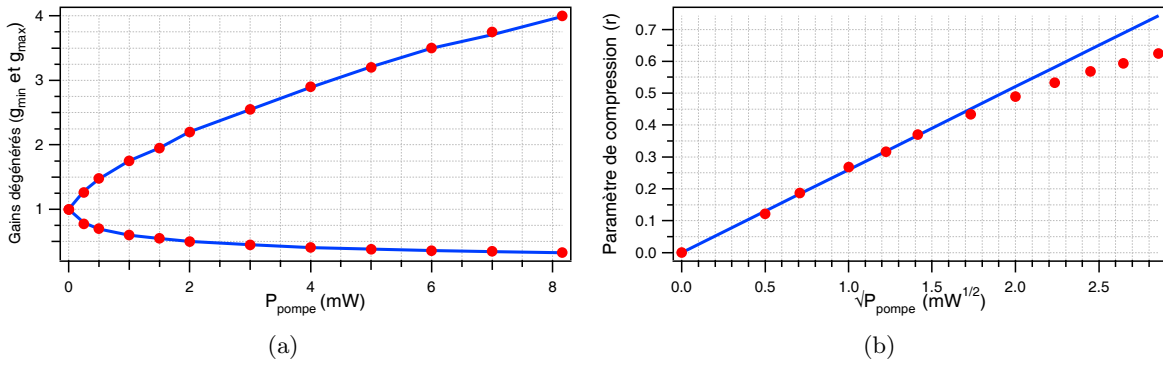


FIGURE 3.16: Paramètres de l'OPA en fonction de la puissance de pompe (a) gains minimum et maximum; les courbes sont obtenues à partir du modèle pour $\gamma = 0,6$ (b) paramètre de compression en fonction de la racine carré de la puissance; la pente de la droite vaut $0,26 \text{ mW}^{-1/2}$ et est obtenue à partir des 6 premiers points

qui semble être essentiellement de l'absorption multiphotonique, à forte puissance : il passe en dessous de la droite théorique correspondant à $r \propto \sqrt{\langle \hat{a}_p^\dagger \hat{a}_d \rangle} \propto \sqrt{P_{\text{pompe}}}$. La table 3.3 résume les caractéristiques de l'OPA pour les deux configurations avec des puissances de pompe respectivement maximale et optimisée.

Paramètre	Puissance maximale	Optimisé
Puissance de pompe	$P_{\text{pompe}} = 8,5 \text{ mW}$	$P_{\text{pompe}} \approx 3,5 \text{ mW}$
Gain paramétrique (non dégénéré)	$G = 1,44$ ($g = 1,3$)	$G = 1,05 - 1,20$ ($g \approx 1,12$)
Gain minimum (dégénéré)	$g_{\min} = 0,33$ (-4,8 dB)	$g_{\min} = 0,51$ (-2,9 dB)
Gain maximum (dégénéré)	$g_{\max} = 4,0$ (6 dB)	$g_{\max} = 2$ (3,1 dB)
Produit des gains (dégénéré)	$g_{\min} g_{\max} = 1,3$	$g_{\min} g_{\max} = 1,04$
Compression	$r = 0,53$ ($s = 0,35$)	$r \approx 0,35$ ($s \approx 0,5$)
Gain parasite	$\gamma = 0,6$ ($h \approx 1,05$)	$\gamma \approx 0,4$ ($h \approx 1,02$)
BLIIRA	$\approx 5\%$ d'absorption	$\approx 2,5\%$ d'absorption

TABLE 3.3: Caractéristiques de l'OPA

En pratique pour effectuer les réglages on commence d'abord par la configuration dégénérée (il est en effet bien plus facile de superposer le faisceau pompe et le faisceau sonde s'ils sont colinéaires). Une fois cette configuration optimisée un miroir placé sur une platine de translation permet de déplacer le faisceau sonde et ainsi de passer en configuration dégénérée.

3.4 Détection et mesures projectives

Maintenant que nous pouvons produire des états gaussiens et leur faire subir un certain nombre d'opérations elles-aussi gaussiennes, il ne nous reste plus que deux éléments : les opérations non gaussiennes et la mesure des états. Les premières seront effectuées par des mesures projectives utilisant des détecteurs de photons et la mesure finale par une ou deux détections homodynes.

3.4.1 La détection homodyne

3.4.1.1 Principe

Pour mesurer les quadratures du champ, il faut pouvoir avoir accès à sa nature ondulatoire, ce qui n'est pas directement possible avec une simple photodiode, qui ne peut pas suivre la fréquence élevée du rayonnement lumineux. Par contre, l'interférence entre deux faisceaux fait apparaître un terme proportionnel aux quadratures. Il faut donc trouver un moyen d'isoler ce terme. De plus la phase est une grandeur relative, il nous faudra donc une référence de phase afin de définir les quadratures considérées.

La détection homodyne répond à tous ces critères. Le principe consiste à faire interférer sur une lame séparatrice 50/50 le signal à mesurer (E_s) avec une référence appelée oscillateur local (E_{OL}). On mesure ensuite les faisceaux en sortie (E_+ et E_-) avec deux photodiodes. Le faisceau de référence est un état cohérent considéré comme classique auquel on applique un déphasage θ . Les mesures donnent alors :

$$I_{\pm} \propto I_{\text{OL}} + I_s \pm \sqrt{2I_{\text{OL}}} X_{\theta} \quad (3.19)$$

En soustrayant les deux signaux d'interférence nous obtenons donc un signal proportionnel à une quadrature précise ; de plus l'intensité de l'oscillateur local fait office d'amplificateur, permettant ainsi de mesurer des champs extrêmement faibles.

Dans la plupart des expériences d'optique quantique utilisant des détections homodynes celles-ci fonctionnent dans le domaine fréquentiel [87, 88]. Une telle technique a l'avantage de s'affranchir de nombreux bruits techniques basses fréquences, ainsi que de réduire les effets des problèmes d'équilibrage (qui sont essentiellement basses fréquences) ; malheureusement elle n'est pas compatible avec une expérience en régime impulsif. Il nous faut donc de travailler dans le domaine temporel avec un détecteur large bande.

3.4.1.2 Imperfections

Pertes et bruit électronique Les pertes peuvent être modélisées par une lame semi-réfléchissante prélevant une partie de l'état qui est ensuite perdue. En termes d'opérateur cela donne

$$\hat{a}'_s = \sqrt{\eta}\hat{a}_s + \sqrt{1-\eta}\hat{a}_0 \quad (3.20)$$

où η est l'atténuation en intensité et $\hat{a}_0$ l'opérateur du mode vide couplé à notre signal par la lame semi-réfléchissante du modèle. Ce dernier est un point très important : en optique quantique les pertes couplent le mode avec le vide, du coup en plus de l'atténuation elles entraînent un ajout de bruit (afin de conserver l'inégalité de Heisenberg) ; plus précisément les pertes ont tendance à ramener le bruit vers celui du vide.

L'efficacité homodyne peut se scinder en quatre contributions, chacune modélisée par une lame séparatrice, soit $\eta = \eta_{\text{opt}}\eta_{\text{mod}}\eta_{\text{quant}}\eta_{\text{elec}}$:

- Les pertes optiques (η_{opt}) : elles correspondent aux pertes induites par l'ensemble des éléments optiques placés après la génération de l'état (imperfections des miroirs et traitements anti-reflets, cubes séparateurs de polarisation, ...).
- La mauvaise adaptation des modes (η_{mod}) : la mauvaise adaptation des modes du signal et de l'oscillateur local (ou *mode-matching*) induit qu'une partie du signal ne va pas interférer avec l'oscillateur local ; elle ne sera donc pas vue par la détection homodyne ; ce qui est strictement équivalent à des pertes. Elle peut être mesurée en préparant les deux modes dans des états cohérents de même amplitude : le contraste des interférences entre ceux-ci vaut alors $\sqrt{\eta_{\text{mod}}}$. Pour un traitement plus complet se référer à [89].
- L'efficacité quantique des photodiodes (η_{quant}) : l'efficacité quantique d'un détecteur correspond à la probabilité qu'un photon incident crée un porteur de charge. Une fraction η_{quant} des photons va se répercuter dans le signal électrique tandis que les autres ne donneront rien ; il s'agit donc là encore d'une perte.
- Le bruit électronique (η_{elec}) : ce dernier point paraît plus étrange puisque il ne fait que rajouter du bruit sans atténuer le signal. Ce qu'il faut bien garder en tête c'est que la détection homodyne ne nous donne qu'une tension proportionnelle au signal que l'on veut mesurer, nous devons donc mettre à l'échelle ses valeurs en utilisant une référence dont la valeur est connue. Cette référence est le bruit du vide ; et lorsqu'on mesure ce bruit il faut prendre en compte la contribution du bruit électronique. Dans ces conditions nous avons deux possibilités : soit nous corrigeons le bruit du vide du bruit électronique pour effectuer la mise à l'échelle, auquel cas le bruit électronique doit effectivement être modélisé par une augmentation de la variance (convolution de la fonction de Wigner avec une gaussienne centrée en zéro et dont la variance vaut celle du bruit électronique mise à l'échelle). Soit nous utilisons la valeur non corrigée pour la mise à l'échelle et nous pouvons modéliser le bruit électronique par une perte [90]. En effet dans ce cas nous divisons le signal par une

quantité légèrement trop grande donnant ainsi l'atténuation, quant au bruit ajouté il est équivalent à celui ajouté par le vide (cf. section A.3). Nous adopterons cette méthode qui simplifie les calculs. La contribution du bruit électronique peut être calculée à partir de la variance mesurée du vide (V_{SNL}) et de celle du bruit électronique (V_e) : $\eta_{\text{elec}} = 1 - V_e/V_{\text{SNL}}$.

Les pertes peuvent être mesurées en envoyant un état cohérent connu sur la détection homodyne et en mesurant l'amplitude des oscillations observées en faisant varier la phase. La calibration de l'état cohérent envoyé est effectuée à l'aide d'un puissance-mètre : on le place avant les densités servant à réduire l'état cohérent au bon ordre de grandeur. Ensuite, ayant mesuré la transmission de ces densités on peut calculer l'amplitude de l'état cohérent vu par la détection homodyne, par un *fit* des données, après moyennage sur 400 points afin de réduire le bruit des mesures. Il ne reste alors plus qu'à tracer la courbe représentant l'amplitude mesurée sur la détection homodyne en fonction de l'amplitude de l'état envoyé et à en calculer la pente. Pour notre système l'efficacité homodyne vaut $\eta = 0,68 \pm 0,04$, réparties en $\eta_{\text{opt}} = 0,87$; $\eta_{\text{mod}} = 0,83$; $\eta_{\text{quant}} = 0,95$ et $\eta_{\text{elec}} = 0,99$.

Bruit de l'oscillateur local Si on remplace l'oscillateur local de l'équation 3.19 par un oscillateur local bruité $I_{\text{OL}} + \delta I_{\text{OL}}$, nous voyons que le signal mesuré devient proportionnel à $\sqrt{I_{\text{OL}}} + \delta I_{\text{OL}} \hat{X}_\theta$. Si les fluctuations sont faibles devant l'intensité moyenne nous pourrions donc les négliger.

Mauvais équilibrage Les termes d'intensité dans l'équation 3.19 sont bien plus grands que le terme d'interférence, et il faut donc équilibrer l'intensité des deux voies avec une très grande précision afin que les termes constants s'annulent bien. Pour analyser les effets d'un déséquilibre considérons que la lame séparatrice effectuant l'interférence aie en fait une transmission $T = \frac{1}{2} + \epsilon$; dans ce cas la différence entre les photocourants devient proportionnelle à $\sqrt{I_{\text{OL}}} \hat{X}_\theta + 2\epsilon(I_{\text{OL}} + I_s) \approx \sqrt{I_{\text{OL}}} \hat{X}_\theta + 2\epsilon I_{\text{OL}}$. En pratique $\Delta^2 \hat{X}_\theta \sim \frac{1}{2}$, ce qui implique que nous devons avoir $\epsilon \ll \frac{1}{2\sqrt{I_{\text{OL}}}}$. Dans notre cas, l'oscillateur local a une puissance d'environ 20 μW , ce qui correspond à un ordre de grandeur de 10^8 photons par impulsion : nous devons donc avoir un équilibrage meilleur que 10^{-4} .

3.4.1.3 Montage

Nous pourrions avoir besoin de mesurer des états bimodes, ce qui nécessite deux détections homodynes. Afin de contrôler la différence de phase entre les deux détections homodynes, un système composé d'une lame quart-d'onde et d'une lame demi-onde permet d'obtenir deux oscillateurs locaux ayant la même intensité et un déphasage θ déterminé. Pour cela, en partant d'un faisceau polarisé verticalement, les lames doivent être orienté suivant les angles $\theta_{\lambda/4} = \theta/2$ et $\theta_{\lambda/2} = \pi/8 + \theta/2$. Lors du réglage du déphasage, celui-ci peut être mesuré de deux façons différentes

- en envoyant deux états comprimés sur des quadratures orthogonales, obtenus en recombinant les deux modes d'une paire EPR. On mesure alors la variance des données issues de chaque détection en fonction de la phase moyenne des oscillateurs locaux (balayée à l'aide d'un piézoélectrique). Le déphasage est alors calculé grâce à un fit sinusoïdale de ces données.
- en envoyant un état cohérent que l'on partage entre les deux détections. On regarde alors la moyenne des signaux en fonction de la phase moyenne. Comme pour la mesure de l'efficacité, il convient de moyennner (sur environ 400 points) de manière à réduire le bruit de la mesure. Le déphasage est, cette fois, obtenu directement à partir du fit sinusoïdale.

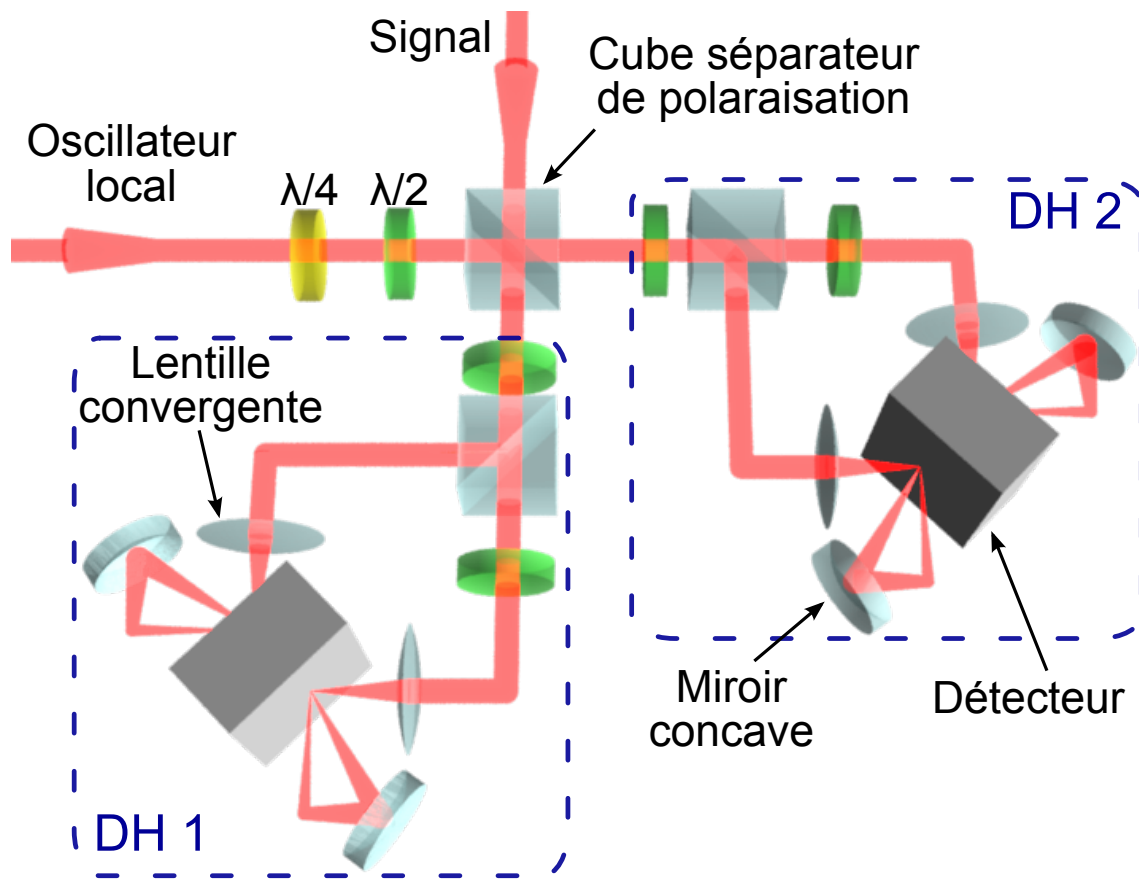


FIGURE 3.17: Montage du système à deux détections homodynes

Afin de pouvoir équilibrer précisément les détections homodynes, les interférences sont effectuées à l'aide d'un couple lame demi-onde – cube séparateur de polarisation. Afin de limiter les pertes, la polarisation du faisceau transmis est remise à la verticale et un jeu de deux miroirs par photodiodes (un plan et un concave) viennent récupérer les réflexions des faisceaux sur les photodiodes et les renvoyer sur celles-ci (ceci permet d'éliminer entre 1 et 3% de pertes).

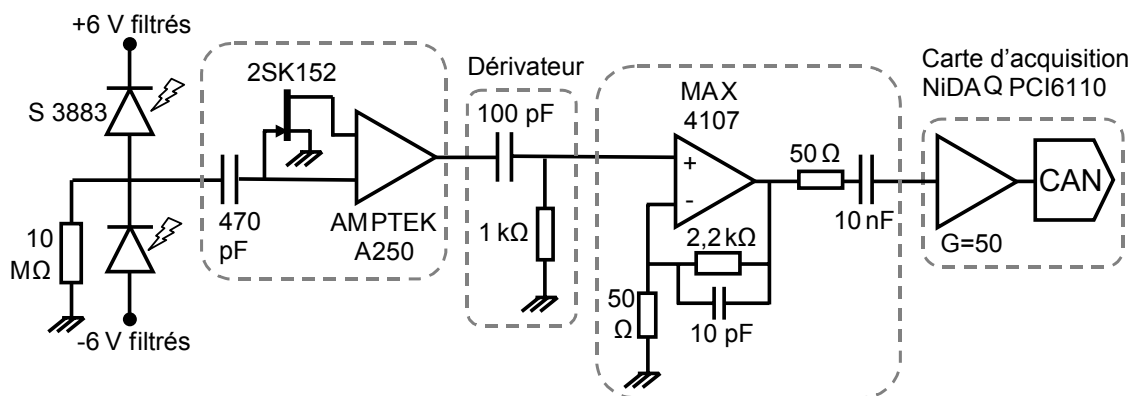


FIGURE 3.18: Schéma électronique de la détection homodyne

La détection des deux faisceaux issus de l'interférence est effectuée par des photodiodes *Hamamatsu S3883* possédant une bonne efficacité quantique (94,5%) et un faible courant d'obscurité (28 pA). Le couple utilisé est choisi parmi un lot d'une quarantaine de photodiodes de manière à ce que leurs réponses soient les plus semblables. Elles sont polarisées en inverse à ± 6 V, les tensions de polarisation étant réglables afin d'égaliser les capacités parasites. Les photocourants sont directement soustraits avant d'être amplifié (afin d'éviter un éventuel déséquilibre dans les amplifications et/ou des saturations). Le schéma électronique précis du montage électronique est montré figure 3.18. L'acquisition est effectuée par une carte *National instrument PCI-6110* permettant d'acquérir jusqu'à 5 millions de points par seconde sur 4 quatre voies avec un bruit négligeable devant celui de la détection homodyne (0,1 mV contre 2,5 mV en écart-type). L'acquisition est déclenchée par un signal TTL provenant du laser et synchronisé avec les impulsions par des lignes à retard. Celles-ci sont réglées pour que la lecture ait lieu au moment où le signal se stabilise après une ou deux oscillations initiales (dues à une légère différence du temps de réponse entre les deux voies). Enfin l'acquisition est contrôlée par des programmes dédiés à chaque expérience et écrits en C++.

La première chose à faire avant d'utiliser les détections homodynes consiste à vérifier qu'elles sont bien équilibrées. Dans ce cas la variance est linéaire en fonction de la puissance de l'oscillateur local, tandis que si la compensation est mauvaise le terme dominant est en I_{OL} (cf. précédemment) ce qui implique une variance quadratique en fonction de cette même puissance. En cas de déséquilibre, il faut d'abord équilibrer les intensités dans les deux voies puis ensuite supprimer le bruit technique du laser, visible en effectuant une transformée de Fourier du signal de la détection homodyne. Ceci peut être fait soit en réglant le *cavity dumper* (position dans la cavité et onde RF) soit en ajustant la position et la focalisation des faisceaux sur les photodiodes, permettant ainsi un réglage plus fin de l'équilibrage.

Malgré toutes ses précautions les signaux issus des détections homodynes présentent un bruit basse fréquence ainsi qu'une dérive de l'équilibrage ; nous pouvons nous affranchir de ceux-ci en effectuant une moyenne glissante sur ces signaux pour autant qu'ils aient une valeur moyenne nulle.

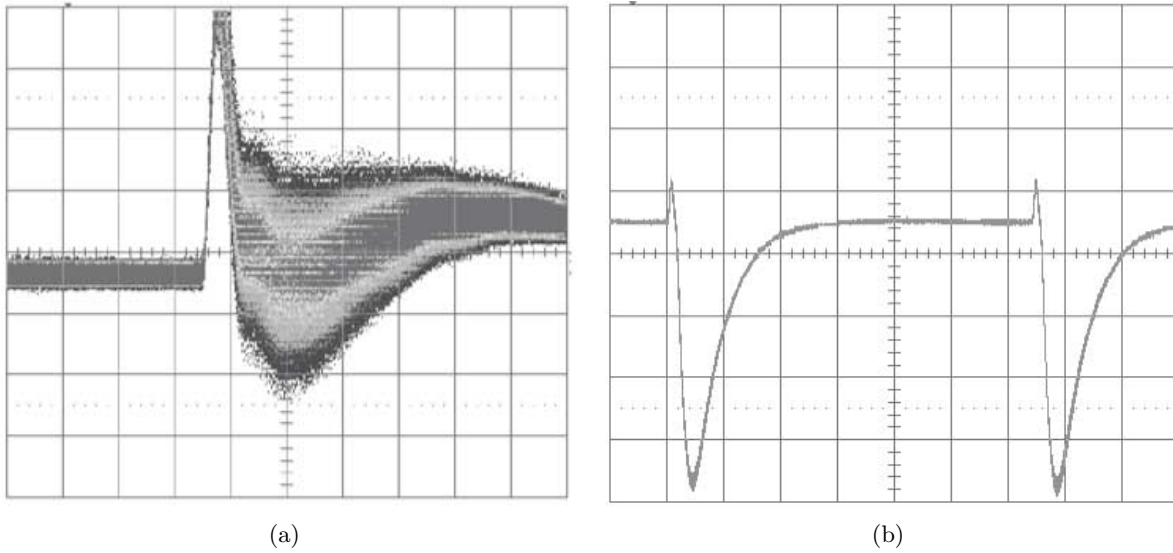


FIGURE 3.19: Signal de la détection homodyne (a) lorsque les voie sont correctement équilibrées (échelle : 50 ns/div et 50 mV/div) (b) lorsque les voie sont déséquilibrées (échelle : 200 ns/div et 500 mV/div)

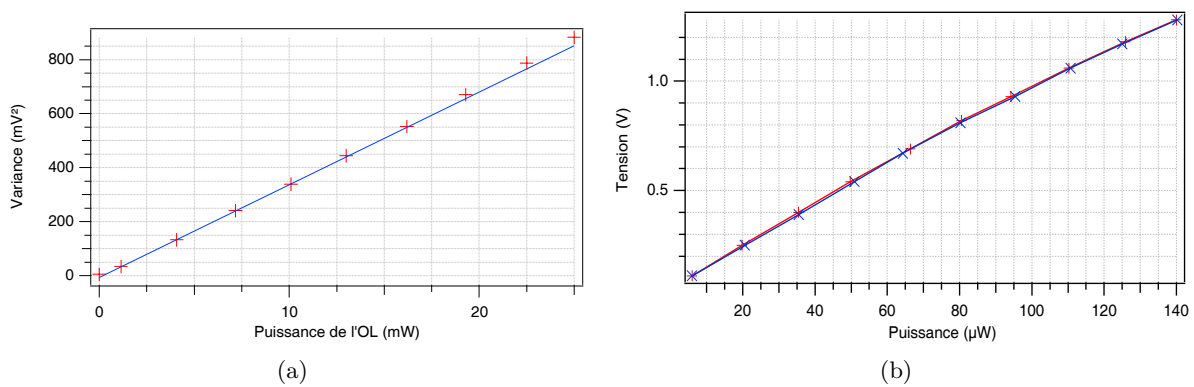


FIGURE 3.20: Préparation de la détection homodyne (a) Test de calibration de la détection homodyne (b) Réponse des photodiodes choisies

3.4.2 La photodiode à avalanche

3.4.2.1 Caractéristiques du dispositif

La photodiode à avalanche est un appareil capable de détecter des photons mais incapable de les compter. Comme son nom l'indique elle fonctionne sur un principe d'avalanche : l'absorption d'un photon crée une paire électron-trou qui est accélérée par une forte tension de polarisation inverse. Si l'électron et le trou acquièrent une énergie cinétique suffisamment grande ils vont alors pouvoir entraîner la création de nouvelles paires électron-trou à leur tour accélérées et ainsi de suite jusqu'à obtenir un courant macroscopique. Celui-ci est ensuite converti en signal TTL. L'impossibilité de compter les photons vient du fait que les avalanches produites s'auto-entretiennent, conduisant ainsi à un courant qui va augmenter jusqu'à être arrêté, à un certain courant seuil, par un dispositif spécialisé. Ce qui donne des résultats sensiblement identiques quel que soit le nombre de photons incidents.

Nous avons dit précédemment que la fluorescence paramétrique (c'est-à-dire les états sortant de l'OPA par émission spontanée) est fortement multimode ; si la détection homodyne réalise un filtrage automatique de par l'interférence avec l'oscillateur local, ce n'est pas le cas des APD qui voient tous les photons. Il faut donc filtrer les impulsions qui leur arrivent dessus. Pour cela nous utilisons une fibre monomode faisant office de filtre spatial suivie d'un réseau de diffraction et d'une fente pour le filtrage spectral. Ces filtres combinés à l'efficacité quantique des APD qui est d'environ 45% donnent de nombreuses pertes. Une mesure plus précise de celles-ci peut être effectuée de manière strictement similaire à la détection homodyne, en envoyant un état cohérent connu ; la seule différence est qu'au lieu de mesurer l'amplitude directement, on compte le nombre de clics, qui dépend de la probabilité de mesurer au moins un photon donnée par l'équation 2.64.

Les APD utilisées sont des *SPCM-AQR-13* de la société *Perkin-Elmer*, elles ont un courant d'obscurité d'environ 200 coups/s ; en regardant uniquement au moment d'arrivée des impulsions nous pouvons réduire celui-ci à moins de 10–15 coups/s. À noter qu'environ la moitié de ces coups restant viennent de réflexions parasites de l'oscillateur local qui parviennent jusqu'aux APD. Nous pouvons enfin remarquer que la détection d'un photon produit une impulsion TTL d'environ 30 ns et que le temps mort pendant lequel le détecteur est aveugle suite à une détection vaut 50 ns. Ces deux temps sont très courts par rapport au délai de 1,3 μ s entre chaque impulsion et donc ne nous posent aucun problème.

3.4.2.2 Mesures projectives et conditionnement

Les projections correspondant à cette mesure sont donc, dans le cas idéal, $|0\rangle\langle 0|$ si aucun photon n'est détecté et $\hat{1} - |0\rangle\langle 0|$ si des photons sont détectés. Les mesures des APD étant destructives, le mode sur lequel est effectuée la projection est ensuite perdu. D'un point de vue de la modélisation mathématique cette caractéristique est plutôt pratique puisqu'elle permet d'utiliser la formule de recouvrement (équation 2.40). Dans notre expérience les fortes pertes sur la voie APD font que la projection sur le vide est très peu fiable ; le projecteur sur la présence de photon est par contre inchangé. Les coups parasites (coup d'obscurité et mauvais filtrage) ont évidemment l'effet inverse : ils laissent inchangé la projection sur le vide et modifient celle sur la présence de photon. Cet effet est cependant bien moins important dans notre dispositif.

Bien que ne pouvant pas compter les photons en soi il est possible de les utiliser pour créer un dispositif ayant des capacités limitées de comptage, en utilisant plusieurs APD et un multiplexage spatial : le signal est divisé entre plusieurs faisceaux qui vont chacun sur une APD différente. Étant donné nos pertes ce dispositif est, encore une fois, surtout utile pour projeter sur des états

contenant au moins n photons où n est le nombre d'APD utilisées [91, 92].

Outre ces projections, les détecteurs et compteurs de photons peuvent aussi être utilisés pour effectuer des opérations de *dégaussification* consistant à soustraire un [101, 94] ou plusieurs photons [95] à un état. Ces opérations sont réalisées en prélevant une partie du faisceau grâce à une lame semi-réfléchissante et en comptant le nombre de photons dans la partie prélevée. Dans le cas idéal, le résultat du comptage donne le nombre de photons soustraits à l'état ; dans le cas réel avec des APD, en plus des considérations déjà évoquées, nous devons définir la réflexion de la lame de manière à ce que la probabilité de prélever plus de photons qu'on ne le veut soit négligeable.

3.5 Production de photons uniques

L'étude des photons uniques en termes de variables continues est assez récente [96, 97] ; elle a néanmoins bien évolué puisqu'elle peut maintenant être réalisée avec de forts taux de répétition [280] permettant ainsi de les utiliser comme une ressource pour d'autres expériences ou pour ajuster le dispositif expérimental.

Ce chapitre a pour but de présenter la génération de photons uniques avec notre dispositif, qui pourront ensuite servir dans d'autres protocoles ou plus généralement afin d'optimiser tous les réglages nécessaires pour faire fonctionner correctement le dispositif expérimental. Nous en profiterons pour introduire les techniques utilisées pour la modélisation des expériences présentées dans ce manuscrit.

3.5.1 Production

Le dispositif est relativement simple (figure 3.21). Le mode complémentaire est injecté dans une fibre monomode de 2 m de long qui sélectionne le mode spatial TEM_{00} ; l'efficacité de couplage varie entre 73 et 83%. Il est ensuite envoyé sur un réseau ayant une efficacité de diffraction de 90%, suivie d'une fente réglable placée au foyer d'une lentille de focale 100 mm afin de sélectionner le mode spectral voulu. La résolution de ce système correspond à un déplacement de la fenêtre spectrale de 7 nm pour un décalage de la fente de 1 mm. La largeur de la fente est réglée pour laisser passer 30% du faisceau issu de la sonde servant aux réglages (un petit peu plus de 1 nm de largeur spectrale). Le mode signal est quant à lui envoyé sur la détection homodyne où il interfère avec l'oscillateur local qui est créé en prélevant une petite partie du faisceau laser avant les cristaux non linéaires. La superposition des deux faisceaux est réglée en regardant les interférences classiques entre l'oscillateur local et la sonde, leur contraste vaut typiquement 92%.

Le taux de production des photons uniques varie entre 5 000 et 10 000 par seconde. Les données sont prises sur 5 000 points. Le photon unique étant invariant par rotation de la phase nous n'avons pas besoin de contrôler celle-ci ni de mesurer plusieurs quadratures.

3.5.2 Modélisation

Les deux chapitres précédents nous ont permis de nous constituer une « boîte à outils » contenant tout ce dont nous avons besoin pour modéliser les expériences et ainsi calculer les formules analytiques des états produits. Pour cela nous allons partir de l'expression de l'état de départ, ajouter les divers imperfections dues aux opérations gaussiennes puis appliquer l'opération — non-gaussienne — de conditionnement ainsi que les imperfections liées à celles-ci.

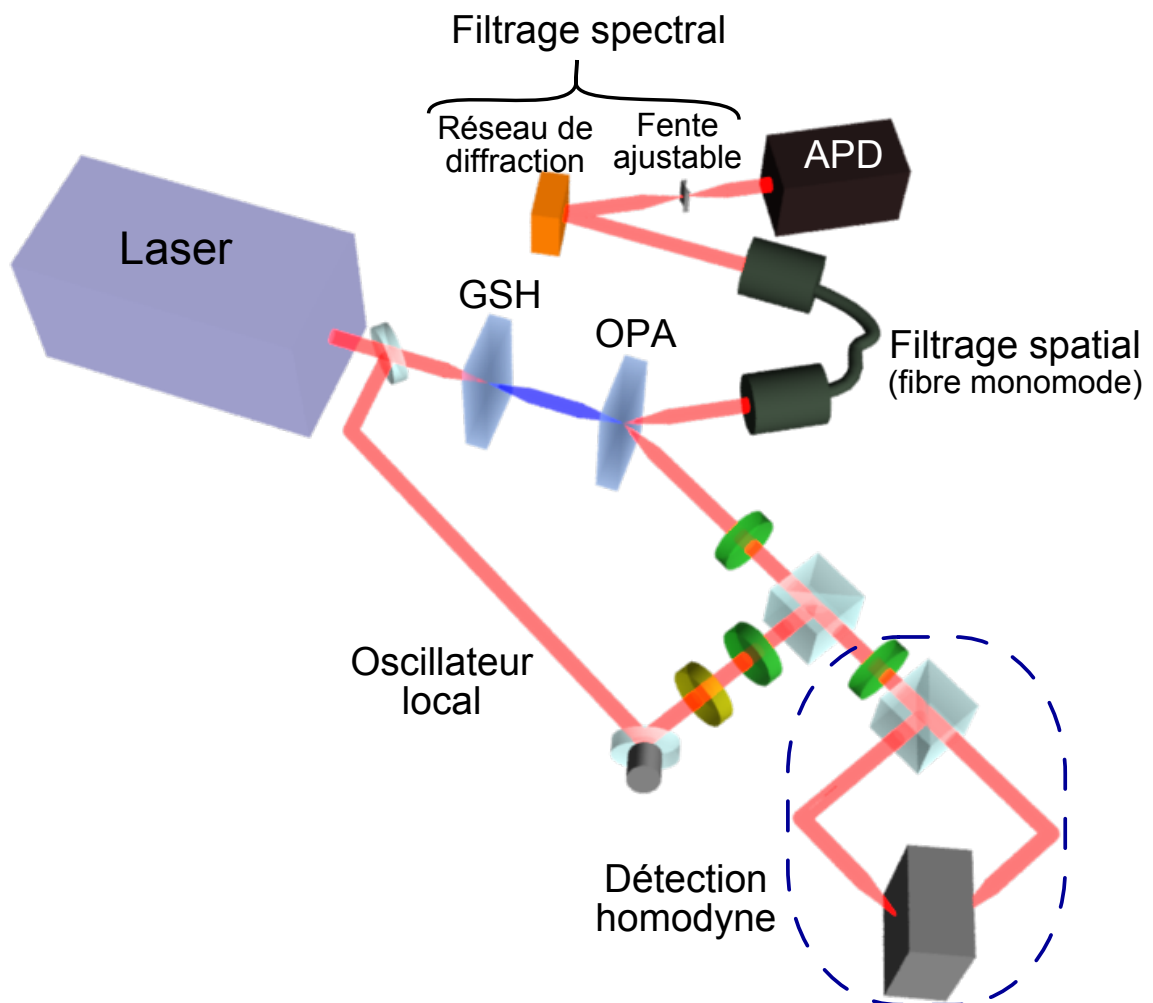


FIGURE 3.21: *Protocole expérimental pour le production de photons uniques*

3.5.2.1 Amplification paramétrique

Partons d'un état EPR qui serait produit par un OPA idéal.

$$W_{\text{EPR}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2} e^{-\frac{(x_1-x_2)^2}{2s} - \frac{(x_1+x_2)^2}{2/s} - \frac{(p_1-p_2)^2}{2/s} - \frac{(p_1+p_2)^2}{2s}} \quad (3.21)$$

Nous avons vu que l'OPA réel peut être modélisé par un OPA idéal suivi de deux autres amplificateurs dont les modes complémentaires sont perdus dans l'environnement. Pour obtenir l'état réellement produit il nous suffit donc de coupler les deux modes de notre état idéal à deux vides suivant la transformation

$$\begin{pmatrix} \sqrt{h} & 0 & -\sqrt{h-1} & 0 \\ 0 & \sqrt{h} & 0 & \sqrt{h-1} \\ -\sqrt{h-1} & 0 & \sqrt{h} & 0 \\ 0 & \sqrt{h-1} & 0 & \sqrt{h} \end{pmatrix} \quad (3.22)$$

puis à tracer sur ces deux modes de vide. Ceci nous donne, d'après les formules de section A.2, la fonction de Wigner suivante

$$W_{\text{OPA}}(x_1, p_1, x_2, p_2) = \frac{e^{-\frac{(x_1-x_2)^2}{2(hs+h-1)} - \frac{(x_1+x_2)^2}{2(\frac{h}{s}+h-1)} - \frac{(p_1-p_2)^2}{2(\frac{h}{s}+h-1)} - \frac{(p_1+p_2)^2}{2(hs+h-1)}}}{\pi(hs+h-1)(\frac{h}{s}+h-1)} \quad (3.23)$$

3.5.2.2 Pertes homodynes et APD

Une fois l'état produit, l'un des faisceaux sera envoyé sur une détection homodyne et l'autre sur une APD ; ils vont donc subir des pertes différentes. En fait il n'est pas nécessaire de prendre en compte les pertes de la voie APD puisqu'elles ne modifient pas la projection. Cependant si nous voulons simplifier les calculs en utilisant les mêmes formules que précédemment alors nous devons avoir les mêmes pertes dans les deux voies. Pour cela nous pouvons profiter du fait que l'efficacité de la voie APD μ est bien plus faible que celle de la détection homodyne η , pour la décomposer en une première transmission égale à l'efficacité homodyne et une deuxième valant $\nu = \mu/\eta$ et qui ne sera pas prise en compte.

Maintenant que nous devons appliquer la même transformation aux deux modes de notre état nous pouvons en calculer le résultat. Afin de simplifier les notations nommons a et b les paramètres issus de la transformation, c'est-à-dire

$$a = \eta(hs + h - 1) + 1 - \eta \quad (3.24)$$

$$b = \eta\left(\frac{h}{s} + h - 1\right) + 1 - \eta \quad (3.25)$$

La fonction de Wigner avec les pertes vaut donc

$$W_{\eta}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 ab} e^{-\frac{(x_1-x_2)^2}{2a} - \frac{(p_1-p_2)^2}{2b} - \frac{(x_1+x_2)^2}{2b} - \frac{(p_1+p_2)^2}{2a}} \quad (3.26)$$

3.5.2.3 Conditionnement

Étant donné notre dispositif (détecteur incapable de compter et faible efficacité APD due aux filtres et à l'efficacité quantique des APD), pour produire un photon unique avec une bonne qualité nous devons faire en sorte que la probabilité d'en avoir plus d'un dans chaque faisceau

de la paire EPR soit faible. Dans ce cas il devient extrêmement peu probable que deux photons ou plus parviennent jusqu'à l'APD. Nous pouvons alors remarquer que la modélisation des pertes par une lame séparatrice nous donne exactement le protocole de soustraction de photon évoqué plus haut dans lequel le mode soustrait serait ensuite perdu (figure 3.22). L'opération

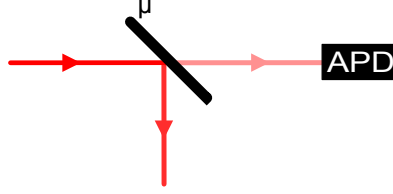


FIGURE 3.22: Détection par une APD avec de fortes pertes.

à effectuer correspond donc à $\text{Tr}_2 \left(\hat{a}_2 \hat{\rho} \hat{a}_2^\dagger \right)$ suivi d'une renormalisation puisque les opérateurs $\hat{a}$ et $\hat{a}^\dagger$ ne sont pas unitaires. Or lorsque l'on effectue la trace d'un produit on peut remplacer celui-ci par une de ses permutation circulaires, ce qui donne $\text{Tr}_2 \left(\hat{a}_2 \hat{\rho} \hat{a}_2^\dagger \right) = \text{Tr}_2 \left(\hat{\rho} \hat{a}_2^\dagger \hat{a}_2 \right) = \langle \hat{n}_2 \rangle_2$. Mathématiquement cette opération est donc équivalente à mesurer le nombre de photon dans le mode de conditionnement. On peut voir cela de la façon suivante : la mesure va décomposer l'état suivant le nombre de photons dans le mode de conditionnement et donner, dans le cas général, un poids égal à ce nombre à chaque contribution. Avec notre approximation cela revient à éliminer la contribution correspondant à une absence de photon et ne garder que celle où il y en a un de détecté.

L'expression de l'opérateur nombre en termes de quadratures est $\hat{n} = (\hat{X}^2 + \hat{P}^2 - 1)/2$; en utilisant l'équation 2.37 nous pouvons donc calculer sa fonction de Wigner :

$$W_n(x, p) = \frac{x^2 + p^2 - 1}{4\pi} \quad (3.27)$$

ce qui nous permet ensuite de calculer la fonction de Wigner obtenue lors d'un bon conditionnement :

$$W_{\text{cond}}(x_1, p_1) = \frac{\iint W_\eta(x_1, p_1, x_2, p_2) W_n(x_2, p_2) dx_2 dp_2}{\iiint W_\eta(x_1, p_1, x_2, p_2) W_n(x_2, p_2) dx_1 dp_1 dx_2 dp_2} \quad (3.28)$$

Il ne nous reste plus qu'à prendre en compte les coups parasites. Nous les modéliserons en considérant que le résultat est un mélange statistique entre l'état issu du bon conditionnement et l'état non conditionné (correspondant donc à ce qui est produit lors d'un mauvais conditionnement, c'est-à-dire un état thermique). Nous appellerons ξ la probabilité que le conditionnement soit bien issu du bon mode ; celle-ci est en quelque sorte la « pureté modale » de notre faisceau de conditionnement après filtrage. Nous avons ainsi

$$W_{\text{phot}}(x, p) = \xi W_{\text{cond}}(x, p) + (1 - \xi) W_{\text{th}}(x, p) \quad (3.29)$$

Au final nous obtenons donc

$$W_{\text{phot}} = \frac{1}{\pi\sigma} \left(1 - \delta + \delta \frac{x^2 + p^2}{\sigma^2} \right) e^{-\frac{x^2 + p^2}{\sigma^2}} \quad (3.30)$$

avec

$$\delta = \frac{2\xi\eta h^2 g (g - 1)}{\sigma^2 (hg - 1)} \quad (3.31)$$

$$\sigma = 2\eta (hg - 1) + 1 \quad (3.32)$$

Le calcul avec les vrais projecteurs donne des résultats très similaires [82] mais est bien plus difficilement exploitable numériquement car il fait intervenir une différence de deux gaussiennes qui pour les valeurs expérimentales des paramètres devient une petite différence entre deux grands nombres.

3.5.3 Utilisations

3.5.3.1 Taux de production et matrice densité

Afin de voir les contributions des différents états de Fock dans notre mélange réellement produit il est plus pratique de revenir à la matrice densité. Ses éléments peuvent être calculés à partir de l'équation 2.49, ce qui nous donne

$$\langle m | \hat{\rho}_1 | n \rangle = 2 \frac{(\sigma^2 - 1)^{n-1}}{(\sigma^2 + 1)^{n+2}} ((1 - \delta)\sigma^4 + \delta(1 + 2n)\sigma^2 - 1) \delta_{m,n} \quad (3.33)$$

nous pouvons remarquer que l'opérateur densité est diagonal : l'état réellement produit est donc un mélange d'états de Fock (le passage en coordonnées polaires de l'équation 2.49 montre que c'est d'ailleurs le cas pour tout état indépendant de la phase). Nous pouvons alors calculer la contribution de ceux-ci à l'état produit (figure 3.23(a)), ce qui nous sera utile pour les expériences utilisant des photons unique.

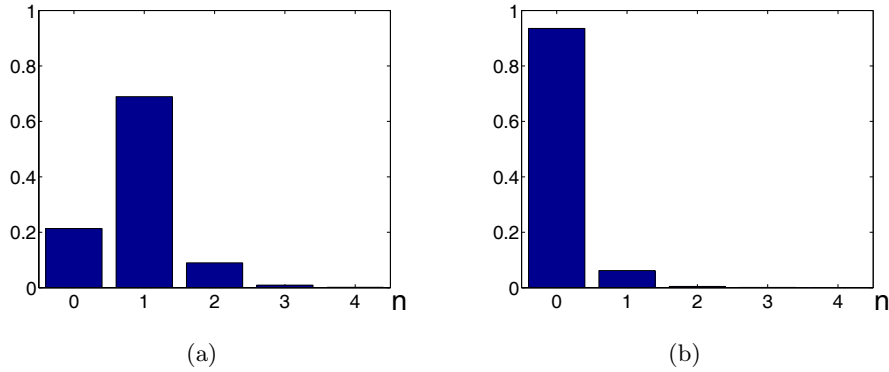


FIGURE 3.23: Coefficients diagonaux des matrices densité (a) photon unique (b) état thermique

Formellement, si nous faisons tendre σ vers 0, nous obtenons un mélange de vide et de photon unique. σ est donc relié aux contributions des états de plus de 1 photon, et correspond donc à l'élargissement de la fonction de Wigner. Ce n'est en réalité pas étonnant puisque qu'il correspond à l'écart-type de l'état thermique obtenu sans conditionnement, et que celui-ci augmente avec le nombre de photons de la paire EPR. Il va essentiellement dépendre de la puissance de pompe, et peut donc être réglé. Le paramètre δ , quant à lui, peut être relié au mauvais conditionnement et aux pertes du photons. Il dépend donc des imperfections du dispositif et est directement responsable de la qualité de l'état produit. Il varie entre 0, correspondant à un état non-conditionné, et 2, pour le cas idéal.

Par conséquent, en prenant $\delta = 0$ dans l'équation 3.33, nous obtenons la matrice densité d'un état thermique :

$$\langle m | \hat{\rho}_{\text{th}} | n \rangle = 2 \frac{(\sigma^2 - 1)^n}{(\sigma^2 + 1)^{n+1}} \delta_{m,n} \quad (3.34)$$

Cette dernière permet de calculer le taux de production des photons uniques : l'APD reçoit en effet un état thermique, en remplaçant l'efficacité homodyne η par l'efficacité de la voie APD μ , nous obtenons donc l'état vu par celle-ci. Ce qui nous donne le taux de production : $P_{\text{phot}} = 1 - \langle 0 | \hat{\rho}_{\text{th}} | 0 \rangle$

3.5.3.2 Caractérisation des imperfections

Nous l'avons dit en introduction de ce chapitre, la production de photons uniques peut-être utilisée pour ajuster les réglages du dispositif expérimental. En plus de la cadence de production déjà évoquée, les photons uniques présentent en effet deux autres propriétés qui en font un bon état pour quantifier l'importance des imperfections :

- Sa réalisation fait appel à tous éléments présents dans le dispositif, donnant ainsi un aperçu de son fonctionnement.
- Étant un état non-gaussien, sa fonction de Wigner présente une négativité quand il est pur, et celle-ci est en plus très sensible aux imperfections. Elle constitue donc un bon critère pour juger de la qualité de l'état produit et par là même du dispositif.

Nous avons donc besoin d'une grandeur quantifiant cette négativité et qui soit facile à retrouver à partir des données. En regardant l'expression analytique de la fonction de Wigner on remarque qu'elle est négative uniquement si $\delta > 1$. De plus, la négativité croît avec δ , le cas idéal étant $\delta = 2$. C'est donc un bon paramètre pour juger de la qualité d'un état quantique.

Les deux paramètres δ et σ sont d'autant plus pratiques, qu'ils ne nécessitent pas d'effectuer une tomographie complète pour obtenir leur valeur ; ils peuvent en effet être déterminés à partir des moments des distributions de probabilités mesurées (on rappelle que les quadratures ont une moyenne nulle) :

$$\sigma = 2 \langle x^2 \rangle \left(1 - \sqrt{1 - \frac{\langle x^4 \rangle}{3 \langle x^2 \rangle}} \right) \quad (3.35)$$

$$\delta = 2 \langle x^2 \rangle \left(1 - \sqrt{\frac{4}{3} - \frac{4 \langle x^4 \rangle}{9 \langle x^2 \rangle}} \right) \quad (3.36)$$

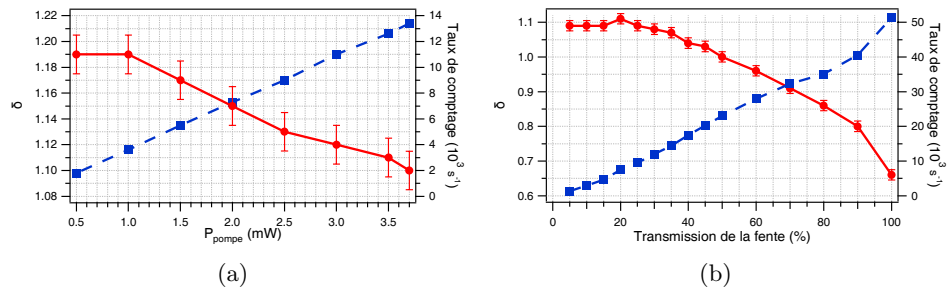


FIGURE 3.24: Influence des degrés de liberté sur les états produits (les disques correspondent au paramètre δ et les carrés au taux de comptage) (a) Puissance de pompe de l'OPA (pour $T_{\text{fente}} = 30\%$) (b) Transmission de la fente utilisée pour le filtrage spectral (pour $P_{\text{pompe}} = 3,7 \text{ mW}$)

Ceci nous donne donc un excellent moyen de caractériser l'impact des divers degrés de liberté du dispositif sur les performances. Les deux principaux sont la puissance de pompe de l'OPA et la transmission de la fente placée devant l'APD ; ils ont comme point commun d'entraîner un compromis entre qualité de l'état en sortie et taux de comptage, illustré par la figure 3.24. Ainsi si

l'on augmente la puissance du faisceau pompe qui alimente l'OPA, le taux d'émission spontanée dans les faisceaux signal et complémentaire augmente. Comme on conditionne les mesures sur les photons détectés par l'APD, cela augmente le taux de comptage des événements. Par contre, lorsque la puissance de pompe augmente il y a plus souvent des états à plus d'un photon, qui ne sont pas voulus. De plus, la distorsion du mode produit dans l'OPA est aussi augmentée. Tout ceci diminue donc la qualité des états produits. De même augmenter la transmission de la fente diminue le filtrage et donc la qualité des états ; toutefois plus on filtre moins il y a de photons qui parviennent à l'APD et donc plus le taux de comptage est bas.

Enfin l'un des plus grands défis de l'information quantique est que les états quantiques sont extrêmement sensibles aux imperfections. Or le bon fonctionnement de notre dispositif, comme beaucoup d'autres, demande sans cesse de trouver le meilleur compromis entre les différentes imperfections. Apporter une amélioration sur un point précis peut en dégrader plusieurs autres et au final empirer les performances globales. Cette caractéristique peut s'avérer déroutante puisqu'une optimisation d'un élément semblant bonne lorsqu'on ne fait attention qu'à celui-ci peut s'avérer préjudiciable quand on en revient au système complet. De tels paramètres permettant de juger rapidement de l'impact global des réglages sur les états produit se trouvent donc être des atouts capitaux non seulement pour l'optimisation initiale du système mais aussi pour contrôler le déroulement d'une expérience notamment en permettant de vérifier la stabilité du système et le cas échéant de compenser les dérives. Bien que cela fasse gagner beaucoup de temps, il ne faut pas non plus s'attendre à des miracles : la mesure du paramètre δ nous indique la présence de problèmes mais ne nous dit pas leur nature exacte et encore moins comment les régler ! En résumé cette méthode si elle ne nous épargne pas les longs mois de réglages, n'en est pas moins un outil capital aussi bien pour la modification et l'amélioration du dispositif que pour la réalisation des expériences.

3.6 Conclusion

Au final nous possédons un système polyvalent permettant de réaliser de nombreuses expériences d'optique quantique à variables continues utilisant des états non-gaussiens. Ce dispositif n'a cessé d'évoluer depuis les thèses de Jérôme Wenger [81, 99, 100, 101] et d'Alexei Ourjoumtsev [82, 91, 94, 92, 102] et des solutions sont constamment envisagées afin d'apporter de nouvelles améliorations comme le montrera la Partie III.

Le principal ajout à ce dispositif effectué et utilisé lors de ce travail de thèse est sans conteste l'insertion d'états cohérents pour les expériences. Bien que pouvant paraître anodin cette introduction n'est pas sans effet : le fonctionnement des expériences, et dans une certaine mesure leur analyse, était jusqu'ici adapté à des états de valeur moyenne nulle qui présentent un certain nombre d'avantages (utilisation de moyenne glissantes, symétries supplémentaires, ...). De plus ils introduisent aussi une nouvelle phase que nous devons contrôler.

Nous disposons donc à présent d'un ensemble d'outils aussi bien expérimentaux que théoriques pour la réalisation d'expériences d'optique quantique à variables continues. Nous savons générer aussi bien des états gaussiens que non gaussiens ainsi que les opérations, elles aussi gaussienne ou non, permettant de les manipuler. Nous avons aussi tous les outils pour mesurer et caractériser les états finaux de même que pour modéliser toutes ces opérations. Ces modélisations nous servent aussi bien pour comparer prédictions et résultats expérimentaux, permettant ainsi de comprendre l'origine des imperfections des états de sorties, que pour la préparation des expériences, nous indiquant ce que nous devons nous attendre à mesurer et quels sont les paramètres expérimentaux qui influent le plus sur la qualité du protocole et donc sur lesquels nous devrions

le plus nous attarder.

Enfin nous possédons un moyen de quantifier la qualité des réglages et l'importance des imperfections, ce qui pourra être bien utile lors des expériences pour les réglages initiaux ainsi que les optimisations nécessaires entre les prises de données afin de compenser les dérives expérimentales. Tout ceci s'est révélé très utile pour les diverses modifications et améliorations que l'on est amené à faire tout au long d'une thèse, à commencer par le remontage du dispositif qui à eu lieu au début de ce travail de thèse suite au déménagement du laboratoire.

Deuxième partie

Résultats expérimentaux

Chapitre 4

Génération de superpositions non locales d'états cohérents

Sommaire

4.1 Introduction	75
4.1.1 L'intrication	75
4.1.1.1 L'intrication et l'information quantique	75
4.1.1.2 Mesure de l'intrication	77
4.1.1.3 États de Bell	78
4.1.2 Principe de l'expérience	79
4.2 Réalisation expérimentale	81
4.2.1 Montage	81
4.2.2 Modélisation	82
4.2.3 Résultats	84
4.3 Comparaison avec un état de Bell discret	86
4.3.1 Montage	86
4.3.2 Modélisation	87
4.3.3 Résultats	87
4.4 Discussion	89
4.4.1 Communications à longues distances	89
4.4.2 Calcul quantique et systèmes hybrides	90
4.4.3 Conclusion	94

4.1 Introduction

4.1.1 L'intrication

4.1.1.1 L'intrication et l'information quantique

Nous pouvons dégager quatre particularités principales à la physique quantique : la dualité onde corpuscule, le principe de superposition, l'intrication et l'incertitude sur les mesures. Nous avons déjà évoqués les rôles importants tenus par la superposition d'états et le principe d'incertitude dans l'information quantique. La dualité onde-corpuscule est au cœur même de notre

approche qui marie les deux visions, permettant ainsi la réalisation de certaines tâches de l'information quantique impossible avec les états gaussiens d'une approche purement ondulatoire. Mais quel est celui de l'intrication ? Jusqu'ici nous l'avons uniquement utilisé pour produire des états de Fock mais nous n'avons par encore mentionné ses liens avec l'information quantique.

Communication quantique Les systèmes de communication quantique mis au point actuellement, et même commercialisés pour certains, n'utilisent pas l'intrication. Il leur suffit en effet de préparer un état quantique bien défini et de le transmettre. Mais les états quantiques sont très fragiles et deviennent donc rapidement un mélange statistique lors des transmissions, ce qui implique que ces systèmes ont des portées limitées (au mieux la centaine de kilomètre [103]).

Afin de palier à ce problème, il faudrait des *répéteurs quantiques* qui permettraient de transférer les états quantiques sur de longues distances avec une dégradation la plus faible possible. Leur principe repose sur la téléportation quantique [104, 52], qui nécessite de l'intrication. Celle-ci consiste à partager un état intriqué entre les deux interlocuteurs, Alice et Bob. Alice va ensuite effectuer une mesure conjointe entre sa « partie » de l'état intriqué et l'état qu'elle veut transmettre. Suivant le résultat de cette mesure, transmis par un canal classique, Bob va pouvoir modifier l'autre partie de l'état intriqué afin qu'elle reproduise l'état à transmettre [105, 106].

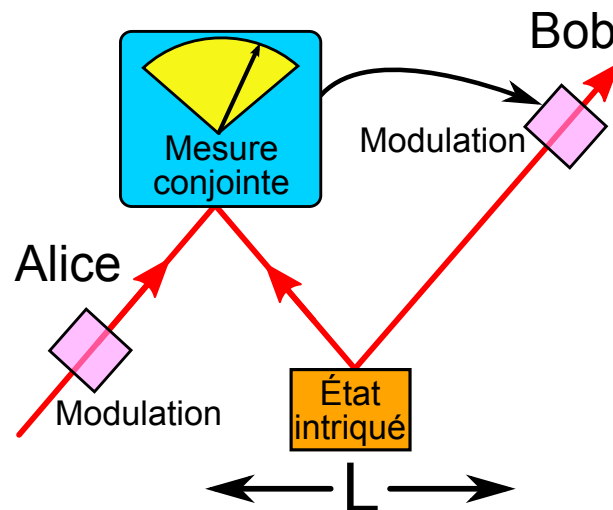


FIGURE 4.1: Téléportation quantique

En pratique l'intrication est elle aussi dégradée, conduisant à une détérioration de l'état téléporté, mais cette dégradation peut être contrée par un ensemble de techniques. Les répéteurs quantiques sont des intermédiaires qui, d'une manière analogue aux répéteurs classiques (qui amplifient le signal afin de compenser les pertes), vont permettre d'augmenter encore la portée de la transmission. Commençons donc par introduire l'un de ces intermédiaires, que l'on nommera Charlie. Au lieu de partager un état intriqué entre Alice et Bob, nous allons en partager un premier entre Alice et Charlie et un second entre Charlie et Bob. En effectuant une mesure conjointe entre les deux parties qu'il possède Charlie va réaliser un *transfert d'intrication* (ou *entanglement swapping*) : l'intrication qu'il possédait avec chacun de ses partenaires et transférée entre ceux-ci [107, 108]. Alice et Bob, qui avaient initialement des états indépendants, partagent alors un état intriqué. Il est bien entendu possible d'étendre ce principe à autant d'intermédiaires que l'on veut. Si la longueur de ces tronçons est correctement choisie, le temps pour effectuer une

communication passe d'une croissance exponentielle en fonction de la distance à une croissance polynomiale [109, 110]. Maintenant que nous avons séparé notre canal de transmission en tronçons ayant des pertes raisonnables, il convient, tout comme pour les communications classiques, de les corriger. Ceci est réalisé par un processus appelé *distillation d'intrication* [111, 112, 113] qui transforme un grand nombre d'états faiblement intriqués en un petit nombre d'états fortement intriqués : l'intrication est concentrée dans un sous-ensemble des états de départ [114, 102, 115]. Ce processus peut être réalisé en utilisant uniquement des opérations locales et des communications classiques (LOCC). Enfin, au moins une partie de ces opérations ne réussissent pas à tous les coups : c'est le cas de la mesure conjointe pour les variables discrètes et de la distillation d'intrication (qui nécessite des états et/ou des opérations non gaussiennes) pour le cas continu. Il est donc indispensable de pouvoir stocker les états quantiques [116, 117, 118, 119] correspondant aux réussites si l'on ne veut pas avoir un taux de succès ridiculement faible, et c'est actuellement le point qui présente le plus de difficultés.

L'ordinateur quantique L'intrication est aussi liée au calcul quantique : en effet les portes à plusieurs qubit permettent d'intriquer des qubits indépendants. Mais elle peut aussi y avoir un rôle plus direct. Nous savons que ces portes sont très difficiles à réaliser avec la lumière, et une idée astucieuse consiste alors à inverser les relations entre ces portes et l'intrication : au lieu que les portes créent l'intrication nous pouvons utiliser cette intrication pour réaliser les portes. L'idée consiste à partir d'un très grand état intriqué, appelé état *cluster*, et à réaliser les portes à l'aide de mesures projectives correctement choisies en fonction du calcul désiré. Cette technique est nommée *one-way quantum computing* (car les mesures projectives sont irréversibles contrairement aux opérations unitaires normalement utilisées) [120].

4.1.1.2 Mesure de l'intrication

Un état est intriqué s'il n'est pas séparable, c'est-à-dire qu'on ne peut pas décrire séparément les objets qui le composent. Une définition plus stricte pour les états bipartites est que deux systèmes sont intriqués si leurs corrélations ne peuvent pas être obtenues par des LOCC. Il est bien plus difficile de définir l'intrication pour des états contenant plus de deux objets, aussi nous nous restreindrons au cas bipartite.

Pour un état pur l'intrication peut-être quantifiée par l'entropie après réduction à un seul mode. C'est-à-dire, pour un état $\hat{\rho}_{AB}$ avec $\hat{\rho}_A = \text{Tr}_B(\hat{\rho}_{AB})$ et $\hat{\rho}_B = \text{Tr}_A(\hat{\rho}_{AB})$:

$$E = -\text{Tr}(\hat{\rho}_A \log_2(\hat{\rho}_A)) = -\text{Tr}(\hat{\rho}_B \log_2(\hat{\rho}_B)) \quad (4.1)$$

L'entropie mesure notre méconnaissance de l'objet ; dans ce cas particulier elle provient de l'intrication entre les deux objets : si on ne regarde qu'une partie du système, plus cette partie sera indépendante de la seconde plus on sera sûr de son état. Au contraire s'il y a une grande dépendance, c'est-à-dire une forte intrication, alors le fait de ne pas connaître la seconde partie « propage » cette méconnaissance sur la première partie. Cette mesure est appelée *entropie d'intrication*. Pour deux qubits maximalement intriqués nous avons $E = 1$, cette quantité d'intrication est appelée ebit.

L'affaire se complique pour les états mixtes. La mesure précédente ne peut pas être utilisée car elle ne fait pas la différence entre la méconnaissance due à l'intrication et celle due au mélange statistique. Comment alors quantifier l'intrication ? On pourrait penser à utiliser des processus nécessitant l'intrication : plus ils sont réalisés efficacement plus l'intrication serait grande. Cette méthode conduit à de nombreuses mesures de l'intrication (comme par exemple le critère de Reid-EPR lié au « paradoxe EPR » [121, 122] ou la fidélité de téléportation quantique

[123, 124]). L'ennui est que ces mesures donnent des résultats différents : un état $|\psi\rangle$ peut être plus intriqué qu'un autre état $|\varphi\rangle$ pour une mesure donnée alors qu'une seconde mesure donnera le résultat inverse.

Il existe en réalité un très grand nombre de mesures différentes, leur sens physique n'étant pas toujours aussi évident. Et parmi celles-ci beaucoup nécessitent des procédures d'optimisation complexes les rendant peu, voire pas du tout, utilisables en pratique. Dans ce chapitre nous utiliserons une des rares mesures calculables : la négativité [125]. Elle est définie par

$$\mathcal{N} = \frac{\left\| \hat{\rho}_{AB}^{T_A} \right\|_1 - 1}{2} \quad (4.2)$$

$\hat{\rho}_{AB}^{T_A}$ étant la matrice densité obtenue par transposée partielle dans le sous-espace A de l'opérateur densité de l'état bipartite $\hat{\rho}_{AB}$:

$$\langle \psi_A, \psi_B | \hat{\rho}_{AB}^{T_A} | \varphi_A, \varphi_B \rangle = \langle \varphi_A, \psi_B | \hat{\rho}_{AB} | \psi_A, \varphi_B \rangle \quad (4.3)$$

La norme $\|\hat{\rho}\|_1$ vaut

$$\|\hat{\rho}\|_1 = \text{Tr} \left(\sqrt{\hat{\rho}^\dagger \hat{\rho}} \right) \quad (4.4)$$

La fonction racine n'est pas analytique nous devons donc, comme pour l'entropie, revenir à la matrice densité et la diagonaliser afin d'obtenir la négativité. Celle-ci est en fait égale à la somme des valeurs propres négatives de $\hat{\rho}_{AB}^{T_A}$ (en valeur absolue). Nous pouvons alors vérifier que pour un état séparable la négativité est nulle : en effet nous avons dans ce cas $\hat{\rho}_{AB}^{T_A} = \hat{\rho}_{AB}$, et les valeurs propres restent donc positives. La négativité des deux qubits maximale intriqués évoqués précédemment, dont l'entropie d'intrication était de 1ebit, est de 1/2.

4.1.1.3 États de Bell

Les états de Bell sont les états maximalelement intriqués entre deux qubits. Dans le cas discret ils valent :

$$|\phi_{\pm}\rangle = \frac{|0\rangle|0\rangle \pm |1\rangle|1\rangle}{\sqrt{2}} \quad (4.5)$$

$$|\psi_{\pm}\rangle = \frac{|0\rangle|1\rangle \pm |1\rangle|0\rangle}{\sqrt{2}} \quad (4.6)$$

L'équivalent en terme de variables continues avec un codage sur des états cohérents de phases opposés donne

$$|\phi_{\pm, \text{coher}}\rangle = \frac{|\alpha\rangle|\alpha\rangle \pm |-\alpha\rangle|-\alpha\rangle}{\sqrt{2(1 \pm e^{-4|\alpha|^2})}} \quad (4.7)$$

$$|\psi_{\pm, \text{coher}}\rangle = \frac{|\alpha\rangle|-\alpha\rangle \pm |-\alpha\rangle|\alpha\rangle}{\sqrt{2(1 \pm e^{-4|\alpha|^2})}} \quad (4.8)$$

Il est aussi possible de coder les qubit sur la parité d'états chats de Schrödinger, les états de Bell étant alors uniquement une permutation des précédents

$$|\phi_{+,chat}\rangle = \frac{|C_+\rangle|C_+\rangle + |C_-\rangle|C_-\rangle}{\sqrt{2}} = |\phi_{+,coher}\rangle \quad (4.9)$$

$$|\phi_{-,chat}\rangle = \frac{|C_+\rangle|C_+\rangle - |C_-\rangle|C_-\rangle}{\sqrt{2}} = |\psi_{+,coher}\rangle \quad (4.10)$$

$$|\psi_{+,chat}\rangle = \frac{|C_+\rangle|C_-\rangle + |C_-\rangle|C_+\rangle}{\sqrt{2}} = |\phi_{-,coher}\rangle \quad (4.11)$$

$$|\psi_{-,chat}\rangle = \frac{|C_+\rangle|C_-\rangle - |C_-\rangle|C_+\rangle}{\sqrt{2}} = |\psi_{-,coher}\rangle \quad (4.12)$$

Si nous voulons réaliser des protocoles d'information quantique avec des qubits à variables continues, c'est donc ces états que nous devons utiliser, mais comment les produire ? Le plus simple consiste à envoyer un chat de Schrödinger monomode $(|\sqrt{2}\alpha\rangle \pm |-\sqrt{2}\alpha\rangle)/\sqrt{2}$ sur une lame séparatrice 50/50 : suivant la parité du chat et la face d'entrée on obtient alors les quatre combinaisons possibles (l'intrication venant du fait que la phase des deux faisceaux en sortie dépend de la phase en entrée).

Cependant ces états se dégradent très rapidement avec les pertes, et même d'autant plus vite que les chats sont grands. Nous l'avons dit la distillation d'intrication permet d'y remédier, mais ce n'est pas le seul moyen. Il est aussi possible de contourner le problème en créant l'intrication à distance, et c'est cette dernière méthode que nous nous proposons d'étudier.

4.1.2 Principe de l'expérience

La distillation ne peut être réalisée uniquement avec des états et opérations gaussiennes, il en va certainement de même de l'intrication à distance. À contrario, nous allons voir qu'il est relativement facile, en partant de chats de même parité, de réaliser les états de Bell $|\psi_{\pm,chat}\rangle$. Il suffit en effet pour cela de changer la parité d'un des deux états de départ. Or nous pouvons remarquer que c'est exactement ce que fait la procédure de dégaussification présentée au chapitre précédant, en effet

$$\hat{a}(|\alpha\rangle \pm |-\alpha\rangle) = \alpha(|\alpha\rangle \mp |-\alpha\rangle) \quad (4.13)$$

Son mode de fonctionnement correspond de plus exactement à ce dont nous avons besoin : le fait qu'il faille envoyer une partie de faisceau dans la voie APD et que la qualité de l'état soit indépendant des pertes dans celle-ci permet d'y intégrer un canal de communication. Il nous reste encore un point à régler : afin de créer l'intrication il ne suffit pas de retirer un photon à l'un des modes, il faut aussi que l'on ne puisse pas savoir à quel mode ce photon a été retiré. Le but est ainsi d'appliquer l'opérateur $(\hat{a}_1 \pm \hat{a}_2)/\sqrt{2}$ à l'ensemble de nos deux modes, c'est en effet cet opérateur qui va les intriquer.

À partir de là le principe est assez simple (figure 4.2) :

1. On prélève une petite fraction de chacun des deux chats, de façon à ce que la probabilité de prélever plus d'un photon soit négligeable par rapport à celle de prélever un photon. La probabilité qu'un photon soit prélevé dans chacun des modes doit aussi être faible.
2. La phase de l'un des prélèvements est modifiée et ils sont tous les deux envoyés à travers un canal de communication.

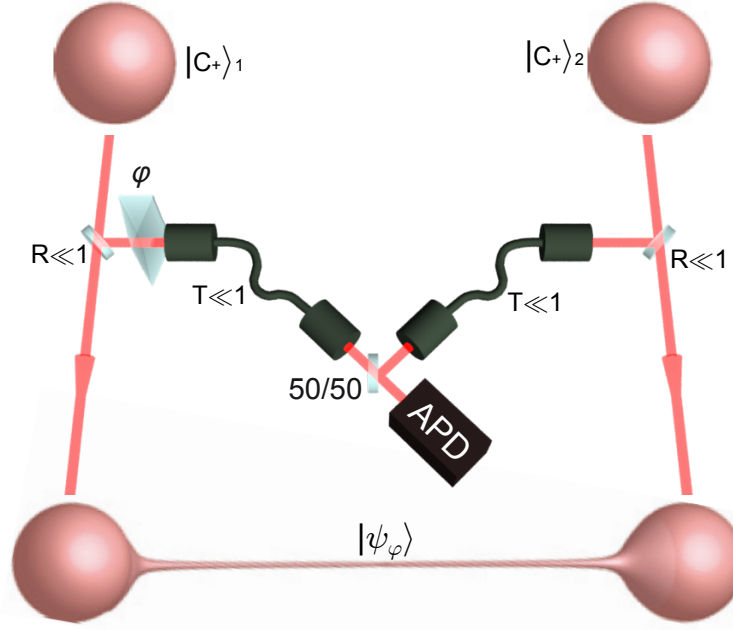


FIGURE 4.2: Principe de l'expérience

3. Ils sont enfin mélangés sur une lame semi-réfléchissante 50/50 et un détecteur de photons est placé à l'une des sorties. Il est aussi possible de doubler le taux de répétition en plaçant un compteur de photon sur chacune des sorties, mais il faut alors prendre en compte le fait que la phase peut être inversée suivant le compteur ayant détecté le photon.

Ainsi au lieu d'envoyer des faisceaux intriqués aux deux partenaires ce sont ceux-ci qui envoient une partie de leurs faisceaux afin de les intriquer. De même au lieu que les pertes dégradent l'intrication elles baissent le taux de succès. Dans la pratique ce dispositif ne permet tout de même pas des communications aussi longues que l'on veut sans perte de l'intrication : en effet les coups parasites la détruisent et il faut donc que ceux-ci restent négligeables devant la cadence de répétition. Cette contrainte n'est cependant pas spécifique à notre protocole, elle est en effet aussi valable pour la distillation d'intrication.

L'opérateur appliqué avec ce schéma est $(\hat{a}_1 - e^{i\varphi}\hat{a}_2)/\sqrt{2}$, et l'état produit après conditionnement est donc

$$|\psi_\varphi\rangle = -i \sin\left(\frac{\varphi}{2}\right) \frac{|C_+\rangle_1 |C_-\rangle_2 + |C_-\rangle_1 |C_+\rangle_2}{\sqrt{2}} + \cos\left(\frac{\varphi}{2}\right) \frac{|C_+\rangle_1 |C_-\rangle_2 - |C_-\rangle_1 |C_+\rangle_2}{\sqrt{2}} \quad (4.14)$$

Nous retrouvons donc bien l'état $|\psi_{-, \text{chat}}\rangle$ pour $\varphi = 0$ et $|\psi_{+, \text{chat}}\rangle$ pour $\varphi = \pi/2$. Nous pouvons aussi remarquer qu'il est également possible de créer les états de Bell $|\phi_{\pm, \text{chat}}\rangle$ en partant de deux chats de parités opposées.

Pour simplifier nous nous contenterons d'effectuer l'expérience avec de très petits chats de Schrödinger ($|\alpha| \lesssim 1$) que nous pourrions alors approximer par des vides comprimés pour les chats pairs, et par des états « chatons » [94, 71, 126] pour les chats impairs. Cette version n'est pas uniquement une approximation des états de Bell avec des chats ou des états cohérents. En effet les états chatons sont en réalité des photons uniques comprimés, et l'état produit est donc un état de Bell discret comprimé. La compression étant une opération réversible et locale, nous

pouvons vérifier que l'intrication correspond bien à un ebit. En réalité ce processus d'intrication peut être appliqué à tout état non classique (c'est-à-dire qui ne soit pas un état cohérent ou un état thermique). Par exemple sur des états de Fock contenant chacun n photons il produit l'état intriqué $(|n-1\rangle|n\rangle \pm |n\rangle|n-1\rangle)/\sqrt{2}$; il peut donc aussi être utilisé pour produire à distance des états de Bell discrets.

Enfin nous pouvons remarquer que ce protocole est en réalité un cas particulier de transfert d'intrication : en prélevant une partie du faisceau nous produisons un état intriqué entre le faisceau principal et le prélèvement. La recombinaison suivie de la détection de photon fait alors office de mesure conjointe qui va transférer l'intrication aux deux faisceaux principaux.

4.2 Réalisation expérimentale

4.2.1 Montage

Le but n'est pas de réaliser un vrai système d'intrication à distance mais uniquement une démonstration de principe. Nous nous passerons donc des fibres optiques servant de canal de transmissions, les pertes de la voie APD suffiront à en mimer les effets. De plus nous n'agissons pas sur la phase, produisant ainsi uniquement l'état $|\psi_{\varphi=0}\rangle = |\psi_{-, \text{chat}}\rangle$.

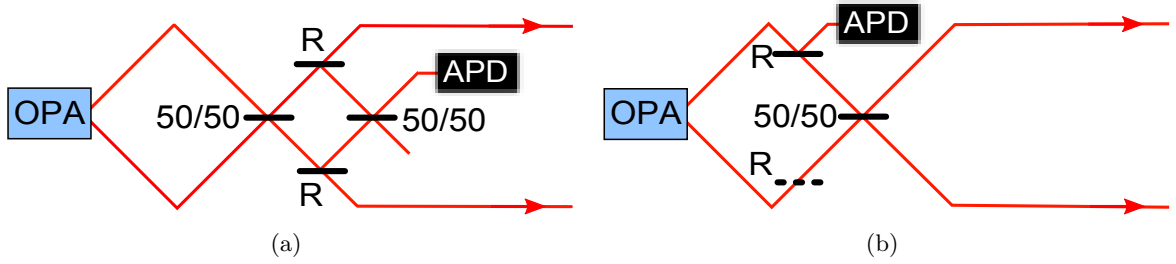


FIGURE 4.3: Schéma de l'expérience (a) schéma réel (b) schéma simplifié

La première chose dont nous avons besoin pour réaliser cette expérience est une paire d'états comprimés, ils seront produits par recombinaison d'un état EPR comme présenté au chapitre précédent. La figure 4.3(a) montre alors que le dispositif de soustraction forme un interféromètre de Mach-Zender, avec une différence de marche nulle puisque $\varphi = 0$. Or dans ce cas, la sortie de cet interféromètre est identique à l'entrée, de sorte que cet interféromètre est inutile. Il peut donc être enlevé, ce qui permet de simplifier le schéma en celui de la figure 4.3(b). Le principe de cette simplification est le suivant : si nous retirons un photon avant la transformation de l'état EPR en deux vides comprimés alors nous ne savons pas lequel de ceux-ci se retrouvera avec un photon en moins après la recombinaison. Dans le cas idéal, nous pouvons facilement retrouver cette simplification mathématiquement :

$$\frac{\hat{a}_1 - \hat{a}_2}{2} \hat{U}_{\text{BS}} = \hat{U}_{\text{BS}} \hat{U}_{\text{BS}}^\dagger \frac{\hat{a}_1 - \hat{a}_2}{2} \hat{U}_{\text{BS}} = \hat{U}_{\text{BS}} \hat{a}_2 \quad (4.15)$$

où $\hat{U}_{\text{BS}}$ est l'opérateur associé à la lame séparatrice. On peut pousser cette simplification en remarquant que la réflexion R des lames de prélèvement est faible, et que l'action de celle qui ne mène vers aucune APD est donc négligeable ; par conséquent, elle peut être retirée. Un second avantage de cette simplification est qu'en n'ayant pas de recombinaison après le prélèvement nous n'avons pas besoin d'un système actif de contrôle de la phase.

Le schéma final de l'expérience est illustré figure 4.4 : l'OPA produit une paire EPR avec un gain $r = 1,18$, ce qui correspond à une compression de 3,6 dB. Les faisceaux sont ensuite superposés par un PBS. Une lame demi-onde placée sur le faisceau complémentaire permet le prélèvement de $R = 5\%$ du faisceau : ce prélèvement sort par l'autre voie du cube et est envoyée vers l'APD dont l'efficacité vaut $\mu = 7\%$. Une lame demi-onde ayant un angle de 45° permet le mélange de la paire EPR pour former, sans conditionnement, les deux vides comprimés. Ils sont ensuite séparés et superposés aux oscillateurs locaux sur un second PBS et envoyés chacun sur une détection homodyne. Nous avons déjà précisé que la phase relative entre les deux faisceaux de la paire EPR ne joue aucun rôle, nous contrôlerons donc la différence de phase entre les quadratures mesurées à l'aide du duo de lames d'ondes placées sur les oscillateurs locaux. Il ne reste plus qu'à contrôler la phase globale grâce à une cale piézo-électrique placée sur l'oscillateur local. Au lieu de la stabiliser nous la balayerons et tirons les quadratures en s'aidant des variances de l'état non conditionné. Le taux de succès est de 500 cps/s, ce qui est largement au dessus des coups d'obscurité (10 cps/s).

4.2.2 Modélisation

Pour la modélisation nous pouvons là aussi déplacer les pertes homodynes juste après l'OPA (en modifiant en conséquence celles de la voie APD). Nous partons donc du même état que pour le photon unique (équation 3.26) :

$$W_\eta(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 ab} e^{-\frac{(x_1-x_2)^2}{2a} - \frac{(p_1-p_2)^2}{2b} - \frac{(x_1+x_2)^2}{2b} - \frac{(p_1+p_2)^2}{2a}} \quad (4.16)$$

avec

$$a = \eta(hs + h - 1) + 1 - \eta \quad (4.17)$$

$$b = \eta\left(\frac{h}{s} + h - 1\right) + 1 - \eta \quad (4.18)$$

On applique alors la réflexion d'une partie du mode 1 dans le mode de conditionnement C, sans oublier le vide $W_0(x, p)$ à l'autre entrée

$$\begin{aligned} W_{\text{mel}}(x_1, p_1, x_2, p_2, x_C, p_C) &= W_\eta(tx_1 + rx_C, tp_1 + rp_C, x_2, p_2) \\ &\times W_0(tx_C - rx_1, tp_C - rp_1) \end{aligned} \quad (4.19)$$

L'étape suivante consiste à appliquer le conditionnement (cf. section 3.5.2.3)

$$W_{\text{cond}}(x_1, p_1, x_2, p_2) = \xi \frac{\text{Tr}_C(W_{\text{mel}} W_{n,C})}{\text{Tr}(W_{\text{mel}} W_{n,C})} + (1 - \xi) \frac{\text{Tr}_C(W_{\text{mel}})}{\text{Tr}(W_{\text{mel}})} \quad (4.20)$$

Il ne nous reste alors plus qu'à recombinaer les deux modes restant sur une 50/50

$$W(x_1, p_1, x_2, p_2) = W_{\text{cond}}\left(\frac{x_1 + x_2}{\sqrt{2}}, \frac{p_1 + p_2}{\sqrt{2}}, \frac{x_2 - x_1}{\sqrt{2}}, \frac{p_2 - p_1}{\sqrt{2}}\right) \quad (4.21)$$

Ce qui donne

$$\begin{aligned} W(x_1, p_1, x_2, p_2) &= \frac{1}{\pi^2(a'b' - c^2)} e^{-\frac{b'(x_1^2 + p_2^2) + a'(p_1^2 + x_2^2) + 2c(x_1 x_2 + p_1 p_2)}{a'b' - c^2}} \\ &\times (1 - a'\beta^2 - b'\alpha^2 + 2c\alpha\beta + (\beta x_1 + \alpha x_2)^2 + (\alpha p_1 + \beta p_2)^2) \end{aligned} \quad (4.22)$$

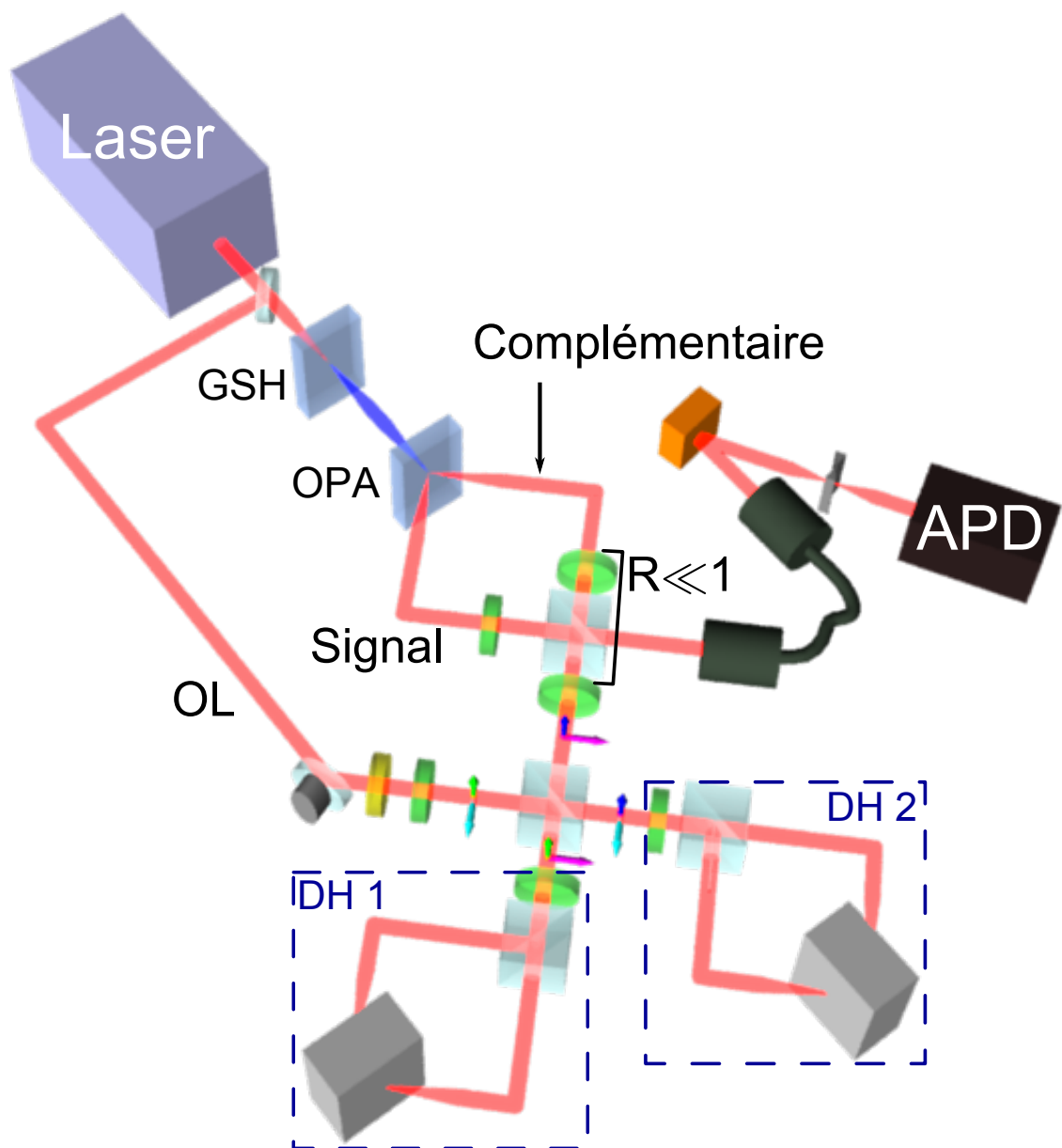


FIGURE 4.4: Schéma du dispositif expérimental

avec

$$a' = \frac{(1+t)^2 a + (1-t)^2 b + 2(1-t^2)}{4} \quad (4.23)$$

$$b' = \frac{(1-t)^2 a + (1+t)^2 b + 2(1-t^2)}{4} \quad (4.24)$$

$$c = (1-T) \frac{a+b-2}{4} \quad (4.25)$$

$$\alpha = \sqrt{\xi} \frac{2ab - (1+t)a - (1-t)b}{2(a'b' - c^2)\sqrt{a+b-2}} \quad (4.26)$$

$$\beta = \sqrt{\xi} \frac{2ab - (1-t)a - (1+t)b}{2(a'b' - c^2)\sqrt{a+b-2}} \quad (4.27)$$

L'état non conditionné correspond à $\alpha = \beta = 0$ ($\xi = 0$), ce qui revient à enlever la partie polynomiale (le conditionnement est la seule opération qui ne conserve pas le caractère gaussien). Bien que cela ne saute pas aux yeux, il est bien séparable.

Cette fonction de Wigner n'est équivalente à celle obtenue lors d'une intrication à distance que dans le cas $R \ll 1$; une modélisation similaire à celle que l'on vient de faire mais pour le protocole réel donne une fonction de Wigner ayant la même forme mais avec des paramètres différents [82], ceux-ci coïncident néanmoins dans la limite $R \rightarrow 1$

$$a' = \eta T(hs + h - 2) + 1 \quad (4.28)$$

$$b' = \eta T\left(\frac{h}{s} + h - 2\right) + 1 \quad (4.29)$$

$$c = 0 \quad (4.30)$$

$$\alpha = \frac{h/s + h - 2}{b'} \sqrt{\frac{\eta \xi T}{h(s + 1/s) + 2h - 4}} \quad (4.31)$$

$$\beta = \frac{hs + h - 2}{a'} \sqrt{\frac{\eta \xi T}{h(s + 1/s) + 2h - 4}} \quad (4.32)$$

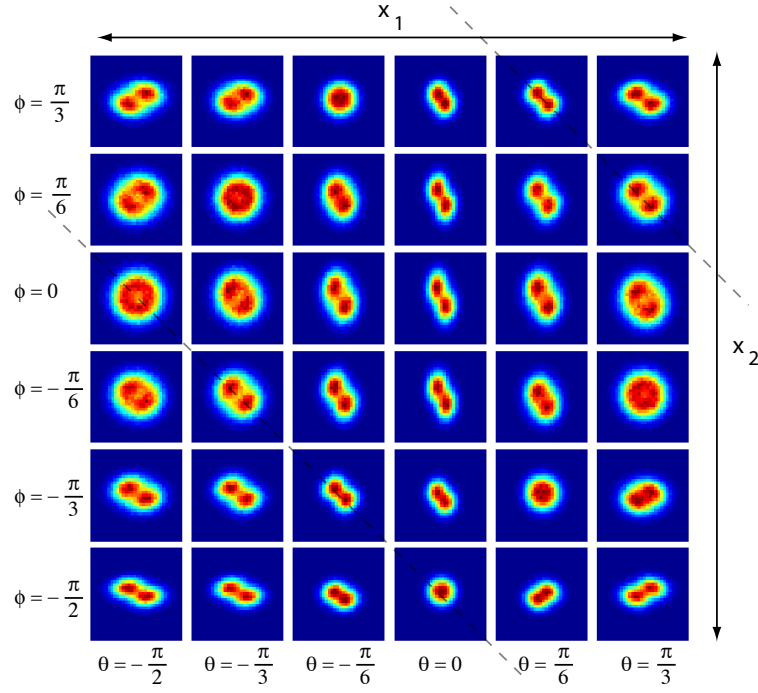
4.2.3 Résultats

Nous avons réalisé la tomographie bimode de l'état produit à partir de 36 distributions de probabilités jointes $P(x_{1,\theta}, x_{2,\phi})$ pour des phases θ et ϕ prenant leurs valeurs parmi $\{-\frac{\pi}{2}, -\frac{\pi}{3}, -\frac{\pi}{6}, 0, \frac{\pi}{6}, \frac{\pi}{3}\}$ [127]. Chaque distribution est calculée à partir de 160 000 mesures et est représentée sous forme d'histogramme de 40×40 bins; ces histogrammes sont présentés figure 4.5.

Le modèle permet, à partir des moments des distributions, de calculer la valeurs des paramètres expérimentaux manquant, à savoir l'excès de gain de l'OPA $\gamma = 0,34$ et la pureté modale $\xi = 0,88$. Nous pouvons alors voir que les distributions sont très similaires à celles obtenus par le modèle; l'écart normalisé entre les deux vaut, en moyenne,

$$\iint |P_{\text{mod}}(x_{1,\theta}, x_{2,\phi}) - P_{\text{exp}}(x_{1,\theta}, x_{2,\phi})| dx_{1,\theta} dx_{2,\phi} = (6,3 \pm 0,5)\%$$

On peut voir que les distributions de probabilités des quadratures forment deux lobes qui fusionnent en un « volcan » pour des quadratures orthogonales. Ces distributions mettent enfin en évidence la symétrie $P(x_{1,\frac{\pi}{2}-\phi}, x_{2,\frac{\pi}{2}-\theta}) = P(x_{2,\phi}, x_{1,\theta})$ particulière à cet état (matérialisée par les tirets dans la figure 4.5).

FIGURE 4.5: *Distribution de probabilité des quadratures*

Ces distributions permettent de reconstruire la fonction de Wigner de l'état préparé. La figure 4.6 présente les coupes $W(x_1, p_1, 0, 0)$ et $W(x_1, 0, x_2, 0)$, avec et sans corrections des pertes homodynes. La version sans correction est obtenue par transformée inverse de Radon et celle avec correction est évidemment obtenue grâce à l'algorithme de maximum de vraisemblance. Il est difficile de dire si la fonction de Wigner non corrigée est négative à cause des fluctuations expérimentales ; la fonction corrigée ne laisse par contre aucun doute sur la question : à l'origine nous avons $W(0, 0, 0, 0) = -0,04 \pm 0,01$ (le cas idéal correspond à une valeur de $1/\pi^2 \approx 0,10$). La fidélité avec un état idéal $|\psi_{\varphi=0}\rangle$ (pour $|\alpha|^2 = 0,65$) est de $(64 \pm 5)\%$ et celle avec deux vides comprimés indépendant de $(23 \pm 3)\%$.

À partir de la matrice densité de l'état corrigé des pertes homodynes nous pouvons calculer l'intrication de l'état : $\mathcal{N} = 0,25 \pm 0,04$. Les principaux défauts expérimentaux conduisant à cette réduction de moitié de l'intrication par rapport à la valeur idéale concernent la production des vides comprimés. Il s'agit de l'excès de gain de l'OPA et de la pureté modale liée à la largeur de la fluorescence paramétrique. Le modèle nous donne pour un état EPR pur (correspondant à $\gamma \approx 0$ et $\xi = 0,98$ uniquement dû aux coups d'obscurité) une intrication de $\mathcal{N} = 0,47$. Si on employait le vrai protocole avec des vides comprimés produits de manière indépendante nous aurions une intrication de $\mathcal{N} = 0,41$ pour des états purs et $\mathcal{N} = 0,22$ pour les paramètres réels de notre expérience.

Nous pouvons donc créer une forte intrication, moitié de l'intrication maximale pour des pertes de 7%. Celles-ci correspondent à environ 60 km de fibre optiques aux longueurs d'onde des télécoms optiques ($\lambda = 1,55 \mu\text{m}$, pertes de 0,2 dB/km) soit une distance totale de 120 km entre les deux sites.

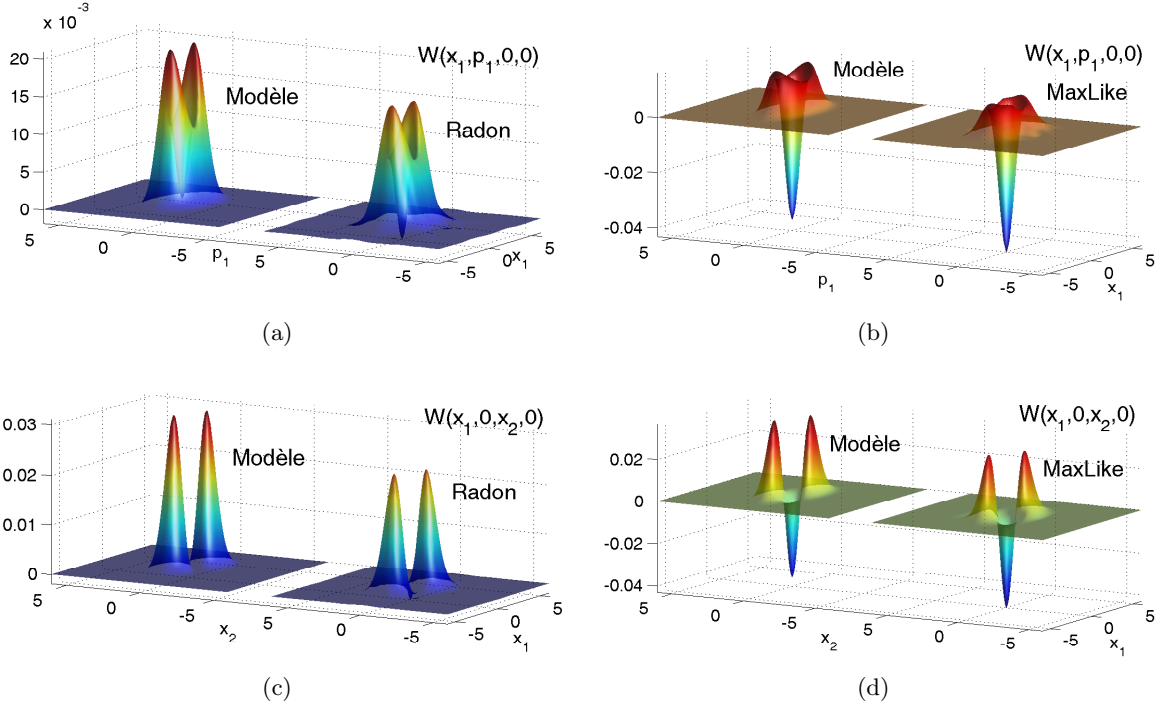


FIGURE 4.6: Coupes de la fonction de Wigner (a) $W(x_1, p_1, 0, 0)$ sans correction (b) $W(x_1, p_1, 0, 0)$ corrigée des pertes homodynes (c) $W(x_1, 0, x_2, 0)$ sans correction (d) $W(x_1, 0, x_2, 0)$ corrigée des pertes homodynes

4.3 Comparaison avec un état de Bell discret

4.3.1 Montage

Puisque nous avons produit un état de Bell continu il serait intéressant de le comparer à un état de Bell discret. La tomographie de cet état a déjà été réalisée par Babichev *et al.*, toutefois nous aurons besoin de reconstruire cet état afin d'effectuer la comparaison avec les mêmes paramètres expérimentaux. Bien qu'il soit possible de produire des états de Bell discrets à distance avec notre protocole de soustraction cohérente, cela nécessiterait de produire deux photons uniques; nous nous contenterons donc de les réaliser par la méthode « habituelle », c'est-à-dire en envoyant un photon unique sur une lame séparatrice 50/50.

L'utilisation des couples lame demi-onde et PBS pour réaliser des lames demi-onde variables permet de passer très facilement de la production d'états de Bell continus à celle d'états de Bell discrets. En effet si nous tournons la lame demi-onde située sur le complémentaire afin de réfléchir l'intégralité du faisceau (figure 4.7) alors nous revenons à la production d'un photon unique dans le mode signal. La lame de recombinaison joue alors le rôle de celle qui crée l'intrication en « séparant » le photon unique en deux modes qui sont mesurés par les détecteurs homodynes.

Les paramètres expérimentaux sont donc les mêmes excepté la pureté modale qui monte à $\xi = 0,94$; en effet dans la version précédente elle était diminuée par l'imperfection du *mode-matching* entre les deux faisceaux issus de l'OPA. Le taux de succès est bien plus important puisqu'on envoie 20 fois plus de puissance vers l'APD; il correspond au taux de production du photon unique.

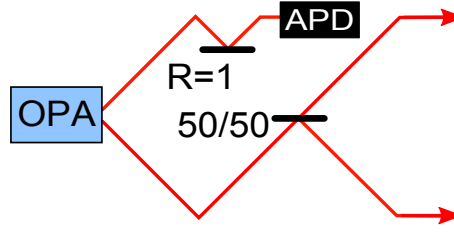


FIGURE 4.7: Schéma du dispositif pour produire un état de Bell discret

4.3.2 Modélisation

Le modèle est très simple à calculer : nous pouvons soit prendre le modèle précédent avec $t = 0$, soit directement mélanger le photon unique dont on connaît la fonction de Wigner (section 3.5) avec le vide :

$$W_{\text{discr}}(x_1, p_1, x_2, p_2) = W_0\left(\frac{x_1 + x_2}{\sqrt{2}}, \frac{p_1 + p_2}{\sqrt{2}}\right) W_{\text{phot}}\left(\frac{x_2 - x_1}{\sqrt{2}}, \frac{p_2 - p_1}{\sqrt{2}}\right) \quad (4.33)$$

Ce qui donne

$$W_{\text{discr}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 \sigma^2} \left(1 - \delta + \delta \frac{(x_1 - x_2)^2 + (p_1 - p_2)^2}{2\sigma^2} \right) \times e^{-\frac{(x_1 - x_2)^2}{2\sigma^2} - \frac{(x_1 + x_2)^2}{2} - \frac{(p_1 - p_2)^2}{2\sigma^2} - \frac{(p_1 + p_2)^2}{2}} \quad (4.34)$$

où δ et σ sont les paramètres du photon uniques donnés aux équations 3.31 et 3.32.

Puisqu'il s'agit d'un état discret, calculons aussi sa matrice densité à l'aide de l'équation 2.49

$$\begin{aligned} \langle n_1, n_2 | \rho_{\text{discr}} | m_1, m_2 \rangle &= 2^{1-n_1-n_2} \frac{(\sigma^2 - 1)^{n_1+n_2-1}}{(\sigma^2 + 1)^{n_1+n_2+2}} \frac{((n_1 + n_2)!)!}{n_1! m_1! n_2! m_2!} \\ &\times (\sigma^4(1 - \delta) + \sigma^2 \delta (1 + 2n_1 + 2n_2) - 1) \delta_{n_1+n_2, m_1+m_2} \end{aligned} \quad (4.35)$$

4.3.3 Résultats

Les états de Fock étant invariant par rotation de la phase, l'état mesuré ne dépend que de la différence entre les phases des quadratures mesurées. Nous pouvons donc nous contenter des six distributions de quadratures illustrées figure 4.8. Chacune d'entre elle a été reconstruite à partir de 960 000 points expérimentaux puis répartis en histogrammes. L'accord avec le modèle est encore une fois très bon (écart normalisé de 4%).

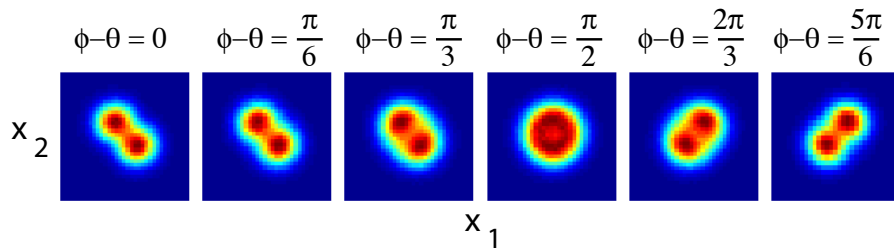


FIGURE 4.8: Distribution de probabilité des quadratures

Comme on s'y attendait, les fonctions de Wigner (figure 4.9) sont similaires à celles de l'état de Bell continu mais sans la compression. La négativité à l'origine est cette fois bien visible sans correction ($W(0,0,0,0) = -0,003$) et celle de la version corrigée des pertes homodynes est relativement proche du cas idéal ($W(0,0,0,0) = -0,032$ contre $-0,10$).

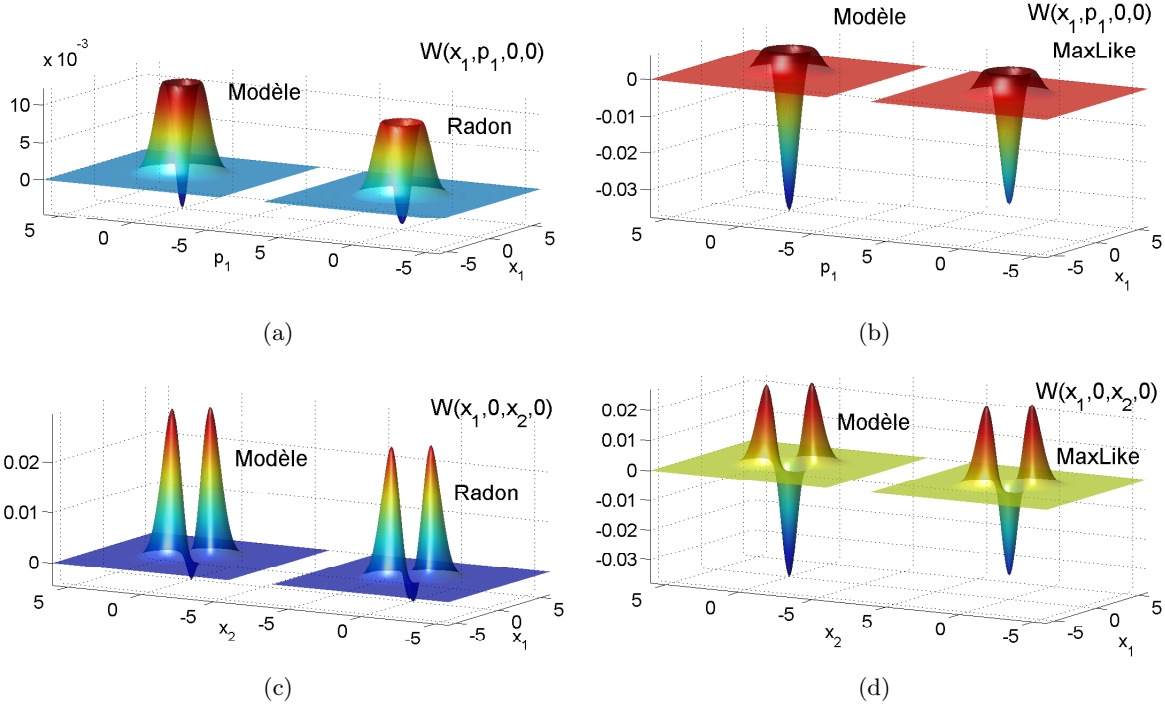


FIGURE 4.9: Coupes de la fonction de Wigner (a) $W(x_1, p_1, 0, 0)$ sans correction (b) $W(x_1, p_1, 0, 0)$ corrigée des pertes homodynes (c) $W(x_1, 0, x_2, 0)$ sans correction (d) $W(x_1, 0, x_2, 0)$ corrigée des pertes homodynes

À partir de la matrice densité on peut à nouveau quantifier l'intrication de l'état corrigé : $\mathcal{N} = 0,19 \pm 0,02$. On peut ainsi remarquer que malgré la réflectivité non nulle de la lame semi-réfléchissante et la moins bonne pureté homodyne dans le cas continu ainsi que la négativité plus marquée de la fonction de Wigner de la version discrète, l'état de Bell continu possède une meilleure intrication que son homologue discret pour le même dispositif de production. Nous devons toutefois noter que les paramètres optimaux ne sont pas les mêmes dans les deux cas : en effet si un fort gain paramétrique n'a d'autre influence sur l'état continu que le taux de succès, ce n'est pas le cas de la version discrète. Pour celle-ci un fort gain comme celui utilisé pour ces deux expériences entraîne une proportion non négligeable d'états à deux photons (environ 15%), ce qui implique une dégradation de l'état produit et donc de son intrication. Si nous avions gardé constant, entre les deux expériences, le taux de succès au lieu du gain nous aurions eu une intrication de $\mathcal{N} = 0,34$.

4.4 Discussion

4.4.1 Communications à longues distances

Afin d'étudier l'efficacité de ce protocole considérons un schéma simple de transmission sur une distance de 100 km (soit une transmission de 10% pour chacune des deux fibres optiques, de 50 km, à longueur télécoms) entre Alice et Bob sans répéteurs ; ils seront uniquement assistés de Charlie, situé à mi-chemin et chargé de créer l'intrication (directement ou à distance). Pour juger de cette efficacité nous pouvons le comparer à d'autres protocoles réalisés avec les mêmes paramètres expérimentaux, comme nous l'avons fait dans la section précédente.

Il existe de nombreux protocoles possibles qui combinent les différents états que l'on peut utiliser (états EPR, états de Bell discrets, états de Bell continus, ...), et les différents moyens de lutter contre les pertes (distillation, intrication à distance, transfert d'intrication, ...). Ils sont bien trop nombreux pour que nous les regardions tous, et nous nous contenterons donc de partir des protocoles impliqués dans ces deux expériences¹. Les paramètres expérimentaux utilisés seront donc les mêmes que ceux déjà présentés dans ce chapitre, avec un taux de répétition du laser de 800 kHz. Pour notre protocole d'intrication à distance nous utiliserons par contre une APD sur chacune des voies afin d'optimiser le taux de succès. Enfin nous étudierons les capacités des divers protocoles considérés avec et sans les imperfections afin d'en déterminer le rôle ; dans ce dernier cas les seuls paramètres utilisés seront le gain de l'OPA ($g = 1, 18$) et la réflexion de la lame de prélèvement ($R = 5\%$) pour l'intrication à distance.

Dans nos expériences nous retrouvons deux types d'états intriqués produits localement : les paires EPR, utilisées comme base dans les deux expériences, et les états de Bell discrets, étudiés dans la deuxième. L'intrication est alors extrêmement sensible aux pertes du canal de communication : dans le cas idéal elle passe de $\mathcal{N} = 0,64$ lors de la production par Charlie à $\mathcal{N} = 0,03$ lors de la réception par Alice et Bob pour l'état EPR et de $\mathcal{N} = 0,31$ à $\mathcal{N} = 0,003$ pour l'état de Bell ($\mathcal{N} = 0,002$ avec toutes les imperfections). Ceci est évidemment bien moins bon que le protocole d'intrication à distance qui dans le cas idéal garde une intrication de $\mathcal{N} = 0,43$ qu'il soit utilisé localement ou à distance. En considérant toutes les imperfections il reste $\mathcal{N} = 0,11$, ce qui est encore bien supérieur à la transmission directe.

Pour être vraiment honnête il ne faut cependant pas comparer l'intrication à distance à une simple transmission directe mais plutôt à cette dernière suivie d'une opération de distillation de l'intrication. Les états EPR donnant de meilleurs résultats que les états de Bell discret, y compris pour le taux de répétition qui est de 800 kHz pour les premiers (avec ou sans imperfections) et pour les seconds de 120 kHz sans imperfections et 8,1 kHz avec, nous ne nous étendrons pas d'avantages sur la version discrète. En considérant la procédure d'augmentation d'intrication déjà réalisée dans le groupe sur des états EPR [102, 82] sans ses imperfections nous pourrions faire monter l'intrication de l'état EPR transmis jusqu'à $\mathcal{N} = 0,038$ avec un taux de répétition qui descend à 4 Hz, ce qui reste bien inférieur aux capacités de notre protocole d'intrication à distance dont le taux de répétitions avec les imperfections est de 81 Hz.

Nous pouvons néanmoins noter que la distillation d'intrication peut être utilisée de manière itérative afin d'augmenter toujours plus l'intrication, ce n'est pas le cas de l'intrication à distance qui du coup ne pourra être utilisée seule que pour des distances moyennes (de l'ordre de quelques centaines de kilomètres).

Pour finir, nous avons dit (4.1.2) que notre protocole est en fait un cas particulier de transfert d'intrication ; nous pouvons donc le comparer à un dispositif où Alice et Bob préparent

1. Les modélisations supplémenaires nécessaires à cette comparaison peuvent être trouvées en appendice C.

chacun un état EPR et en envoient une partie à Charlie qui se charge d'effectuer le transfert d'intrication. Celui-ci produit un état que l'on peut modéliser par la fonction de Wigner suivante (cf. section C.3) :

$$W_{\text{transf}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 \sigma^4} \left(1 - \delta + \delta \frac{(x_1 - x_2)^2 + (p_1 - p_2)^2}{2\sigma^2} \right) e^{-\frac{x_1^2 + p_1^2 + x_2^2 + p_2^2}{\sigma^2}} \quad (4.36)$$

où les paramètres δ et σ sont les mêmes que pour le photon unique. L'état obtenu est d'ailleurs identique à celui que l'on obtiendrait en mélangeant notre photon unique avec un état thermique produit par le même OPA. Nous pouvons remarquer qu'à faible gain ($\sigma \rightarrow 1$ dans l'expression mathématique) la probabilité de produire plus d'un photon par faisceau est négligeable ; les deux branches correspondent alors au protocole utilisé pour produire des photons uniques, mais le conditionnement est délocalisé entre les deux modes conservés par Alice et Bob ; nous retrouvons donc les états de Bell discrets $\psi_{\pm}$, le signe dépendant du bras dans lequel est détecté le photon. D'ailleurs, ce dernier point montre la différence entre le transfert d'intrication et la distribution directe d'états de Bell discrets : les pertes agissent sur le taux de répétition au lieu de l'intrication, comme pour notre protocole. La première constatation est évidente : le transfert d'intrication utilise la totalité d'un mode de sortie de l'OPA là où l'intrication à distance en utilise 5% la cadence sera donc 20 fois plus grande. Sans les imperfections l'intrication à distance est plus efficace ($\mathcal{N} = 0,43$ contre $\mathcal{N} = 0,26$ pour le transfert d'intrication), le choix se fait donc entre l'intrication et la cadence. Par contre lorsque l'on applique toutes les imperfections le protocole de transfert d'intrication prend clairement le pas sur l'intrication à distance : intrication très légèrement meilleure ($\mathcal{N} = 0,12$ contre $\mathcal{N} = 0,11$) pour un taux de répétition toujours 20 fois supérieur.

Protocole	Sans pertes ni imperfections	Avec pertes, sans imperfections	Avec pertes et imperfections
Intrication à distance	0,43	0,43	0,11
État EPR direct	0,64	0,03	0,03
État de Bell discret direct	0,31	0,003	0,002
État EPR distillé	1,03	$\sim 0,04$	0,038
Transfert d'intrication	0,26	0,26	0,11

TABLE 4.1: Intrication des divers protocoles de communication

Nous pouvons donc en conclure que les protocoles où Alice et Bob envoient un faisceau à Charlie qui se charge d'effectuer l'intrication sur les modes restants par une mesure conjointe sont meilleurs que ceux où Charlie distribue des états intriqués y compris avec distillation d'intrication. Le choix s'effectue alors entre des états continus utilisant un codage emprunté aux variables discrètes (états de Bell), et un état profitant pleinement de son caractère continu (états EPR), la préférence allant à ce dernier pour les raisons que nous venons de détailler .

4.4.2 Calcul quantique et systèmes hybrides

Pour le calcul quantique il est bien plus efficace de créer des états de Bell continus en envoyant un chat de Schrödinger sur une lame semi-réfléchissante 50/50. Notre protocole a en effet été

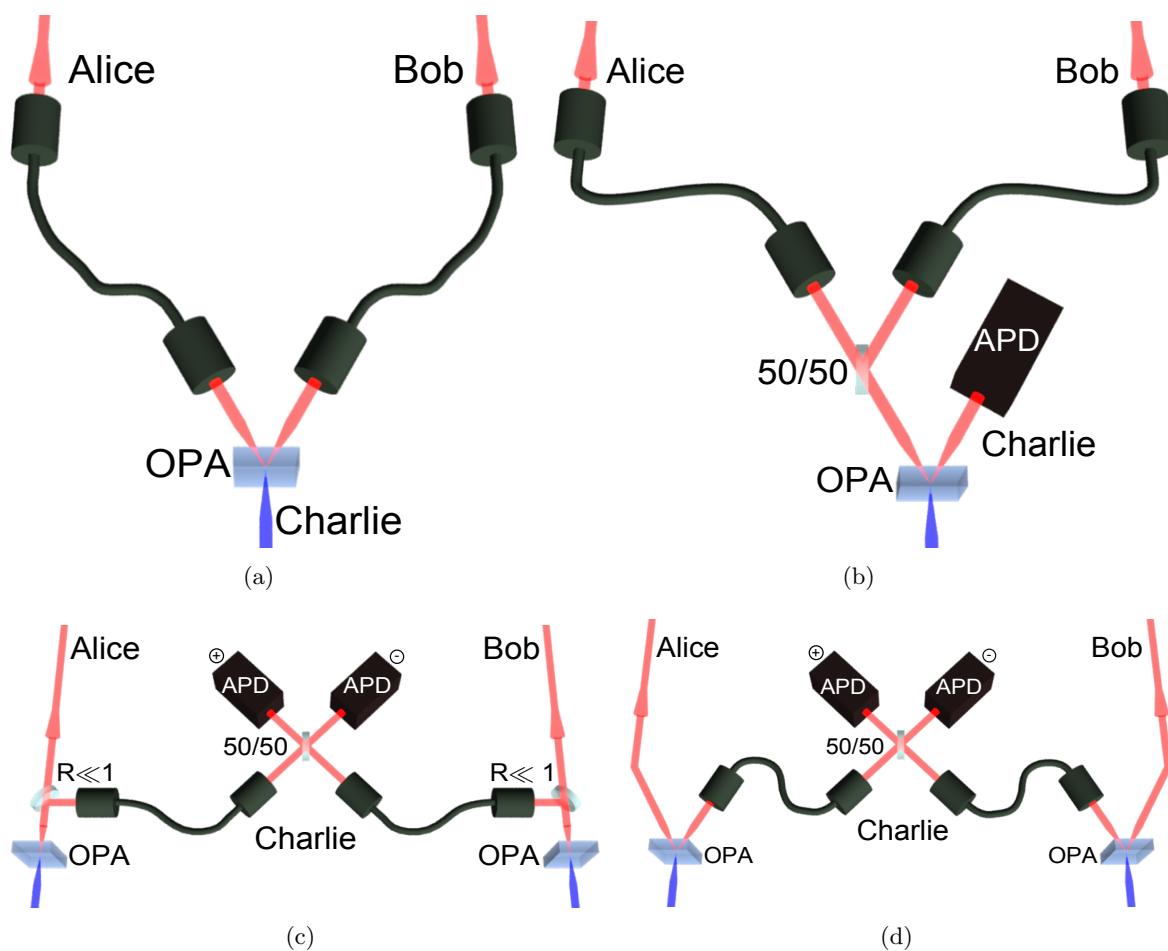


FIGURE 4.10: Différents protocoles de communication quantique envisagés (a) Envoi d'un état EPR, éventuellement avec distillation de l'intrication (b) Envoi d'un état de Bell discret (c) Intrication à distance (d) Transfert d'intrication

conçu pour palier aux pertes dans les canaux de communications à longues distances or ces canaux n'existent pas dans les ordinateurs quantiques.

Cependant le fait de créer l'intrication par une opération non gaussienne faisant intervenir l'approche discrète ouvre la voie à de nouvelles applications. Les domaines des variables discrètes et celui des variables continues se sont longtemps développés en parallèle. Pour le calcul quantique ils ont ainsi menés à des systèmes ayant chacun leurs inconvénients et leurs avantages. Depuis quelques années de nombreux travaux, dont ceux de cette thèse font partie, visent à mélanger ces deux approches afin de permettre un certain nombre de processus qui ne sont pas accessibles avec une approche uniquement continue pour laquelle les états restent, pour l'heure actuelle, gaussiens. Nous pouvons toutefois aller plus loin dans ce rapprochement en mélangeant les deux approches au sein même de l'ordinateur quantique. L'idée est alors de mélanger qubits discrets et continus afin de bénéficier des avantages des deux approches [128, 129]. Pour cela un élément essentiel est la transformation d'un qubit discret en qubit continu, et c'est là que notre méthode d'intrication entre en jeu.

Notre protocole repose sur le fait que la parité d'un chat de Schrödinger est inversée par la soustraction d'un photon, réalisant ainsi une porte NOT pour un encodage sur la parité des chats. Puisque le conditionnement s'effectue sur la détection d'un photon nous pouvons mélanger le prélèvement avec un qubit encodé sur les états de Fock $|0\rangle$ et $|1\rangle$ et placer des compteurs de photons sur les deux sorties afin de s'assurer qu'un seul photon est détecté (figure 4.11). Avec un tel dispositif l'inversion de la parité du chat est commandée par la présence ou l'absence du photon dans le mode discret, plus précisément un qubit de la forme $\mu|0\rangle + \nu|1\rangle$ correspond, en choisissant correctement les transmissions des deux lames et en conditionnant sur le bon détecteur, à l'opération $\nu\hat{a}\hat{1} + \mu\hat{a}$ sur l'état chat. Pour un chat pair en entrée nous avons les correspondances suivantes :

$$|0\rangle \leftrightarrow \frac{|\alpha\rangle - |-\alpha\rangle}{\sqrt{2}} \quad (4.37)$$

$$|1\rangle \leftrightarrow \frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2}} \quad (4.38)$$

Un chat impair donne évidemment la correspondance inverse. Ce dispositif peut être vu comme une téléportation d'un état discret sur un état continu : comme pour l'intrication à distance nous avons création d'un état intriqué par la première lame semi-réfléchissante puis ensuite mesure conjointe, ici avec le qubit discret.

Cet encodage du qubit discret n'est cependant pas le plus pratique, il serait préférable d'employer un encodage en polarisation sur un photon unique. Le codage se fait alors sur un état bimode et les superpositions à poids égaux sont des états de Bell discrets $|\psi_{\pm}\rangle$, les équivalents continus sont alors tout naturellement les états de Bell $|\psi_{\pm,\text{chat}}\rangle$ préparés par notre système. Il convient maintenant de modifier celui-ci pour y incorporer l'état de Bell discret à transférer. Pour cela nous utiliserons deux chats dans des polarisations différentes (qui peuvent être superposés spatialement) ; pour chaque polarisation le mélange sera effectué avec le mode correspondant de l'état discret et nous devrons conditionner sur la présence d'exactly un photon dans chaque polarisation (figure 4.12). Le chat qui verra sa parité changée sera donc celui dont la polarisation ne contient pas le photon unique. Pour des chats pairs en entrée et un état discret de la forme $\mu|1,0\rangle + \nu|0,1\rangle$ la transformation correspondante, est $\mu\hat{a}_2 - \nu\hat{a}_1$ (à un facteur de phase près suivant les deux combinaisons de conditionnement possibles). La correspondance effectuée est

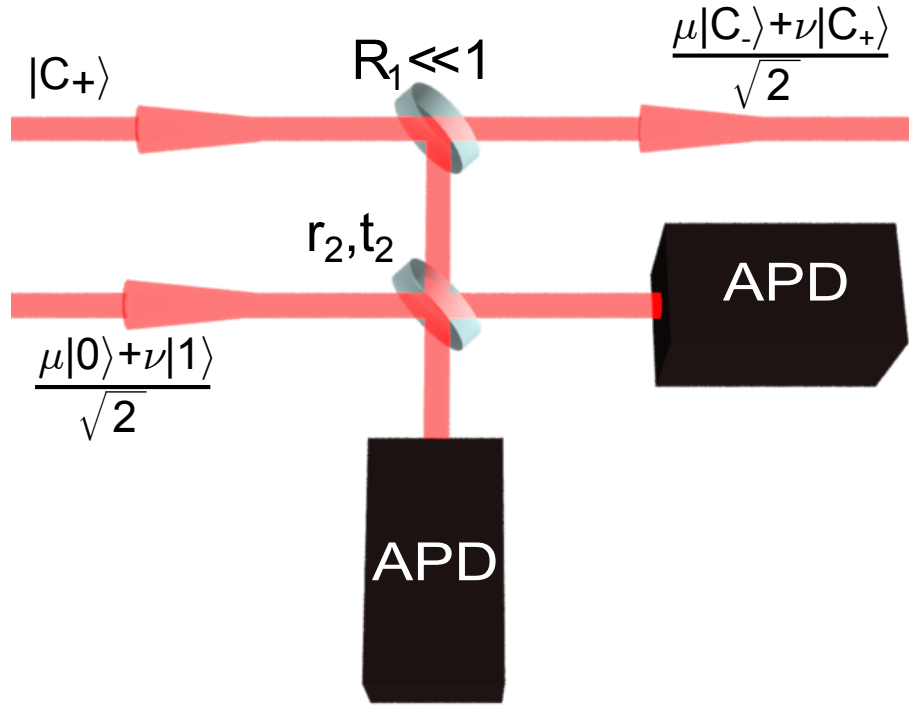


FIGURE 4.11: Protocole transformant un qubit discret encodé sur les états de Fock $|0\rangle$ et $|1\rangle$ en un qubit continu encodé sur la parité d'un chat de Schrödinger

alors la suivante :

$$|1, 0\rangle \leftrightarrow \left(\frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2}} \right) \left(\frac{|\alpha\rangle - |-\alpha\rangle}{\sqrt{2}} \right) \quad (4.39)$$

$$|0, 1\rangle \leftrightarrow - \left(\frac{|\alpha\rangle - |-\alpha\rangle}{\sqrt{2}} \right) \left(\frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2}} \right) \quad (4.40)$$

Nous pouvons remarquer que ce protocole modifie la phase quantique d'un des qubits, ce qui peut être compensé en ajoutant une porte de phase. Notons aussi que nous gardons l'information concernant l'éventuelle perte du qubit discret sur son homologue continu : la perte du qubit donne des chats de même parité alors qu'un qubit valide correspond à des chats de parités différentes.

Cette méthode de codage des qubits continus semble assez lourde. Nous pouvons alors nous demander s'il est possible de transformer un qubit discret codé sur la polarisation en un qubit continu codé sur la parité d'un seul chat. Puisqu'il y a obligatoirement un photon dans le qubit discret nous devons forcément conditionner sur deux photons, dans ce cas il nous faut apporter un deuxième photon de manière à ne pas changer systématiquement la parité du chat. Mais ce photon ne doit pas être tout le temps présent sinon c'est l'effet inverse qui se produit : le chat ne voit jamais sa parité changer. L'idée est alors de remplacer dans le montage précédent le prélèvement du deuxième chat par une superposition de vide et de photon unique (figure 4.13). En ajustant correctement les transmissions des lames et en conditionnant sur le bon couple (toujours avec des polarisations différentes) on obtient à nouveau la transformation $\nu\alpha\hat{1} + \mu\hat{a}$ mais cette fois pour l'état discret $\mu|1, 0\rangle + \nu|0, 1\rangle$. L'inconvénient de cette méthode est que l'on ne peut pas détecter la perte du photon, et c'est bien normal : la perte doit correspondre à un état bien distinct ce qui demande d'effectuer le codage sur une base de dimension supérieure à 2. Ce dernier protocole est donc plus simple mais bien moins efficace que le précédent, au final on

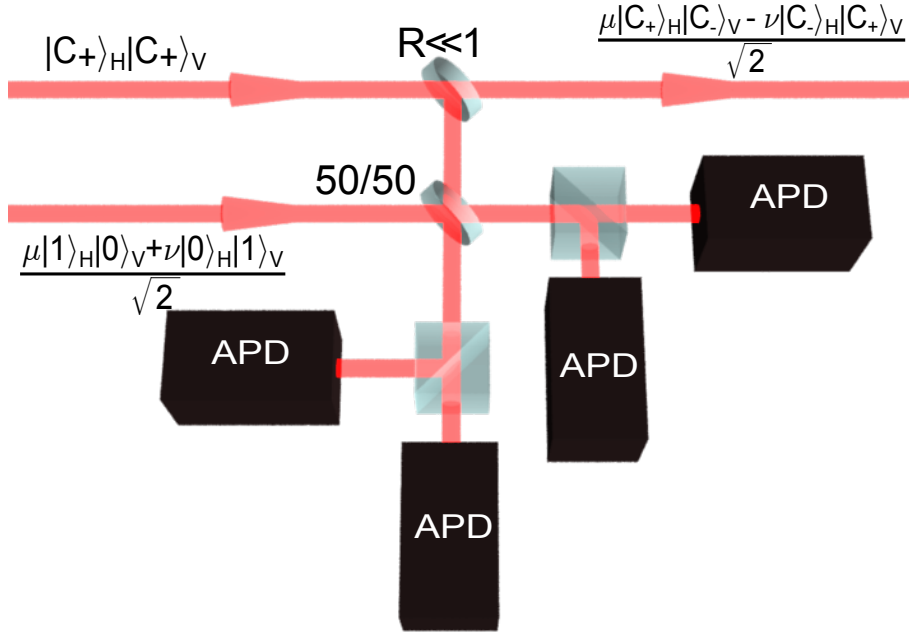


FIGURE 4.12: Protocole transformant un qubit discret encodée sur la polarisation d'un photon en un qubit continu encodé sur les parités de deux chats de Schrödinger

en revient toujours au même facteur limitant : la dimension de l'espace utilisé pour le codage.

4.4.3 Conclusion

Nous avons démontré qu'il est possible d'intriquer à distance deux états initialement indépendants et que cette intrication n'est pas affectée par les pertes du canal de transmission. Comparé à d'autres protocoles de communication utilisant les mêmes outils, cette méthode s'avère tout-à-fait efficace. Enfin nous avons aussi montré qu'en modifiant légèrement notre dispositif il est possible d'en faire un convertisseur de qubit discret vers continu utile pour des approches hybrides du calcul quantique optique.

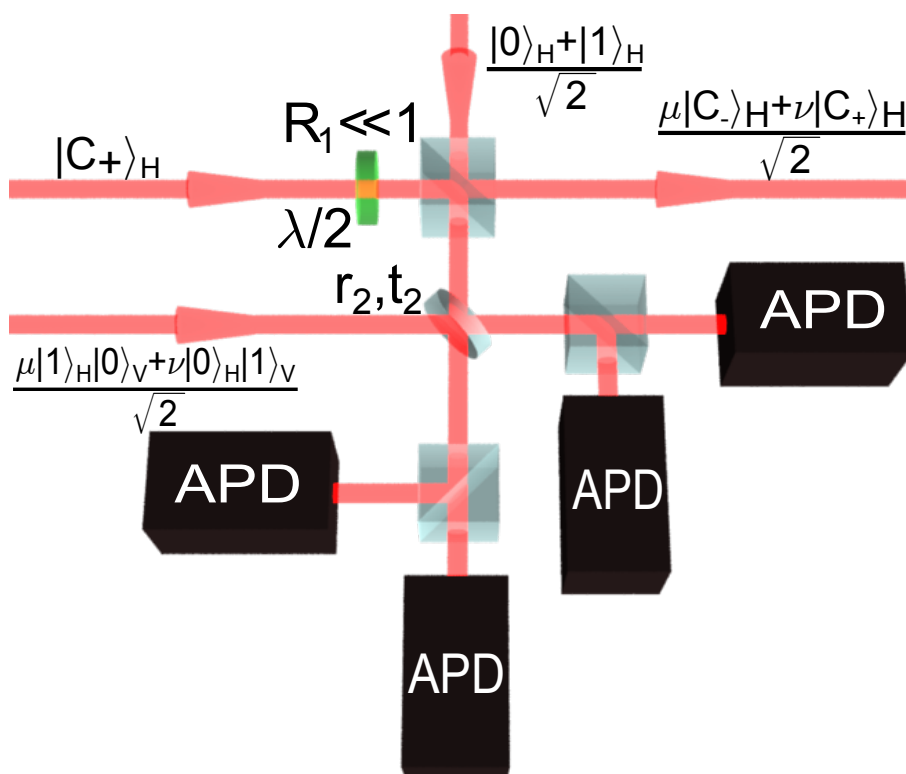


FIGURE 4.13: Protocole transformant un qubit discret encodée sur la polarisation d'un photon en un qubit continu encodé sur la parité d'un seul chat de Schrödinger

Chapitre 5

L'amplificateur sans bruit non déterministe

Sommaire

5.1 Amplification d'un signal quantique	97
5.1.1 Amplification déterministe	98
5.1.2 Principe de l'amplificateur sans bruit non déterministe	100
5.2 Réalisation expérimentale	106
5.2.1 Montage	106
5.2.2 Modélisation de l'expérience	108
5.2.2.1 Considérations générales	108
5.2.2.2 Modèle numérique	109
5.2.2.3 Modèle analytique	111
5.2.3 Résultats	115
5.3 Discussion	119
5.3.1 Cohérence avec les lois de la physique quantique	119
5.3.2 Autres expériences	120
5.3.3 Utilisations diverses	121
5.3.4 Retour sur le calcul quantique hybride	122
5.3.5 Conclusion	124

5.1 Amplification d'un signal quantique

À l'heure de la technologie moderne il est des processus physiques indispensables que l'on utilise tous les jours sans même en être conscient. L'amplification, dont la miniaturisation de la version électrique sous la forme d'un petit élément appelé transistor a permis l'essor de l'électronique, est de ceux-là. Cet élément indispensable à la plupart des systèmes de mesure et de communication se trouve non seulement partout dans nos laboratoires, et ce manuscrit en fait bien état, mais aussi dans la vie quotidienne : microphones, CCD, antennes, répéteurs dans les systèmes de communications (optique ou électrique), ... Le monde classique regorge de tels systèmes, mais qu'en est-il de son homologue quantique ?

Toutes les mesures d'un état quantique sont entachées de bruit. Pour des états cohérents ce bruit reste constant quelle que soit l'amplitude de ces états. Un signal fort est donc lisible de manière bien plus fiable qu'un signal faible, donnant ainsi toute son importance à l'amplification d'un signal quantique. Mais comment réaliser une amplification qui garderait le bruit quantique identique ? Un tel processus revient à vouloir transformer un état cohérent en un autre de plus grande amplitude :

$$\alpha \mapsto g\alpha \quad (5.1)$$

Il est appelée *amplification sans bruit*¹, et son action sur des états cohérents est représentée figure 5.1.

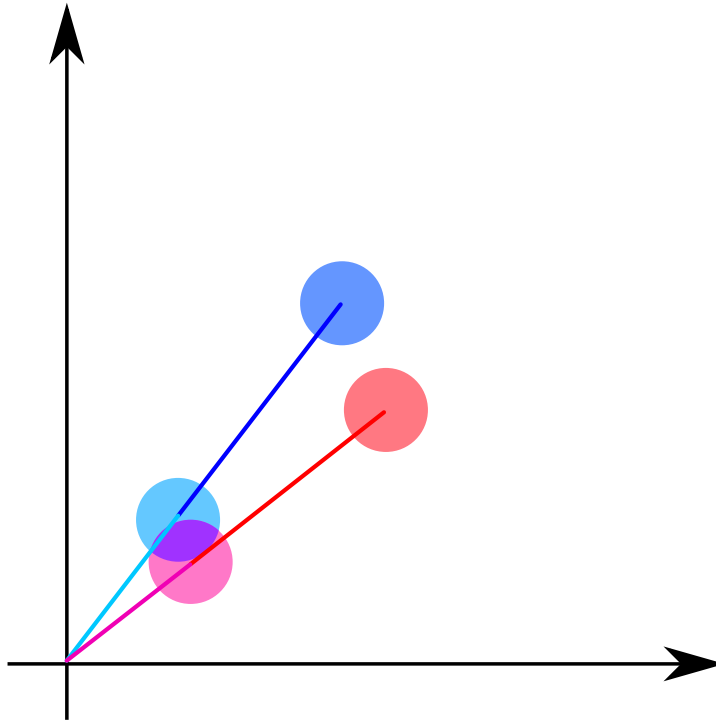


FIGURE 5.1: Action d'un amplificateur sans bruit sur des états cohérents.

5.1.1 Amplification déterministe

Un tel dispositif pose tout-de-même un important problème fondamental : imaginons que l'on ait un tel amplificateur avec un gain de deux en intensité, c'est-à-dire $g = \sqrt{2}$, et qu'en sortie on place une lame semi-réfléchissante 50/50 ; nous aurions alors en sortie deux copies identiques de l'état d'entrée. En d'autres termes nous aurions cloné parfaitement notre état, ce qui est interdit par la physique quantique [130, 131]. Il semble donc impossible d'effectuer une telle transformation, du moins dans le cadre « habituel » de la physique quantique où les évolutions se font à l'aide d'opérateurs unitaires et sont donc déterministes.

Cette impossibilité peut être démontrée de manière plus rigoureuse. Pour cela nous allons utiliser un raisonnement développé par Tim Ralph et Austin Lund [156] :

1. Cette dénomination a aussi parfois été donnée de manière moins stricte aux amplificateurs qui se contentent d'amplifier le bruit sans dégrader le rapport signal sur bruit, dans ce manuscrit nous réserverons cette appellation aux amplificateurs qui laissent le bruit identique.

1. Supposons qu'il existe une transformation unitaire $\hat{T}$ permettant d'effectuer cette amplification

$$\hat{T}|\alpha\rangle = c|g\alpha\rangle \quad (5.2)$$

où c est un nombre complexe de module 1. Il est introduit ici afin de rester le plus général possible, tout état quantique étant défini à une phase près.

2. Définissons maintenant l'opérateur d'annihilation $\hat{b} = \hat{T}\hat{a}\hat{T}^\dagger$. Si nous l'appliquons à l'état amplifié $|g\alpha\rangle$, nous obtenons :

$$\hat{b}|g\alpha\rangle = \frac{1}{c}(\hat{T}\hat{a}\hat{T}^\dagger)\hat{T}|\alpha\rangle \quad (5.3a)$$

$$= \frac{1}{c}\hat{T}\hat{a}|\alpha\rangle \quad (5.3b)$$

$$= \frac{1}{c}\alpha\hat{T}|\alpha\rangle \quad (5.3c)$$

$$= \alpha|g\alpha\rangle \quad (5.3d)$$

Les états cohérents sont donc états propres de cet opérateur avec comme valeur propre $\frac{1}{g}\alpha$. Nous pouvons donc en conclure que cet opérateur vaut

$$\hat{b} = \frac{1}{g}\hat{a} \quad (5.4)$$

3. De l'équation 5.4 nous pouvons déduire que

$$[\hat{b}, \hat{b}^\dagger] = \frac{1}{g^2} \quad (5.5)$$

4. Or à partir de la définition de l'opérateur $\hat{b}$ nous obtenons

$$[\hat{b}, \hat{b}^\dagger] = \hat{T}[\hat{a}, \hat{a}^\dagger]\hat{T}^\dagger = \hat{T}\hat{T}^\dagger = 1 \quad (5.6)$$

5. En combinant les deux équations précédentes nous arrivons donc à la contrainte $g = 1$: une telle amplification n'existe que dans le cas trivial où elle n'amplifie pas.

Nous avons donc bien prouvé que l'on ne pouvait pas amplifier un état cohérent avec une transformation unitaire, tout en conservant un état cohérent en sortie.

Pour autant il existe des amplificateurs déterministes, comme nous l'avons vu dans les chapitres précédents. Mais quelle est leur action sur le bruit ? De manière générale on peut exprimer l'action d'un amplificateur sur les quadratures par [133, 134]

$$\hat{X}_{\text{out}} = g_X \hat{X}_{\text{in}} + \hat{B}_X \quad (5.7)$$

$$\hat{P}_{\text{out}} = g_P \hat{P}_{\text{in}} + \hat{B}_P \quad (5.8)$$

g_X et g_P sont les gains pour chacune des quadratures, et $\hat{B}_X$ et $\hat{B}_P$ sont deux opérateurs hermitiens de valeur moyenne nulle qui représentent le bruit ajouté par l'amplificateur. Ils sont nécessaire afin de conserver le commutateur

$$[\hat{X}_{\text{out}}, \hat{P}_{\text{out}}] = i \quad (5.9)$$

Les bruits étant, par définition, indépendants du signal, ils commutent avec les opérateurs quadratures ; nous avons donc

$$[\hat{X}_{\text{out}}, \hat{P}_{\text{out}}] = g_X g_P [\hat{X}_{\text{in}}, \hat{P}_{\text{in}}] + [\hat{B}_X, \hat{B}_P] \quad (5.10)$$

Ce qui donne

$$[\hat{B}_X, \hat{B}_P] = -i(g_X g_P - 1) \quad (5.11)$$

On en déduit donc que ces opérateurs ajoutent des bruits dont les variances vérifient

$$\Delta \hat{B}_X \Delta \hat{B}_P = \frac{|g_X g_P - 1|}{2} \quad (5.12)$$

Ceci nous donne deux types d'amplificateurs possibles.

L'amplificateur indépendant de la phase Il s'agit du cas classique où l'amplification est la même pour les deux quadratures $g_X = g_P = g$. La réalisation la plus commune pour des états quantiques utilise un OPA non dégénéré, vu en section 3.3.3.1 : le mode signal contient le mode amplifié tandis que le mode complémentaire est ignoré. L'application du changement de variable de l'équation 3.9 modélisant l'action de celui-ci à un état cohérent, couplé à un mode vide sur lequel nous traçons ensuite, donne (cf. section A.2) la fonction de Wigner de l'état amplifié :

$$W_{\text{indep}} = \frac{1}{\pi(2g^2 - 1)} e^{-\frac{(x - \sqrt{2}g \text{Re}(\alpha))^2}{2g^2 - 1} - \frac{(p - \sqrt{2}g \text{Im}(\alpha))^2}{2g^2 - 1}} \quad (5.13)$$

Nous pouvons alors séparer la variance de cet état amplifié, qui vaut en l'occurrence $g^2 - \frac{1}{2}$, en deux contributions : $\frac{g^2}{2}$ et $\frac{g^2 - 1}{2}$. La première correspond à l'amplification du bruit initial de l'état cohérent par un facteur identique à celui du signal ; la deuxième est l'excès de bruit ajouté par l'amplificateur, qui correspond bien à celui donné par l'équation 5.12. Comme le montre la figure 5.2, il n'améliore pas la résolution de l'état quantique, au contraire il l'empire même.

L'amplificateur dépendant de la phase L'équation 5.12 nous montre qu'il est tout-de-même possible de ne pas ajouter de bruit à condition d'avoir des gains différents pour les deux quadratures [135, 136]. Plus précisément nous devons amplifier une quadrature et dé-amplifier l'autre d'un même facteur, c'est-à-dire $g_X = 1/g_P$. Autrement dit nous devons comprimer l'état d'entrée. Comme nous l'avons déjà dit, ceci peut-être réalisé, par exemple, à l'aide d'un OPA dégénéré. L'état en sortie étant alors un état comprimé dont la fonction de Wigner est donnée par l'équation 2.66. L'inconvénient est qu'il est nécessaire de connaître la phase de l'état d'entrée afin d'amplifier la bonne quadrature. Concernant la résolution, si elle n'est plus empirée, elle n'en est pas améliorée pour autant. En effet le bruit est amplifié / comprimé de la même manière que le signal (figure 5.3), conservant ainsi le même rapport signal sur bruit .

5.1.2 Principe de l'amplificateur sans bruit non déterministe

L'impossibilité d'amplifier sans bruit est due à la linéarité et à l'unitarité de l'évolution quantique, il nous faudrait donc un moyen de briser celles-ci. Or il existe justement un élément de la physique quantique qui entraîne de nombreux questionnements et donne naissance à des théories le plus diverses sur son interprétation pour la raison qu'il ne s'inscrit pas dans cette évolution linéaire et unitaire. Cet élément c'est la mesure. Plus précisément la partie de la mesure

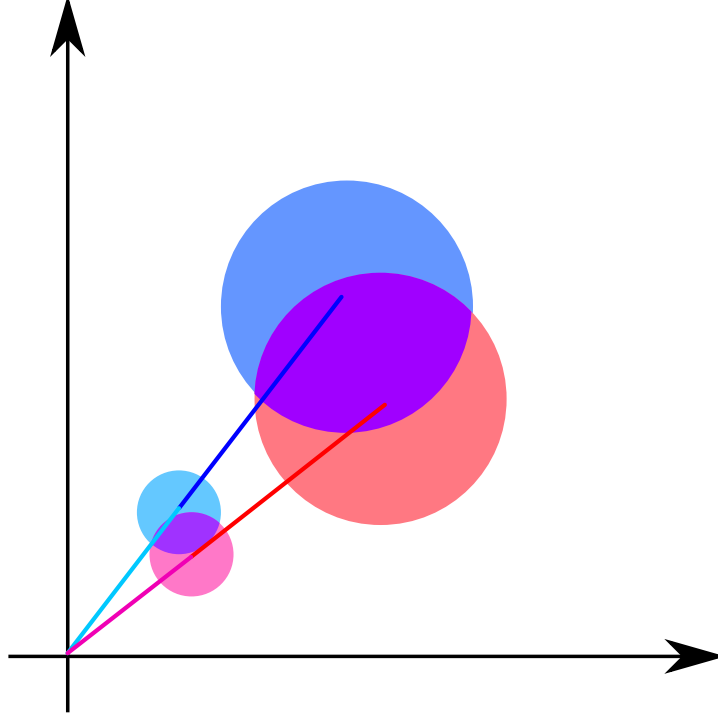


FIGURE 5.2: Action d'un amplificateur déterministe indépendant de la phase sur des états cohérents.

qui brise cette évolution est la projection sur l'état mesuré. Encore faut-il que celle-ci projette sur le bon état, ce qui est aléatoire ; nous perdons donc le côté déterministe.

Il a déjà été démontré que certains processus interdits de manière déterministe étaient réalisables dans une version probabiliste ; c'est notamment le cas du clonage [137, 138] et de l'augmentation locale d'intrication [139]. Afin d'étudier cette possibilité nous allons suivre un raisonnement inspiré de celui proposé par David Menzies et Sarah Croke [140]

1. Considérons maintenant la transformation conditionnée

$$|\alpha\rangle \langle\alpha| \mapsto \mathcal{T}_C(|\alpha\rangle \langle\alpha|) = |g\alpha\rangle \langle g\alpha| \quad (5.14)$$

qui, sans le conditionnement, correspond à la transformation déterministe

$$|\alpha\rangle \langle\alpha| \mapsto \mathcal{T}(|\alpha\rangle \langle\alpha|) = p_\alpha |g\alpha\rangle \langle g\alpha| + (1 - p_\alpha) \hat{\rho}_{\text{autres}} \quad (5.15)$$

où p_α est la probabilité de réussite et $\hat{\rho}_{\text{autres}}$ l'état obtenu lorsque la mesure ne donne pas le bon résultat. De telles transformations qui agissent sur les matrices densités sont appelées des *superopérateurs*.

2. Ces superopérateurs sont *complètement positifs* (i.e. les opérateurs issus des transformations sont tous complètement positifs), ce qui permet d'effectuer une *décomposition de Kraus* [141, 142] sur ces derniers. Celle-ci consiste à écrire l'action du superopérateur $\mathcal{T}_C$ sous la forme

$$\mathcal{T}_C(\hat{\rho}) = \sum_{k=0}^{\mathcal{N}} \hat{M}_k \hat{\rho} \hat{M}_k^\dagger \quad (5.16)$$

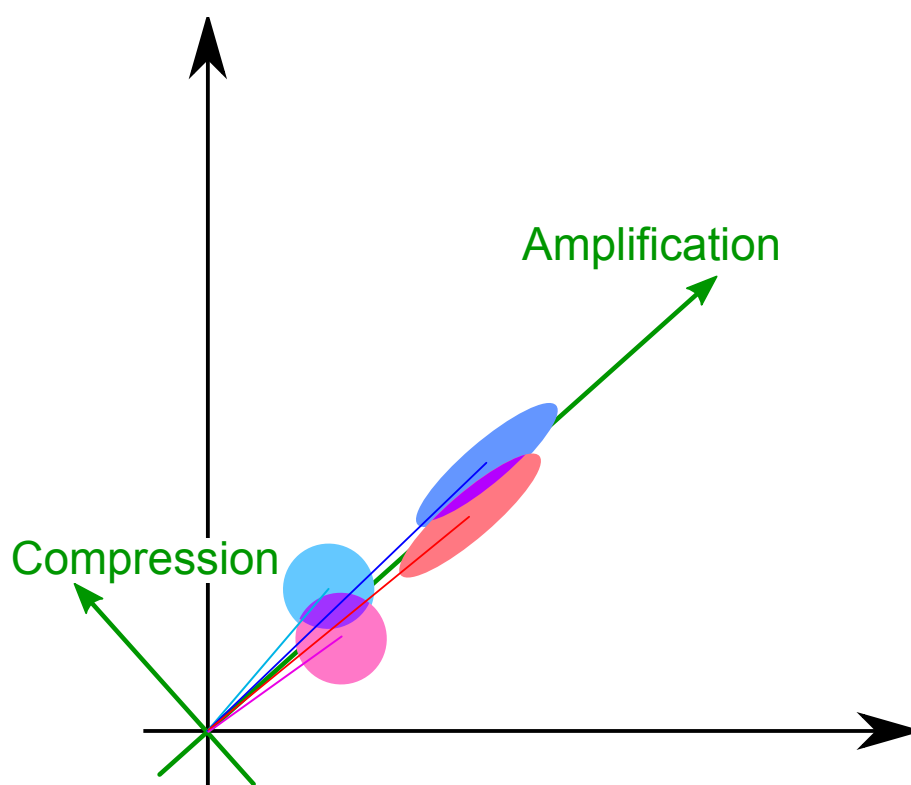


FIGURE 5.3: Action d'un amplificateur déterministe dépendant de la phase sur des états cohérents.

Les $\hat{M}_k$ sont des opérateurs non unitaires qui vérifient

$$\sum_{k=0}^N \hat{M}_k \hat{M}_k^\dagger \leq \hat{\mathbb{I}} \quad (5.17)$$

Le superopérateur $\mathcal{T}$ se décompose de la même façon sur un ensemble plus grand d'opérateurs, liés, dans notre cas, aux différentes projections de la mesure qui nous sert de conditionnement. L'équation 5.17 devient alors une égalité, ce qui correspond au fait que la somme des probabilités des résultats de la mesure vaut 1. L'équation 5.16 peut ainsi être vue comme une restriction de la décomposition de $\mathcal{T}$ aux seuls opérateurs correspondant au bon conditionnement.

3. Dans notre cas, il n'y qu'un seul opérateur correspondant au bon conditionnement et donc qui nous donne l'amplification. Une décomposition sur la base de Fock montre que cet opérateur vaut

$$\hat{M}_0 = d g^{\hat{n}} \quad (5.18)$$

où d est un facteur de proportionnalité lié à la probabilité de succès par $p_\alpha = |d|^2 e^{-(g^2-1)|\alpha|^2}$

4. L'équation 5.17 nous donne la contrainte $\hat{M}_0 \hat{M}_0^\dagger \leq \hat{\mathbb{I}}$, c'est-à-dire

$$|d|^2 \sum_{n=0}^{\infty} g^{2n} |n\rangle \langle n| \leq \sum_{n=0}^{\infty} |n\rangle \langle n| \quad (5.19)$$

Ce qui, pour un gain $g > 1$, n'est possible que si $d = 0$, autrement dit si la probabilité de réussir est nulle !

5. L'amplification sans bruit n'est toujours pas possible certes, mais nous n'avons pas dit notre dernier mot. Si nous ne pouvons pas la réaliser de manière exacte, alors essayons d'en effectuer une approximation. Pour cela nous allons tronquer l'espace de Hilbert à au plus N photons, négligeant ainsi les termes à plus grand nombre de photons dans la composition des états cohérents. L'opérateur s'écrit alors

$$\hat{M}_0 = d_N \sum_{n=0}^N g^n |n\rangle \langle n| \quad (5.20)$$

Nous pouvons alors nous limiter à la contrainte $|d_N|^2 < g^{-2N}$ qui autorise une probabilité de réussite non nulle.

Il est donc bel est bien possible de réaliser une amplification sans bruit, à condition que ce ne soit qu'une approximation non déterministe. Nous noterons que le taux succès décroît quand :

- le gain de l'amplificateur augmente.
- la dimension de l'espace de Hilbert augmente, c'est-à-dire quand on veut amplifier des états plus grands.

Amplification d'un petit état cohérent Tim Ralph et Austin Lund ont proposé un montage permettant de réaliser une telle amplification [143]. Intéressons-nous dans un premier temps à la version pour les petits états cohérent de leur proposition. L'approximation consiste à se placer dans un espace de Hilbert de dimension 2, et à écrire les états cohérents comme

une superposition du vide et d'un photon unique, ce qui nécessite $g\alpha \ll 1$. Dans ce cas la transformation à réaliser est (à un facteur de normalisation près) :

$$|0\rangle + \alpha |1\rangle \mapsto |0\rangle + g\alpha |1\rangle \quad (5.21)$$

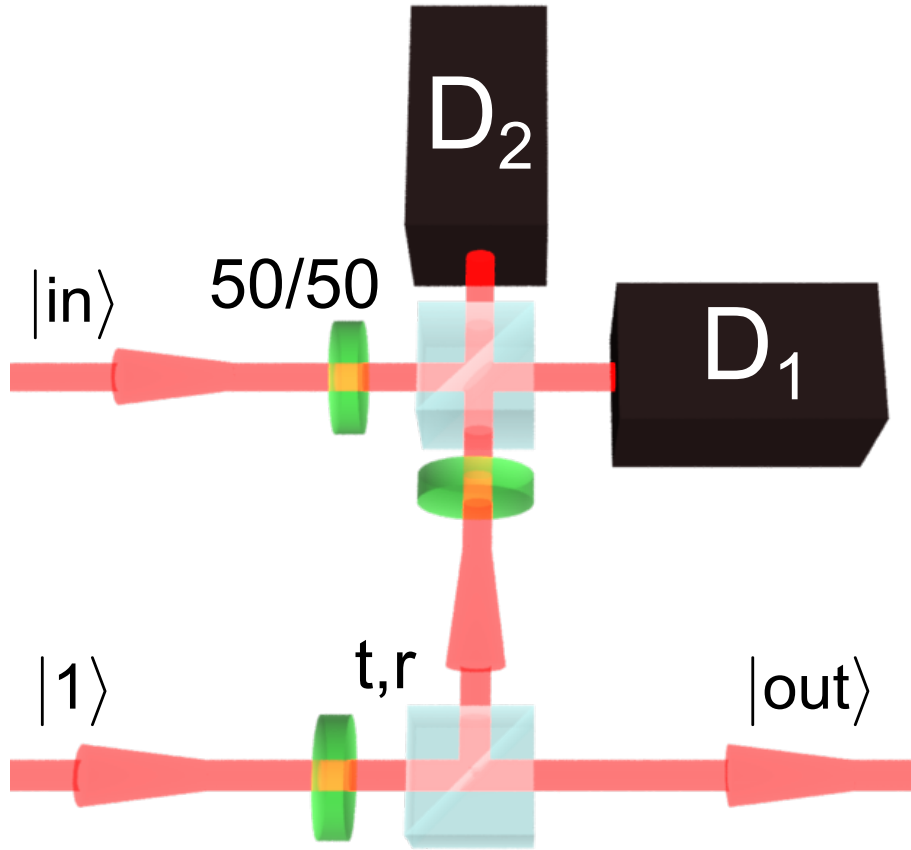


FIGURE 5.4: Principe de l'amplificateur sans bruit pour de petits états cohérents

La figure 5.4 montre comment réaliser cette transformation :

1. Un photon unique est envoyé sur une lame semi-réfléchissante de réflexivité r , la partie transmise constituant la sortie de notre amplificateur.
2. La partie réfléchie va ensuite interférer avec l'état cohérent à amplifier sur une lame semi-réfléchissante 50/50 ; les deux faisceaux en sortie sont alors envoyés sur deux compteurs de photons.
3. Lorsqu'un seul photon est détecté par l'ensemble des deux détecteurs, alors nous avons produit le bon état en sortie, moyennant éventuellement une inversion de la phase sur l'état en sortie suivant le bras dans lequel photon a été détecté.

On peut remarquer que le principe est similaire à celui d'une téléportation quantique [104, 105, 106, 52] : production d'un état intriqué, mesure conjointe d'un des modes de celui-ci avec l'état d'entrée et action « classique » sur l'autre mode en fonction du résultat de la mesure pour obtenir l'état de sortie. Il est aussi proche du protocole de troncature d'état appelé « ciseaux quantiques » [144, 145]

Il ne paraît pas évident au premier abord qu'un tel dispositif puisse sortir une version amplifiée de l'état d'entrée, même si la similarité avec la téléportation quantique tend à aller dans ce sens. Attardons-nous donc sur le fonctionnement. Lorsque l'on regarde différentes expériences de physique quantique on retrouve un certain nombre de processus « standards » ; parmi ceux-ci figure l'indiscernabilité du chemin suivi par un photon avant sa détection, déjà à la base du protocole présenté au chapitre précédant. Ce processus est à nouveau au cœur de notre amplificateur : en effet lorsqu'un photon est détecté il nous est impossible de savoir s'il vient de l'état cohérent ou du photon unique. Ce qui nous mène à une superposition quantique de deux possibilités :

- Soit il n'y avait pas de photon dans l'état cohérent, ce qui veut dire que le photon unique a été réfléchi : c'est lui que nous avons détecté, et donc on retrouve le vide en sortie.
- Soit il y avait un photon dans l'état cohérent, dans ce cas c'est forcément lui qui a été détecté et le photon unique a donc été transmis et se retrouve dans l'état de sortie.

Nous voyons donc que l'on a bien en sortie une superposition de vide et de photon unique qui est liée à la superposition en entrée. Cependant le rapport entre le vide et le photon unique dans la superposition est modifié : il dépend non seulement du rapport d'origine mais aussi de la réflexion de la lame, permettant ainsi de l'amplifier.

Cette explication en termes de photons montre bien l'amplification de l'amplitude, mais elle ne nous dit pas comment la phase est conservée, pour cela il faut évidemment revenir à une description ondulatoire. Tout est dans l'interférence entre l'état cohérent et le photon unique. Lorsqu'il n'y a qu'un seul photon, un seul des détecteurs capte un signal. En termes d'ondes, pour que l'on ait ce cas de figure, il faut que l'interférence soit constructive d'un côté et destructive de l'autre. Autrement dit les deux états qui interfèrent doivent soit avoir la même phase soit avoir des phases opposées en fonction du bras dans lequel le photon est détecté. Du coup tout se passe comme si le photon unique « choisissait une phase » correspondant à ce critère, c'est ce que Franck Laloë et William Mullin appellent *l'appariement spontané de la phase* [146], phase que l'on retrouve donc dans l'état en sortie.

Après cette présentation intuitive, effectuons le calcul. Nous avons donc un état $r|0\rangle_1|1\rangle_2 + t|1\rangle_1|0\rangle_2$ dont le mode 2 interfère avec $|0\rangle_3 + \alpha|1\rangle_3$, puis on projette soit sur $|0\rangle_2|1\rangle_3$, soit sur $|1\rangle_2|0\rangle_3$. En inversant le signe pour l'une des projections nous obtenons le même résultat pour les deux, à savoir :

$$e^{-\frac{|\alpha|^2}{2}} \frac{r}{\sqrt{2}} \left(|0\rangle + \frac{\sqrt{1-r^2}}{r} \alpha |1\rangle \right) \quad (5.22)$$

Nous avons donc bien l'état voulu en posant

$$g = \frac{\sqrt{1-r^2}}{r} \quad (5.23)$$

Le dispositif amplifie bel est bien pour une réflexion inférieure à 50% ; de plus il fonctionne pour n'importe quel gain, dans la limite $g\alpha \ll 1$. La probabilité de réussite de l'amplification vaut

$$P = e^{-|\alpha|^2} r^2 \left(1 + g^2 |\alpha|^2 \right) \approx r^2 \quad (5.24)$$

On retrouve aussi le comportement déjà évoqué lorsque le gain augmente : il faut alors baisser la réflexion, ce qui diminue le nombre de photons allant vers les détecteurs de conditionnement et donc le taux de succès.

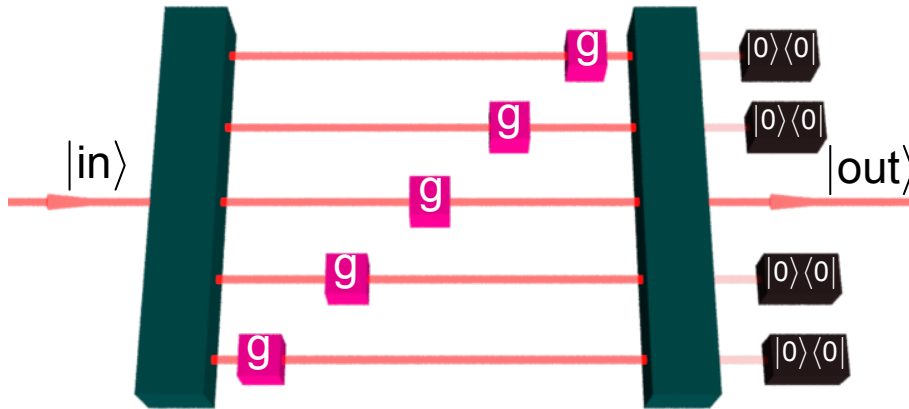


FIGURE 5.5: Principe de l'amplificateur sans bruit pour des états cohérents plus grands

Amplification d'un état cohérent plus grand Le passage à des états cohérents plus grands se fait en divisant l'état cohérent en N états suffisamment petits pour être amplifiés par le dispositif précédent (figure 5.5) [143]. Une fois que chacun de ces états a été amplifié, ils sont alors recombinaés de manière conditionnelle. Les états en sortie n'étant en effet que des approximations d'états cohérents, il est nécessaire de vérifier que tous les photons vont bien dans la même voie en conditionnant sur l'absence de photons dans les mauvaises voies. On retrouve là la deuxième partie de notre réflexion précédente sur le taux de succès : plus on divise l'état d'entrée, et donc plus on augmente la taille de l'espace de Hilbert, plus il faut de coïncidences et donc plus la probabilité de réussite diminue.

5.2 Réalisation expérimentale

5.2.1 Montage

L'amplification d'un petit état cohérent peut donc être vue comme un élément de base pour l'amplification sans bruit, et c'est cet élément que nous avons cherché à mettre en oeuvre pour démontrer expérimentalement la faisabilité de l'amplification sans bruit. Rappelons le principe de cet élément de base, décrit sur la figure 5.4 : un photon unique, dont la génération a déjà été exposée section 3.5, est séparé sur une séparatrice asymétrique ; l'une des voies de sortie de cette séparatrice constitue la sortie de l'amplificateur, que nous analyserons avec une détection homodyne, tandis que l'autre est mélangée sur une séparatrice 50/50 avec un état cohérent fortement atténué. Le succès de l'amplification est conditionné à un événement de détection sur les voies de sortie de cette dernière séparatrice.

Comme nous l'avons vu précédemment, la faible transmission des filtres situés avant les détecteurs implique une probabilité négligeable de détecter plus de un photon, nous permettant ainsi de remplacer les compteurs de photons par de simples APD. Mais dans ce cas précis nous pouvons même aller plus loin : il est en effet tout aussi peu probable de détecter un photon sur chacun des deux détecteurs ; par conséquent il nous est possible d'utiliser un seul détecteur sans que cela ne modifie de manière significative le résultat. Nous nous contenterons donc d'utiliser un seul détecteur en sortie de la lame 50/50, sur la voie qui ne nécessite pas d'inverser la phase. Si cela simplifie le dispositif nous devons néanmoins noter que le taux de succès est divisé par deux.

Nous pouvons tout de même nous demander si notre incapacité à résoudre le nombre de

photon ne détériore pas trop l'état en sortie. Pour cela voyons dans quels cas le conditionnement peut se faire sur plus d'un photon. Si l'approximation des petits états cohérent est valide, la probabilité d'avoir plus d'un photon venant de l'état cohérent est négligeable. De même nous pouvons régler le gain de l'OPA de manière à avoir suffisamment peu d'états à plus de un photon dans notre « photon unique ». La seule possibilité à prendre en compte est celle où nous avons un photon dans chacun des deux. Ce cas de figure arrive avec une probabilité proportionnelle à $r\alpha$, ce qui n'est pas négligeable devant celle du photon venant de l'état cohérent ($\propto \alpha$). Seulement dans le premier cas de figure nous obtenons en sortie le vide et dans le deuxième cas un photon unique. Et l'approximation des petits états cohérent implique que le vide soit bien plus présent ; dans le cas idéal, sa probabilité est proportionnelle à r . Ainsi en reprenant les calculs de la section précédente et en y ajoutant un conditionnement sur ce nouveau cas de figure nous obtenons en sortie la matrice densité suivante :

$$\hat{\rho} \propto \begin{pmatrix} 1 + |\alpha|^2 & g\alpha \\ g\alpha^* & g^2 |\alpha|^2 \end{pmatrix} \approx \begin{pmatrix} 1 & g\alpha \\ g\alpha^* & g^2 |\alpha|^2 \end{pmatrix} \quad (5.25)$$

qui reste une bonne approximation de celle d'un état cohérent d'amplitude $g\alpha$.

Il nous reste encore un problème à régler, à savoir la stabilité de la phase entre les divers modes. Le schéma de principe de la figure 5.4 nous montre que trois modes entrent en jeu, nous devons donc contrôler la différence de phase entre ces trois modes et l'oscillateur local. Nous connaissons déjà l'astuce consistant à encoder deux modes sur des polarisations orthogonales d'un même mode spatial. Nous pouvons utiliser cette méthode sur deux couples de modes, mais il nous faut une autre technique pour contrôler la phase entre ces deux couples. La solution retenue consiste à utiliser, en jouant sur la polarisation, une partie de l'état cohérent d'entrée comme référence de phase. Pour reprendre en détails le fonctionnement de cette expérience, dont le montage complet est montré figure 5.6 :

- Nous commençons par un photon unique qui n'a pas de phase définie et ne nous pose donc pas de problème jusqu'à sa séparation.
- En utilisant des lames demi-ondes et un cube séparateur pour la lame semi-réfléchissante asymétrique nous pouvons superposer spatialement (mais avec des polarisations orthogonales) l'état à amplifier avec la partie réfléchie du photon unique². La phase de cette dernière reste alors constante par rapport à la référence.
- Toujours grâce au même cube la partie transmise de l'état cohérent, qui sert donc de référence de phase, est quant à elle superposée avec la partie transmise du photon. Cette dernière a donc elle aussi une phase constante par rapport à la référence.
- Ces états transmis sont alors redirigés sur le cube séparateur de polarisation situé à l'entrée des deux détecteurs homodynes où chaque mode est envoyé sur l'une des détecteurs. La différence de phase entre les deux oscillateurs locaux est réglée pour être nulle (cf. section 3.4.1.3), et on peut alors se servir des mesures homodynes effectuées sur l'état cohérent de référence pour fixer la phase de l'oscillateur local utilisé par l'autre détection homodyne et analyser ainsi l'état de sortie.

Cette solution expérimentale soulève néanmoins un nouveau problème : celui de la mesure de l'état d'entrée. À l'origine nous espérions le mesurer en l'envoyant sur la détection homodyne, permettant ainsi d'effectuer la mesure dans les mêmes conditions que pour l'état amplifié, mais la présence du faisceau de référence envoyé sur la deuxième détection rend cette technique bien plus compliquée. En effet il faut inverser parfaitement le partage de l'état cohérent entre APD

2. Cette technique permet en outre de pouvoir régler le gain de l'amplificateur en tournant la lame demi-onde situé sur le trajet du photon unique

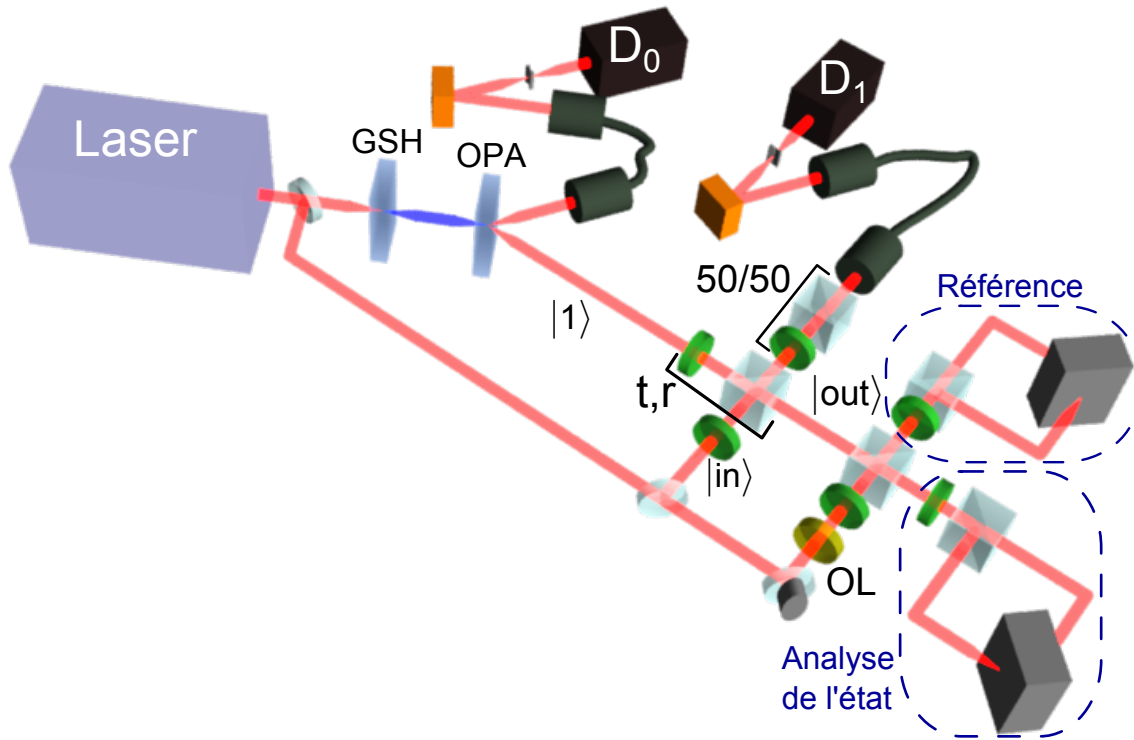


FIGURE 5.6: Montage expérimental de l'amplificateur sans bruit pour de petits états cohérents.

et deuxième détection homodyne, ce qui demande à chaque fois de repasser par des faisceaux non atténués et donc d'effectuer un certain nombre de réglages supplémentaires. Il s'est avéré que la méthode la plus simple consiste à mesurer l'état cohérent avec l'APD et à prendre en compte les différentes pertes par le calcul pour comparer les états sans et avec l'amplificateur. À partir du moment où nous connaissons les pertes de la voie APD (section 3.4.2) et les pertes homodynes (section 3.4.1), cette méthode ne pose aucun problème.

5.2.2 Modélisation de l'expérience

Nous connaissons le comportement de l'amplificateur dans le cas idéal, mais que se passe-t-il dans le cas réel, avec les imperfections expérimentales, et qu'arrive-t-il lorsque l'approximation des petits états cohérents n'est plus valable ? L'expérience nous le dira, mais avoir un modèle qui nous le prédit s'avère très utile, que se soit pour les réglages ou pour l'analyse des données.

5.2.2.1 Considérations générales

Afin d'élaborer un modèle la première étape consiste à se demander quelles sont les imperfections qu'il faut prendre en compte :

- La première imperfection réside dans la production du photon unique ; celle-ci a déjà été analysée section 3.5, et nous pourrions donc en utiliser les résultats. Les paramètres δ et σ qui le caractérise peuvent être déterminé soit en envoyant l'intégralité du faisceau contenant le photon unique sur une détection homodyne, soit en regardant directement le photon unique atténué par la lame séparatrice asymétrique, ce dernier est alors caractérisé

par les paramètres

$$\sigma_t^2 = t^2 \sigma^2 + 1 - t^2 \quad (5.26)$$

$$\delta_t = t^2 \delta \frac{\sigma^2}{\sigma_t^2} \quad (5.27)$$

- Ensuite vient l’adaptation spatio-temporelle des deux modes devant interférer. Nous pouvons la modéliser par un coefficient de recouvrement (ou *mode-matching*) ϵ . L’équation 3.30 nous donnant le « photon unique » dans le mode détecté par la détection homodyne et l’APD, c’est l’état cohérent que nous décomposons en un mode $\hat{a}_{//}$ interférant avec le photon unique et un mode orthogonal $\hat{a}_{\perp}$:

$$\hat{a}_{\alpha} = \sqrt{\epsilon} \hat{a}_{//} + \sqrt{1 - \epsilon} \hat{a}_{\perp} \quad (5.28)$$

Tout comme ceux que nous avons vu précédemment, notamment pour la modélisation de l’efficacité homodyne, ce recouvrement pourra être estimé à l’aide du contraste des interférences. Cela nécessite cependant de ne pas atténuer le faisceau et donc de faire l’hypothèse qu’une fois les atténuateurs mis en place le mode-matching reste identique.

- Les dernières imperfections viennent des systèmes de détection, à commencer par la transmission μ de la voie APD. Son effet consiste essentiellement à réduire le taux de succès, comme nous le verrons par la suite.
- Nous l’avons déjà vu, l’imperfections de l’APD qui modifie le plus l’état en sortie est la pureté modale ξ' . Seulement, dans le cas qui nous intéresse, le recouvrement ϵ peut également être vu comme une source de coups parasites : la détection d’un photon dans le mode $\hat{a}_{\perp}$ est en effet complètement décorrélée de la sortie de l’amplificateur, et a la même incidence qu’un coup parasite. La contribution du recouvrement peut être modélisée de la même façon, et nous définirons une pureté modale effective ξ_{eff} prenant en compte toutes les sources de coups parasites. L’état obtenu lors d’un mauvais conditionnement est l’état non-conditionné ; dans le cas présent il s’agit du photon unique avec des pertes correspondant à la réflexion de la première lame semi-réfléchissante.
- Enfin la dernière imperfection à prendre en compte est l’efficacité homodyne η déjà mentionnée dans les chapitres précédents.

5.2.2.2 Modèle numérique

À partir de ces considérations nous pouvons, dans un premier temps, réaliser un modèle numérique de l’expérience. Ce modèle nous a permis d’avoir une première idée de ce que nous allions obtenir, ainsi que les valeurs nécessaires et/ou suffisantes des paramètres quantifiant les diverses imperfections afin d’obtenir un résultat convenable. Il nous a de plus servi à faire les premières analyses des résultats [147].

Deux modèles numériques ont en réalité été utilisés pour cette expérience. Le premier repose sur un certain nombre de calculs en base de Fock effectués par Rémi Blandino et Simon Fossier et donnant les formules et algorithmes à appliquer pour obtenir la matrice densité de l’état amplifié. C’est ce modèle qui a permis de prédire les résultats attendus et de définir quelles optimisations du dispositif expérimental étaient nécessaires. Son côté « hybride » composé à la fois de calculs numériques et de formules issues de calculs analytiques donne un très bon compromis entre le temps mis pour développer le modèle et le temps pris par son utilisation, spécialement lors du travail préparatoire d’une expérience nécessitant d’identifier l’impact des différentes imperfections. Il a en outre permis de poser les jalons pour les autres modèles.

Les étapes constituant ce modèle sont les suivantes (cf. figure 5.7) :

1. Lors d'un calcul en base de Fock, la première chose à faire consiste à définir la taille de l'espace de Hilbert, c'est-à-dire le nombre maximal de photons que l'on va considérer. Dans le cas présent la limite est fixée à 2.
2. Nous partons d'une part de la matrice densité du « photon unique imparfait », qui est en réalité un mélange de vide, d'état à un photon et d'état à deux photons (puisque l'on s'arrête à ce nombre de photons). Puis nous effectuons le mélange avec le vide sur la lame semi-réfléchissante asymétrique.
3. D'autre part nous avons un état cohérent, tronqué lui aussi à 2 photons, que nous allons également mélanger sur une lame semi-réfléchissante asymétrique avec le vide afin de modéliser le mode-matching.
4. Nous effectuons ensuite le mélange sur la séparatrice 50/50; puis enfin la détection, consistant à mesurer au moins un photon dans les modes incidents sur une APD, en prenant en compte la transmission imparfaite du canal.

Nous obtenons alors la matrice densité de l'état produit, et il ne reste plus qu'à prendre en compte les pertes homodynes pour obtenir l'état mesuré. Concernant les mauvais clics du conditionnement, on ne considère dans ce modèle que la partie due au mauvais mode-matching.

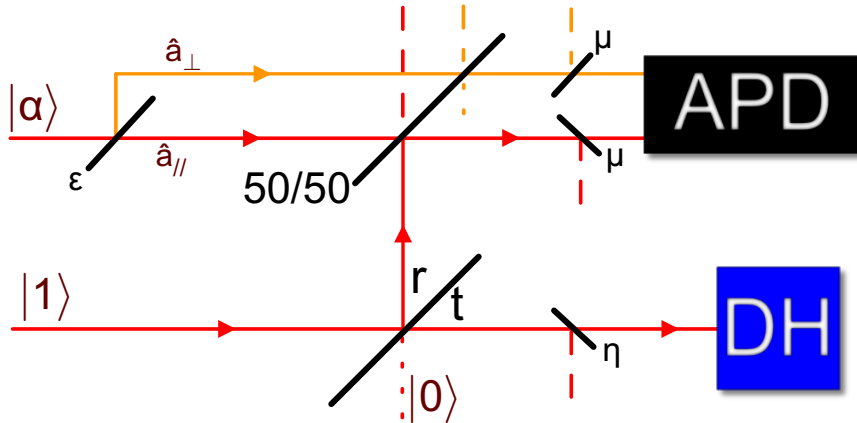


FIGURE 5.7: *Modélisation de l'amplificateur sans bruit*

Le deuxième modèle numérique utilisé est cette fois entièrement numérique : il consiste simplement à effectuer le processus décrit précédemment à l'aide d'outils numériques [148] sans effectuer de calculs préalables. Le temps nécessaire au modèle pour donner un résultat est par conséquent plus long, par contre sa complexification, par exemple en augmentant la taille de l'espace de Hilbert, ne demande généralement que de modifier quelques lignes de code. Ce modèle a été utilisé pour analyser les données expérimentales en les comparant au résultat du modèle. Un tel modèle est en effet idéal dans ce cas dans la mesure où nous pouvons déterminer de manière indépendante les valeurs des imperfections (cf. chapitre 3) ; le principal degré de liberté sur le modèle est alors la complexité du calcul, et nous avons d'ailleurs augmenté la taille de l'espace de Hilbert jusqu'à 5 photons.

5.2.2.3 Modèle analytique

Si les modèles numériques sont très pratiques, pour une analyse complète une formule analytique de la fonction de Wigner de l'état produit est bien plus parlante. Un tel calcul est cependant plus complexe et demande par conséquent plus de temps. En contrepartie une fois les formules obtenues leur application est extrêmement rapide, sans compter que dans bien des cas il est possible de remarquer un certain nombre de propriétés « à l'œil ».

Le modèle analytique développé lors de ce travail de thèse suit la même logique que le modèle numérique à quelques exceptions près. La première d'entre elles est qu'il n'est pas nécessaire de couper l'espace de Hilbert en restreignant le nombre de photons.

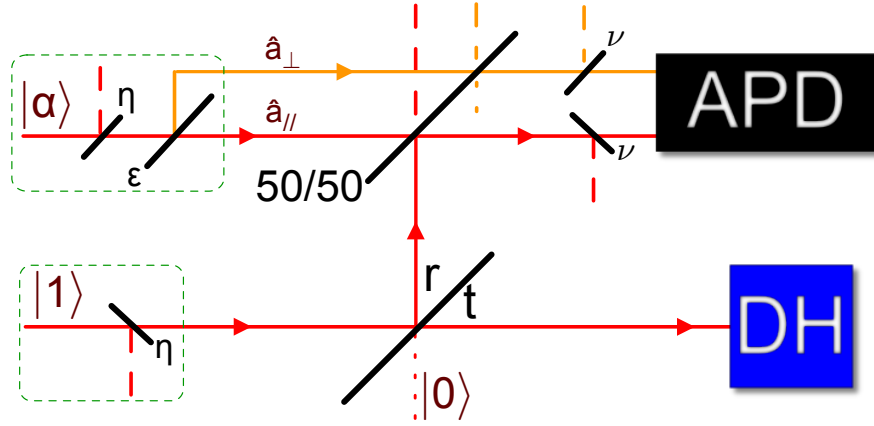


FIGURE 5.8: Modélisation équivalente de l'amplificateur sans bruit

Fonction de Wigner Avant de s'attaquer aux calculs du modèle il convient de faire quelques modifications afin de rendre ceux-ci un peu plus simple sans pour autant en changer le résultat. Ces modifications consistent à « déplacer les fuites » de la même manière que nous l'avons fait aux chapitres précédents. On peut en effet remarquer que le schéma du modèle utilisé jusqu'ici, figure 5.7, est identique à celui de la figure 5.8 où les pertes homodynes ont été déplacées en entrée et les pertes de la voie APD ont été remplacées par $\nu = \frac{\mu}{\eta}$ afin de compenser. De plus nous utiliserons cette fois la pureté modale effective ξ_{eff} définie plus haut, ce qui nous permet, dans un premier temps, de considérer le mauvais mode-matching uniquement comme une perte.

Une fois ces considérations prises en compte nous avons deux états en entrée :

- Le photon unique dont la fonction de Wigner est donnée par l'équation 3.30, les deux paramètres δ et σ (dont les expressions sont données aux équations 3.31 et 3.32.) qui la caractérisent pouvant être facilement mesurés expérimentalement.
- Un état cohérent ayant une amplitude effective $\alpha_{\text{eff}} = \sqrt{\eta}\sqrt{\epsilon}\alpha$ où α est l'amplitude de l'état à amplifier.

À partir de là nous pouvons dérouler les calculs du modèle :

1. La première étape consiste à mélanger le photon unique avec le vide sur une lame semi-réfléchissante de réflectivité r et de transmittivité t . Ce qui donne une fonction de Wigner bi-mode

$$W_{\text{asym}}(x, p, x_1, p_1) = W_{\text{phot}}(tx - rx_1, tp - rp_1) W_0(rx + tx_1, rp + tp_1) \quad (5.29)$$

2. Nous pouvons désormais faire interférer le mode réfléchi, avec l'état cohérent puis tracer sur la voie de sortie qui ne contient pas d'APD. Ceci nous donne alors une nouvelle fonction de Wigner bi-mode plus complexe.

$$\begin{aligned}
W_{\text{sym}}(x, p, x_{\text{APD}}, p_{\text{APD}}) &= \frac{1}{\pi^2(\sigma^2 + \sigma_t^2)} \left(2 \left(1 - \frac{(1+t^2)\delta\sigma^2}{\sigma^2 + \sigma_t^2} \right) \right. \\
&\quad \left. + \frac{4\delta\sigma^2}{(\sigma^2 + \sigma_t^2)^2} \left(\left(\sqrt{2}tx + r(x_{\text{APD}} - \alpha_{\text{eff}}) \right)^2 + (\sqrt{2}tp + rp_{\text{APD}})^2 \right) \right) \\
&\quad \times e^{-\frac{x^2 + t^2(x - \sqrt{2}g(x_{\text{APD}} - \alpha_{\text{eff}}))^2 + t^2\sigma^2(gx + \sqrt{2}(x_{\text{APD}} - \alpha_{\text{eff}}))^2}{\sigma^2 + \sigma_t^2}} \\
&\quad \times e^{-\frac{p^2 + t^2(p - \sqrt{2}gp_{\text{APD}})^2 + t^2\sigma^2(gp + \sqrt{2}p_{\text{APD}})^2}{\sigma^2 + \sigma_t^2}}
\end{aligned} \tag{5.30}$$

3. Revenons maintenant à un état monomode en appliquant le conditionnement sur le mode de l'APD. Nous effectuerons celui-ci de la même manière que pour le photon unique (section 3.5) en appliquant l'opérateur d'annihilation puis en traçant. Nous obtenons ainsi la fonction de Wigner de l'état pour un bon conditionnement.
4. Il ne nous reste plus qu'à ajouter le mauvais conditionnement. L'état final est un mélange statistique de l'état bien conditionné avec une probabilité ξ_{eff} et d'un photon unique atténué avec une probabilité $1 - \xi_{\text{eff}}$.

Au final nous avons une fonction de Wigner de la forme

$$W_{\text{ampl}}(x, p) = \frac{1}{\pi\sigma_t^2} \left(c_0 + c_1x + (c_2 + c_3x)(x^2 + p^2) + c_4(x^2 + p^2)^2 \right) e^{-\frac{x^2 + p^2}{\sigma_t^2}} \tag{5.31}$$

où les coefficients valent :

$$\begin{aligned}
c_0 &= -\xi_{\text{eff}} \frac{1 - 2\delta_t + (1 + 2g^2\alpha_{\text{eff}}^2)(\delta_t - 1)\sigma_t^2}{\sigma_t^2(2g^2\alpha_{\text{eff}}^2 + (1 + \delta_t)\sigma_t^2 - 1)} \\
&\quad + (1 - \xi_{\text{eff}})(1 - \delta_t)
\end{aligned} \tag{5.32}$$

$$c_1 = \xi_{\text{eff}} \frac{2\sqrt{2}g\alpha_{\text{eff}}(\sigma_t^2 - 1 + \delta_t(2 - \sigma_t^2))}{\sigma_t^2(2g^2\alpha_{\text{eff}}^2 + (1 + \delta_t)\sigma_t^2 - 1)} \tag{5.33}$$

$$\begin{aligned}
c_2 &= \xi_{\text{eff}} \frac{(\sigma_t^2 - 1)^2 - \delta_t(4 - (5 + 2g^2\alpha_{\text{eff}}^2))\sigma_t^2 + \sigma_t^4}{\sigma_t^4(2g^2\alpha_{\text{eff}}^2 + (1 + \delta_t)\sigma_t^2 - 1)} \\
&\quad + (1 - \xi_{\text{eff}}) \frac{\delta_t}{\sigma_t^2}
\end{aligned} \tag{5.34}$$

$$c_3 = \xi_{\text{eff}} \frac{2\sqrt{2}g\alpha_{\text{eff}}\delta_t(\sigma_t^2 - 1)}{\sigma_t^4(2g^2\alpha_{\text{eff}}^2 + (1 + \delta_t)\sigma_t^2 - 1)} \tag{5.35}$$

$$c_4 = \xi_{\text{eff}} \frac{\delta_t(\sigma_t^2 - 1)^2}{\sigma_t^6(2g^2\alpha_{\text{eff}}^2 + (1 + \delta_t)\sigma_t^2 - 1)} \tag{5.36}$$

g étant le gain théorique de l'amplificateur dont l'expression, en fonction de la réflexion de la lame séparatrice asymétrique, est donnée à l'équation 5.23.

Gain effectif À partir de l'équation 5.31 il est possible de calculer le gain effectif de l'amplificateur, c'est-à-dire le gain que l'on trouve sur la valeur moyenne des quadratures en mesurant celles-ci avec le même dispositif (et donc les mêmes pertes).

$$g_{\text{eff}} = \frac{\langle \hat{X}_{\text{out}} \rangle}{\langle \hat{X}_{\text{in}} \rangle} \quad (5.37)$$

avec $\langle \hat{X}_{\text{in}} \rangle = \sqrt{2\eta}\alpha$. Ce qui nous donne

$$g_{\text{eff}} = \frac{\sigma_t^2 (c_1 + 2c_3\sigma_t^2)}{2\sqrt{2}\sqrt{\eta}\alpha} \quad (5.38)$$

$$g_{\text{eff}} = \frac{g\sqrt{\epsilon}\xi_{\text{eff}}((1+\delta_t)\sigma_t^2 - 1)}{2g^2\alpha_{\text{eff}}^2 + (1+\delta_t)\sigma_t^2 - 1} \quad (5.39)$$

Variance De la même manière nous pouvons calculer la variance des différentes quadratures³

$$\Delta^2 X_\theta = V_{\text{const}} + V_{\text{cos}} \cos^2(\theta) \quad (5.40)$$

$$V_{\text{const}} = \frac{\sigma_t^2}{2} (c_0 + 2c_2\sigma_t^2 + 6c_4\sigma_t^4) \quad (5.41)$$

$$V_{\text{cos}} = -\frac{\sigma_t^4}{4} (c_1 + 2c_3\sigma_t^2)^2 \quad (5.42)$$

ou en fonction des paramètres

$$V_{\text{const}} = \frac{\xi_{\text{eff}} + (1+\delta_t)(2g^2\alpha_{\text{eff}}^2 - 2\xi_{\text{eff}} - 1)\sigma_t^2 + ((1+\delta_t)^2 + (1-(\delta_t-2)\delta_t)\xi_{\text{eff}})\sigma_t^2}{2(2g^2\alpha_{\text{eff}}^2 + (1+\delta_t)\sigma_t^2 - 1)} \quad (5.43)$$

$$V_{\text{cos}} = -\frac{2g^2\alpha_{\text{eff}}^2\xi_{\text{eff}}^2((1+\delta_t)\sigma_t^2 - 1)^2}{(2g^2\alpha_{\text{eff}}^2 + (1+\delta_t)\sigma_t^2 - 1)^2} \quad (5.44)$$

Taux de succès Pour obtenir la probabilité de réussite nous devons revenir à l'équation 5.30 et tracer sur le mode de sortie. En appliquant les pertes ν nous obtenons alors la fonction de Wigner de l'état vu par l'APD

$$W_{\text{APD}}(x_{\text{APD}}, p_{\text{APD}}) = \frac{1}{\pi\sigma_{\text{APD}}^2} \left(1 - \delta_{\text{APD}} + \delta_{\text{APD}} \frac{(x_{\text{APD}} - \sqrt{\frac{\nu}{2}}\alpha_{\text{eff}})^2 + p_{\text{APD}}^2}{\sigma_{\text{APD}}^2} \right) \times e^{-\frac{(x_{\text{APD}} - \sqrt{\frac{\nu}{2}}\alpha_{\text{eff}})^2 + p_{\text{APD}}^2}{\sigma_{\text{APD}}^2}} \quad (5.45)$$

où

$$\sigma_{\text{APD}}^2 = \nu \frac{r^2}{2} (\sigma^2 - 1) + 1 \quad (5.46)$$

$$\delta_{\text{APD}} = \nu \frac{r^2}{2} \delta \frac{\sigma^2}{\sigma_{\text{APD}}^2} \quad (5.47)$$

3. qui, rappelons le, vaut $\frac{1}{2}$ pour un état cohérent avec nos conventions.

Cette fois nous pouvons laisser de côté l'approximation considérant qu'un seul photon est détecté et calculer la probabilité de détecter au moins un photon à partir de l'équation 2.49

$$P_{\text{id}} = 1 - 2 \frac{1 + (2 + \sigma_{\text{APD}}^2 + \delta_{\text{APD}}(\mu\epsilon\alpha^2 - \sigma_{\text{APD}}^2 - 1)) \sigma_{\text{APD}}^2}{(1 + \sigma_{\text{APD}}^2)^3} e^{-\frac{\mu\epsilon\alpha^2}{1 + \sigma_{\text{APD}}^2}} \quad (5.48)$$

Cette probabilité ne prend pas en compte les mauvais conditionnements, la vraie probabilité de détecter un photon s'obtient simplement en remarquant que

$$P_{\text{ampl}} = \frac{P_{\text{id}}}{\xi_{\text{eff}}} \quad (5.49)$$

Pureté modale effective Revenons maintenant sur la pureté modale effective ; en première approximation nous considérerons que le mode de l'état cohérent qui n'interfère pas avec le photon unique passe néanmoins à travers les filtres de manière identique au mode issu de l'interférence. En rappelant que, contrairement à ξ_{eff} , la probabilité ξ' ne distingue pas les photons en provenance des modes $\hat{a}_{\parallel}$ et $\hat{a}_{\perp}$ (voir équation 5.28 et figures 5.7 et 5.8), et en appelant P_{coher} la probabilité de détecter un photon dans le mode $\hat{a}_{\perp}$, ceci se traduit par :

$$P_{\text{id}} + P_{\text{coher}} = \xi' P_{\text{ampl}} \quad (5.50)$$

en négligeant la probabilité de détecter un photon dans chacun des deux modes . La probabilité P_{coher} peut se déduire de l'équation 2.64

$$P_{\text{coher}} = 1 - e^{-\frac{1}{2}\mu(1-\epsilon)\alpha^2} \quad (5.51)$$

En reliant les équations 5.49 et 5.50 on obtient la valeur de la pureté modale effective

$$\xi_{\text{eff}} = \xi' \frac{P_{\text{id}}}{P_{\text{id}} + P_{\text{coher}}} \quad (5.52)$$

Comme on pouvait s'y attendre, celle-ci ne dépend pas uniquement des réglages mais aussi de l'état d'entrée. Sa variation est cependant très faible : avec les paramètres utilisés pour l'expérience nous passons de $\xi_{\text{eff}} = 1$ pour $\alpha = 0$ à $\xi_{\text{eff}} = .98$ pour $\alpha = 1$.

Ces valeurs montrent bien que, comme nous l'avions dit plus tôt , les mauvais conditionnements viennent quasi-exclusivement du mode n'ayant pas interféré. Afin de s'en rendre compte analysons ce qu'il se passe lors d'un mauvais conditionnement. Un tel événement correspond à un clic alors qu'il n'y a pas de photon dans l'état cohérent et que le photon unique a été transmis, ce qui nous laisse deux possibilités

- Il s'agit d'un coup d'obscurité : si on compare le taux de succès aux coups d'obscurité de nos APD (cf. section 3.4.2) on remarque que la probabilité est belle et bien négligeable.
- Le photon vient du photon unique. Les filtres des deux voies APD étant sensiblement identiques, il y a une très forte probabilité pour que ce photon soit déjà pris en compte dans la description du photon unique (contribution parasite d'un état à plus d'un photon), dans ce cas il ne doit pas être pris en compte dans la pureté homodyne. En fait seul nous intéresse le cas où le photon serait dans un mode filtré par la première voie APD mais pas par la deuxième et qui n'aurait pas été perdu malgré les pertes importantes. Bien que difficile à quantifier il est aisé de se convaincre qu'un tel cas de figure est extrêmement rare.

Au final nous pourrions prendre $\xi' = 1$ sans que cela ne modifie les prédictions de manière significative.

5.2.3 Résultats

Lors de la réalisation expérimentale nous disposons d'un photon unique de mauvaise qualité (forte présence de vide, caractérisée par le paramètre δ). Les prédictions du premier modèle numérique indiquent cependant que cela influe peu sur les caractéristiques de l'amplificateur, les effets se faisant surtout sentir sur le taux de succès ; tout comme pour l'efficacité de la voie APD. Les paramètres critiques de l'expérience sont la présence de plus d'un photon dans le « photon unique » (paramètre σ) et surtout l'adaptation des modes ϵ . Les valeurs des différents paramètres, mesurés indépendamment de l'expérience en elle-même mais au même moment, sont données dans la table 5.1.

Paramètres	δ	σ	ϵ	μ	ξ_{eff}	η
Valeurs	0,78	1,078	0,96	0,11	1 à 0,98	0,68

TABLE 5.1: Valeurs des paramètres caractérisant les imperfections expérimentales

Nous avons réalisé l'expérience pour des états cohérents d'une amplitude allant de $\alpha = 0,55$ à $\alpha \approx 1$ ainsi que pour le vide avec un gain en amplitude $g = 2$, c'est-à-dire 6 dB en intensité. Pour chaque état en entrée nous avons effectué la tomographie de l'état en sortie à partir d'un

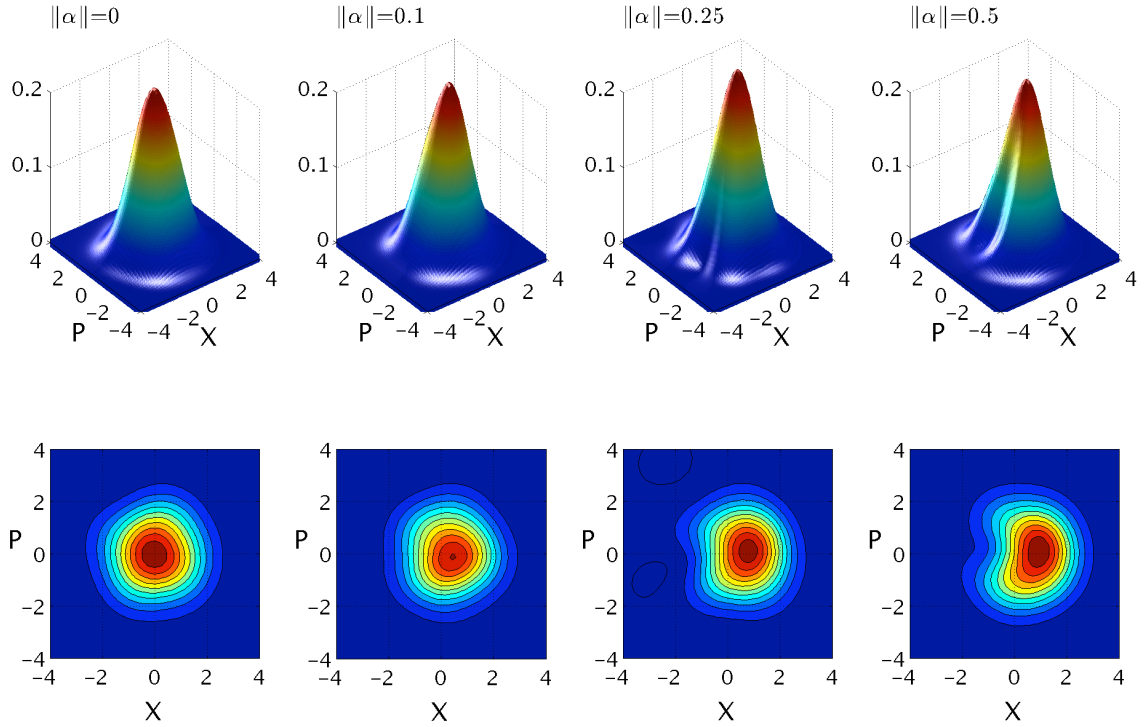


FIGURE 5.9: Reconstruction expérimental des fonctions de Wigner des états amplifiés

ensemble de 200 000 points répartis en 12 histogrammes en fonction de la quadrature mesurée. La figure 5.9 montre quelques unes de ces tomographies. Les fonctions de Wigner permettent d'évaluer l'ampleur des distorsions et ainsi de savoir jusqu'à quelle amplitude l'approximation des petits états cohérents reste valable. Lorsque l'amplitude augmente nous voyons deux effets dont l'importance augmente aussi :

- La distorsion du bruit qui n'est plus circulaire, ce qui est traduit dans le modèle analytique par le terme en cosinus dans la variance. Elle est néanmoins très différente de celle de

l'amplificateur déterministe dépendant de la phase : outre qu'elle est bien plus faible, elle ne dépend pas d'une phase définie par l'amplificateur : le bruit est toujours plus faible sur l'amplitude que sur la phase. Rappelons-nous en effet que nous avons pris comme référence la phase de notre état d'entrée : ainsi dans l'expression de la variance de notre modèle analytique l'origine des phases $\theta = 0$ correspond à la phase de l'état cohérent.

- La non gaussiannité de l'état, qui est directement liée à la troncature de notre état cohérent en une superposition de vide et de photon unique, et donc à la pertinence de l'approximation des petits états cohérents. Une telle superposition n'est en effet pas un état gaussien, mais lorsque l'approximation est valable le caractère non gaussien reste négligeable. Une autre façon de le voir consiste à dire que dans ce dernier cas l'état en sortie est essentiellement constitué de vide qui est bien un état gaussien. Par ailleurs les fonctions de Wigner sur la figure 5.9, même celles qui apparaissent clairement non gaussiennes, ne présentent pas de parties négatives : cette négativité est en effet très faible ici, même en considérant théoriquement des états purs, et disparaît rapidement avec les imperfections.

À partir des histogrammes des mesures nous pouvons aussi analyser les caractéristiques de l'amplificateur, à commencer par le gain effectif défini dans l'équation 5.37. Comme énoncé précédemment, l'état cohérent en entrée est mesuré à l'aide de l'APD de conditionnement (en coupant le faisceau contenant le photon unique de manière à ce que seul l'état cohérent arrive sur celle-ci). D'après l'équation 2.64, le taux de comptage observé vaut alors $f_\alpha = f_{\text{rep}} \left(1 - e^{-\mu|\alpha|^2}\right)$, où f_{rep} est la fréquence de répétition du laser. Ce qui nous donne :

$$\alpha = \sqrt{\frac{1}{\mu} \ln \left(\frac{f_{\text{rep}}}{f_{\text{rep}} - f_\alpha} \right)} \quad (5.53)$$

Nous pouvons alors appliquer les pertes homodynes afin de comparer des valeurs mesurées de manière identique dans le calcul du gain. Il est donc nécessaire de bien connaître l'efficacité des deux systèmes de détection (APD et détection homodyne). Rappelons qu'elles sont déterminées en envoyant un faisceau classique d'amplitude connue, atténué par des densités neutres elles aussi connues (cf. section 3.4.1.2 et 3.4.2.1) ; la précision de cette méthode, qui est d'environ 6 % en valeur relative, est suffisante.

La figure 5.10 montre les résultats expérimentaux comparés aux données du modèle. Les carrés donnent les points expérimentaux, la courbe en trait plein les résultats du modèle analytique et la courbe en tirets ceux du modèle numérique. Cette figure montre une bonne concordance entre les modèles et les données expérimentales (la fidélité moyenne entre les états expérimentaux et ceux donnés par les modèles est de 99%) et une concordance quasi-parfaite entre les deux modèles ; seuls les états avec une plus grande amplitude montrent une légère différence à cause des approximations effectuées dans les modèles qui ne sont valables qu'à faible amplitude, et sont de plus différentes pour les deux modèles (troncature de l'espace de Fock pour le modèle numérique et taux négligeable d'états à plus de un photon arrivant sur l'APD pour le modèle analytique). Pour $\alpha < 0,1$ le gain reste proche de la valeur théorique de 2, démontrant ainsi une bonne robustesse vis-à-vis des imperfections. Il décroît néanmoins très vite une fois passée cette limite, ce qui est dû au fait que l'état cohérent amplifié est tronqué à un photon. Cet effet n'a cependant rien de particulier à notre système et peut-être comparé à la saturation des amplificateurs classiques.

La première caractéristique nous a permis de vérifier que l'on avait bien un amplificateur, la deuxième s'attache à l'autre point important du « cahier des charges » de l'amplificateur, à savoir l'évolution du bruit. Pour la quantifier nous avons décidé d'utiliser un outil usuel quand

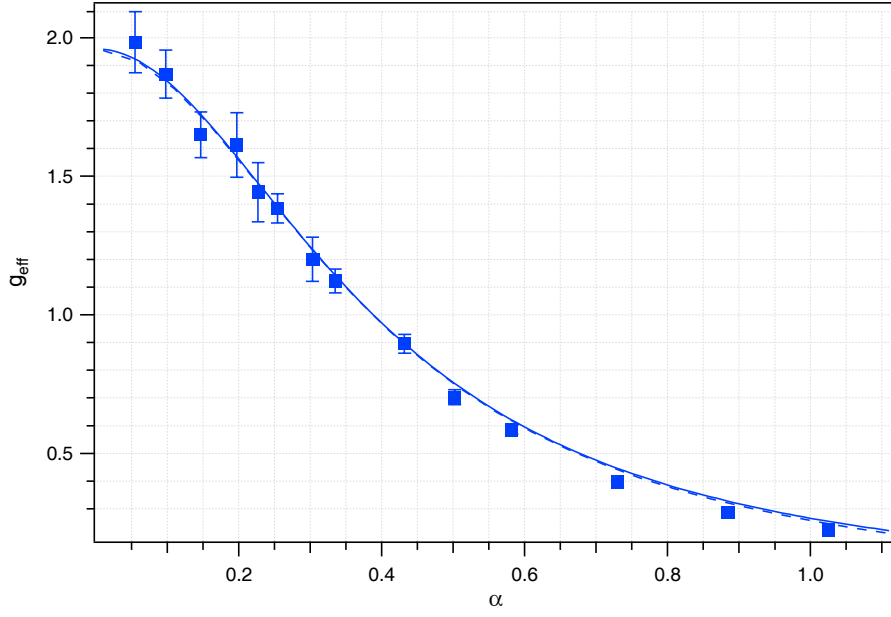


FIGURE 5.10: Gain effectif de l'amplificateur sans bruit

il est question d'amplificateurs, notamment en électronique : le bruit équivalent en entrée

$$N_{\text{eq}} = \frac{\Delta^2 \hat{X}_{\text{out}}}{g_{\text{eff}}^2} - \Delta^2 \hat{X}_{\text{in}} \quad (5.54)$$

Il représente le bruit qu'il faudrait rajouter à l'entrée pour reproduire le bruit à la sortie si l'amplificateur se contentait de l'amplifier. La figure 5.11 montre celui-ci : étant donnée la distorsion observée sur les tomographies nous avons représenté la valeur moyenne du bruit (carrés bleus), sa valeur minimale prise sur la quadrature $\hat{X}$ (disques rouges) et sa valeur maximale prise sur la quadrature $\hat{P}$ (triangles verts). À des fins de comparaisons nous avons ajouté, en pointillé, le bruit équivalent en entrée d'un amplificateur déterministe idéal indépendant de la phase⁴ (à noter que la version dépendante de la phase a un bruit équivalent en entrée strictement nul). Alors que les amplificateurs ordinaires ont un bruit équivalent en entrée toujours positif, puisqu'au mieux ils amplifient le bruit, celui de notre amplificateur est négatif dans sa zone de bon fonctionnement. Ceci est le signe que notre système n'amplifie pas le bruit et donc augmente le rapport signal sur bruit

$$\frac{\text{SNR}_{\text{out}}}{\text{SNR}_{\text{in}}} = \frac{\Delta^2 \hat{X}_{\text{in}}}{\Delta^2 \hat{X}_{\text{in}} + N_{\text{eq}}} > 1 \quad (5.55)$$

Notre dispositif ajoute tout de même un peu de bruit à cause des imperfections expérimentales, cet ajout est essentiellement dû à la présence d'état à plus de un photon dans le « photon unique ».

Au final nous pouvons déduire de tous ces résultats que la zone de fonctionnement de notre amplificateur s'étend au maximum jusqu'à une amplitude d'entrée de $\alpha = 0, 1$.

Enfin le taux de succès varie de 1 % à environ 6 % en fonction de l'état d'entrée comme le montre la figure 5.12 (hors prise en compte de la probabilité de production du photon unique).

4. Remplacé par une simple lame séparatrice pour les gains $g < 1$.

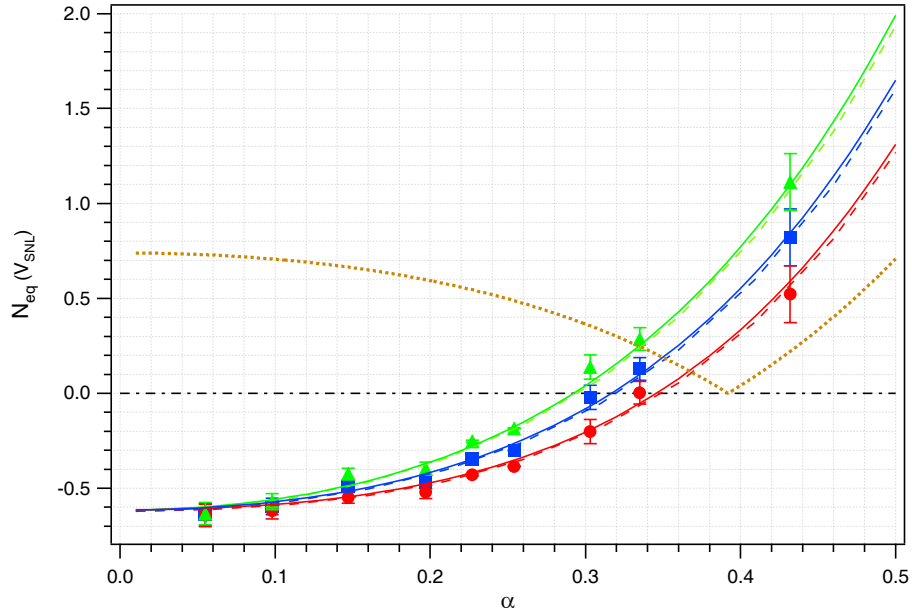


FIGURE 5.11: Bruit équivalent en entrée (minimal, moyen et maximal) de l'amplificateur sans bruit (en unité de bruit du vide). Les marqueurs représentent les données expérimentales, les tirets le modèle numérique et les traits plein le modèle analytique. La courbe en pointillés correspond à un amplificateur déterministe indépendant de la phase pour le même gain (remplacé par une lame séparatrice lorsqu'il y a déamplification).

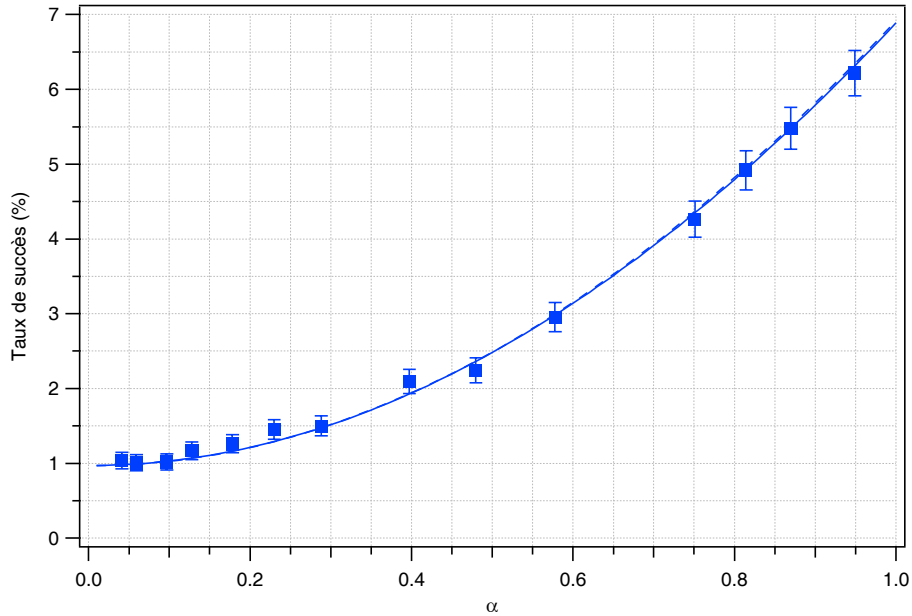


FIGURE 5.12: Taux de succès de l'amplificateur sans bruit

Nous pouvons voir que les données expérimentales correspondent bien aux deux modèles (courbes en trait plein et en tirets), et notamment aux formules 5.48 à 5.52. Néanmoins, pour la zone intéressante cela correspond à un facteur 10 avec la probabilité idéale donnée par l'équation 5.24 : c'est en effet sur ce facteur que se font le plus ressentir les effets des pertes.

5.3 Discussion

5.3.1 Cohérence avec les lois de la physique quantique

Nous avons bel et bien réussi à contourner les lois de la physique en réalisant un processus qui est normalement interdit. Nous avons même réussi à faire mieux que le cas classique puisque nous avons augmenté le rapport signal sur bruit. Nous sommes donc en droit de nous poser la question de la cohérence avec les lois de la physique : bien que nous les ayons contournées, nous devons tout de même les retrouver d'une façon ou d'une autre. La question qui se pose alors est : comment les retrouver ? Si nous jetons un coup d'œil aux autres expériences qui semblaient violer les lois de la physique, comme par exemple en faisant apparaître des vitesses plus grande que celle de la lumière, nous remarquons qu'à chaque fois on se tourne vers l'information. C'est en effet celle-ci qui respecte toujours les lois de la physique.



FIGURE 5.13: Communication entre Alice et Bob en utilisant l'amplificateur sans bruit

Analysons donc ce qu'il se passe si nous utilisons notre amplificateur dans un dispositif de communication. Prenons les éternels interlocuteurs des protocoles de communication, à savoir Alice et Bob. Dans un protocole ordinaire de communication quantique à variables continues Alice envoie un message à Bob en modulant en amplitude et en phase un faisceau laser, et ce dernier lit le message grâce à une détection homodyne. L'information transmise par ce canal entre Alice et Bob est quantifiée par l'information mutuelle I_{AB} ; cette grandeur issue de la théorie de l'information représente, pour un système bipartite, l'information que l'on a acquis sur une des parties en déterminant l'autre. Elle est symétrique : Alice en sait autant sur la partie de Bob que ce dernier sur la partie d'Alice. Elle correspond donc bien à l'information partagée par Alice et Bob grâce à la communication. Dans ce cas de figure les deux parties correspondent aux impulsions lumineuses avant et après la transmission.

Pour des états quantiques gaussiens l'information mutuelle est donnée par [149]

$$I_{AB} = \frac{1}{2} \log_2 \left(1 + \frac{V_A}{V_{\text{SNL}}} \right) \quad (5.56)$$

où V_A est la variance de la modulation et V_{SNL} le bruit du vide. Voyons ce qu'il se passe si maintenant Bob ajoute l'amplificateur sans bruit afin d'augmenter le rapport signal sur bruit. Lorsque l'amplification a lieu le signal d'Alice est amplifié, et, en se plaçant dans une hypothèse d'impulsions fortement atténuées (avec $V_A \ll V_{\text{SNL}}$) afin d'amplifier des états cohérents de faible

amplitude, l'information mutuelle devient alors :

$$I_{AB}^{\text{ampl}} = \frac{1}{2} \log_2(1 + g^2 \frac{V_A}{V_{\text{SNL}}}) \approx g^2 I_{AB} \quad (5.57)$$

À première vue il semble bien que la transmission d'information ait augmenté, seulement cela serait oublier la contrepartie de l'amplification, c'est-à-dire le côté probabiliste. Nous devons en effet prendre aussi en compte les événements où l'amplification n'a pas marché, entraînant la destruction du signal. L'information mutuelle totale vaut donc

$$I_{AB}^{\text{tot}} = P I_{AB}^{\text{ampl}} \quad (5.58a)$$

$$\approx r^2 g^2 I_{AB} \quad (5.58b)$$

$$\approx (1 - r^2) I_{AB} \quad (5.58c)$$

Au final l'amplificateur ne permet donc pas d'augmenter la transmission d'information, gardant ainsi la cohérence avec les lois de la physique interdisant l'amplification de la quantité d'information transmise. Son action consiste uniquement à concentrer l'information dans certains événements.

5.3.2 Autres expériences

Le principe d'amplification sans bruit a eu un certain succès dans la communauté des variables continues, voire même au delà, menant à la proposition d'autres protocoles et la réalisation d'autres expériences.

La première d'entre elles est la démonstration, a peu près au même moment que nous, du même dispositif par le groupe de Geoff Pryde[156]. Ils utilisent cependant une approche différente : ils n'amplifient pas un vrai état cohérent mais uniquement un mélange statistique de vide et de photon unique, et analysent l'état en sortie avec uniquement un compteur de photon. Un tel dispositif procure l'avantage de pouvoir travailler bien plus efficacement et facilement avec des états de très faibles amplitudes, ce qui leur a permis de démontrer la linéarité du gain là où nous ne faisons que l'effleurer. Par contre leur caractérisation des états de sortie reste très limitée : ils ne peuvent en effet pas effectuer de tomographie permettant de montrer que l'état reste bien approximativement un état cohérent, ni analyser le bruit de l'état amplifié. De même ils ne peuvent montrer de manière non ambiguë la zone de fonctionnement de l'amplificateur de par leur méthode de production de l'état d'entrée.

Une autre méthode pour amplifier sans bruit a d'abord été évoquée par Petr Marek et Radim Filip[150] puis généralisée et développée par Jaromír Fiurášek[151] et démontrée expérimentalement dans le groupe de Marco Bellini[152]. Elle est toujours basée sur l'approximation des petits états cohérents comme une superposition de vide et de photon unique mais consiste à approximer l'opérateur théorique d'amplification $g^{\hat{n}}$ par son développement polynomial au premier ordre, c'est à dire $g^{\hat{n}} \approx 1 + (g - 1)\hat{n}$. Pour un gain $g = 2$ cela revient, de par les relations de commutation, à appliquer l'opérateur de création suivi de l'opérateur d'annihilation sur l'état à amplifier. L'avantage de cette méthode est qu'elle ne reproduit pas l'état en entrée mais le modifie ; les états en sortie présentent toujours une contribution des états à plus de un photon, bien que celle-ci ne soit pas avec le bon poids. La transformation réalisée étant $|n\rangle \mapsto (n+1)|n\rangle$, cela permet néanmoins de pousser plus loin l'approximation des petits états cohérents et donc d'avoir une zone de fonctionnement plus grande. De plus lorsqu'aucune des opérations n'a réussi l'état n'est pas détruit. À contrario cette méthode possède deux inconvénients : premièrement il n'est

évidemment possible d'avoir que des gains en amplitude entiers, donnant ainsi une moins bonne variabilité au système. Deuxièmement, augmenter le gain demande d'effectuer plus d'ajouts et de soustractions de photons, et donc demande plus de coïncidences entre des événements aléatoires. Là où notre système voit son taux de succès décroître hyperboliquement avec le gain la décroissance est ici exponentielle. Il est possible d'améliorer celui-ci en effectuant une série de N additions de photon suivies d'autant de soustractions⁵ afin d'obtenir un gain $g = (N + 1)!$ au prix donc d'une nouvelle baisse de la variabilité du gain, mais la décroissance reste bien plus importante que pour notre dispositif.

C'est justement à cause du taux de succès que Petr Marek et Radim Filip ne se sont pas attardés sur ce système et ont proposé de remplacer l'ajout de photons par l'ajout de bruit thermique [150]. Ce dispositif est le plus simple et celui qui possède un taux de succès le plus grand. Malheureusement comme toujours il y a un prix à payer pour cette amélioration ; en l'occurrence il ne s'agit plus réellement d'une amplification mais d'une *concentration de la phase* : l'état en sortie est déformé et ne peut donc plus guère être utilisé autrement que par un détecteur ; par contre le rapport signal sur bruit pour une mesure de la phase a bel et bien été augmenté. Ce protocole a été démontré expérimentalement dans le groupe de Ulrik Andersen[153].

5.3.3 Utilisations diverses

Les premières utilisations de l'amplificateur sans bruit qui viennent à l'esprit consistent à augmenter la résolution d'un faible signal dans un système de détection ou de compenser les pertes des canaux de transmission dans les systèmes de communications. C'est d'ailleurs dans ce but que nous avons commencé à nous y intéresser, en relation avec un système de distribution quantique de clés à variables continues développé dans le groupe en collaboration avec Thalès [154, 155]. Une analyse de sécurité est actuellement en cours d'étude à ce sujet. D'un point de vue plus général, l'amplificateur peut servir à augmenter la résolution d'une mesure, que ce soit pour des applications de communication ou de métrologie. Seulement nous avons vu section 5.3.1 que l'information transmise était globalement diminuée, la question se pose alors de savoir s'il n'est pas préférable d'effectuer une série de mesures afin d'augmenter statistiquement la résolution (grâce au théorème central limite). Cette dernière méthode est en effet plus simple et efficace, puisqu'elle n'introduit pas d'imperfection supplémentaire⁶ ; mais elle demande de pouvoir mesurer plusieurs fois le même état. Une telle contrainte n'est par exemple pas réalisable dans un système de cryptographie quantique où les états transmis doivent être aléatoires.

Une autre utilisation immédiate, et déjà évoquée en section 5.1.1, est le clonage [143]. Bien que l'amplification sans bruit permette de réaliser le clonage d'un état cohérent, les choses se gâtent dans la pratique. En effet une application réaliste ne viserait pas à cloner un seul état mais plutôt un ensemble d'états répartis selon une certaine distribution. Dans ce cas, même en considérant que l'on peut se contenter d'un extrait des données de par le caractère probabiliste, la distribution clonée sera différente de celle d'origine. La cause en est la dépendance du taux de succès envers l'état d'entrée : certains états seront plus souvent clonés que les autres, modifiant ainsi la répartition en amplitude des états.

Nous avons défini l'action de l'amplificateur sans bruit à l'aide des états cohérents, mais il

5. Il s'agit d'ailleurs en fait du protocole évoqué par Petr Marek et Radim Filip, la version exposée précédemment étant plus exactement celle de Jaromír Fiurášek

6. À partir des équations 5.23 et 5.24 nous trouvons que la précision (pour une mesure en amplitude) est, sans imperfections, améliorée d'un facteur $\sqrt{N} - 1$ en fonction du nombre moyen de tentatives nécessaires N , ce qui reste très proche du facteur $\sqrt{N}$ obtenu pour une accumulation statistique.

peut très bien être appliqué à d'autres états. Appliqué par exemple à un mode d'un état EPR il permet d'en augmenter l'intrication [143, 156]. En effet si nous appliquons l'amplificateur à la décomposition en base de Fock d'un état EPR, donnée par l'équation 2.69, nous obtenons

$$\sum_{n=0}^{\infty} \tanh^n(r) |n\rangle |n\rangle \mapsto \sum_{n=0}^{\infty} \tanh^n(r) g^n |n\rangle |n\rangle \quad (5.59)$$

donnant ainsi un nouveau paramètre de compression plus grand ; nous devons tout de même avoir $\tanh(r)g \leq 1$ afin de garder un état physique, c'est-à-dire normalisable. D'un point de vue pratique, nous devons utiliser une approximation de l'opérateur $g^{\hat{n}}$, qui mènerait bel et bien à un état physique ; ce qui se voit bien d'un point de vue expérimental puisque un tel réglage de gain y est possible. Seulement l'état en sortie ne serait plus du tout un état EPR, et ce quelque soit l'approximation utilisée, contrairement à la saturation que nous avons vu sur l'amplification d'un état cohérent⁷. Un tel dispositif est bien plus compliqué à mettre en place que celui de notre expérience, de plus les états EPR sont bien plus sensibles aux imperfections que les états cohérents, conduisant ainsi à un gain moins important.

D'autres états auxquels nous pourrions appliquer l'amplificateur sont les états Chats de Schrödinger. Puisque nous savons produire de petits chats [94, 92] avec des taux de succès raisonnable, cela pourrait être un moyen intéressant d'en produire de plus grands en faisant grandir des petits chats. Il resterait tout de même à faire une étude plus détaillée pour en vérifier la faisabilité, notamment en considérant que les états préparés ne sont en réalité que des approximations de chats de Schrödinger, et voir s'il l'on gagnerait réellement quelque chose. Cette expérience ne pourrait pas être menée avec notre dispositif puisque les chats ne peuvent être approximés par une superposition de vide et de photon unique. Il faudrait alors soit une version à plusieurs étages, soit l'amplificateur à ajout et soustraction de photon, qui peut aussi fonctionner sur les bases $(|0\rangle, |2\rangle)$ et $(|1\rangle, |3\rangle)$.

Nous pouvons encore citer une dernière application proposée par Nicolas Gisin *et al.* [157] : il ne s'agit pas d'une amplification d'un signal continu mais d'une modification du dispositif afin de compenser une partie des pertes subies par un qubit discret. L'idée est alors de coder les qubits sur la polarisation d'un photon, les pertes venant alors ajouter une contribution du vide à notre qubit. Pour diminuer l'effet de celle-ci il faut donc augmenter la part de photon unique au détriment du vide, et c'est exactement ce que fait notre système. Le dispositif est cependant plus complexe de par la présence de deux modes pour le photon unique : il faut donc remplacer le photon unique en entrée par deux photons uniques, un dans chaque mode, et conditionner sur la détection d'un photon dans chacun des deux modes.

5.3.4 Retour sur le calcul quantique hybride

Notre dispositif à un seul étage ne peut pas amplifier un état chat, mais il ne reste pourtant pas sans effet dessus : quel est alors celui-ci ? Nous avons évoqué la ressemblance avec le protocole des « ciseaux quantiques » ; un étage de l'amplificateur est plus précisément une généralisation de ce protocole (qui correspond à $g = 1$). Notre dispositif a donc une action similaire : en l'occurrence il tronque l'état en entrée pour n'en laisser que les composantes à 0 et 1 photons (en modifiant leur rapport). En d'autres termes il transforme un chat pair en vide et un chat impair en photon unique, ce qui réalise l'opération inverse de celle présentée à la fin du chapitre

7. Ce phénomène serait d'ailleurs lui aussi présent dans cette application.

précédent : la conversion d'un qubit continu en qubit discret. Plus précisément l'action de l'amplificateur sans bruit à un étage sur un qubit continu codé sur la parité d'un chat de Schrödinger correspond à la transformation

$$a \frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2}} + b \frac{|\alpha\rangle - |-\alpha\rangle}{\sqrt{2}} \mapsto a|0\rangle + bg\alpha|1\rangle \quad (5.60)$$

En prenant $g = 1/\alpha$ nous avons donc bien l'opération désirée. Tout comme pour le protocole inverse ce comportement peut être vu par le lien avec la téléportation : ce système correspond lui aussi à une téléportation hybride mais cette fois c'est un état continu qui est téléporté sur un mode discret.

L'idéal serait toutefois de pouvoir effectuer la même transformation avec un codage en polarisation. Dans ce cas le chapitre précédent nous a appris que le qubit continu devait être codé lui aussi sur deux chats séparés en polarisation. Nous avons déjà évoqué une version de l'amplificateur fonctionnant pour un codage en polarisation ; cette version a été conçue pour fonctionner uniquement avec des qubits discrets en entrée comme en sortie mais si nous y envoyons un qubit continu à la place alors nous obtenons le comportement voulu (figure 5.14). En effet la troncature de notre qubit continu codé en polarisation donne exactement un qubit discret.

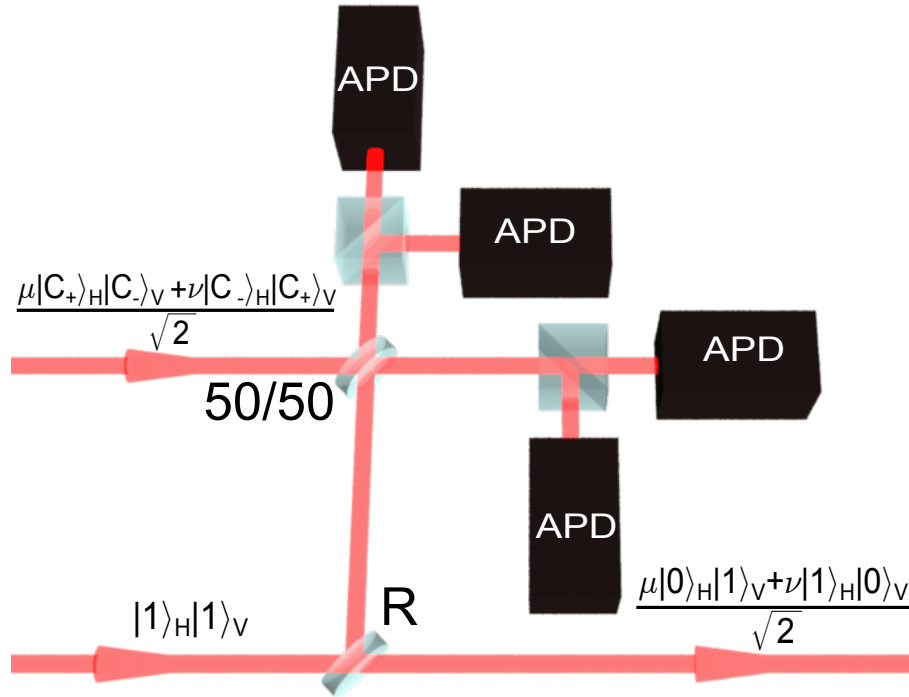


FIGURE 5.14: Protocole transformant un qubit encodé sur les parités de deux chats de Schrödinger en un qubit continu discret encodée sur la polarisation d'un photon

Voyons ce dispositif plus en détails. La troncature donne les correspondances suivantes, pour chaque polarisation,

$$|C_-, C_+\rangle \leftrightarrow |1, 0\rangle \quad (5.61)$$

$$|C_+, C_-\rangle \leftrightarrow |0, 1\rangle \quad (5.62)$$

De manière générale la transformation réalisée par cette version de l'amplificateur est

$$\begin{aligned} & a |C_+, C_-\rangle + b |C_-, C_+\rangle + c |C_+, C_+\rangle + d |C_-, C_-\rangle \\ & \mapsto g\alpha (a |1, 0\rangle + b |0, 1\rangle) + c |0, 0\rangle + g^2\alpha^2 d |1, 1\rangle \end{aligned} \quad (5.63)$$

Dans le cas d'un qubit pur $c = d = 0$, le gain joue alors uniquement un rôle dans le taux de succès, la valeur optimale est $g = 1$ c'est-à-dire pour le protocole des ciseaux quantiques. Cependant nous avons vu que dans le protocole inverse, transformant un qubit discret en qubit continu, la perte du photon crée l'un des deux états avec des chats de même parité. Dans ce cas nous pouvons régler le gain du dispositif de manière à réduire la contribution de cet état dans l'état final, purifiant ainsi le qubit.

5.3.5 Conclusion

Au final nous avons démontré qu'il est possible de contourner les restrictions imposées par l'unitarité et la linéarité de l'évolution quantique. Nous avons ainsi pu réaliser un dispositif amplifiant un signal optique sans en amplifier le bruit quantique. L'augmentation du rapport signal sur bruit induite par ce processus s'est même révélée être très robuste face aux imperfections.

Nous avons aussi passé en revue un certain nombre d'applications potentielles de l'amplificateur sans bruit. Mais l'histoire des sciences nous enseigne que les applications finales des recherches fondamentales sont souvent très loin de ce que l'on attendait à l'origine. Qui sait alors à quoi il servira ? D'autant plus que son comportement pour le moins atypique est encore loin d'avoir livré tout ses secrets. Prenons un exemple pour illustrer ce fait : lorsqu'on l'utilise sur un état cohérent il en augmente l'amplitude et donc le nombre moyen de photon ; pourtant si on l'utilise sur un état de Fock il le laissera inchangé et ce quel qu'en soit le gain. La physique quantique n'est déjà pas toujours très intuitive mais l'amplificateur sans bruit fait sans doute partie de ses éléments les moins intuitifs.

Chapitre 6

Mesure de la non-gaussianité d'un état

Sommaire

6.1 La non-gaussianité	125
6.1.1 États gaussiens et états non-gaussiens	125
6.1.2 Les mesures	127
6.1.3 Non gaussianité et non-classicité	130
6.2 Réalisation expérimentale	131
6.2.1 Montage	131
6.2.2 Modélisation	133
6.2.3 Résultats	135
6.3 Discussion	141
6.3.1 Non-gaussianité de quelques états courants	141
6.3.2 Non-gaussianité des processus quantiques	143
6.3.3 Non-gaussianité et information	144
6.3.4 Conclusion	145

6.1 La non-gaussianité

6.1.1 États gaussiens et états non-gaussiens

Les états gaussiens ont une place de choix dans le domaine des variables continues. Outre le fait qu'ils sont faciles à produire, ces états sont, à variance fixée, extremums pour de nombreux critères. Nous avons ainsi déjà mentionné le fait que les états d'incertitude minimale sont gaussiens, mais cela ne s'arrête pas là. Les états gaussiens sont aussi ceux qui maximisent l'entropie de von Neuman [158] et l'entropie conditionnelle [159, 160]. De même, pour un système de distribution de clés quantique à variables continues, les attaques individuelles optimales que peut effectuer un espion sont gaussiennes [159]. À contrario ils minimisent l'information mutuelle [162], l'intrication et le taux de clés secrètes transmises [163].

Nous avons aussi pu voir que les opérations possibles avec les états gaussiens sont limitées ; nous avons en effet déjà mentionné dans les chapitres précédents un certain nombre d'opérations

qui ne peuvent être réalisées que par l'intermédiaire d'états ou de processus non gaussiens. À cela nous pouvons ajouter une autre catégorie d'opérations qui sont possibles avec des états gaussiens mais que la non-gaussianité permet aussi d'améliorer. Nous pouvons citer comme exemple la téléportation [164, 165, 166], le clonage quantique [167], le stockage d'états quantiques [168] et l'estimation de paramètres [169, 170].

Ainsi sont apparus de nombreux protocoles permettant de dégaussifier des états [101, 94, 71, 171] ou de créer directement des états non-gaussiens [172, 173, 92]. On peut aussi noter la présence de protocoles de gaussification [174] permettant notamment de retrouver un état (approximativement) gaussien après une opération non-gaussienne. Les méthodes utilisées sont tout aussi nombreuses : en plus des techniques étudiées dans ce manuscrit on peut citer l'effet Kerr [175] et les autres processus non-linéaires d'ordre supérieur à deux [176, 177], les cavités actives avec effet Kerr croisé [178, 62], la synthèse d'opérateur unitaire arbitraire [179], l'interférométrie conditionnelle [180], ... On peut aussi retrouver des états gaussiens dans d'autres domaines relevant de l'information quantique comme la CQED (*Cavity Quantum Electrodynamics*) [181] ou les circuits supraconducteurs [182].

Les états non-gaussiens sont donc devenus une ressource capitale pour l'information quantique. À ce point arrive la question de caractérisation de cette non-gaussianité. La quantification de celle-ci permet, par exemple, d'étudier plus profondément l'apport de la non-gaussianité à la téléportation [183] ou à l'intrication [184]. Elle peut aussi être utilisée pour quantifier la validité de l'approximation utilisée dans l'amplificateur sans bruit du chapitre précédant (bien que dans ce dernier cas la fidélité est tout aussi adaptée).

Mais avant de s'attaquer à la mesure de la non-gaussianité, disons quelques mots sur le formalisme gaussien. Une distribution gaussienne est entièrement définie par le vecteur de ses valeurs moyennes et sa matrice de covariance

$$K = \begin{pmatrix} \langle \delta \hat{X}_1^2 \rangle & \langle \delta \hat{X}_1 \delta \hat{P}_1 \rangle & \langle \delta \hat{X}_1 \delta \hat{X}_2 \rangle & \cdots & \langle \delta \hat{X}_1 \delta \hat{P}_n \rangle \\ \langle \delta \hat{P}_1 \delta \hat{X}_1 \rangle & \langle \delta \hat{P}_1^2 \rangle & \langle \delta \hat{P}_1 \delta \hat{X}_2 \rangle & \cdots & \langle \delta \hat{P}_1 \delta \hat{P}_n \rangle \\ \langle \delta \hat{X}_2 \delta \hat{X}_1 \rangle & \langle \delta \hat{X}_2 \delta \hat{P}_1 \rangle & \langle \delta \hat{X}_2^2 \rangle & \cdots & \langle \delta \hat{X}_2 \delta \hat{P}_n \rangle \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \langle \delta \hat{P}_n \delta \hat{X}_1 \rangle & \langle \delta \hat{P}_n \delta \hat{P}_1 \rangle & \langle \delta \hat{P}_n \delta \hat{X}_2 \rangle & \cdots & \langle \delta \hat{P}_n^2 \rangle \end{pmatrix} \quad (6.1)$$

avec $\delta \hat{X}_i = \hat{X}_i - \langle \hat{X}_i \rangle$ et $\delta \hat{P}_i = \hat{P}_i - \langle \hat{P}_i \rangle$. Par des opérations locales (déphasage et compression) il est possible de changer de base de façon à séparer les quadratures $\hat{X}_i$ et $\hat{P}_j$. La matrice de covariance se simplifie alors en

$$K = \begin{pmatrix} \langle \delta \hat{X}_1^2 \rangle & 0 & \langle \delta \hat{X}_1 \delta \hat{X}_2 \rangle & \cdots & 0 \\ 0 & \langle \delta \hat{P}_1^2 \rangle & 0 & \cdots & \langle \delta \hat{P}_1 \delta \hat{P}_n \rangle \\ \langle \delta \hat{X}_2 \delta \hat{X}_1 \rangle & 0 & \langle \delta \hat{P}_1^2 \rangle & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & \langle \delta \hat{P}_n \delta \hat{P}_1 \rangle & 0 & \cdots & \langle \delta \hat{P}_n^2 \rangle \end{pmatrix} \quad (6.2)$$

Elle permet de calculer les divers grandeurs intéressantes comme la pureté et l'entropie :

$$\mathcal{P} = \frac{1}{2^N \sqrt{\det(K)}} \quad (6.3)$$

$$S = f(\sqrt{\det(K)}) \quad (6.4)$$

ou N est le nombre de modes et avec

$$f(x) = (x + \frac{1}{2}) \log_2 \left(x + \frac{1}{2} \right) - (x - \frac{1}{2}) \log_2 \left(x - \frac{1}{2} \right) \quad (6.5)$$

6.1.2 Les mesures

Les distributions gaussiennes sont aussi très présentes en statistique classique, commençons donc par nous pencher sur la caractérisation classique du caractère non-gaussien. Nous venons de dire que les distributions gaussiennes sont entièrement caractérisées par leur moments d'ordre 1 et 2 : c'est ainsi qu'en statistique le caractère non-gaussien d'une distribution est caractérisé par les moments d'ordre 3 et 4, en les comparant à leurs valeurs pour des distribution gaussiennes. À partir de ceux-ci, on définit le *coefficient de dissymétrie* γ_1 (ou *skewness*), le *kurtosis* β_2 (ou *coefficient d'aplatissement*) et l'*excès de kurtosis* γ_2 . Pour une distribution de probabilité $P(x)$ ils valent [185]

$$\gamma_1 = \frac{\langle (x - \langle x \rangle)^3 \rangle}{(\Delta^2 x)^{3/2}} \quad (6.6)$$

$$\beta_2 = \frac{\langle (x - \langle x \rangle)^4 \rangle}{(\Delta^2 x)^2} \quad (6.7)$$

$$\gamma_2 = \beta_2 - 3 \quad (6.8)$$

Une distribution avec un excès de kurtosis positif est appelée *super-gaussienne* ou *leptokurtique* et est plus pointue ; à l'inverse si l'excès de kurtosis est négatif la distribution est plus aplatie elle est alors appelée *sub-gaussienne* ou *platikurtique*. Le coefficient de dissymétrie indique quant à lui si la distribution est plus étalée vers la gauche ($\gamma_1 > 0$) ou vers la droite ($\gamma_1 < 0$). Ces caractérisations précises sont cependant valables uniquement pour une distribution de probabilité uni-dimensionnelle [186], ce qui empêche leur adaptation à la fonction de Wigner qui est caractérisée par les deux dimensions x et p . Le fait que l'excès de kurtosis est uniquement nul pour des distributions gaussiennes reste cependant valable quelle que soit le nombre de dimensions. On pourrait donc utiliser la valeur absolue de l'excès de kurtosis pour caractériser la non-gaussianité. Malheureusement d'un point de vue expérimental le kurtosis est très sensible aux données aberrantes, ce n'est donc pas une mesure robuste ; de plus il n'existe pas de moyen pour l'estimer sans avoir de biais expérimental qui soit indépendant de la distribution [187].

Pour trouver une meilleure mesure posons-nous la question suivante : à quoi correspond une mesure de non-gaussianité d'un état ? Ce que nous voulons savoir c'est à quel point il s'éloigne d'un état gaussien qui lui ressemblerait. À partir de cette constatation nous pouvons définir la non-gaussianité comme la « distance » entre l'état $\hat{\rho}$ à caractériser et une référence gaussienne $\hat{\tau}$ (figure 6.1). Il nous reste alors deux points à éclaircir : comment choisir cette référence et comment mesurer la distance ? Le premier point ne pose pas réellement de problème : l'état gaussien qui ressemble le plus à l'état $\hat{\rho}$ est celui qui possède le même vecteur de valeurs moyennes et la même matrice de covariance. Le deuxième point amène plus de difficultés, puisqu'il n'y a pas de mesure unique, et plusieurs méthodes ont été proposées.

La première méthode proposée [188, 189] utilise la distance d'Hilbert-Schmidt définie par

$$D_{\text{HS}}(\hat{\rho}, \hat{\tau}) = \sqrt{\frac{\text{Tr}((\hat{\rho} - \hat{\tau})^2)}{2}} \quad (6.9)$$

$$= \sqrt{\frac{\mathcal{P}(\hat{\rho}) + \mathcal{P}(\hat{\tau}) - 2\text{Tr}(\hat{\rho}\hat{\tau})}{2}} \quad (6.10)$$

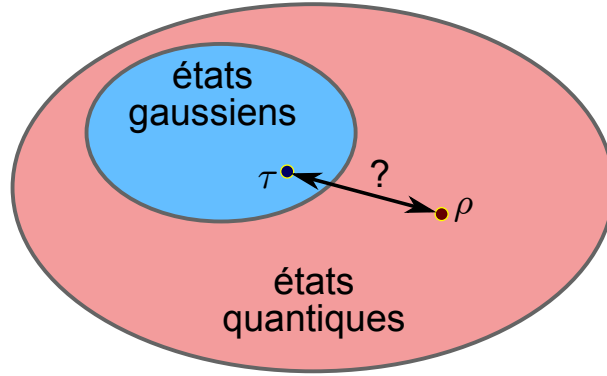


FIGURE 6.1: Mesure de non-gaussianité par la distance avec une référence

La mesure de non-gaussianité δ_1 vaut alors

$$\delta_1(\hat{\rho}) = \frac{D_{\text{HS}}^2(\hat{\rho}, \hat{\tau})}{\mathcal{P}(\hat{\rho})} = \frac{\text{Tr}((\hat{\rho} - \hat{\tau})^2)}{2\text{Tr}(\hat{\rho}^2)} \quad (6.11)$$

D'après cette définition cette mesure est toujours positive, et elle est nulle uniquement pour un état gaussien. De plus elle possède les propriétés suivantes :

- Elle est invariante par une transformation unitaire symplectique de l'opérateur densité, c'est-à-dire de la forme $e^{i\hat{H}}$ avec $\hat{H}$ au plus bilinéaire en fonction des opérateurs de champs. Ceci permet aux transformations telles que le déplacement, la compression, l'amplification paramétrique ou le mélange sur une lame séparatrice de ne pas modifier la non-gaussianité.
- Pour un état bipartite de type $\hat{\rho} = \hat{\rho}_A \otimes \hat{\rho}_G$ où $\hat{\rho}_G$ est gaussien, $\delta_1(\hat{\rho}) = \delta_1(\hat{\rho}_A)$. Il n'y a par contre pas additivité de la non-gaussianité pour un produit tensoriel quelconque.
- La non-gaussianité est bornée : $\delta_1(\hat{\rho}) \leq 1/2$ [190].

Si on revient aux statistiques classiques on observe que ce principe de distance avec un état gaussien de référence est aussi utilisé. On trouve alors une autre mesure, notamment utilisée pour l'*analyse en composantes indépendantes*¹ : la néguentropie [191]. La distance entre les distributions est alors mesurée par l'entropie relative. Cette quantité, issue de la théorie de l'information, peut-être reliée à une action très commune en science expérimentale : la représentation d'un phénomène physique par un modèle. Lorsque que nous établissons un modèle, il décrit en partie le phénomène et nous donne donc des informations sur celui-ci. Néanmoins les modèles sont rarement, si ce n'est jamais, parfaits. Il nous manque donc une certaine quantité d'information pour connaître parfaitement le système observé. Ces informations sont quantifiées par l'entropie relative. Si elle est nulle alors notre modèle décrit parfaitement le système (en termes de probabilités les distributions réelle et du modèle sont les mêmes), si elle est égale à l'entropie du système alors c'est que le modèle ne nous apprend strictement rien. Mathématiquement si les p_i et la distribution de probabilités réelle et q_i celle du modèle, l'entropie relative vaut alors :

$$S(p||q) = \sum_{i=1}^N p_i \log_2(p_i/q_i) \quad (6.12)$$

Elle peut ainsi être vue comme une distance entre les deux distributions, mais nous devons garder en mémoire que ce n'en n'est pas une au sens mathématique du terme, et elle n'est, en

1. il s'agit d'une méthode issue de la séparation aveugle de sources.

autres, pas symétrique. La version quantique de l'entropie relative vaut

$$S(\hat{\rho}||\hat{\tau}) = \text{Tr}(\hat{\rho}(\log_2(\hat{\rho}) - \log_2(\hat{\tau}))) \quad (6.13)$$

et permet de définir une seconde mesure de non-gaussianité [192, 189]

$$\delta_2(\hat{\rho}) = S(\hat{\rho}||\hat{\tau}) \quad (6.14)$$

L'entropie relative a d'ailleurs déjà été utilisée en information quantique comme mesure de l'indiscernabilité statistique [193]. Dans le cas qui nous intéresse le choix de l'état de référence implique² $\text{Tr}(\hat{\rho} \log_2(\hat{\tau})) = \text{Tr}(\hat{\tau} \log_2(\hat{\rho}))$, ce qui permet de réécrire la mesure

$$\delta_2(\hat{\rho}) = S(\hat{\tau}) - S(\hat{\rho}) \quad (6.15)$$

Les états gaussiens maximisant l'entropie, la mesure est toujours positive (c'est d'ailleurs toujours le cas pour l'entropie relative); de plus elle est bien nulle uniquement pour un état gaussien. Elle possède les propriétés suivantes :

- Elle est invariante par une transformation unitaire symplectique, tout comme la précédente mesure.
- Elle est additive pour un état séparable : $\delta_2(\hat{\rho}_A \otimes \hat{\rho}_B) = \delta_2(\hat{\rho}_A) + \delta_2(\hat{\rho}_B)$, et pour un état bipartite générique nous avons $\delta_2(\hat{\rho}_{AB}) \geq \delta_2(\hat{\rho}_A) + \delta_2(\hat{\rho}_B)$.

Nous pouvons remarquer que dans le cas d'un état pur la non gaussianité est égale à l'entropie de l'état de référence; cette mesure est alors indépendante des moments d'ordre supérieur à deux.

Enfin une troisième méthode a été récemment proposée [200], basée cette fois sur la fonction Q. Celle-ci est une autre représentation des états quantique dans l'espace des phases, correspondant à la fonction caractéristique [41]

$$\left\langle e^{i\hat{a}\frac{\hat{X}-i\hat{P}}{\sqrt{2}}} e^{i\hat{a}^\dagger\frac{\hat{X}+i\hat{P}}{\sqrt{2}}} \right\rangle \quad (6.16)$$

Mais elle peut aussi être vue d'une autre manière : s'il n'est pas possible de définir une vraie distribution de probabilité dans l'espace des phases, il est par contre possible de définir la probabilité qu'un état puisse être mesuré, i.e. projeté, comme étant un état quasi-classique ayant une position bien définie dans l'espace des phases. C'est ce à quoi correspond la fonction Q :

$$Q(\alpha) = \frac{1}{2\pi} \langle \alpha | \hat{\rho} | \alpha \rangle \quad (6.17)$$

avec $\alpha = (x + ip)/\sqrt{2}$. Elle peut donc être facilement obtenue à partir de la fonction de Wigner par la formule de recouvrement

$$Q(\alpha) = \iint W(x', p') W_\alpha(x', p') dx' dp' \quad (6.18)$$

L'avantage de cette fonction est qu'elle est une vraie distribution de probabilité, on peut donc lui appliquer les outils de la statistique classique. De plus, le passage de la fonction de Wigner à la fonction Q se fait par convolution avec une gaussienne; ainsi les états gaussiens ont une fonction Q gaussienne et les états non-gaussiens une fonction Q non-gaussienne. Il est donc

2. En effet $\log_2(\hat{\tau})$ est un polynôme d'ordre au plus 2.

possible d'estimer la distance entre les états $\hat{\rho}$ et $\hat{\tau}$ à partir de celle-ci. C'est ce que fait cette mesure en utilisant pour cela la différence entre les entropies de Wehrl définies par

$$H(\hat{\rho}) = - \int \cdots \int Q_{\hat{\rho}}(\alpha_1, \dots, \alpha_N) \log_2(Q_{\hat{\rho}}(\alpha_1, \dots, \alpha_N)) d^2\alpha_1 \dots d^2\alpha_N \quad (6.19)$$

La mesure de non-gaussianité vaut alors

$$\delta_3(\hat{\rho}) = H(\hat{\tau}) - H(\hat{\rho}) \quad (6.20)$$

L'entropie de Wehrl est aussi maximale pour un état gaussien, à variance fixe, cette mesure est donc elle aussi positive et nulle uniquement pour un état gaussien. Elle possède de plus comme autres propriétés :

- Elle est invariante par déplacement dans l'espace des phases et par un passage à travers un système linéaire passif.
- Elle est invariante par changement d'échelle $Q(\alpha) \mapsto \lambda^{2N} Q(\lambda\alpha)$.
- Elle est additive pour un état séparable.

La table 6.1 résume les principales caractéristiques de ces trois mesures, la deuxième semble, à première vue, meilleure car elle est à la fois additive pour un produit tensoriel et invariante par opération gaussienne. Toutes ces mesures sont calculables numériquement à partir des résultats d'une tomographie ; la plus facilement calculable est évidemment la première et la plus compliquée la deuxième puisqu'elle demande de diagonaliser la matrice densité, mais cela reste néanmoins largement à la portée des outils numériques actuels. Pour un modèle analytique c'est bien plus compliqué : les deux dernières sont difficilement calculables dans le cas général et la première, bien que calculable, mène généralement à des formules compliquées et peu utilisables en pratique. Le seul cas réellement calculable analytiquement est la deuxième mesure pour un état pur puisqu'elle ne dépend alors que de la matrice de covariance.

Caractéristiques	δ_1	δ_2	δ_3
invariance par opération unitaire symplectique	oui	oui	seulement déplacements et opérations correspondant à un système linéaire passif
invariance par changement d'échelle	non	non	oui
additivité en cas de produit tensoriel	non	oui	oui
bornée	oui (1/2)	non	non

TABLE 6.1: Principales caractéristiques des différentes mesures de non-gaussianité

6.1.3 Non gaussianité et non-classicité

On peut sentir qu'il existe un certain lien entre la non-gaussianité et la non-classicité des états ; c'est notamment le cas des états purs non gaussiens qui sont aussi non classiques, dans le sens où leur fonction de Wigner est négative. Il serait donc intéressant de comparer les deux, mais pour cela il nous faut également une mesure de non-classicité. Ces mesures sont elles aussi généralement basées sur des distances ; le problème est qu'il n'y a pas de référence classique bien définie comme pour la référence gaussienne. Il faut alors prendre la distance avec l'ensemble des états considérés comme classiques, ce qui demande alors une minimisation rendant la mesure peu calculable en pratique. De plus il existe deux définitions possibles des états classiques.

La première définition consiste à considérer uniquement les états cohérents et les états thermiques déplacés (correspondant à un état cohérent amplifié par un amplificateur paramétrique indépendant de la phase). C'est certainement la plus utilisée ; tout comme pour la non-gaussianité on trouve des mesures basées sur la distance de Hilbert-Schmidt [195, 196] et sur l'entropie relative [197]. Cette définition implique que des états gaussiens peuvent être non-classiques (états comprimés).

La deuxième, que nous avons utilisée jusqu'ici, considère tous les états à fonction de Wigner positive comme classiques. Cette définition a notamment donné une mesure basée sur l'entropie relative [198]. Avec cette mesure, les états purs non-gaussiens sont tous non-classiques et vice-versa. Par contre il peut y avoir des états mixtes non-gaussiens mais classiques. Nous continuerons d'utiliser cette définition pour notre comparaison ; de plus afin de ne pas devoir minimiser sur un ensemble d'état nous réutiliserons l'estimation déjà présente dans les chapitres précédents, bien qu'elle ne constitue pas une vraie mesure, à savoir la valeur de la fonction de Wigner en son point le plus négatif³

$$\nu(\hat{\rho}) = \frac{\min_{x,p} (W(x,p))}{\min_{x,p} (W_{\text{phot}}(x,p))} \quad (6.21)$$

où nous avons choisi le photon unique comme référence car c'est l'état ayant la plus forte négativité parmi tous ceux qui sont présentés dans ce manuscrit. Nous avons donc $0 \leq \nu(\hat{\rho}) \leq 1$, avec $\nu(\hat{\rho}) = 0$ uniquement pour un état classique.

6.2 Réalisation expérimentale

6.2.1 Montage

Pour apporter une démonstration expérimentale de la mesure de non gaussianité, la première question à se poser est : quel état utiliser ? Les états non gaussiens les plus « courants » sont les états de Fock ou encore des vides comprimés dégaussifiés. Seulement l'idéal serait de pouvoir passer continûment d'un état fortement non gaussien à un état quasiment gaussien, et ce n'est pas possible avec ces états qui sont toujours fortement non-gaussiens. On pourrait aussi penser à utiliser la sortie de l'amplificateur sans bruit mais on tombe alors dans le défaut inverse : la non-gaussianité reste assez peu marquée et c'est encore plus vrai pour la non-classicité. Conservons l'idée de dégaussifier un état cohérent : quelle autre méthode existe-t-il ? Si la technique habituelle de suppression d'un photon n'a aucun effet, ce n'est par contre pas le cas de l'opération inverse, à savoir l'ajout d'un photon. Cette opération a en plus toutes les caractéristiques voulues : sur un état d'amplitude nulle elle donne un photon unique, à la fois fortement non gaussien et non classique, tandis que sur un état de grande amplitude son action est négligeable, permettant ainsi à l'état de garder un caractère quasiment gaussien et classique.

Le montage est un mélange entre l'amplification paramétrique d'un faisceau cohérent et la génération de photon unique (figure 6.2) : comme pour la première un état cohérent est introduit dans le mode signal de l'OPA et de manière identique à la seconde on conditionne sur la mesure d'un photon unique dans le mode complémentaire. Le principe est similaire à celui qui permet de créer un état de Fock à partir d'un état EPR : s'il n'y a qu'un photon dans le mode complémentaire, c'est que nous avons ajouté un seul photon à l'état cohérent. Nous devons bien entendu veiller à ce que la probabilité d'avoir plus d'un photon dans le mode complémentaire soit

3. Notons que jusqu'ici nous avons uniquement besoin de regarder la valeur à l'origine car c'était à chaque fois le point le plus négatif.

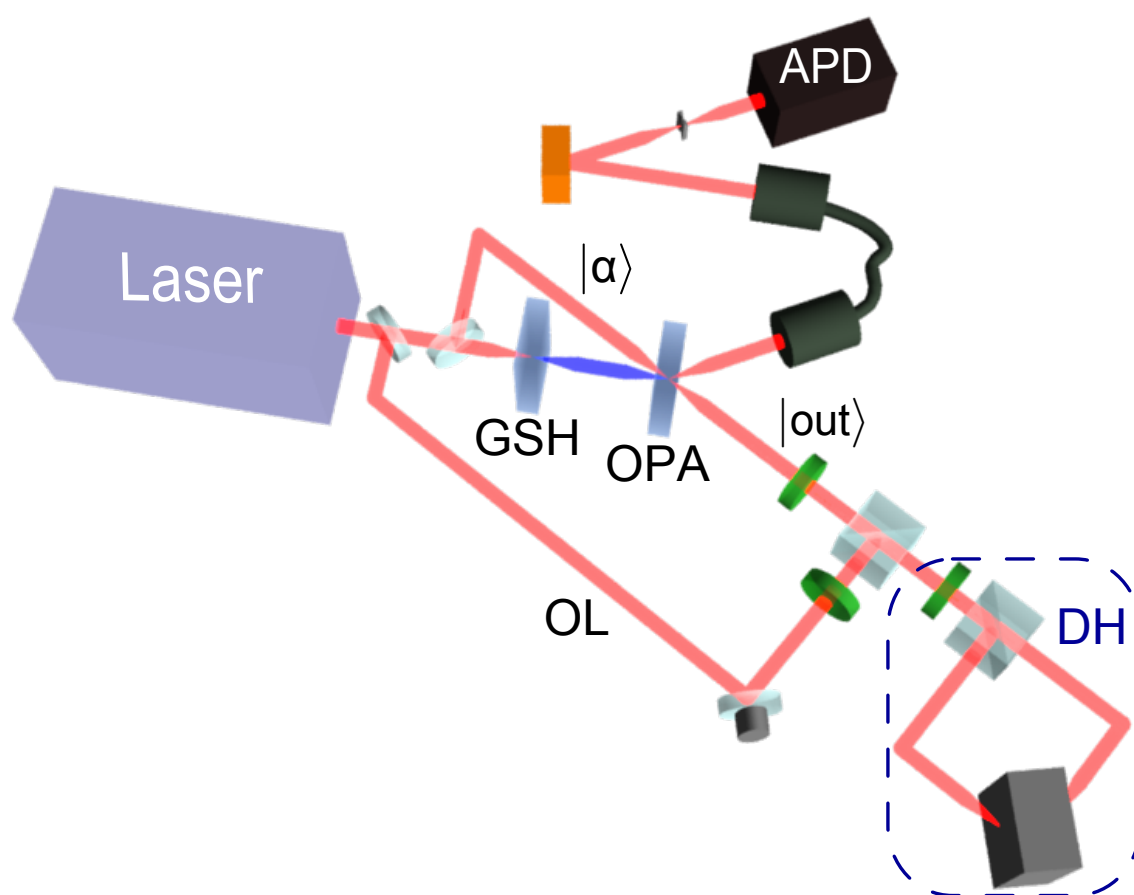


FIGURE 6.2: Schéma expérimental pour l'ajout d'un photon à un état cohérent

négligeable, ce qui implique les mêmes restrictions qu'au chapitre précédent : faible compression pour l'OPA et faible amplitude de l'état cohérent qui devra donc être fortement atténué.

Pour le contrôle de la phase nous utiliserons uniquement une référence qui sera cette fois constituée par le faisceau lui-même en prenant les données non conditionnées qui correspondent à l'état cohérent en entrée amplifié par l'OPA. Les paramètres expérimentaux de l'expérience (obtenus à partir des données expérimentales) sont donnés table 6.2. Ils correspondent aux paramètres du photon unique (obtenu pour $\alpha = 0$) $\delta = 1, 15$ et $\sigma = 1, 002$.

Paramètres	r	γ	ξ	η
Valeurs	0,105	0,425	0,96	0,71

TABLE 6.2: Valeurs des paramètres caractérisant les imperfections expérimentales

6.2.2 Modélisation

La modélisation de l'expérience suit le même principe que les précédentes, à nouveau nous allons déplacer les pertes homodynes en sortie de l'OPA et remplacer celles de la voie APD (μ) par $\nu = \mu/\eta$.

Fonction de Wigner Nous allons partir d'un état bimode correspondant à un état cohérent dans le mode signal et le vide dans le mode complémentaire

$$W_{\text{in}}(x, p, x_c, x_c) = \frac{1}{\pi^2} e^{-(x-\sqrt{2}\alpha)^2 - p^2 - x_c^2 - p_c^2} \quad (6.22)$$

auquel nous appliquons la transformation correspondant à l'amplification paramétrique (équation 3.9) avec un gain $g = \cosh^2(r)$ et un facteur de compression $s = e^{-2r}$.

$$W_g(x, p, x_c, x_c) = \frac{1}{\pi^2} e^{-\frac{(x+x_c-\sqrt{2}\frac{\alpha}{\sqrt{s}})^2}{2/\sqrt{s}} - \frac{(x-x_c-\sqrt{2}\sqrt{s}\alpha)^2}{2\sqrt{s}} - \frac{(p-p_c-\sqrt{2}\frac{\alpha}{\sqrt{s}})^2}{2/\sqrt{s}} - \frac{(p+p_c-\sqrt{2}\sqrt{s}\alpha)^2}{2\sqrt{s}}} \quad (6.23)$$

Nous pouvons alors ajouter l'amplification parasite puis les pertes homodynes sur chacun des deux modes en utilisant les formules de la section A.2 ce qui donne

$$W_\eta(x, p, x_c, x_c) = \frac{1}{\pi^2} e^{-\frac{(x+x_c-\sqrt{2}g_{\text{max}}\alpha)^2}{2b} - \frac{(x-x_c-\sqrt{2}g_{\text{min}}\alpha)^2}{a} - \frac{(p-p_c-\sqrt{2}g_{\text{max}}\alpha)^2}{2b} - \frac{(p+p_c-\sqrt{2}g_{\text{min}}\alpha)^2}{2a}} \quad (6.24)$$

avec

$$g_{\text{max}} = \sqrt{\frac{\eta h}{s}} \quad (6.25)$$

$$g_{\text{min}} = \sqrt{\eta h s} \quad (6.26)$$

$$a = \eta(hs + h - 1) + 1 - \eta \quad (6.27)$$

$$b = \eta\left(\frac{h}{s} + h - 1\right) + 1 - \eta \quad (6.28)$$

Il ne nous reste plus qu'à appliquer le conditionnement

$$W_{\text{cond}}(x, p) = \frac{\iint W_\eta(x, p, x_c, p_c) W_n(x_c, p_c) dx_c dp_c}{\iiint W_\eta(x, p, x_c, p_c) W_n(x_c, p_c) dx dp dx_c dp_c} \quad (6.29)$$

et à effectuer le mélange avec l'état non conditionné correspondant à l'amplification paramétrique d'un état cohérent

$$W_{\text{aj}}(x, p) = \xi W_{\text{cond}}(x, p) + (1 - \xi) W_{\text{n.c.}}(x, p) \quad (6.30)$$

avec

$$W_{\text{n.c.}}(x, p) = \iint W_{\eta}(x, p, x_c, x_c) dx_c dp_c \quad (6.31)$$

Ce qui donne

$$W_{\text{aj}}(x, p) = \frac{1}{\pi\sigma^2} \left(1 - \delta_{\alpha} - \zeta_{\alpha} + \delta_{\alpha} \frac{(x - \sqrt{2}\kappa\alpha)^2 + p^2}{\sigma^2} \right) e^{-\frac{(x - \sqrt{2}\sqrt{g_{\text{eff}}}\alpha)^2}{\sigma^2} - \frac{p^2}{\sigma^2}} \quad (6.32)$$

avec les paramètres qui valent

$$\delta_{\alpha} = \frac{\delta}{1 + \frac{h(g-1)}{hg-1}\alpha^2} \quad (6.33)$$

$$\zeta_{\alpha} = \frac{\delta_{\alpha}\sigma^2\alpha^2}{2g_{\text{eff}}} \quad (6.34)$$

$$\kappa = \frac{2\eta - 1}{2\sqrt{g_{\text{eff}}}} \quad (6.35)$$

$$g_{\text{eff}} = \eta gh \quad (6.36)$$

où δ et σ sont les paramètres du photon unique.

Référence gaussienne À partir de la fonction de Wigner il est possible de calculer les moments d'ordre 1 et 2 permettant de définir la référence gaussienne :

$$\langle x \rangle = \sqrt{2} \frac{\delta_{\alpha}\sigma^2 + 2g_{\text{eff}}}{2\sqrt{g_{\text{eff}}}} \alpha \quad (6.37)$$

$$\langle p \rangle = 0 \quad (6.38)$$

$$\langle \delta x^2 \rangle = (1 + \delta_{\alpha}(1 - 2\zeta_{\alpha})) \frac{\sigma^2}{2} \quad (6.39)$$

$$\langle \delta p^2 \rangle = (1 + \delta_{\alpha}) \frac{\sigma^2}{2} \quad (6.40)$$

$$\langle \delta x \delta p \rangle = 0 \quad (6.41)$$

La référence gaussienne vaut alors

$$W_{\text{ref}}(x, p) = \frac{1}{2\pi\sqrt{\langle \delta x^2 \rangle \langle \delta p^2 \rangle}} e^{-\frac{(x - \langle x \rangle)^2}{2\langle \delta x^2 \rangle} - \frac{p^2}{2\langle \delta p^2 \rangle}} \quad (6.42)$$

Pureté et mesure δ_1 Nous pouvons aussi calculer la pureté de l'état

$$\mathcal{P} = \frac{2 - 2(1 - \zeta_{\alpha})\delta_{\alpha} + \delta_{\alpha}^2}{2\sigma^2} \quad (6.43)$$

La distance de Hilbert-Schmidt pourrait aussi être calculée analytiquement, mais cela conduit à une formule longue et peu commode à utiliser, il est donc préférable d'utiliser un modèle numérique (autant du point de vue de la simplicité que du temps utilisé et de la fiabilité vis-à-vis des erreurs de calculs). Nous pourrions néanmoins effectuer la décomposition de celle-ci en fonction des puretés et du recouvrement et utiliser les équations 6.3 et 6.43 pour les puretés.

Entropie de la référence gaussienne et mesure δ_2 Pour la deuxième mesure les choses se compliquent ; nous ne pouvons en effet calculer analytiquement que l'entropie de la référence gaussienne qui vaut

$$S(W_{\text{ref}}) = f \left(\sqrt{(1 + \delta_\alpha)(1 + \delta_\alpha(1 - 2\zeta_\alpha))} \frac{\sigma^2}{2} \right) \quad (6.44)$$

Ceci nous permet cependant d'obtenir une formule dans le cas idéal $g, h, \xi \rightarrow 1$ où l'état obtenu est pur :

$$\delta_2(W_{\text{aj}}) = f \left(\frac{1}{2} \sqrt{1 + \frac{4\eta(1 + \eta - \alpha^2(\eta - 1))}{(1 + \alpha^2)^3}} \right) \quad (6.45)$$

Pour le cas réel nous devons utiliser un modèle numérique.

Fonction Q et mesure δ_3 Enfin nous pouvons aussi utiliser l'équation 6.18 afin d'obtenir l'expression de la fonction Q

$$Q_{\text{aj}}(x, p) = \frac{1}{\pi(1 + \sigma^2)^2} \left(1 - \delta_\alpha \sigma^2 - \zeta_\alpha(1 + \sigma^2) + \sigma^2 + \delta_\alpha \sigma^2 \frac{(x - \sqrt{2}\kappa_Q \alpha)^2 + p^2}{1 + \sigma^2} \right) \times e^{-\frac{(x - \sqrt{2}\sqrt{g_{\text{eff}}}\alpha)^2}{1 + \sigma^2} - \frac{p^2}{1 + \sigma^2}} \quad (6.46)$$

avec

$$\kappa_Q = \frac{\kappa(1 + \sigma^2) - \sqrt{g_{\text{eff}}}}{\sigma^2} \quad (6.47)$$

La référence gaussienne peut être trouvée de la même manière

$$Q_{\text{ref}}(x, p) = \frac{1}{\pi \sqrt{(2\langle \delta x^2 \rangle + 1)(2\langle \delta p^2 \rangle + 1)}} e^{-\frac{(x - \langle x \rangle)^2}{2\langle \delta x^2 \rangle + 1} - \frac{p^2}{2\langle \delta p^2 \rangle + 1}} \quad (6.48)$$

Nous pouvons alors calculer l'entropie de la référence gaussienne

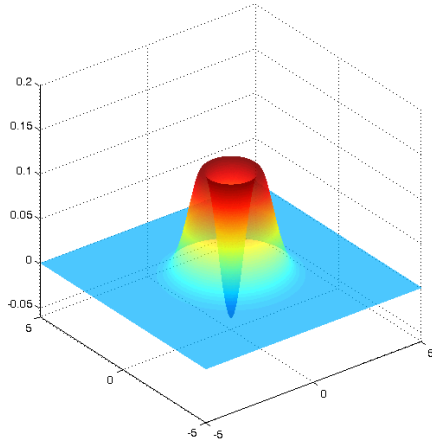
$$H(Q_{\text{ref}}) = \log_2(2\pi e) + \frac{\log_2(2\langle \delta x^2 \rangle + 1) + \log_2(2\langle \delta p^2 \rangle + 1)}{2} \quad (6.49)$$

L'entropie de l'état produit nécessitera par contre elle aussi l'utilisation d'un calcul numérique.

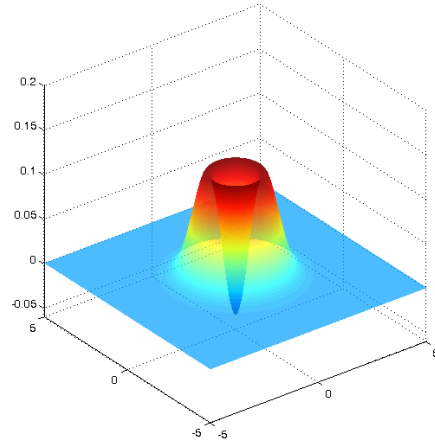
6.2.3 Résultats

Les données ont été prises en collaboration avec Nicolò Spagnolo venu passer quelques mois chez nous. Pour chaque valeur de l'état cohérent en entrée, 800 000 points ont été pris et répartis en 12 histogrammes. Le tri s'effectuant sur l'état cohérent légèrement amplifié et non sur une référence auxiliaire, nous ne pouvons pas descendre en deçà de $\alpha \approx 0,5$: en dessous de cette limite la résolution n'est plus suffisante pour trier correctement les données. Ceci n'est évidemment pas valable pour $\alpha = 0$, correspondant au photon unique, qui ne nécessite pas de tri.

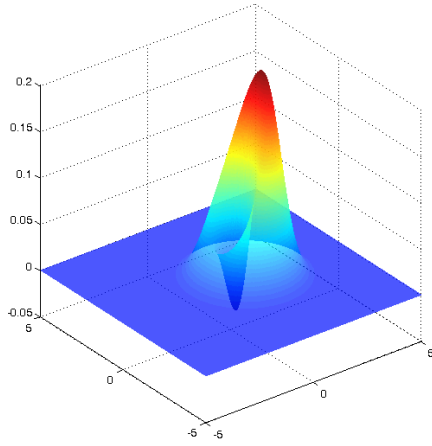
La figure 6.3 montre les distributions de Wigner expérimentales [199], reconstruites par l'algorithme de maximum de vraisemblance, et théoriques. La fidélité entre les deux vaut $\mathcal{F} = (98,9 \pm 0,6)\%$. Nous pouvons nettement voir que plus l'état cohérent en entrée est grand, plus la distribution se rapproche d'une gaussienne. De plus la négativité des fonctions de Wigner est bien visible (sans correction des pertes homodynes), ce qui permet à notre critère de non-classicité d'être pertinent.



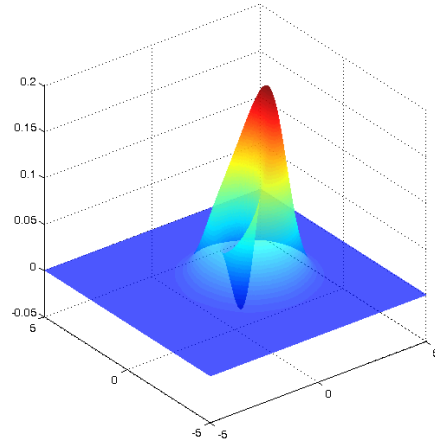
(a)



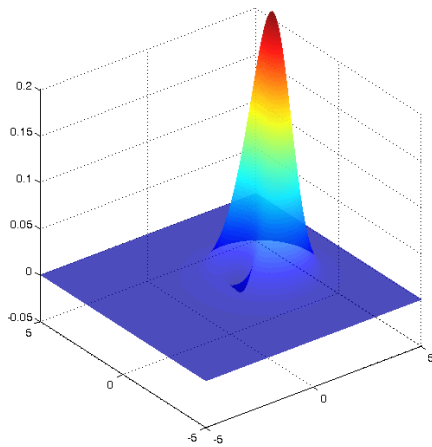
(b)



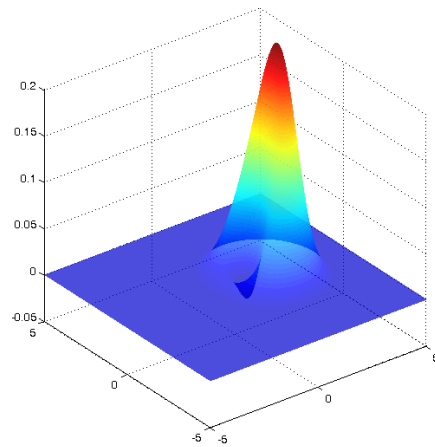
(c)



(d)



(e)



(f)

FIGURE 6.3: Fonctions de Wigner (a) théorique pour $\alpha = 0$ (b) expérimentale pour $\alpha = 0$ (c) théorique pour $\alpha = 0,48$ (d) expérimentale pour $\alpha = 0,48$ (e) théorique pour $\alpha = 1,02$ (f) expérimentale pour $\alpha = 1,02$

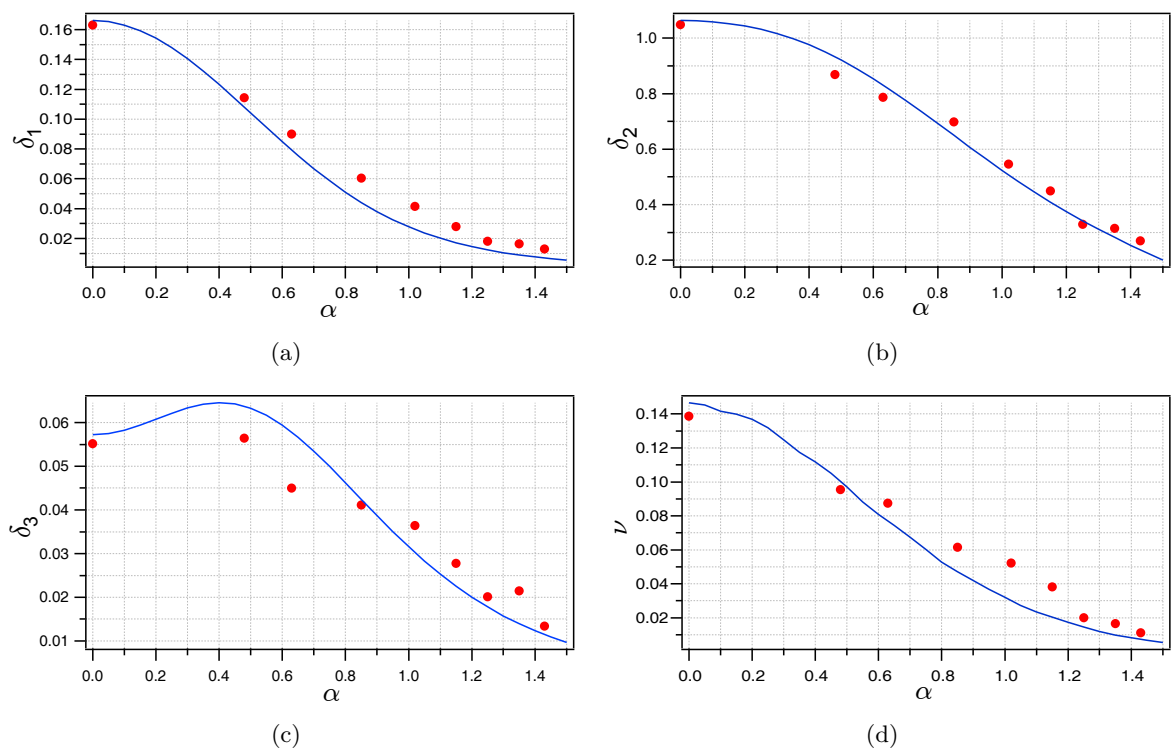


FIGURE 6.4: Différentes mesures de non-gaussianité et estimation de la non-classicité. (a) avec la distance d'Hilbert-Schmidt (b) avec la néguentropie (c) avec l'entropie de Wehrl (d) non-classicité

À partir de ces données nous pouvons comparer les différentes mesures de non-gaussianité et effectuer le parallèle avec notre estimation de la non-classicité. Ces résultats sont illustrés figure 6.4. Les courbes en trait plein correspondent au modèle et les points aux données expérimentales. Les différentes mesures ont globalement bien le comportement attendu. La troisième, calculée à partir de la fonction Q , présente cependant un comportement singulier pour les plus petits états cohérents, qui sont malheureusement dans la zone qui ne peut être vérifiée expérimentalement par notre dispositif ; nous reviendrons plus en détail sur ce point un peu plus loin. De plus la mesure expérimentale est moins « régulière » que pour les deux premières, laissant penser à une plus forte dépendance vis-à-vis du bruit expérimental.

La première mesure, basée sur la distance d'Hilbert-Schmidt, présente clairement un biais entre le modèle et les données expérimentales. Ceci est intrinsèque à la mesure. Pour s'en convaincre regardons ce qu'il se passe si on veut utiliser expérimentalement cette mesure sur un état gaussien $\hat{\rho}$: sa fonction de Wigner correspond alors à celle de l'état de référence auquel vient s'ajouter, en plus du bruit dû aux imperfections qui sont modelisables, le bruit que l'on observe sur toute mesure expérimentale (ce bruit est essentiellement dû au fait que toute mesure est constitué d'un échantillon fini de points alors qu'une reconstruction parfaite demande une infinité de points), soit $W_{\rho}(x, p) = W_{\tau}(x, p) + \delta W(x, p)$ avec $\langle \delta W \rangle = 0$ (notons bien que $\delta W(x, p)$ n'est pas une distribution mais un ensemble de variables aléatoires paramétrisées par x et p), le carré de la distance d'Hilbert-Schmidt vaut alors

$$D_{\text{HS}}(\hat{\rho}, \hat{\tau}) = \pi \iint \delta W(x, p)^2 dx dp \propto \langle \delta W^2 \rangle \quad (6.50)$$

Au lieu d'être nulle, cette distance est donc proportionnelle à la variance du bruit expérimental. Dans le cas général il est facile de voir que celui-ci va entraîner une surestimation de la non-gaussianité.

Un raisonnement similaire peut-être tenu pour notre estimation de la non-classicité : en prenant le minimum nous allons généralement prendre un point artificiellement abaissé par ce bruit, ce qui entraîne en moyenne une surestimation. C'est en fait même pire car cette estimation est en plus sensible aux données aberrantes : imaginons que l'on ai un point non significatif qui soit anormalement bas c'est alors lui qui sera utilisé pour estimer la non-classicité, faussant ainsi la mesure. Nous devons cependant garder en mémoire que la reconstruction de la fonction de Wigner lisse celle-ci, réduisant ainsi ce bruit expérimental, il est du coup d'autant plus difficile d'en déterminer l'importance indépendamment des effets que nous venons de citer.

Puisque nous avons démontré expérimentalement la validité de notre modèle nous pouvons utiliser celui-ci afin d'analyser le comportement des différentes mesures en fonction des valeurs de paramètres expérimentaux. Le premier paramètre sur lequel nous allons nous pencher est le paramètre de compression r de l'OPA. La figure 6.5 montre que les deux premières mesures de la non-gaussianités et notre estimation de la non-classicité donnent des résultats qui décroissent lorsque la compression, et donc le gain paramétrique, augmente. Ce phénomène vient du fait que la probabilité d'ajouter plus d'un photon doit être négligeable puisque nous ne sommes pas capables de la séparer du cas où un seul photon est ajouté. Or plus le gain paramétrique augmente, moins c'est le cas. L'état en sortie devient un mélange statistique de plus en plus important et qui a tendance à nous ramener à un état gaussien. La troisième mesure de non-gaussianité a un comportement contraire, mais celui-ci dépend uniquement des autres imperfections (excès de gain, pureté modale et efficacité homodyne) : sans elles cette mesure resterait la même pour toutes les valeurs de r , tandis que les autres mesures présentent toujours le même comportement.

Intéressons-nous maintenant à l'impact de ces autres imperfections, à savoir l'excès de gain représenté par le rapport γ , la pureté modale ξ et l'efficacité homodyne η . Ses deux derniers

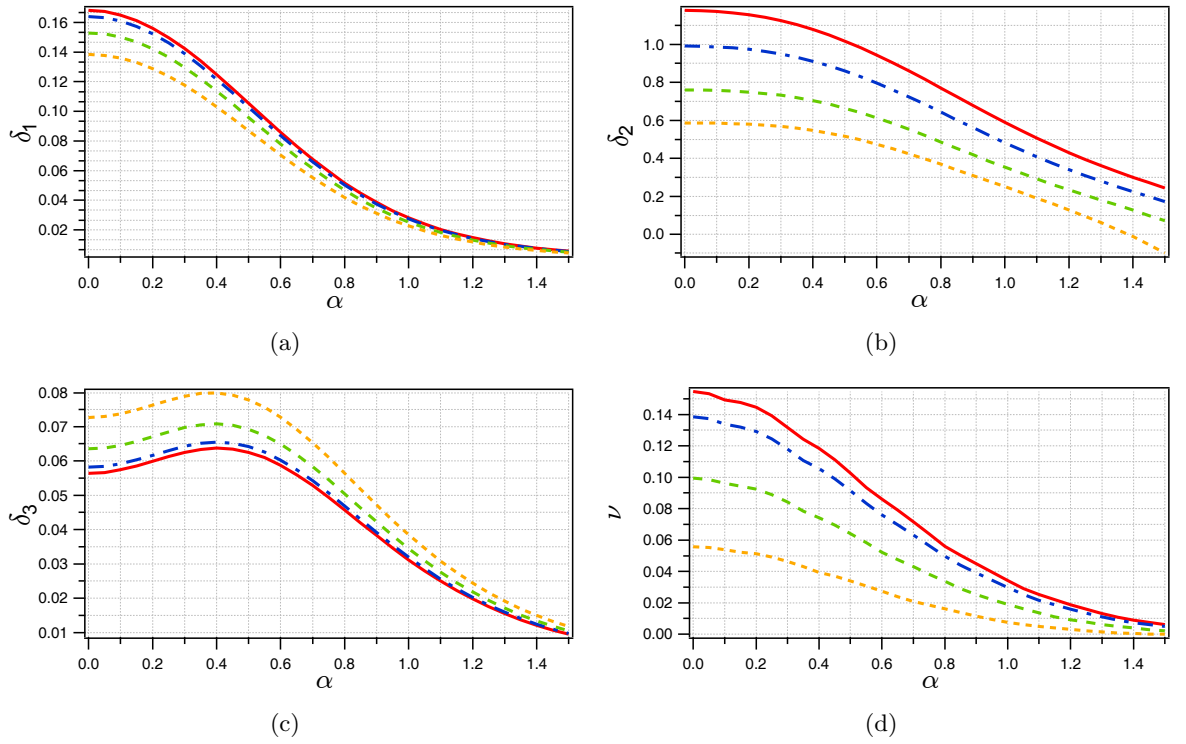


FIGURE 6.5: Influence du paramètre de compression, les courbes correspondent à $r = \approx 0; 0, 15; 0, 30; 0, 45$. (a) non-gaussianité avec la distance d'Hilbert-Schmidt (b) non-gaussianité avec la néguentropie (c) non-gaussianité avec l'entropie de Wehrl (d) non-classicité

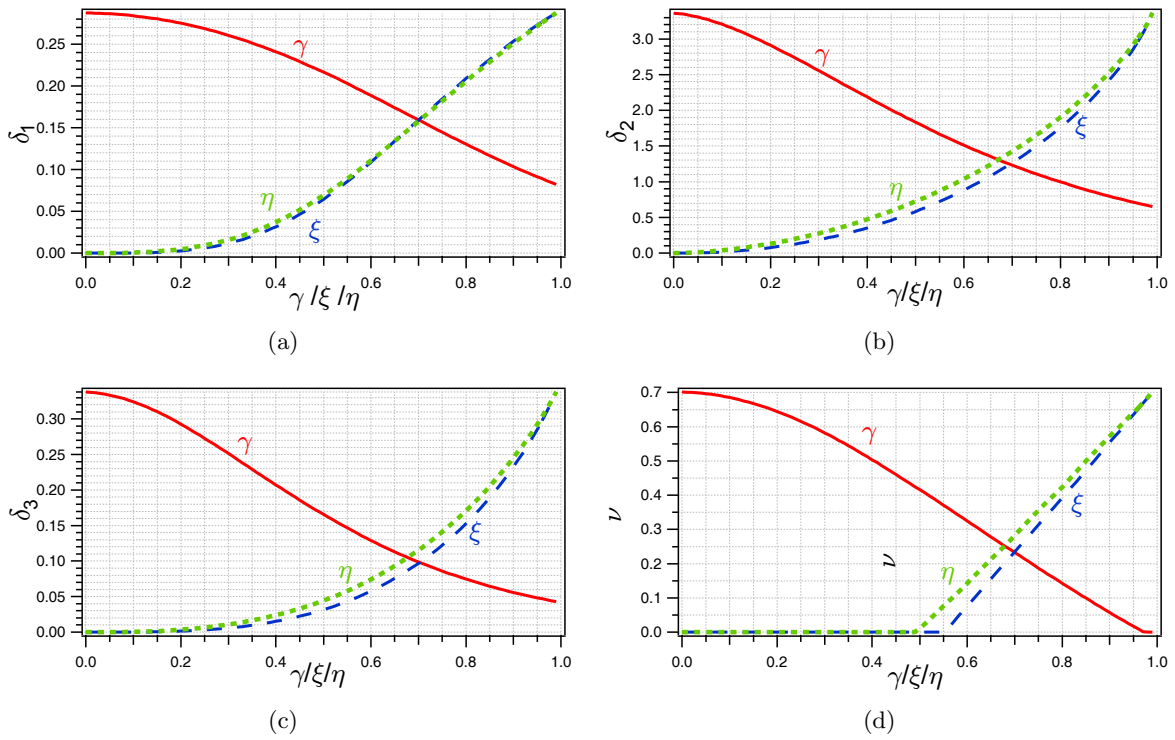


FIGURE 6.6: Influence des imperfections (a) non-gaussianité avec la distance d'Hilbert-Schmidt (b) non-gaussianité avec la néguentropie (c) non-gaussianité avec l'entropie de Wehrl (d) non-classicité

induisent un mélange avec un état gaussien et classique (respectivement un état cohérent amplifié et le vide) qui est d'autant plus à l'avantage de ce dernier que le paramètre est proche de 0 ; la non-gaussianité décroît donc avec ceux-ci. Cette similitude entre leurs effets se voit même dans l'importance de ceux-ci puisque les courbes sont très proche. Ils sont néanmoins définis de manières totalement différentes, ce qui explique que l'efficacité homodyne a un effet légèrement moins important. À l'inverse l'augmentation du paramètre γ diminue la non gaussianité et la non classicité. Son effet est double : il peut d'une part ajouter des photons supplémentaires à l'état cohérent, rapprochant l'état en sortie d'une simple amplification, mais il peut aussi rajouter des photons dans le mode complémentaire, déclenchant ainsi un clic sur l'APD alors qu'aucun photon n'est ajouté à l'état cohérent. Dans les deux cas le résultat est le même que précédemment : on trouve en sortie un mélange avec un état gaussien ; cette fois la part de ce dernier est d'autant plus importante que l'excès de gain est grand. Contrairement au cas précédent toutes les courbes ont le même comportement, nous pouvons juste noter l'effet de seuil dans notre estimation de la non-classicité qui est entièrement due à la définition utilisée, puisque toutes les fonctions de Wigner positives ont le même minimum (0).

Il reste encore un point que nous n'avons pas approfondi : le comportement de la troisième mesure de non-gaussianité, basée sur l'entropie de Wehrl, parfois contraire à celui des autres mesures et à ce que l'on s'attendrait à obtenir. Nous avons pu voir deux points qui posent problèmes :

- L'augmentation de la non-gaussianité avec le paramètre de compression alors que l'on s'attendrait plutôt à l'inverse. Nous pouvons relier ceci au fait que cette mesure de non-gaussianité n'est pas conservée lors d'une compression ; l'explication la plus probable est que les effets de cette non-conservation prédominent sur le comportement attendu produit par l'augmentation du mélange statistique.
- Une légère montée de la non-gaussianité avec l'amplitude de l'état cohérent lorsque celui-ci est faible. Pour éclairer ce phénomène nous devons d'abord rappeler que la troisième mesure (tout comme la deuxième) ne repose pas sur une vraie distance, au sens mathématique, mais sur une quantité liée à l'information. Rien ne nous garantit alors que la mesure se comporte toujours comme une distance ; des effets liés à l'information peuvent tout à fait venir modifier le comportement attendu. Le modèle montre d'ailleurs une forte évolution de l'entropie de Wehrl de l'état non gaussien et de sa référence gaussienne (figure 6.7). Nous pouvons donc émettre l'hypothèse que ces effets sont la cause de cette montée de la non-gaussianité. Cette hypothèse est renforcée par le fait que cette montée dépend des imperfections qui modifient le degré de pureté et donc l'information contenue dans les états observés.

Enfin, la figure 6.7 nous montre que la troisième mesure de non gaussianité correspond à une petite différence entre deux grandeurs ; ce qui amplifie l'importance du bruit expérimental. Cet aspect permet d'expliquer la plus faible régularité des résultats expérimentaux pour cette mesure.

6.3 Discussion

6.3.1 Non-gaussianité de quelques états courants

Nous pouvons nous demander quelle est la non-gaussianité des états non-gaussiens courants, notamment ceux présentés au chapitre 2.

Nous pouvons trouver dans la littérature l'expression des trois mesures de non-gaussianité

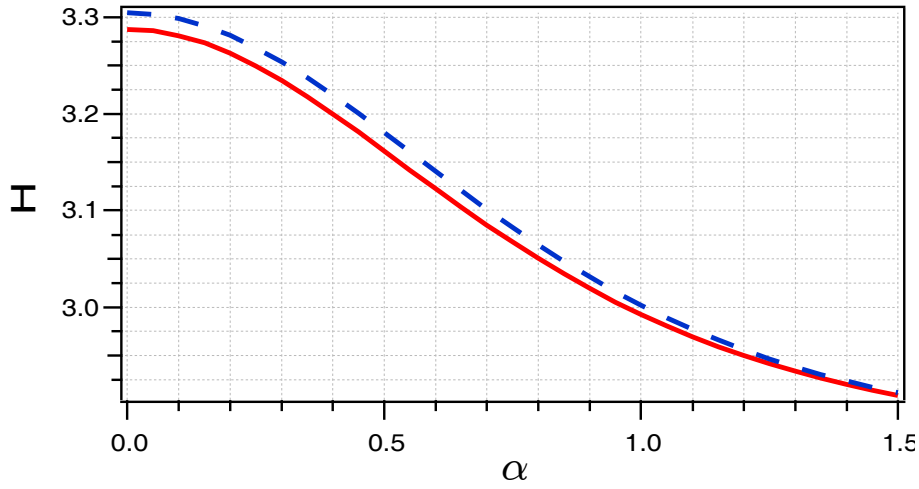


FIGURE 6.7: Entropie de Werhl de l'état non gaussien (trait plein) et de la référence gaussienne (tirets) d'après le modèle.

pour des états de Fock [189, 200] :

$$\delta_1(n) = \frac{1}{2} \left(1 + \frac{1}{2n} - 2 \frac{n^n}{(n+1)^{n+1}} \right) \quad (6.51)$$

$$\delta_2(n) = f \left(n + \frac{1}{2} \right) \quad (6.52)$$

$$\delta_3(n) = \ln(n+1) - m - \ln(n!) + n\psi(n+1) \quad (6.53)$$

où la fonction f est celle de l'équation 6.5 et ψ est la fonction *digamma* dont l'expression, comprenant la constante d'Euler $\gamma = 0,5772\dots$, vaut

$$\psi(n) = \sum_{k=1}^{n-1} \frac{1}{k} - \gamma \quad (6.54)$$

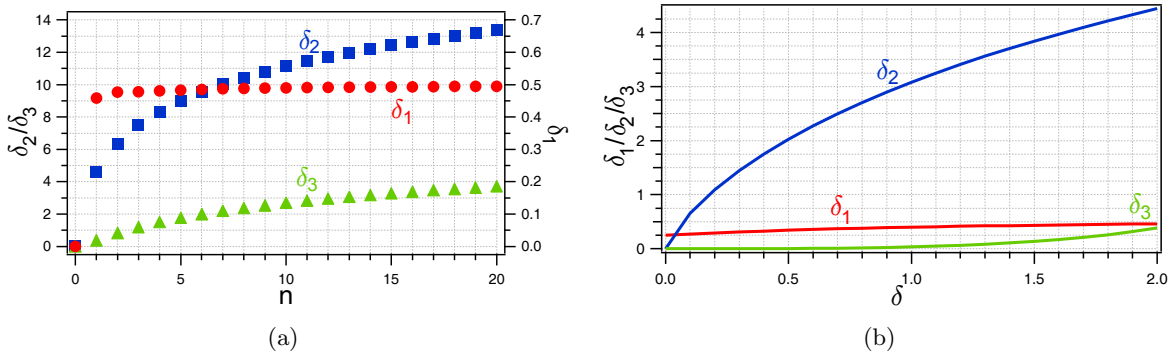


FIGURE 6.8: Non-gaussianité des états de Fock (a) États de Fock idéaux (b) Photon unique réel produit par notre dispositif

La figure 6.8(a) illustre la non-gaussianité des états de Fock suivant les trois mesures. Nous pouvons remarquer que la non-gaussianité augmente bien avec le nombre de photons, ce qui est

pleinement cohérent avec la structure de la fonction de Wigner qui fait intervenir des polynômes de degré de plus en plus grand. L'augmentation de la première mesure, basée sur la distance de Hilbert-Schmidt, reste cependant « faible » de par le fait que cette mesure est bornée, contrairement aux autres mesures qui tendent vers l'infini quand on accroît le nombre de photon. Notons aussi que pour la deuxième mesure, basée sur la néguentropie, les états de Fock maximisent la non-gaussianité à nombre moyen de photons fixé (ce qui paraît cohérent étant donné leur lien avec les variables discrètes) [189]. La première mesure semble avoir le même comportement (mais cela n'a pas encore pu être prouvé théoriquement), ce n'est par contre pas le cas de la troisième mesure de non-gaussianité. En revenant à un cas plus réel nous pouvons aussi calculer numériquement la non-gaussianité des photons uniques produits par notre système en fonction de la valeur du paramètre δ pour $\sigma \rightarrow 1$ (figure 6.8(b)). Rappelons que σ correspond à l'écart-type de l'état thermique obtenu sans le conditionnement, et qu'il peut être rendu aussi faible que l'on veut en baissant le gain de l'OPA (au prix d'une baisse du taux de production).

Les autres états non-gaussiens dont nous avons beaucoup parlé sont les états chats de Schrödinger, la figure 6.9 montre leur non gaussianité. Là aussi la non-gaussianité augmente avec la taille, ce qui est tout à fait en accord avec les jugements « à l'œil » que l'on peut faire en regardant la fonction de Wigner. Nous retrouvons aussi la même différence que précédemment entre la première mesure, qui est bornée, et les deux autres qui tendent vers l'infini. Enfin nous pouvons remarquer la différence de non-gaussianité suivant la parité pour les petits chats. Nous rappellerons pour l'interpréter que pour une faible amplitude le chat pair se rapproche d'un vide comprimé (donc d'un état gaussien), tandis qu'un chat impair se rapproche d'un état chaton (ou photon unique comprimé) qui n'est pas gaussien. Nous pouvons clairement voir pour les deux premières mesures avec un chat impair un palier à faible amplitude où la non-gaussianité est égale à celle d'un photon unique (comprimé ou non puisque ces mesures n'y sont pas sensibles), ce palier peut alors nous donner une indication sur la zone de validité de l'approximation d'un chat impair de petite taille par un photon unique comprimé.

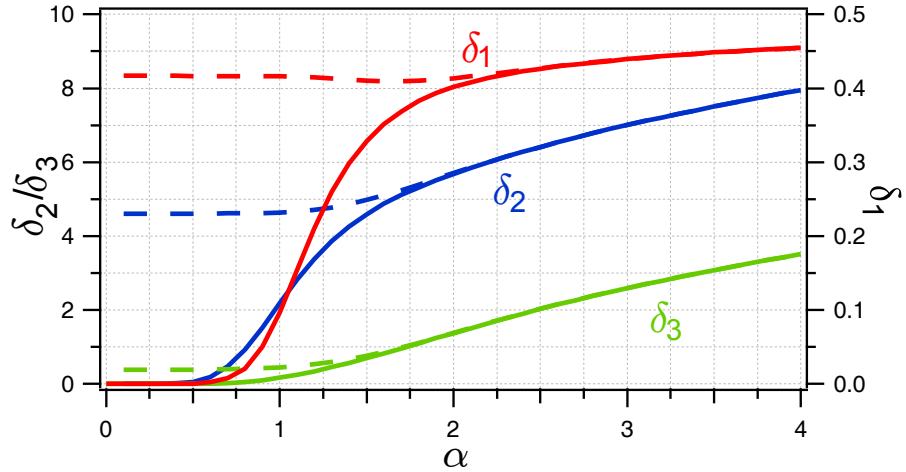


FIGURE 6.9: Non-gaussianité des états chats de Schrödinger pairs (traits pleins) et impairs (tirets)

6.3.2 Non-gaussianité des processus quantiques

À partir de ces mesures de non-gaussianité des états, nous pouvons aller plus loin et nous intéresser à la non-gaussianité des processus. nous avons notamment parlé en introduction de

ce chapitre d'opérations de gaussification et de dégaussification. Une telle mesure a été proposée par Marco Genoni *et al.* (déjà auteurs des deux premières mesures) [189] ; son expression est

$$\delta_n(\mathcal{E}) = \max_{\hat{\rho} \in \mathcal{G}} (\delta_n(\mathcal{E}(\hat{\rho}))) \quad (6.55)$$

où δ_n correspond aux différentes mesures de non-gaussianité étudiées, $\mathcal{E}$ est le superopérateur décrivant l'opération et $\mathcal{G}$ l'ensemble des états gaussiens.

Cette mesure n'est bien nulle que pour une opération gaussienne de par la définition de celle-ci (une opération gaussienne étant, rappelons-le, une opération qui transforme tout état gaussien en un autre état gaussien). Dans le cas de la deuxième mesure de non-gaussianité, nous pouvons prouver que pour un canal gaussien $\delta_2(\mathcal{E}(\hat{\rho})) \leq \delta_2(\hat{\rho})$ [192], ce qui correspond bien au comportement attendu, puisqu'un canal gaussien, outre qu'il s'agit d'une opération gaussienne, gaussifie l'état d'entrée (la gaussification étant due aux pertes du canal).

Malheureusement cette définition est loin d'être facilement utilisable en pratique de par la maximisation, et il n'existe pas à l'heure actuelle d'autres propositions pour quantifier la non-gaussianité d'un processus quantique.

6.3.3 Non-gaussianité et information

Le fait que la deuxième mesure de non-gaussianité utilise l'entropie de Von Neumann permet d'effectuer un lien entre la non-gaussianité et l'information utilisable. La troisième mesure utilise elle aussi une entropie, seulement celle-ci tient aussi compte de la méconnaissance due à l'incertitude sur la mesure qui ne peut être utilisée pour coder des informations.

La première quantité que nous pouvons relier à la non-gaussianité est l'entropie d'Holevo χ , qui définit le nombre de bits classique que l'on peut coder dans un état quantique, et vaut $\chi(\hat{\rho}) = S(\hat{\rho}) - \sum_i p_i \hat{\rho}_i$ où les $\hat{\rho}_i$ sont les états utilisés pour encoder le message. En considérant que ces derniers sont des états purs, nous obtenons immédiatement

$$\chi(\hat{\rho}) = S(\hat{\tau}) - \delta_2(\hat{\rho}) \quad (6.56)$$

L'entropie d'Holevo constitue une limite de l'information qui peut être transmise par un état quantique ; nous retrouvons donc le fait, énoncé en introduction de ce chapitre, qu'un état gaussien maximise l'entropie pour une matrice de covariance fixée. La « perte d'entropie » est même égale à la non-gaussianité de l'état.

Nous pouvons aussi établir un lien avec l'information mutuelle I_{AB} . La différence entre l'information mutuelle d'un état bipartite non-gaussien $\hat{\rho}_{AB}$ et son homologue gaussien $\Delta I_{AB} = I_{AB}(\hat{\rho}_{AB}) - I_{AB}(\hat{\tau}_{AB})$ vaut

$$\Delta I_{AB} = \delta_2(\hat{\rho}_{AB}) - \delta_2(\hat{\rho}_A \otimes \hat{\rho}_B) \geq 0 \quad (6.57)$$

Nous retrouvons cette fois le fait que l'information mutuelle est minimale pour un état gaussien [162] ; la non-gaussianité permet donc de calculer l'apport d'un état non-gaussien à l'information mutuelle.

Enfin la dernière quantité que l'on peut lier à la non-gaussianité est l'entropie conditionnelle. Nous pouvons définir, de même que précédemment, la différence $\Delta S(A|B) = S_{\hat{\rho}}(A|B) - S_{\hat{\tau}}(A|B)$. Celle-ci vaut

$$\Delta S(A|B) = \delta_2(\hat{\rho}_B) - \delta_2(\hat{\rho}_{AB}) \leq 0 \quad (6.58)$$

Or nous avons vu, en introduction de ce manuscrit, que cette entropie conditionnelle peut être strictement négative, et que cette négativité est importante pour les communications quantiques.

6.3.4 Conclusion

Nous avons donc pu tester trois mesures différentes de la non-gaussianité d'un état quantique. L'une de celles-ci, basée sur la néguentropie, s'est révélée meilleure que les deux autres pour le cas étudié. La première basée sur la distance d'Hilbert-Schmidt présente un biais lors des mesures expérimentales, de plus le fait qu'elle est bornée diminue la résolution de la mesure pour des états fortement non-gaussiens. Quant à la troisième, basée sur l'entropie de Wehrl, elle souffre de n'être pas conservée par toutes les opérations gaussiennes et surtout de présenter des comportements contraires à celui attendu dans certains cas en présence de pertes. Nous devons cependant modérer ces propos en rappelant que nous n'avons étudié que quelques cas particuliers, ce qui permet uniquement de donner des pistes pour le comportement général mais pas d'en tirer des conclusions. Il est en effet possible que d'autres cas particuliers fassent apparaître de nouveaux effets rendant une autre mesure plus intéressante. La mesure δ_2 basée sur la néguentropie possède néanmoins un avantage considérable par rapport aux autres qui n'est pas lié à un cas particulier : les liens avec l'information qui permettent de facilement quantifier l'apport de la non-gaussianité à des protocoles d'information quantique.

Enfin nous avons pu voir une forte concordance entre la non-gaussianité et la non-classicité pour le cas étudié, ce qui suggère que la non-gaussianité observée est essentiellement d'origine non-classique. Ce point met en évidence l'intérêt qu'il y aurait à disposer de mesures traduisant à la fois la non-gaussianité et la non-classicité.

Troisième partie

Amélioration du protocole
expérimental

Chapitre 7

L'amplificateur femtoseconde

Sommaire

7.1 Amplification d'impulsions femtosecondes	149
7.1.1 Intérêt d'un amplificateur	149
7.1.2 Différents amplificateurs possibles	150
7.1.2.1 Amplification paramétrique	150
7.1.2.2 Amplification régénérative à base de saphir dopé au titane	151
7.1.2.3 Amplification passive à base de saphir dopé au titane	151
7.1.3 Application au Système expérimental	152
7.2 Dispositif	154
7.2.1 Vue d'ensemble	154
7.2.2 Cristaux utilisés	156
7.2.3 Refroidissement cryogénique	157
7.3 Intégration au dispositif expérimental	157
7.3.1 Faisceau amplifié	157
7.3.2 Doublage de fréquence	160
7.3.3 Amplification paramétrique	162
7.4 Travail restant	162

7.1 Amplification d'impulsions femtosecondes

7.1.1 Intérêt d'un amplificateur

Le principal facteur limitant de notre dispositif est la puissance des impulsions femtosecondes. Il n'existe pas de laser plus puissant qui satisferait aux différents critères (principalement le taux de répétition et la qualité du mode) requis pour mener à bien nos expériences (cf. section B.1). La solution pour obtenir des impulsions plus énergétiques consiste alors à amplifier les impulsions délivrées par notre laser. Le but est de disposer d'une plus grande puissance de faisceau bleu, afin de pomper plus efficacement l'OPA, ce qui apporte plusieurs améliorations :

- la possibilité de moins focaliser la pompe dans l'OPA afin de limiter les distorsions induites par le gain, et ainsi d'améliorer la pureté modale des états produits.

- une augmentation de l'intrication des états EPR et de la compression du vide comprimé, ce qui rend possible certaines expériences comme le clonage quantique [202] ou des tests sans échappatoires des inégalités de Bell [203].
- un nombre de photons moyens plus important, permettant de produire des états de Fock et des chats de Schrödinger plus grands et de réaliser en un temps raisonnable des expériences demandant plus de coïncidences.
- la possibilité d'utiliser plusieurs OPA afin de produire plus d'états quantiques, ce qui est nécessaire pour un nombre grandissant d'expériences.

7.1.2 Différents amplificateurs possibles

Il existe plusieurs techniques pour amplifier une impulsion femtoseconde, la plupart reposant sur le principe de l'*amplification à dérive de fréquence* (ou CPA pour *Chirped Pulse Amplification*) [204] : l'inconvénient de l'amplification femtoseconde est qu'elle met en jeu de très grosses puissances crêtes, et les impulsions amplifiées peuvent alors abîmer le matériau utilisé pour l'amplification ; la CPA permet de palier à ce problème. L'amplification a alors lieu en trois temps :

1. L'impulsion est étirée, et l'énergie est ainsi répartie sur un temps plus long ; c'est cette opération qui réduit considérablement les risques de détérioration. L'étirement peut aller jusqu'à 5 ordres de grandeur.
2. L'amplification en elle-même a ensuite lieu.
3. Au final l'impulsion est recomprimée afin de retrouver son profil temporel de départ.

7.1.2.1 Amplification paramétrique

Nous avons déjà un amplificateur dans notre dispositif : l'OPA ; nous pourrions donc utiliser le même principe d'amplification paramétrique. À première vue l'idée peut paraître farfelue : si nous voulons plus de puissance c'est justement pour améliorer les performances de l'OPA, alors comment pourrions-nous en utiliser un dans ce but ? Ce qu'il faut bien voir c'est que le fonctionnement serait très différent. Pour l'OPA dont nous disposons déjà, nous voulons que les deux modes en sortie, signal et complémentaire, aient la même longueur d'onde que le faisceau de référence. Tous les faisceaux (référence, pompe, sonde) doivent donc provenir du même laser, or c'est justement le fait que la puissance soit limitée qui nous restreint. À contrario, l'amplification femtoseconde n'impose aucune condition sur le complémentaire, qui n'est pas utilisé ; il est donc possible d'utiliser un second laser bien plus puissant pour effectuer le pompage.

Un tel dispositif est appelé OPCPA (*Optical Parametric Chirped Pulse Amplification*) [205, 206, 207] ; il présente l'avantage de ne pas générer d'émission spontanée amplifiée (qui viendrait parasiter le faisceau en ajoutant de la lumière incohérente) et possède de très forts gains, de 4 à 11 ordres de grandeurs [208]. Les cristaux non-linéaires ne stockent pas d'énergie, ils se contentent de transférer une partie de celle de la pompe au signal. Ces deux faisceaux doivent donc être synchronisés temporellement. Les pompes utilisées sont généralement des lasers nanosecondes très énergétiques ; dans notre cas nous pourrions utiliser des lasers à 532 nm (ce qui donnerait un complémentaire à 1450 nm).

Toutefois nous retrouvons le problème du choix entre puissance et taux de répétition : les lasers de pompe suffisamment puissants pour obtenir de bons gains ont des cadences au mieux

de 1 kHz, et les gains évoqués plus haut sont réalisés à des cadences de 1–10 Hz. Ces valeurs sont bien trop faibles pour nous.

7.1.2.2 Amplification régénérative à base de saphir dopé au titane

Si nous continuons à regarder ce que nous avons sur notre table optique nous pouvons remarquer que l'OPA n'est pas le seul élément amplificateur présent. En effet nous en avons aussi un dans notre laser : le milieu amplificateur constitué d'un cristal de saphir dopé au titane. Les techniques d'amplification utilisant l'émission stimulée dans un milieu ayant subi une inversion de population sont actuellement les plus communes. On parle d'amplificateurs solides par opposition aux amplificateurs à colorants plus anciens.

Dans de tels systèmes deux effets s'opposent généralement : l'obtention d'un fort gain et celle d'un bon rendement [209], c'est-à-dire une bonne extraction par l'impulsion de l'énergie stockée dans le cristal. Ainsi si l'on veut un bon gain, alors l'impulsion laisse beaucoup d'énergie dans le cristal après son passage. Or ces gains sont déjà bien plus faibles que ceux présentés précédemment. En fait dans ce cas là le faible rendement va nous aider : puisqu'on arrive à avoir un bon gain en laissant une quantité importante d'énergie dans le cristal alors autant en profiter en faisant de multiples passages dans celui-ci.

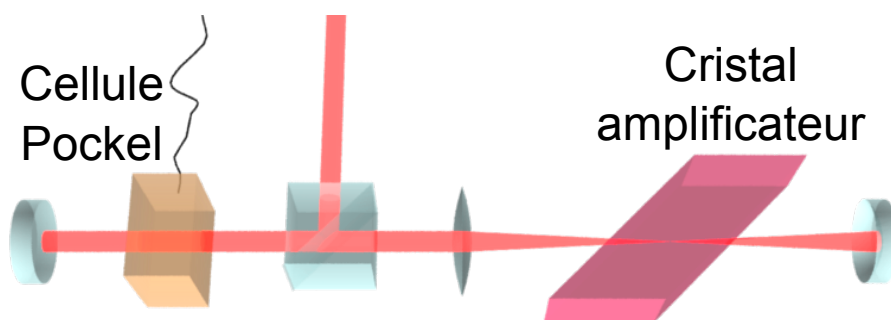


FIGURE 7.1: Amplificateur régénératif

La technique la plus employée pour effectuer une amplification multi-passages s'appelle l'*amplification régénérative* et consiste à piéger l'impulsion dans une cavité contenant le milieu amplificateur puis à l'en sortir une fois le nombre de passages suffisant [210, 211, 212, 213]. La méthode la plus commune pour effectuer ce piégeage consiste à utiliser un polariseur et une cellule de Pockels. L'impulsion reste dans la cavité jusqu'à ce que l'énergie extraite soit inférieure aux pertes engendrées par la traversée de la cavité, ceci permet typiquement des gains allant de 5 à 7 ordres de grandeur. La plupart des amplificateurs commerciaux sont basés sur ce principe, toutefois ce système souffre du même problème que le précédent : la cadence bien que meilleure reste encore trop faible pour nos applications (250 kHz au maximum) [214, 215, 216, 217].

7.1.2.3 Amplification passive à base de saphir dopé au titane

L'alternative à l'amplification régénérative est l'*amplification passive*, qui consiste à utiliser des chemins différents pour chacun des passages ; on parle alors de *multi-passage géométrique*, évitant ainsi l'utilisation d'un système actif pour faire sortir l'impulsion. De part la contrainte géométrique le nombre de passage sera néanmoins plus faible.

Voyons maintenant comment s'effectue le pompage dans le cristal ; il existe deux méthodes :

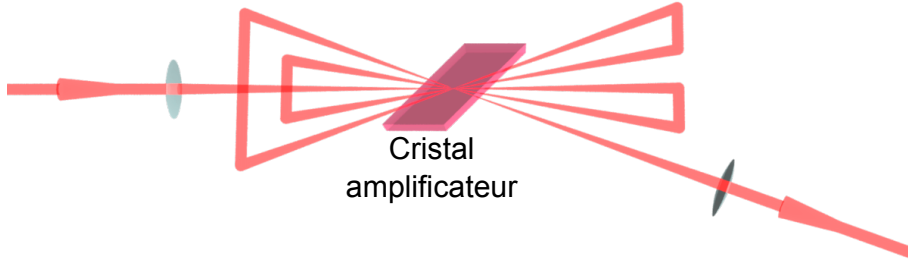


FIGURE 7.2: Amplificateur à multi-passage géométrique

- La première méthode consiste à pomper dans la direction orthogonale à celle de la propagation du laser à amplifier ; elle offre le plus grand volume de zone « active », ce qui implique un alignement moins critique. Le pompage peut aussi être soit impulsionnel (on a alors un bon rendement mais la cadence reste encore une fois trop limitée, de l'ordre du kilohertz), soit continu, mais dans ce cas le rendement est très faible et il faut de nombreux passages pour avoir une amplification efficace, ce qui n'est guère réalisable en multipassage géométrique. De plus la zone de gain est généralement inhomogène : il se forme notamment une lentille thermique conduisant à des aberrations déformant le mode spatial.
- L'autre méthode consiste à pomper dans la même direction que le faisceau à amplifier et à focaliser les différents faisceaux en un point précis. On obtient alors un meilleur gain tout en ayant besoin de moins de passages ; le revers de la médaille est que l'alignement est bien plus critique : lors de chaque passage l'impulsion doit être focalisée exactement au même endroit que le laser de pompe.

Cette dernière solution a déjà été utilisée avec succès à de fortes cadences de répétition [218, 219]. Au final c'est la seule qui peut nous permettre d'obtenir un bon gain sans diminuer le taux de répétition ni trop dégrader le mode spatial. Il n'existe pas dans le commerce d'amplificateur satisfaisant à ces critères, et nous avons donc dû le fabriquer nous même.

7.1.3 Application au Système expérimental

Maintenant que nous savons quel processus d'amplification adopter, voyons comment le mettre en œuvre selon nos besoins. L'amplification dépend des densités d'énergies (ou fluence) J qui s'expriment en J/cm^2 . Le gain est donné par l'équation de Frantz-Nodvick [220]

$$G = \frac{J_{\text{sat}}}{J_{\text{in}}} \ln \left(1 + G_0 \left(e^{\frac{J_{\text{in}}}{J_{\text{sat}}}} - 1 \right) \right) \quad (7.1)$$

avec $G_0 = e^{J_{\text{sto}}/J_{\text{sat}}}$, J_{in} étant la fluence d'entrée du faisceau à amplifier, J_{sat} celle de saturation et J_{sto} celle stockée dans le cristal.

La fluence de saturation vaut $J_{\text{sat}} = \frac{\hbar\omega}{\sigma}$ où σ est la section efficace d'émission stimulée, elle dépend donc du milieu amplificateur et de la longueur d'onde considérée ; pour le titane-saphir à 850 nm elle vaut $1,2 \text{ J}/\text{cm}^2$. Dans un système à quatre niveaux comme le titane-saphir, le temps de relaxation du niveau inférieur de la transition laser est suffisamment court pour qu'habituellement il ne soit quasiment pas peuplé. Seulement dans notre cas ce temps de relaxation (de l'ordre de quelques picosecondes) n'est plus négligeable devant la durée des impulsions¹ : le niveau inférieur n'est plus évacué suffisamment rapidement et le système se comporte comme un système à trois niveaux. Or ces systèmes présentent des saturations deux

1. Étant donné que nous ne pouvons pas nous attendre à des gains énormes, nous n'avons pas besoin d'un

fois plus rapides [221], et par conséquent la fluence de saturation vaut plutôt $0,6 \text{ J/cm}^2$ dans notre cas.

Nous voulons nous placer dans un régime de fort gain et donc de faible rendement (compensé par plusieurs passages) ; celui-ci correspond à $J_{\text{in}} \ll J_{\text{sat}}$. Ce critère est aisément respecté pour des focalisations de l'ordre de la dizaine de micromètre (nous avons même $G_0 \frac{J_{\text{in}}}{J_{\text{sat}}} \ll 1$). L'équation 7.1 peut alors se simplifier en

$$G \approx \frac{J_{\text{sat}}}{J_{\text{in}}} \ln \left(1 + G_0 \frac{J_{\text{in}}}{J_{\text{sat}}} \right) \approx G_0 \quad (7.2)$$

L'énergie stockée dans le cristal vaut $E_{\text{sto}} = E_{\text{abs}} \lambda_p / \lambda$ où E_{abs} est l'énergie absorbée. Mais comment calculer celle-ci à partir de la puissance de pompe P_{pompe} ? L'absorption de l'énergie est liée à l'évolution de la densité de population des états excités dans le cristal, cette évolution suit l'équation

$$\frac{dN(t)}{dt} = \frac{N_0 - N(t)}{\tau} \quad (7.3)$$

où τ est le temps de fluorescence de la transition laser et N_0 la densité de population de l'état excité en régime stationnaire. On en déduit donc que la réponse consiste à faire intervenir ce temps de fluorescence qui, pour le titane-saphir, vaut $\tau = 3,2 \mu\text{s}$. En prenant une absorption de 80%, l'énergie stockée vaut donc

$$E_{\text{sto}} = 0,8 \frac{\lambda_p}{\lambda} \tau P_{\text{pompe}} \quad (7.4)$$

L'énergie n'étant stocké que dans les atomes du cristal éclairés par le laser, l'énergie stockée peut être reliée à la fluence stockée par la surface éclairée : $E_{\text{sto}} = \pi w^2 J_{\text{sto}}$, w étant le rayon du faisceau pompe au niveau du cristal. Ceci nous donne, pour une puissance de pompe de 15 W et une focalisation sur $20 - 25 \mu\text{m}$, un gain allant de 5 à 12 en simple passage. Dans la réalité le gain sera cependant moindre à cause des différentes imperfections, dont la plus importante est le *mode-matching* entre la pompe et le faisceau infrarouge. Nous pourrions alors modéliser le gain par

$$G = e^{c_\lambda P_{\text{pompe}} m(P_{\text{pompe}})} \quad (7.5)$$

où $c_\lambda P_{\text{pompe}}$ correspond à un facteur de gain idéal et $m(P_{\text{pompe}})$ est le *mode-matching*. Les valeurs précédentes nous donnent un paramètre c_λ de l'ordre de $0,1 - 0,2 \text{ W}^{-1}$. Nous reviendrons sur l'influence du mode-matching et la valeur du paramètre m .

Pour finir, nous ne nous sommes pas encore posé la question de la cadence de répétition ; l'expression de G_0 pré-suppose en effet que l'énergie stockée est revenue à son maximum entre chaque impulsion. Or, étant donnée notre cadence élevée, cette hypothèse n'est certainement pas exacte, mais à quel point s'en éloigne-t-on ? Les calculs mènent alors à effectuer le remplacement (cf. D)

$$G_0 \mapsto G_0 e^{\kappa(G-1) \frac{J_{\text{in}}}{J_{\text{sat}}}} \quad (7.6)$$

avec $\kappa = \frac{1}{1 - e^{t_{\text{rep}}/\tau}}$, t_{rep} étant le temps entre deux impulsions. Ceci permet de définir une perte de gain η_G en simple passage, telle que $G \approx \eta_G G_0$, et qui vaut

$$\eta_G = \frac{J_{\text{sat}} - \kappa J_{\text{in}}}{J_{\text{sat}} - \kappa J_{\text{in}} e^{J_{\text{sto}}/J_{\text{sat}}}} \quad (7.7)$$

Cette formule n'est théoriquement valable que pour un simple passage mais, comme nous le verrons par la suite, elle aussi modélise bien une amplification en double passage.

grand étirement des impulsions ; au contraire nous devons même limiter cet étirement afin d'avoir le moins de pertes possibles.

7.2 Dispositif

La conception d'un amplificateur adapté à nos besoins a été effectuée par Julien Laurat et sa mise en œuvre (en dehors du dispositif expérimental) par Aurélien Dantan avant ce travail de thèse. Le remontage du dispositif expérimental au début de celui-ci a été l'occasion d'intégrer cet amplificateur à l'expérience.

7.2.1 Vue d'ensemble

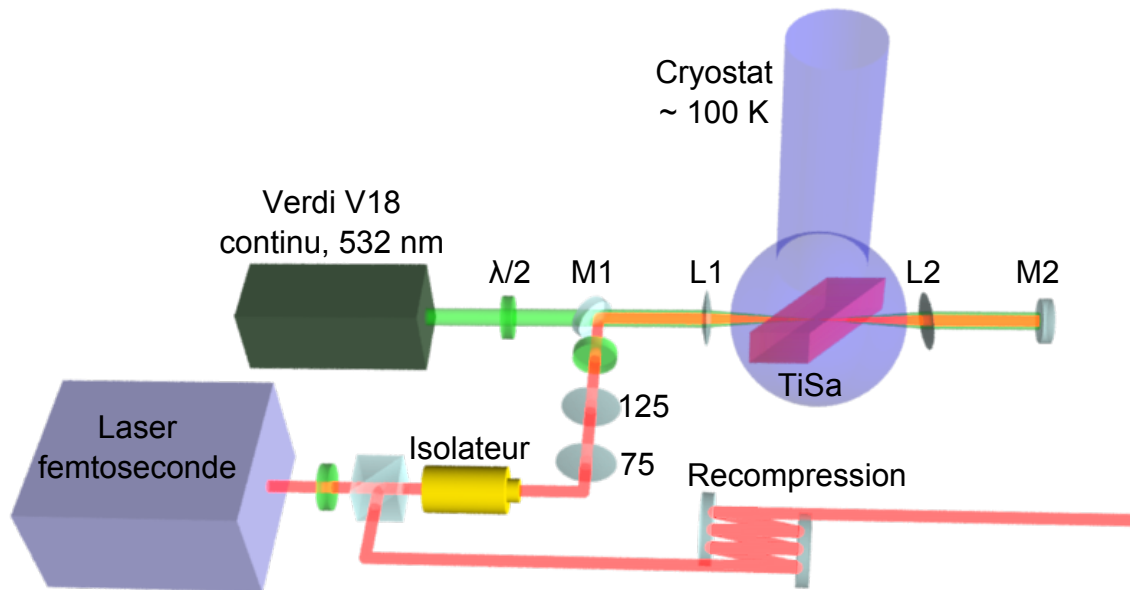


FIGURE 7.3: Montage de l'amplificateur femtoseconde utilisé

Le montage est présenté sur la figure 7.3. Le cristal de titane-saphir est pompé par un faisceau continu à 532 nm généré par un *Verdi V18* de *Coherent*. Il est placé dans un cryostat afin d'être refroidi à la température désirée. Le faisceau infrarouge effectue un aller-retour dans l'amplificateur et est ensuite séparé du trajet d'entrée par un isolateur de Faraday. L'impulsion est uniquement étirée par son passage dans celui-ci avant l'amplification, elle est ensuite recomprimée par des miroirs à dispersion négative de la société *Layertec* (GVD d'environ $-1\,300\text{ fs}^2$). Le principe de ces miroirs est d'organiser les différentes couches diélectriques de manière à ce que les longueurs d'ondes décalées vers le rouge pénètrent plus profondément dans le miroir que celle qui sont décalées vers le bleu, permettant ainsi de compenser la dispersion [222]. Ce système à l'avantage de prendre beaucoup moins de place que le train de prisme présent dans la cavité laser (cf. section B.2.3 et figure 3.2). La recompression est obtenue en effectuant 6 rebonds sur deux miroirs.

Les faisceaux sont focalisés, dans les deux sens, sur $20\text{ }\mu\text{m}$ au niveau de la face avant du cristal à l'aide de deux lentilles $L1$ et $L2$ de focale 75 mm . La superposition des points de focalisation est réalisée en translatant la deuxième lentille du télescope de mise en forme de l'infrarouge. Leur positionnement sur la face avant du cristal, qui donne le meilleur gain, est effectuée en déplaçant le cryostat. Afin de limiter les pertes le cristal taillé à l'angle de Brewster. Ceci introduit toutefois de l'astigmatisme qui peut être compensé en ajustant l'angle de la lentille $L2$.

Les premiers tests effectués en dehors du dispositif expérimental par Aurélien Dantan étaient

plutôt encourageants [223]. La figure 7.4 montre l'évolution du gain en fonction de la température et du taux de répétition. Nous pouvons remarquer que pour chaque puissance de pompe il existe une température optimale. Celle-ci est due à la compétition entre deux effets :

- La détérioration du *mode-matching* $m(P_{\text{pompe}})$ par les effets thermiques [224, 225, 226]. Le principal est la lentille thermique dont la focale vaut

$$f = \frac{\kappa(T)\pi\omega_p}{\rho P_{\text{abs}} dn(T)} \quad (7.8)$$

où $\kappa(T)$ est la conductivité thermique du titane-saphir et $dn(T)$ la variation de son indice, $P_{\text{abs}} \approx 0.88P_{\text{pompe}}$ est la puissance de pompe absorbée et $\rho = 1 - \lambda_p/\lambda$ la fraction de cette dernière qui est convertie en chaleur. Le *mode-matching* pourra alors être modélisé par la formule empirique

$$m(P_{\text{pompe}}) \approx \frac{1}{1 + (f_0/f)^2} \approx \frac{1}{1 + aP_{\text{pompe}}^2 e^{bT - cT^2}} \quad (7.9)$$

- Le spectre d'émission du titane-saphir, qui présente un maximum vers 750 nm à température ambiante et se décale vers le bleu lorsque la température descend [227]. Le facteur de gain idéal $c_\lambda P_{\text{pompe}}$ défini plus haut peut-être modélisé par une lorentzienne

$$c_\lambda P_{\text{pompe}} \approx \frac{dP_{\text{pompe}}}{1 + \left(\frac{T - T_\lambda}{\delta T}\right)^2} \quad (7.10)$$

L'extrapolation des données à faible gain mesurées par Delaigue *et al.*[228] donne $T_\lambda = 210$ K et $\delta T = 200$ K.

Les valeurs expérimentales du gain peuvent être reproduites à l'aide de ce modèle et en utilisant l'équation 7.5 pour $a = 4,7 \cdot 10^{-5} \text{ W}^{-1}$, $b = 0,036 \text{ K}^{-1}$, $c = 7,66 \cdot 10^{-5} \text{ W}^{-2}$ et $d = 0,256 \text{ W}^{-1}$. Enfin nous pouvons noter que $c_\lambda = 0,22 \text{ W}^{-1}$; ce qui correspond bien à l'ordre de grandeur donné plus haut.

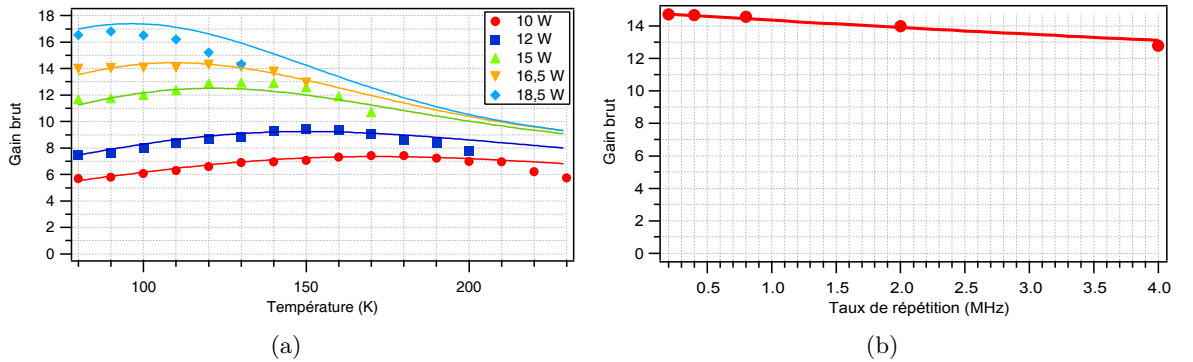


FIGURE 7.4: Résultat des premiers tests (a) En fonction de la température et pour différentes puissances de pompe, les courbes correspondent au modèle (b) En fonction de la cadence de répétition (pour $T = 110$ K et $P_{\text{pompe}} = 15$ W), la courbe correspond au modèle (équation 7.7), seul le décalage vertical est obtenu à partir des données

7.2.2 Cristaux utilisés

Les cristaux utilisés sont fournis par la société *Crystal Systems* ; ils sont fortement dopées ($\alpha_{532} = 7\text{ cm}^{-1}$) et absorbent 88% du faisceau vert. Leur dimension est de 2,5 mm sur 4 mm pour une profondeur de 3 mm. Le contact thermique entre le cristal et le support est effectué par une feuille d'indium enroulée autour du cristal ; un adhésif spécial, bon conducteur thermique et résistant au vide et aux températures cryogéniques, vient en plus renforcer le contact thermique entre les deux parties du support.

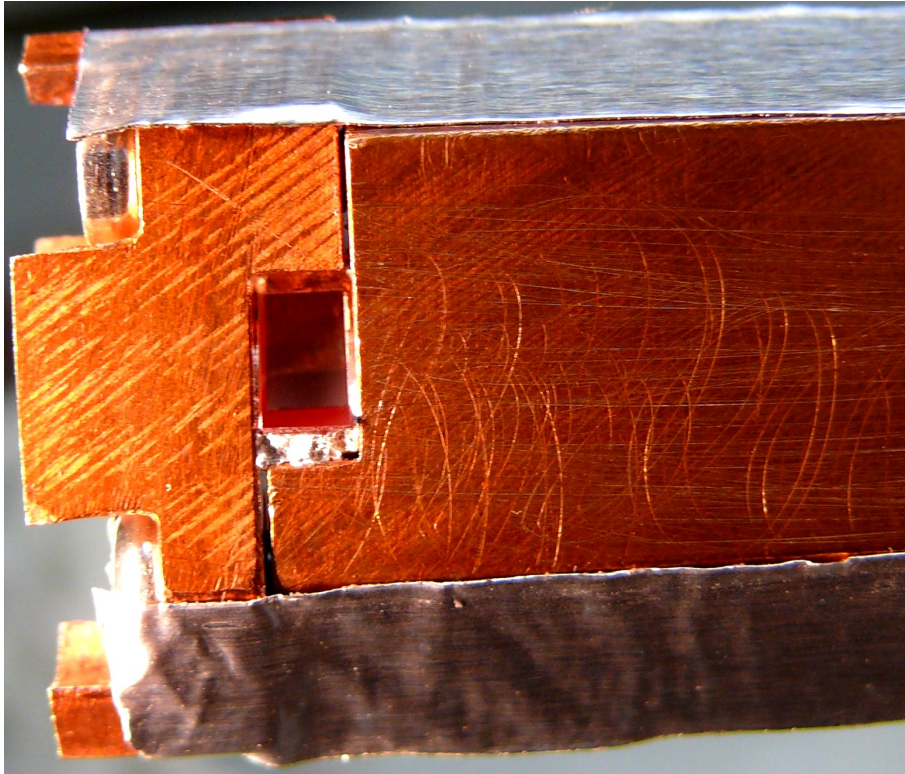


FIGURE 7.5: *Cristal de titane-saphir dans sa monture*

L'étude sur le long terme menée lors de ces travaux de thèse a montré que les cristaux et les joints d'indium finissent inmanquablement par se détériorer. Il apparaît des rayures et des ébréchures sur la face avant du cristal ; cette usure est parfaitement normale et habituelle pour ce type d'amplificateur. Les rayures peuvent notamment être dues au traitement de surface qui finit par ne plus supporter la puissance. En fin de vie certains cristaux finissent par avoir des problèmes de transmission : chute brutale ayant lieu de plus en plus rapidement après le début du pompage ; les causes précises de ce phénomène n'ont toutefois pas pu être déterminées.

Quant à l'indium, il finit par se déposer sur la surface du cristal. Ce dépôt a deux origines. Premièrement le vieillissement qui entraîne un effritement de la feuille d'indium et contre lequel on ne peut rien faire. Deuxièmement un mauvais centrage du faisceau pompe sur le cristal peut entraîner une fonte de l'indium si le point de focalisation est trop proche des bords. La pompe peut aussi « brûler » l'indium et ainsi noircir l'arrête du cristal située à proximité. Ceci peut être évité en centrant le cristal sur le faisceau avec une faible puissance de pompe, puis en ne le déplaçant que légèrement lorsque l'on a toute la puissance de pompe.

La qualité optique des cristaux peut aussi varier d'un échantillon à l'autre entraînant ainsi

une dégradation plus ou moins importante du mode spatial. Enfin, nous avons également testé un cristal avec un dopage moins important ($\alpha_{532} = 4,48 \text{ cm}^{-1}$) dans le but de réduire les effets parasites ; celui-ci ne nous a malheureusement pas permis d'avoir un gain satisfaisant (gain net de 2).

7.2.3 Refroidissement cryogénique

L'utilisation d'un système cryogénique est bien plus complexe qu'une cascade d'éléments Peltier [229], mais elle permet d'atteindre les températures bien plus basses correspondant au maximum de gain.

Le refroidissement s'effectue par un flux continu d'azote liquide à l'intérieur d'un cryostat CFV d'*Oxford Instruments* dans lequel règne un vide de 10^{-4} Pa. Un réglage grossier de la température est d'abord effectué en réglant le débit de sortie de l'azote gazeux, un réglage plus fin est ensuite réalisé à l'aide d'une résistance chauffante régulée par un contrôleur de température *Oxford Instruments ITC*. Si les premiers tests étaient effectués soit en versant directement l'azote dans le cryostat et en laissant remonter la température, soit en envoyant l'azote par un court tube non isolé, l'intégration au dispositif a nécessité l'utilisation d'une canne de transfert sous vide. La conjonction de celle-ci et du contrôleur de température permet de maintenir la température avec une précision de 0,1 K. Il arrive néanmoins que cette canne se bouche. Un dispositif à été conçu avec l'aide d'André Guilbaud afin de pouvoir souffler efficacement de l'air comprimé dans la canne pour éjecter le bouchon.

Un autre phénomène gênant lié aux températures cryogéniques est l'apparition de buée qui se forme à l'extérieur des fenêtres du cryostat, même après avoir mis la canne isolée sous vide. Le refroidissement a d'abord lieu sur la partie haute du cryostat ; les fenêtres refroidissent ensuite à cause de la conduction thermique le long de la paroi extérieure du cryostat. Le refroidissement au niveau des fenêtres, sans le faisceau vert, est généralement d'environ 1°C par heure, mais il est peut monter jusqu'à $2,3^\circ \text{C}$ par heure ; il est aussi contré par le faisceau pompe qui réchauffe les fenêtres. Une meilleure déshumidification a permis de réduire ce phénomène. De plus un système permettant de réchauffer l'extérieur du cryostat a été étudié en cas de persistance.

7.3 Intégration au dispositif expérimental

7.3.1 Faisceau amplifié

La première chose à laquelle il faut faire attention est de bien différencier le gain brut, correspondant au rapport entre les puissances du faisceau en sortie de l'amplificateur avec et sans pompage et le gain net, correspondant au rapport entre les puissances avec et sans l'amplificateur. Le deuxième prend en compte les pertes induites par toutes les optiques de l'amplificateur ; c'est lui qui mesure l'amélioration apportée par l'amplificateur et non le gain brut étudié jusqu'à présent. À cause des nombreux matériaux traversés, les pertes optiques avoisinent les 50% ce qui nous donne donc $G_{\text{net}} = G_{\text{brut}}/2$. Ces pertes viennent essentiellement de l'isolateur optique et du cristal lui-même : ces deux éléments, avec les optiques qui sont autour, ont une transmission d'environ 85 % chacun². Malgré cela le gain net mesuré après intégration au dispositif expérimental reste intéressant comme l'illustre la figure 7.7. Pour des raisons techniques liées à l'usure du laser de pompe il n'a pas été possible d'effectuer des tests poussés à des puissances de pompe plus importantes que 15 W.

2. N'oublions pas qu'ils sont traversés deux fois par le faisceau.



FIGURE 7.6: Cryostat de l'amplificateur. La zone colorée en bleu-cyan correspond à la partie se refroidissant en premier.

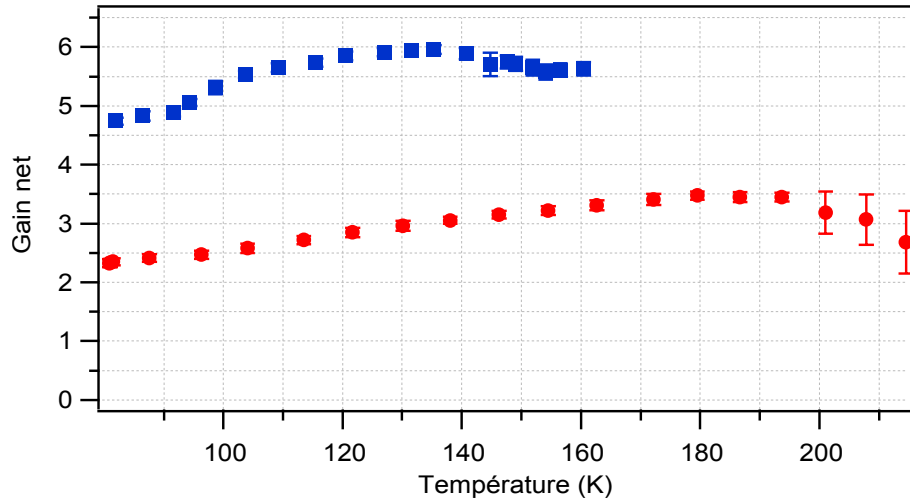


FIGURE 7.7: Gain net de l'amplificateur après intégration au dispositif expérimental pour des puissances de pompe de 10 W (disques) et 15 W (carrés).

Lors des tests l'optimisation du *mode-matching* et de la position du cristal est refaite pour chaque point. Nous noterons cependant que la reproductibilité d'un test à l'autre n'est pas parfaite, signe qu'il reste encore certains facteurs entrant en jeu qu'il reste à identifier. De plus le gain varie en fonction du cristal y compris avec un même dopage ; cette variation semble dépendre de la qualité optique des cristaux.

L'utilisation des outils de diagnostic du laser a permis de démontrer que l'amplificateur n'altérerait en rien les caractéristiques spatio-temporelles des impulsions : le spectre reste identique à celui sans amplification et le signal d'auto-corrélation donne une largeur de l'impulsion sensiblement identique à celle sans amplification (188 fs contre 187 fs lors des tests).

Le profil spatial par contre pose plus de problèmes. L'astigmatisme, qui est correctement compensé par la lentille tiltée, n'est en effet pas la seule dégradation spatiale provoquée par l'amplificateur. Nous avons en effet remarqué une déformation du profil des impulsions, le phénomène le plus important étant l'apparition d'une « seconde bosse » en partie confondue avec la première sur le profil vertical. Cet effet est directement lié aux caractéristiques du cristal puisque le cristal moins fortement dopé entraîne une déformation moins importante (figure 7.8).

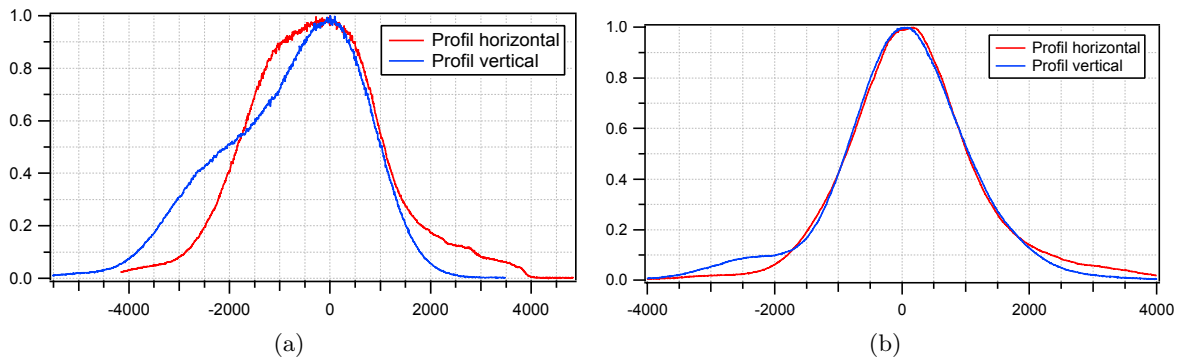


FIGURE 7.8: Profils spatiaux (a) Cristal fortement dopé (qualité optique médiocre) (b) Cristal faiblement dopé

Enfin, le dernier point à étudier concerne le bruit et la stabilité de la puissance. Le bruit relatif sur la puissance moyenne en sortie de l'amplificateur est d'environ 2-3 %, il peut venir d'une part du laser lui-même, celui-ci n'étant pas réglé de manière optimale lors des tests, d'autre part des fluctuations du pointé des lasers, la forte focalisation rendant le système très sensible à celles-ci. La stabilité à long terme a été assez peu étudiée, cependant nous avons pu établir deux causes à une perte de celle-ci :

- Les dérives du dispositif liées aux effets cités précédemment pour le bruit.
- Un découplage des vibrations et mouvements entre la table optique et le cryostat. La première est placée sur des pieds flottants afin d'amortir les vibrations ; le cryostat quant à lui, bien que fixé sur la table optique, est solidaire de la canne de transfert elle-même reliée au sol par l'intermédiaire du réservoir d'azote liquide³. Ceci peut entraîner un désalignement de la position du cristal.

7.3.2 Doublage de fréquence

Après l'intégration au dispositif et cette étude plus poussée du faisceau amplifié, l'étape suivante a été naturellement de doubler la fréquence du faisceau amplifié. Nous avons ainsi pu obtenir une puissance de bleu allant jusqu'à 16,6 mW (correspondant à une efficacité de 25%) sans correction du mode spatial (mais avec un cristal de titane-saphir de bonne qualité). Afin de ne pas endommager le cristal, la taille du faisceau infrarouge est modifiée par un télescope de façon à ce que le *waist* au niveau des cristaux soit plus grand, le but étant de garder la même fluence avec amplification que sans amplification.

Le spectre a été observé à l'aide d'un spectromètre USB (celui utilisé pour le laser ne pouvant pas descendre aussi bas en longueur d'onde) ; ce dernier possède toutefois une moins bonne résolution et sa calibration laisse à désirer. La figure 7.9 illustre les spectres avec et sans amplification et montre que le spectre du faisceau doublé n'est pas modifié par l'amplification.

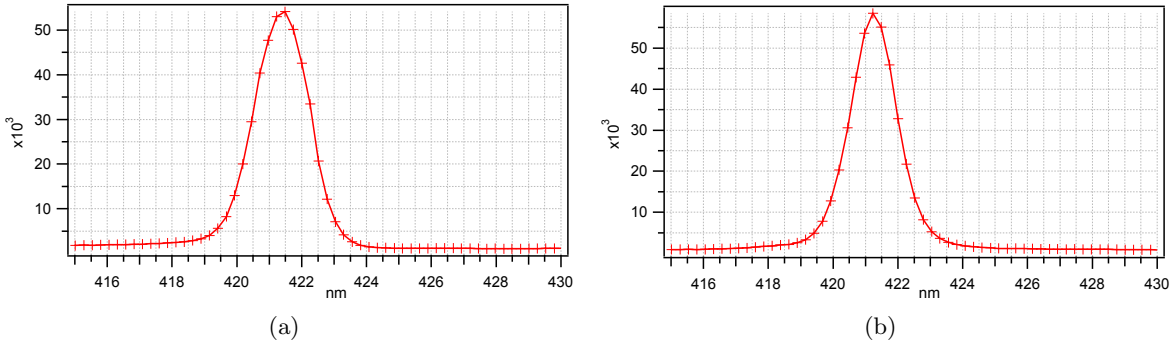


FIGURE 7.9: Spectre du faisceau doublé (résolution : 1 nm) (a) Sans amplification (b) Avec amplification

La durée des impulsions a été obtenue par *cross-corrélation* avec le faisceau infrarouge, et le résultat est montré figure 7.10. Les points représentent les données issues de la *cross-corrélation* et la courbe bleue le fit par une gaussienne d'écart-type σ . La largeur à mi-hauteur vaut alors

$$\Delta\tau_{\text{FWHM}} = 2\sqrt{2\ln(2)}\sigma = 2,35\sigma \quad (7.11)$$

3. La canne de transfert possède bien une partie flexible, mais cette flexibilité est très limitée et non suffisante pour désolidariser complètement le cryostat du réservoir.

La durée des impulsions vaut quant à elle $\Delta t = \Delta\tau_{FWHM}/\sqrt{2}$, ce qui nous donne 223 fs pour le faisceau bleu non amplifié et 275 fs pour le bleu amplifié contre 213 fs pour l'infrarouge⁴. Ces mesures permettent de voir l'étirement des impulsions par la traversée du cristal non-linéaire ; il est plus important pour le faisceau amplifié de par sa puissance. L'impact de cet étirement est cependant difficile à évaluer ; il se fera essentiellement sentir sur la fluorescence paramétrique issue de l'OPA et demande donc d'aller bien plus loin dans l'intégration de l'amplificateur au dispositif expérimental.

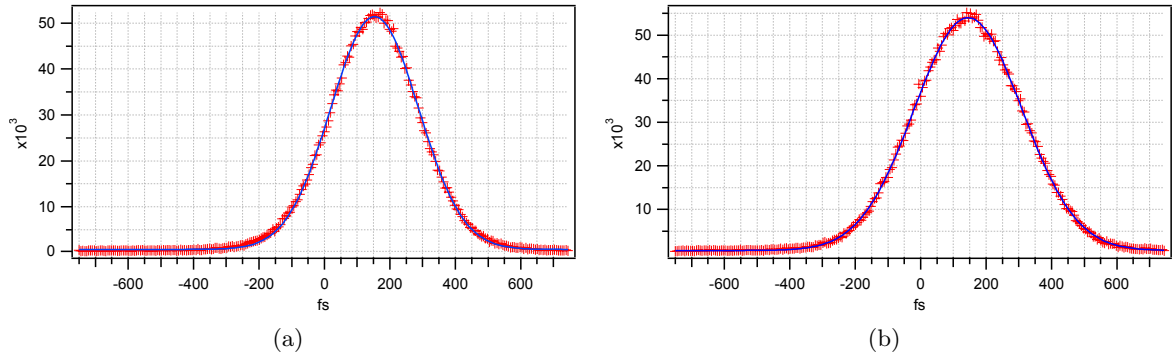


FIGURE 7.10: Signal de cross-correlation du faisceau doublé (a) Sans amplification (b) Avec amplification

Concernant le profil spatial les conclusions sont strictement les mêmes que pour le faisceau amplifié (figure 7.11). Il est à noter que le profil spatial a déjà un impact sur le doublage : ainsi en plaçant un diaphragme avant le GSH nous pouvons observer que, suivant la qualité du faisceau, le maximum de puissance de bleu peut avoir lieu lorsque celui-ci est partiellement fermé (y compris avec un cristal de titane-saphir moins dopé), filtrant ainsi le mode spatial.

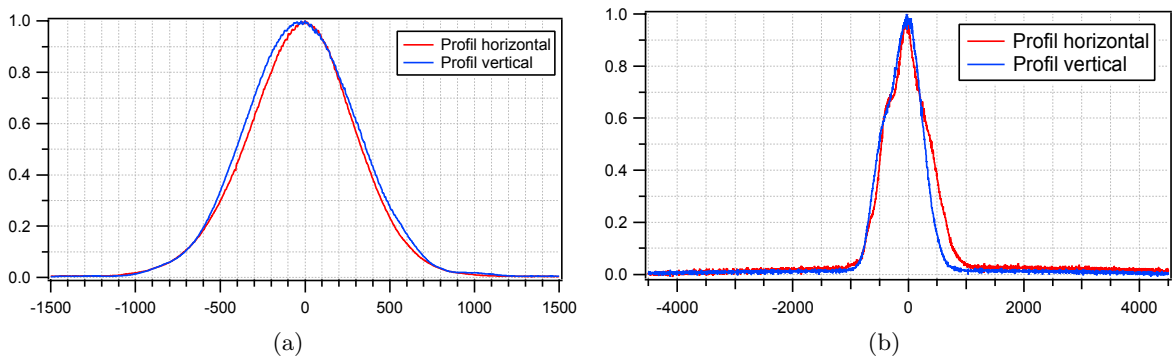


FIGURE 7.11: Profil spatial du faisceau doublé, sans filtrage spatial (a) Sans amplification (b) Avec amplification

4. Rappelons que les caractéristiques des impulsions de notre laser changent d'un jour à l'autre, ce qui explique la différence de durée des impulsions avec les tests évoqués précédemment.

7.3.3 Amplification paramétrique

La troisième étape consiste à optimiser l'OPA avec le faisceau amplifié. À nouveau nous devons faire attention à la focalisation des faisceaux à l'intérieur du cristal. De plus le rapport optimal entre les *waists* des faisceaux pompes et sondes n'est pas forcément le même avec cette puissance que celui trouvé expérimentalement sans l'amplificateur. Il est donc nécessaire d'effectuer une étude approfondie de l'impact de la focalisation des deux faisceaux sur la qualité de l'amplification paramétrique.

Le faisceau sonde utilisé est un faisceau suffisamment intense (environ 100 μW) pour être considéré comme classique : comme pour les réglages sans l'amplificateur femtoseconde, il convient en effet de commencer par améliorer le fonctionnement classique de l'OPA. Cette optimisation n'a malheureusement pas pu être effectuée par manque de temps. Du coup, nous n'avons pas encore pu obtenir une meilleure amplification paramétrique et les résultats peuvent paraître bien faibles : le paramètre de compression r vaut au mieux 0,4 (à comparer aux valeurs des chapitres précédents, comprises entre 0,35 et 0,53) et l'excès de gain est caractérisé par $\gamma = 1$ (contre 0,4 à 0,6). Néanmoins une comparaison de ceux-ci avec des données obtenus sur un OPA non optimisé sans amplification femtoseconde laisse optimiste quant aux résultats que nous pourrions obtenir après optimisation (figure 7.12).

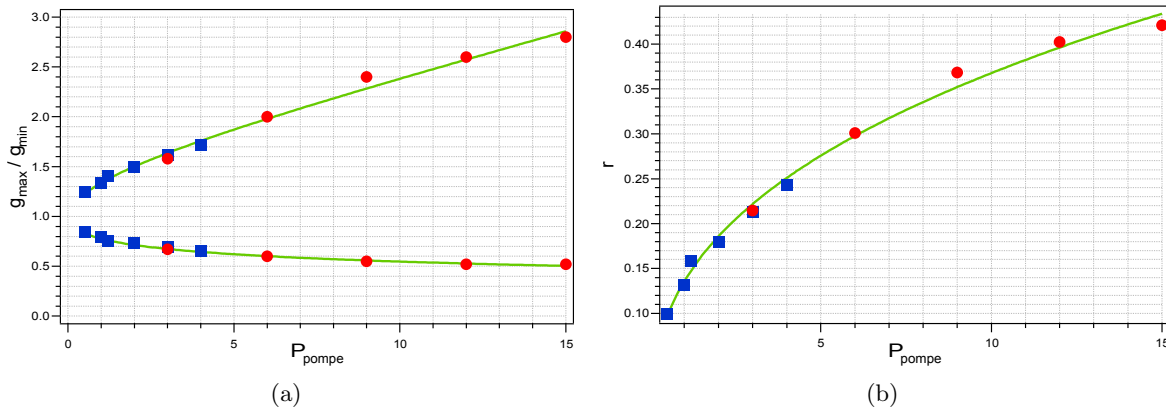


FIGURE 7.12: Paramètre de l'OPA non optimisé en fonction de la puissance de pompe, avec (disques) et sans (carrés) amplification. La courbe verte correspond à $r = 0,17P_{\text{pompe}}^{0,38}$ (obtenu par fit des données) et $\gamma = 1$. (a) gains minimum et maximum (b) paramètre de compression

7.4 Travail restant

Tous ces tests ont montré qu'il y a bien plus de travail à effectuer pour intégrer l'amplificateur au système expérimental qu'on ne le pensait à l'origine. La conséquence en est que cette intégration n'a pu être terminée, les principales étapes restantes sont les suivantes :

1. Terminer l'optimisation de l'OPA avec une amplification paramétrique « classique », ce qui devrait déjà nous donner une première idée de l'importance des améliorations apportées par l'amplificateur. Elle demandera aussi certainement de revenir sur le GSH, notamment pour optimiser le rapport entre puissance de bleu et déformation spatiale.
2. Observer le bruit du vide à la détection homodyne ; celle-ci étant très sensible au bruit, ce test permettra de définir le niveau de bruit acceptable et s'il est nécessaire d'améliorer ce

point ou non.

3. Observer le vide comprimé, ce qui donnera une première vue d'ensemble du fonctionnement du dispositif avec l'amplificateur et permettra d'avoir une idée plus précise de ce qu'il apporte.
4. Enfin tester l'intégralité du dispositif en produisant des photons uniques, ce test permettra de détecter toutes les imperfections introduites par l'amplificateur, permettant ainsi de les corriger.

Il reste aussi plusieurs points noirs à approfondir ; en voici une liste accompagnée de quelques pistes :

- Le profil spatial ; il a jusqu'ici été plutôt mauvais, or il est très important afin de limiter les imperfections de l'OPA et d'obtenir de bon *mode-matching* lors des différentes recombinaisons de faisceaux. Nous avons pu voir que sa déformation dépend beaucoup du cristal mais c'est un paramètre sur lequel nous n'avons qu'un contrôle limité. Deux solutions sont envisageables afin de l'améliorer. La première consiste à filtrer spatialement le faisceau à l'aide d'un sténopé (ou *pinhole*) placé au foyer d'une lentille ; l'inconvénient est que cela entraîne des pertes importantes et ce d'autant plus que l'on voudra une bonne qualité. Nous pouvons cependant espérer compenser ces pertes en améliorant la transmission, plutôt faible, de l'ensemble du dispositif d'amplification. La deuxième solution consiste à corriger le mode spatial à l'aide de techniques d'optique active [230]. Cette méthode est plus compliquée à mettre en œuvre mais les pertes dépendent uniquement du dispositif utilisé ; en fonction de celui-ci nous pouvons donc espérer qu'elles soient plus faibles.
- Nous avons pu observer un plus fort étalement de l'impulsion lors du passage dans le premier cristal non-linéaire, ce qui pourrait entraîner une dégradation du *mode-matching* avec l'oscillateur local. La solution serait alors d'allonger aussi l'impulsion de l'oscillateur local afin qu'il ait la même durée que celles sortant de l'OPA. La méthode la plus simple consiste à introduire un élément dispersif. Une méthode plus poussée et précise serait d'utiliser un dispositif de mise en forme d'impulsions [231] ; là encore cette technique est plus compliquée à mettre en œuvre.
- Le bruit, dont l'effet ne pourra être caractérisé que par la détection homodyne. Nous disposons de photodiodes à quadrants pour mesurer l'importance des fluctuations de pointé, et si elles sont trop importantes la solution consiste à asservir la position du faisceau pompe en utilisant des miroirs commandés par des piézo-électriques. Une autre solution plus générale consiste à asservir l'intensité de l'un des faisceaux, pompe ou infrarouge, à l'aide d'un modulateur acousto-optique.
- Les fluctuations sur le long terme devront être mieux étudiées afin de s'assurer que l'on peut effectuer des expériences sans devoir ré-aligner tout le système trop souvent. Concernant les problèmes de découplage des vibrations et mouvements, ceux-ci peuvent être minimisés en faisant très attention lors des réglages. L'avantage de notre dispositif expérimental est que nous n'avons pas besoin d'être présents lors de la prise de données : une fois tous les réglages effectués il est même possible de lancer l'acquisition et d'effectuer quelques tests de l'extérieur. Ceci réduit les risques de dérèglement du dispositif.

Il reste donc une certaine quantité de travail à effectuer avant d'utiliser l'amplificateur dans des expériences ; ce travail sera néanmoins facilité par les tests déjà réalisés qui ont d'ores et déjà permis d'identifier les différentes sources de problèmes et d'en résoudre une bonne partie.

Chapitre 8

Le VLPC

Sommaire

8.1 Le comptage de photons	165
8.1.1 Introduction	165
8.1.2 Comparaison des différentes méthodes	166
8.1.3 Principe de fonctionnement du VLPC	168
8.2 Montage	170
8.2.1 Vue d'ensemble	170
8.2.2 Montage cryogénique	172
8.2.3 Acheminement du faisceau par fibre optique	173
8.3 Caractérisation	174
8.3.1 Dispositif de test	174
8.3.2 Recherche d'un point de fonctionnement	176
8.3.3 Efficacité en fonction du flux incident	177
8.4 Conclusion	179

8.1 Le comptage de photons

8.1.1 Introduction

Nous avons vu au chapitre précédent que le principal facteur limitant de notre dispositif était la puissance laser. Le second facteur limitant de nos expériences est l'incapacité de nos détecteurs à compter les photons. La quasi-totalité des protocoles utilisant le conditionnement nécessitent que celui-ci s'effectue sur un nombre bien précis de photons. Notre dispositif permet uniquement d'effectuer une approximation de ce conditionnement en s'arrangeant pour que la probabilité que le nombre de photons soit supérieur au nombre désiré soit négligeable. Néanmoins d'une part cela reste une approximation, ce qui dégrade l'état produit, d'autre part il existe des protocoles où cette approximation n'est pas valable. Ces derniers ne peuvent être réalisés avec notre montage.

Comme l'on peut s'en douter les vrais compteurs de photons ne sont pas courants. Il en existe bien quelques uns qui sont commercialisés depuis peu, mais leur performances sont encore loin d'être satisfaisantes [232]. En fait il existe de nombreuses propositions pour réaliser des appareils

capables de discriminer le nombre de photons, mais la plupart ne sont encore que des prototypes. De plus ils font souvent appels à des éléments complexes comme des structures semi-conductrices ou supraconductrices particulières. Il est donc relativement difficile d'en obtenir.

Le choix d'un compteur de photon ne dépend donc pas uniquement de ses caractéristiques et performances, mais aussi de la complexité du dispositif nécessaire à son fonctionnement ainsi que des contacts et collaborations nécessaires pour l'obtention de l'appareil. Dans notre cas nous avons pu récupérer un compteur appelé VLPC (*Visible Light Photon Counter*) ayant déjà été utilisé à Stanford [233, 234]. L'inconvénient de ce détecteur est qu'il est prévu pour fonctionner à des longueurs d'onde plus courtes. Il convient donc de tester ses performances à notre longueur d'onde afin de savoir s'il est intégrable à notre expérience. Ce test est une première, le VLPC n'ayant jusqu'ici été utilisé que dans le visible.

8.1.2 Comparaison des différentes méthodes

Commençons par passer en revue les principales méthodes pour compter les photons. Nous nous restreindrons à celles qui ont fait l'objet de démonstrations expérimentales. Ils en existe cependant d'autres comme celles utilisant des mémoires quantiques qui peuvent théoriquement atteindre des efficacités supérieures à 99 % [235, 236].

La table 8.1, en fin de section, résume les caractéristiques des différents compteurs présentés.

Le multiplexage Il est possible d'utiliser des APD pour approximer un compteur de photons. Le principe consiste à diviser le faisceau en plusieurs modes qui seront lus indépendamment ; le but est que les photons soient séparés afin d'être tous détectés. Dans le cas où nous devons mesurer un nombre n précis de photons nous pouvons nous contenter de n modes si la probabilité d'avoir plus de photons peut être rendue négligeable, mais l'efficacité sera amoindrie car un certain nombre d'événements correspondant au bon nombre de photons ne pourra être détecté (lorsque plusieurs photons se retrouvent dans le même mode). Dans le cas général le nombre de modes utilisés devra résulter d'un compromis : d'un côté il est nécessaire de diviser le faisceau en suffisamment de parties pour que la probabilité que deux photons aillent dans le même mode soit faible, d'un autre côté plus il y a de modes plus les pertes optiques et les coups d'obscurités sont importants. Il existe deux types de multiplexage :

Le multiplexage spatial Il est principalement réalisé à l'aide de lames séparatrices et d'autant de détecteurs qu'il y a de modes [237, 238] ; c'est cette technique qui a été utilisée dans le groupe pour produire des états de Fock à deux photons [91, 92]. Il existe aussi une variante plus compacte consistant à mettre en forme le faisceau pour qu'il soit également réparti sur une matrice de photodiodes à avalanches [239, 240].

Le multiplexage temporel Appelé TMD (pour *time multiplexed detector*), il consiste à décaler temporellement les modes à l'aide de fibres optiques ; il utilise ainsi bien moins de détecteurs. Il peut être réalisé de deux manières différentes. La première est la configuration « balancée » où le faisceau passe par une série d'interféromètres de Mach-Zender dont un des bras est bien plus long et de plus en plus long puis finit sur deux APD [241, 242]. La deuxième est la configuration « bouclée » où il n'y a qu'une séparatrice dont une des sortie renvoie vers l'une des entrées et l'autre sur un seul détecteur [243, 244, 245]. Les comparaisons théoriques montrent que la configuration balancée est la plus efficace [246].

Les détecteurs à base de multiplication La multiplication n'est pas réservée aux détecteurs non discriminants. Il est possible d'obtenir une multiplication qui dépend globalement du nombre de photons incidents.

Le tube photomultiplicateur (PMT) Certainement le plus ancien compteur de photons [247, 248], il fait appel à des technologies rappelant nos vieux écrans cathodiques. La multiplication est effectuée par des dynodes à l'intérieur d'un tube sous vide. Il a un temps de réponse très court, mais souffre d'une très faible efficacité quantique. Il reste cependant utilisé dans bon nombre de dispositifs de spectroscopie et d'appareils de mesure ; c'est d'ailleurs le cas de notre auto-corrélateur.

La photodiode à avalanche Il a été récemment démontré qu'en utilisant une photodiode à avalanche en dessous du régime de saturation à l'aide de techniques de suppression de bruit permettant de détecter de faibles avalanches il était possible de compter les photons [249, 250, 251]. Cette technologie est balbutiante et n'est donc pas encore disponible pour des applications ; elle est néanmoins très prometteuse de par sa simplicité d'utilisation.

Le Visible Light Photon Counter (VLPC) Il fait aussi parti des détecteurs utilisant du multiplexage et un processus d'avalanche, mais le dispositif est encore plus compact puisque le multiplexage a lieu au sein même du détecteur. Il utilise en effet des processus d'avalanche suffisamment bien contrôlés pour que celles-ci restent confinées dans de petites zones du détecteur, permettant ainsi à plusieurs photons de déclencher chacun une avalanche pourvu qu'ils n'arrivent pas dans la même zone active [274, 253]. Il est plus efficace qu'un multiplexage spatial [254], mais fonctionne à température cryogénique (7 K).

Comptage direct des porteurs de charges Lorsque les photons sont absorbés dans un semiconducteur, ils créent des porteurs de charges. Si la multiplication augmente leur nombre afin d'obtenir un signal mesurable avec l'électronique standard, il existe maintenant des techniques pour mesurer directement ces porteurs de charges.

Les détecteurs à base de boîtes quantiques Ils sont constitués d'un transistor à effet de champ contrôlé par une couche de boîtes quantiques. Il a pu être montré que le courant traversant le transistor subit des changements proportionnels aux nombre de porteurs de charge capturés par les boîtes quantiques [255, 256]. Ces détecteurs fonctionnent à température cryogénique pour des efficacités encore faibles, celles-ci étant dues à l'absorption à l'intérieur des couches semiconductrices, nous pouvons cependant espérer de bonnes améliorations dans ce domaine.

Le photodétecteur à intégration de charges Le CIPD (pour *Charge-Integration Photon Detector*) utilise un composant commun : la photodiode PIN. Le but est alors d'utiliser un amplificateur bas bruit suivi d'un circuit intégrateur de charge afin de compter les porteurs de charges. Il fonctionne à température cryogénique (4 K) afin de réduire les coups d'obscurité qui sont pour le coup extrêmement faibles. Les réalisations expérimentales montrent de très bonnes efficacités mais sont centrés sur la longueur d'onde télécom [257, 258] ; l'utilisation de photodiode à base de silicium devrait cependant donner des résultats similaires à nos longueurs d'ondes. Le principal inconvénient est que le taux de répétition maximal qui a pu être démontré est seulement de 40 Hz.

Les détecteurs supraconducteurs Ce sont les détecteurs les plus complexes. Outre l'usage de matériaux supraconducteurs ils requièrent un système cryogénique demandant souvent de refroidir à une température sub-kelvin. Ils ne détectent pas à proprement parler le nombre de photons mais l'énergie absorbée, ce qui est cependant équivalent si on travaille à longueur d'onde fixée. La plupart utilisent le principe du bolomètre et sont capables de détecter de très faibles variations locales de température en se plaçant à proximité de la transition supraconducteur-résistif.

La jonction tunnel supraconductrice (STJ) Ils sont plus proches des détecteurs à avalanches que des bolomètres : l'absorption d'un photon crée une quasi-particule qui va, à son tour, détruire des paires de Cooper, créant ainsi un courant tunnel mesurable et proportionnel à l'énergie absorbée [259, 260]. Leur efficacité est moins bonne que celle des autres détecteurs supraconducteurs, mais il est probable que celle-ci augmente dans les années à venir.

Le détecteur supraconducteur de photons uniques (SSPD) Le SSPD (*Superconducting Single Photon Detector*) utilise un nanofil supraconducteur disposé en méandres sur la zone active et parcouru par un courant proche du seuil de transition supraconducteur-résistif. L'absorption d'un photon crée alors un pont résistif conduisant à un pic de tension mesurable [261, 262]. L'absorption de plusieurs photons à des endroits différents crée plusieurs points chauds, ce qui équivaut à mettre en série les points résistifs [263]. Cette mesure est cependant très peu précise à cause des impédances en jeu. Une autre solution consiste alors à multiplier les fils, soit en gravant par lithographie plusieurs nanofils suivant les même méandres [264], soit en mettant en parallèles plusieurs méandres de nanofils [265, 266].

Le Transition Edge Sensor (TES) Un absorbeur est placé en contact thermique avec un film supraconducteur dont la température est juste en dessous de la température critique et auquel est appliquée une tension constante. Dans cette configuration la résistance, et donc le courant traversant le film, dépend fortement de la température [267]. En utilisant différents absorbeurs il est possible d'utiliser les TES dans toutes les gammes du spectre électromagnétique. Les TES ont d'excellentes performances et ont déjà été utilisés avec succès dans le domaine de l'information quantique [95]. À noter que les photons thermiques peuvent introduire des coups parasites, mais leur taux reste semblable à celui d'une APD. Ils demandent par contre des mesures de courant complexes, généralement à base de SQUID (*Superconducting Quantum Interference Device*).

8.1.3 Principe de fonctionnement du VLPC

Le VLPC est un élément semiconducteur en silicium principalement composé de deux couches [272, 253] : une couche intrinsèque¹ absorbante et une couche de gain faiblement dopée avec de l'arsenic. Le contact électrique s'effectue par un contact transparent du côté de la première couche et une zone fortement dopée du côté de la deuxième ; enfin un traitement anti-reflet est déposé sur le contact transparent (figure 8.1). Les photons incidents peuvent être absorbés soit par la couche absorbante, soit par la couche de gain.

- Lorsqu'un photon est absorbé dans la première couche il crée une paire électron-trou. À cause de la tension appliquée au semiconducteur, l'électron se déplace vers le contact transparent tandis que le trou migre vers la couche de gain. Le faible dopage dans celle-ci induit un niveau d'énergie dû aux impuretés situé à 54 mV en dessous de la bande de conduction. À la température de fonctionnement du VLPC (environ 7 K), les électrons n'ont pas suffisamment d'énergie thermique pour être excités dans la bande de conduction et restent donc dans les impuretés. Le trou va alors apporter l'énergie suffisante pour qu'un électron passe du niveau donneur à la bande de conduction. Cet électron va être accéléré en direction de la face d'entrée et déclenche ainsi le processus d'avalanche sur son sillage.

1. c'est-à-dire qui est semi-conductrice sans dopage.

Comp- teur	Do- maine	Eff.	Obs. (cps/s)	f _{rep}	N _{max}	T
PMT	Visible – UV	7 – 25 %	100	3 GHz	9	Amb.
APD	Visible – Proche IR	10 – 60 %	1 – 250	20 MHz	6	Amb. / 77 K
Multiplexage spatial	Visible – Proche IR	< 50 %	Var.	20 MHz	Var.	Amb.
TMD	Visible – Proche IR	< 50 %	250	1 MHz	8 - 16	Amb.
VLPC	Visible	80 – 90 %	20 000	500 MHz	6	7 K
Boite quan- tique	Visible – Proche IR	10 – 25 %	0,4	> 400 kHz (< 5 MHz ?)	3	4 K
CIPD	900 – 1700 nm	80 %	0,1	40 Hz	20	4 K
STJ	Visible	45 %	~ 0	500 MHz	3	320 mK
SSPD	Visible – Proche IR	20 %	0,01	2 GHz	Var.	4 K
TES	Visible – Proche IR	89 %	~ 0 – 400	> 500 kHz	11	100 – 125 mK

TABLE 8.1: Comparaison des différents compteurs de photons

- Lorsqu'un photon est absorbé par la couche de gain, il va directement fournir l'énergie nécessaire à l'excitation de l'électron déclenchant l'avalanche.

Les électrons issus de l'avalanche, environ 30 000 par photon incident, rejoignent rapidement la surface transparente, générant ainsi un pic de courant d'une durée inférieure à la nanoseconde. À contrario les charges positives se déplacent très lentement dans la zone de gain. Ce déplacement est dû à un recouvrement partiel des états d'énergie des impuretés voisines, réalisé par un choix précis du taux de dopage, induisant un mécanisme de *conduction par sauts* basé sur l'effet tunnel. Cette lenteur fait qu'ils n'acquièrent jamais une énergie cinétique suffisante pour exciter d'autres électrons, ce qui réduit considérablement le bruit de multiplication. De plus ils ralentissent l'avalanche et aident à la contenir dans une zone restreinte (environ 4 μm de diamètre), permettant ainsi la détection de plusieurs photons. Le revers de la médaille est que le temps de relaxation d'une zone, et donc la durée minimale entre deux utilisations de la même zone, est long (environ 3,5 ms) [269].

Une autre caractéristique intéressante du VLPC est son faible bruit de multiplication, qui dépend de la variation du nombre d'avalanches déclenchées par un photon incident. Nous venons déjà de voir deux phénomènes responsables de la bonne qualité du VLPC en termes de bruit de multiplication : si seul les électrons sont responsables des avalanches alors ces variations sont réduites, de même pour le confinement de l'avalanche dans une zone réduite. En outre, il existe une troisième raison : un VLPC ne nécessitant pas une forte tension pour fonctionner, de par le faible écart d'énergie entre la bande donneuse et la bande de conduction, les électrons ont une plus faible énergie cinétique que dans les APD. Il en résulte un délai minimum entre deux déclenchements d'avalanche par le même électron, ce qui réduit le bruit de multiplication [270].

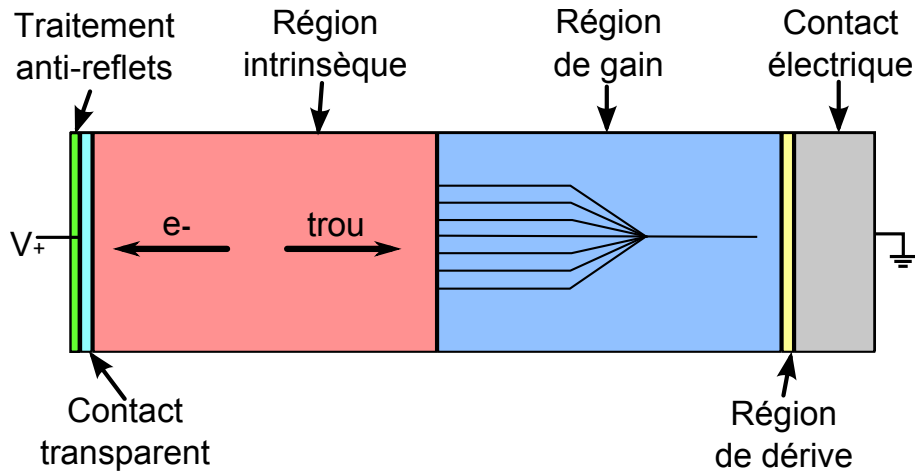


FIGURE 8.1: Structure du VLPC

Un gros inconvénient du VLPC, certainement le plus important, est que les impuretés de la couche de gain sont facilement excitées par les photons thermiques. Ainsi il est sensible aux photons dont la longueur d'onde est comprise entre 1 et 30 μm . L'utilisation des deux couches diminue cette sensibilité, mais elle reste néanmoins suffisante pour induire de forts taux d'obscurité.

Finalement, nous devons aussi noter que le processus de fabrication n'est pas encore parfaitement maîtrisé ; ces détecteurs ne sont en effet pas produits à grande échelle mais fabriqués sur demande par *DRS Technologies, Inc.*, à l'origine pour l'expérience D0 située au *FermiLab*. Par conséquent les performances varient d'un exemplaire à l'autre. Nous devons ajouter qu'une production de VLPC coûte très cher, mais donne de nombreux échantillons ; ainsi la plupart des VLPC utilisés en optique quantique viennent d'exemplaires rejetés pour l'expérience D0 [269].

8.2 Montage

8.2.1 Vue d'ensemble

Le montage, l'optimisation et les tests préliminaires ont été réalisés par Florence Fuchs avec la participation d'André Villing pour le montage électronique.

Le VLPC se présente sous la forme d'une puce en silicium comprenant 8 détecteurs de 1 mm de diamètre. Celui utilisé est monté dans un support lui-même enchâssé dans une monture en cuivre ; ce montage est exactement celui utilisé à Stanford (figure 8.2). Il fonctionne à une température comprise entre 6 et 8 K et doit être polarisé par une tension d'environ 7 V, ces deux paramètres sont réglables afin de trouver le point de fonctionnement optimal.

Étant donné la géométrie du cryostat utilisé, l'acheminement du faisceau lumineux jusqu'au détecteur s'effectue par l'intermédiaire d'une fibre optique. Le signal électrique est quant à lui amplifié une première fois par un pré-amplificateur bas bruit situé à l'intérieur du cryostat (modèle *MGA-81563* d'*Agilent Technologies* offrant un gain de 12,4 dB pour un bruit de 2,8 dB à température ambiante), il est ensuite amplifié une seconde fois par une série de trois amplificateurs en sortie de celui-ci (un modèle *ZFL-1000LN* modifié et deux modèles *ZJL-4HG* de *Mini-Circuits* offrant respectivement des gains de 20 et 17 dB pour des bruits de 2,9 et 3,9 dB).

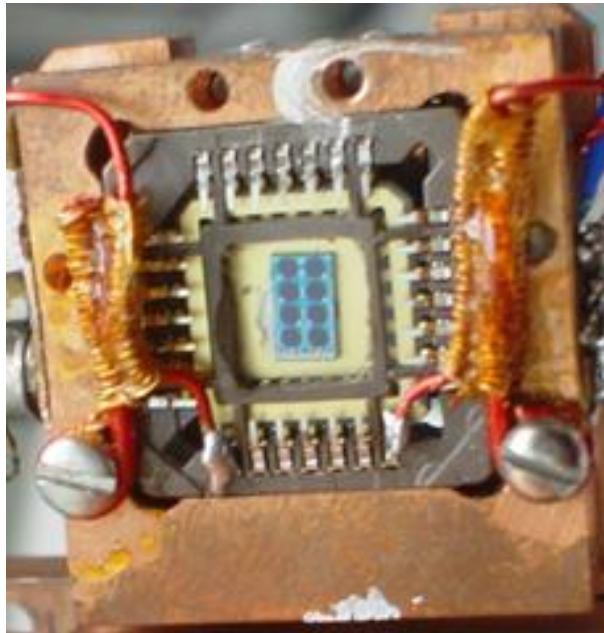


FIGURE 8.2: VLPC dans sa monture

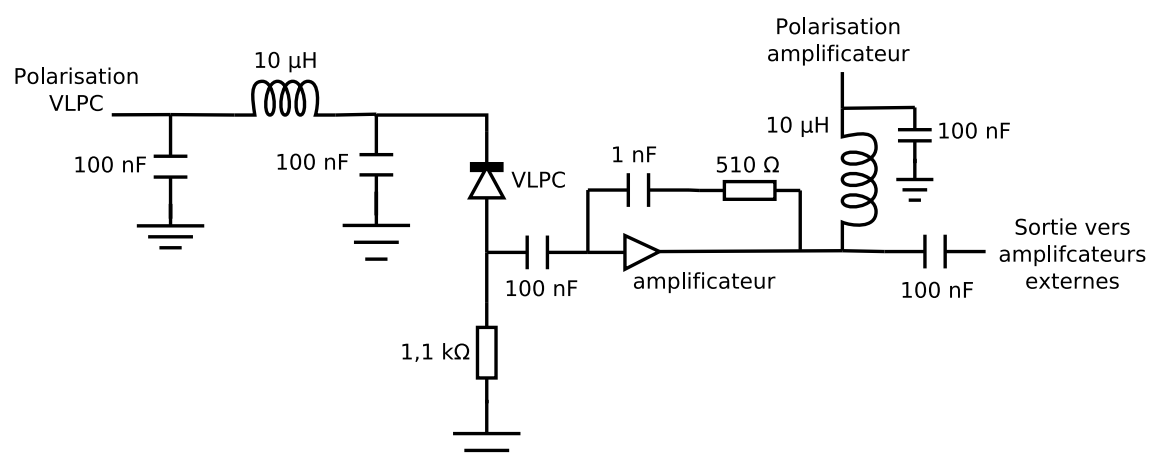


FIGURE 8.3: Circuit électrique de la partie située à l'intérieur du cryostat.

8.2.2 Montage cryogénique

La méthode usuelle pour atteindre la température nécessaire au fonctionnement du VLPC consiste à utiliser un cryostat à bain d'hélium. Cette méthode revient toutefois très cher à cause du prix de l'hélium ; il est alors plus avantageux, notamment pour de petites installations ne bénéficiant pas d'un système de récupération de l'hélium gazeux, d'utiliser un cryostat à cycle fermé. Le cryostat utilisé est un cryostat à tube pulsé *ST405* de *Cryomech*, composé de deux éléments : une tête froide (modèle *PT405*) constituant l'enceinte réfrigérée et un compresseur (modèle *CP950*). Le refroidissement est assuré par un cycle thermodynamique de compression-détente ; la particularité du cryostat à tube pulsé est que ces dernières ne sont pas créées par des pistons mais par des ondes de pression engendrées par le compresseur. L'absence de pièces mobiles évite un certain nombre de difficultés techniques (étanchéité, usure, fiabilité, ...) et réduit les vibrations au niveau de la tête froide.

Le cryostat possède deux étages à 60 et 3,5 K (figure 8.4(a)). Les supports du VLPC et de la fibre optique sont montés sur une équerre positionnée sur le plan froid à 3,5K. Le pré-amplificateur ainsi que le circuit de polarisation sont placés sur l'autre face de ce plan. Le contrôle de la température est effectué par une résistance chauffante et un capteur de température (diode en silicium) placé sur l'équerre.

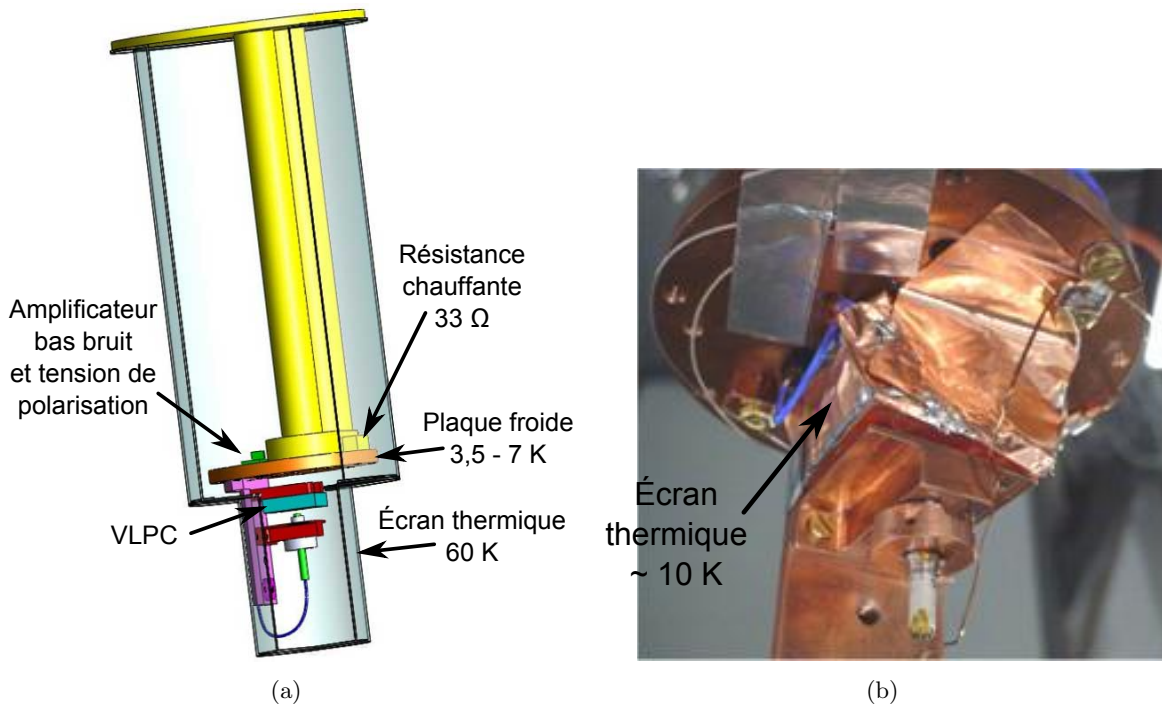


FIGURE 8.4: Refroidissement cryogénique du VLPC (a) Schéma du cryostat (b) Support du VLPC et écran thermique

La gamme de fréquences auxquelles le VLPC est sensible est tellement étendue qu'il détecte les photons thermiques émis par la partie à 60 K. Pour y remédier un écran thermique en cuivre (réalisé à l'aide de feuilles de cuivre adhésives) entoure les supports du VLPC et de la fibre. Cet écran est thermalisé à la température de fonctionnement du VLPC, soit dans les 6–7 K.

8.2.3 Acheminement du faisceau par fibre optique

La fibre menant le faisceau au VLPC à l'intérieur du cryostat doit répondre à plusieurs critères :

- garder ses propriétés aux températures cryogéniques.
- avoir une ouverture numérique et un diamètre de cœur compatibles avec notre dispositif.
- accepter des rayons de courbure courts (< 25 mm) sans que la transmission ne soit trop dégradée.
- transmettre efficacement les photons à 850 nm et très mal les photons thermiques pour $\lambda > 1\mu\text{m}$. Ce dernier point est d'autant plus important que les fibres à température ambiante émettent un rayonnement thermique qui est capté par le VLPC.

Les fibres plastiques standards, en PMMA, absorbent fortement les radiations thermiques, mais leur transmission est bien moins bonne que les fibres en silice. Depuis une vingtaine d'année sont apparues des fibres en polymères perfluorés dont la transmission est bien meilleure. Le choix s'est alors porté vers ces dernières, et plus précisément sur une fibre à gradient d'indice *GigaPOF-60SR* fabriquée par *Chromis Fiber Optics*. L'atténuation de celle-ci est faible à 850 nm ($< 0,06$ dB/m) ; elle n'est par contre pas connue au delà de 1 500 nm, et l'hypothèse effectuée est qu'elle est plus proche d'une fibre en PMMA que d'une fibre en silice. Les tests effectués confirment cette hypothèse et ne montrent pas de baisse significative de la transmission à 850 nm à des températures cryogéniques.

Ces fibres se sont cependant avérées difficiles à polir, induisant ainsi de fortes pertes aux interfaces. Les fibres utilisées, d'une longueur de 1 m, ont alors une transmission comprise entre 20 et 30 % seulement.

Enfin, un dispositif a été conçu afin de centrer le plus précisément possible la fibre optique sur la surface active du VLPC.

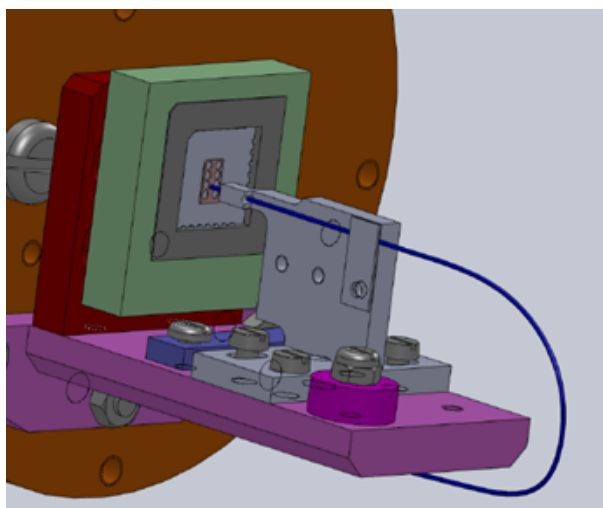


FIGURE 8.5: Monture pour centrer la fibre optique

8.3 Caractérisation

8.3.1 Dispositif de test

La source laser utilisée pour les tests est une diode laser de type VCSEL (*Vertical-Cavity Surface-Emitting Laser*) émettant à 850 nm. Afin de délivrer des impulsions lumineuses elle est contrôlée par un générateur d'impulsions. Celles-ci ont une durée typique de 2,3 ns et sont produites avec un taux de répétition de 10,1 kHz, ce qui donne une puissance moyenne de 14,8 nW. Le laser est ensuite atténué de manière à former de petits états cohérents ayant un nombre moyen de photons de l'ordre de l'unité. Cette atténuation est réalisée par une série de filtres à densité neutre, dont la transmission est connue, qu'il est possible de varier afin de modifier le nombre moyen de photons.

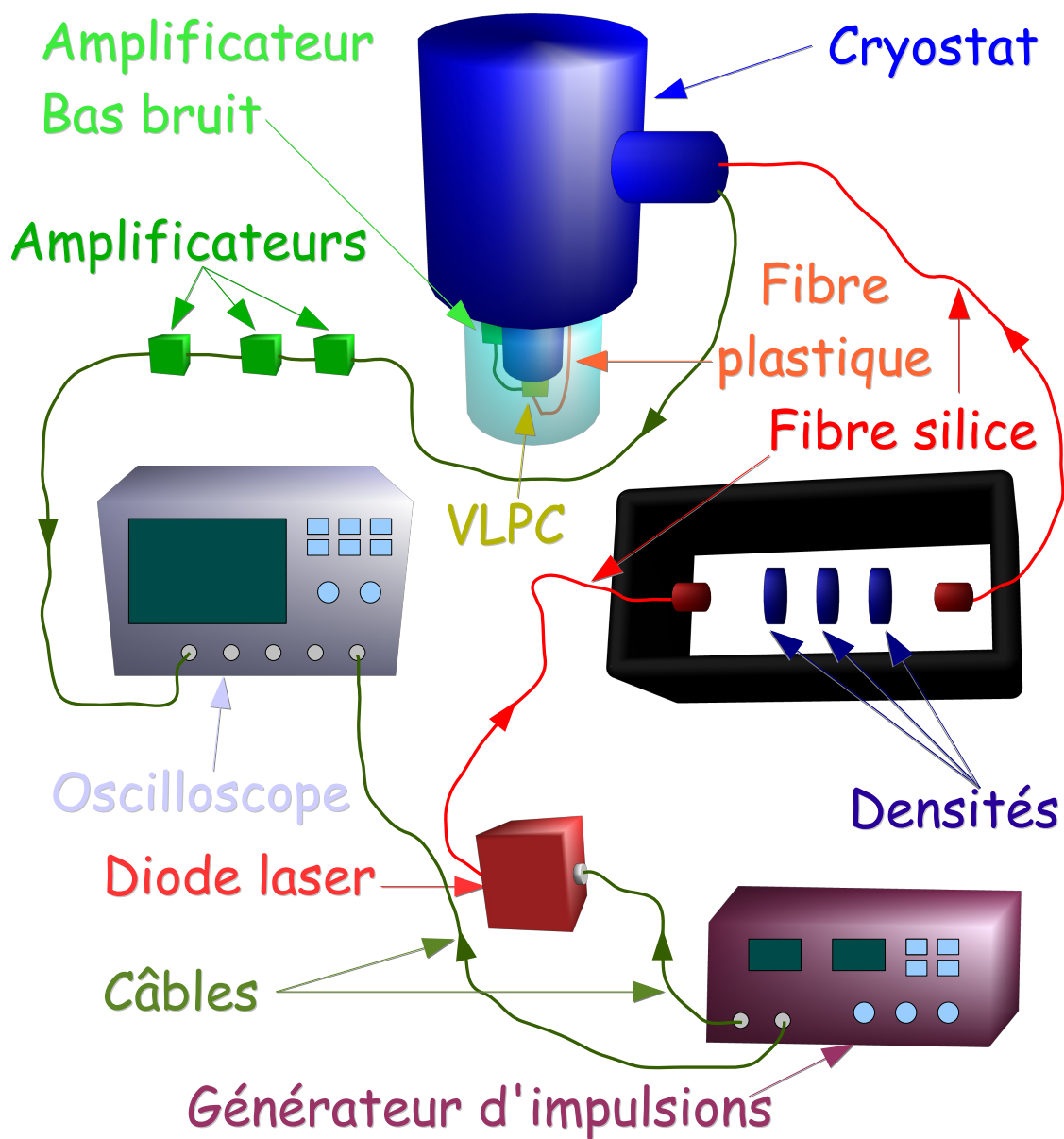


FIGURE 8.6: Dispositif utilisé pour tester le VLPC

L'acquisition est effectuée à l'aide d'un oscilloscope rapide *waveRunner 104 MXi* de *LeCroy* ayant un taux d'échantillonnage de 10 GEch./s. Une mesure correspond à 10 000 acquisitions successives synchronisées avec le générateur d'impulsions ; pour chaque acquisition on mesure l'aire sous la courbe dans une fenêtre de 8 ns centrée sur le pic de tension (cf. figure 8.7). Ces aires sont ensuite réunies dans un histogramme, chaque histogramme correspondant à une mesure.

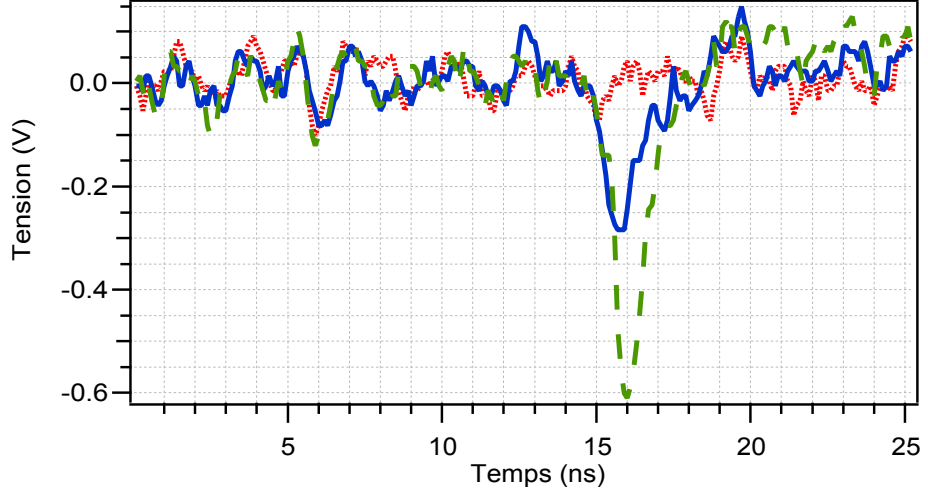


FIGURE 8.7: Signaux en sortie du VLPC, correspondant à la détection de 0 (pointillés), 1 (trait plein) et 2 (tirets) photons.

S'il n'y avait aucun bruit alors le pic de tension mesuré serait toujours le même pour un nombre de photons donné, ce qui se traduirait dans l'histogramme par des pics de largeur nulle. Ce n'est bien évidemment pas le cas : le processus d'avalanche est bruité et il en est de même des amplificateurs électroniques. Dans la pratique les différents pics de l'histogramme, correspondant à des nombres de photons différents, peuvent être modélisés par des gaussiennes [275, 272]. En effet un évènement à plusieurs photons résulte de plusieurs avalanches et est donc la somme d'autant de pics à un photon, ainsi que du bruit électronique. Le fait que la somme de variables aléatoires gaussiennes suit une distribution gaussienne ainsi que le théorème central limite font des gaussiennes une bonne approximation. Pour le pic à un photon les calculs théoriques prédisent une distribution bi-sigmoïdale [270] ; pour un large gain de multiplication, ce qui est le cas du VLPC, elle est cependant proche d'une distribution gaussienne.

Permettre un ajustement indépendant des paramètres des différentes gaussiennes introduirait trop de degrés de liberté, il convient donc d'exprimer ceux-ci selon des lois simples. Puisqu'un évènement à plusieurs photons correspond à la somme d'évènements à un photon nous pouvons considérer l'écart entre les pics comme constant (du moins tant qu'il n'y a pas d'effet de saturation). Ce qui nous donne

$$x_i = x_0 + i\Delta_x \quad (8.1)$$

où Δ_x est l'écart entre les pics, proportionnel au gain de multiplication [273], et x_0 la position du pic à zéro photon. De plus cela entraîne aussi l'additivité des variances des évènements sommés, qui s'écrivent donc

$$\sigma_i^2 = \sigma_0^2 + i\sigma_M^2 \quad (8.2)$$

où σ_0 est la largeur du pic à zéro photon et correspond donc au bruit électronique, σ_M étant le bruit de multiplication d'un évènement à un photon.

L'histogramme peut alors être modélisé par la fonction suivante :

$$h(x) = K \sum_{i=0}^N P_i \frac{1}{\sqrt{2\pi}\sigma_i} e^{-\frac{(x-x_i)^2}{2\sigma_i^2}} \quad (8.3)$$

où N est le nombre de photons maximum, K un coefficient de proportionnalité et P_i la probabilité de mesurer i photons, qui dans notre cas particulier suit la loi de poisson donnée équation 2.64.

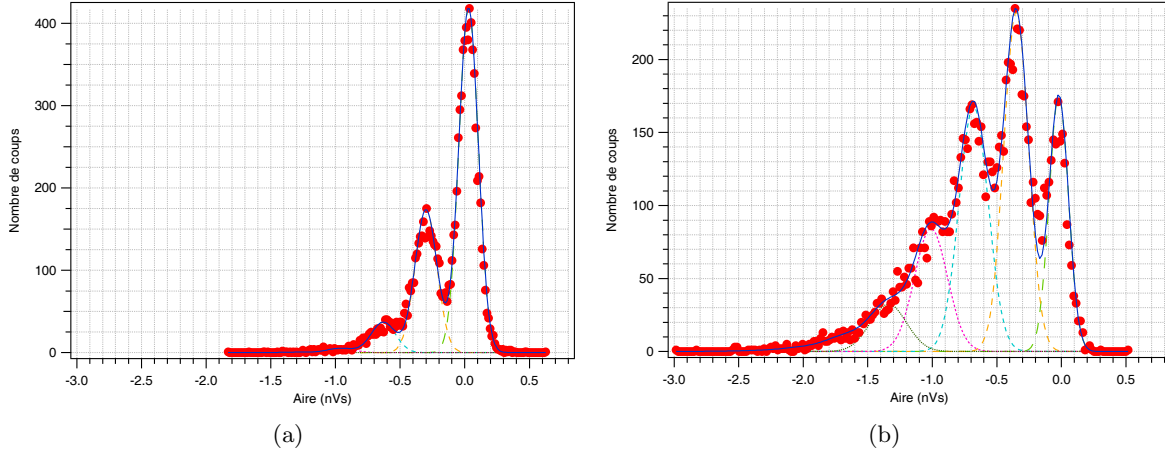


FIGURE 8.8: Histogrammes obtenus à partir du signal du VLPC, la courbe en trait plein correspond au fit et celles en tirets aux divers gaussiennes qui le compose (a) Pour un nombre moyen de photon mesurés de 0,49 (b) Pour un nombre moyen de photons mesurés de 1,7

À partir de ce modèle nous pouvons effectuer un *fit* des histogrammes (cf. figure 8.8) et ainsi déterminer l'amplitude de l'état cohérent vu par le VLPC. Étant donné que nous connaissons la puissance issue de la diode, l'atténuation des densité et la transmission des fibres optiques, nous pouvons aussi calculer l'amplitude de l'état cohérent arrivant sur le détecteur. Tout ceci nous permet alors de mesurer l'efficacité quantique du VLPC.

8.3.2 Recherche d'un point de fonctionnement

Nous pouvons régler deux paramètres qui influent sur la réponse du VLPC : la température et la tension de polarisation. Lorsqu'ils augmentent, l'avalanche est plus facilement déclenchée, et il s'en suit une meilleure efficacité du détecteur mais aussi une plus forte présence de coups d'obscurité.

Or lorsqu'il s'agit de compter les photons, et non seulement de détecter leur présence, ces deux paramètres jouent grandement sur la confiance que l'on peut avoir dans le résultat de la mesure. Une mauvaise efficacité entraîne des déclenchements du conditionnement sur un nombre de photons trop important. Jusqu'ici nous limitons la probabilité qu'il y ait plus de photons que nécessaire afin qu'elle soit négligeable mais nous ne pouvons pas utiliser cette astuce dans tous les cas et l'avantage d'un compteur de photons est justement de rendre ceux-ci accessible à l'expérience. À l'opposé un taux d'obscurité trop important déclenchera le conditionnement sur un nombre de photons trop peu élevé.

En outre ces l'efficacité quantique et les coups d'obscurité peuvent non seulement nuire à la qualité de l'état préparé mais ils peuvent aussi réduire les chances de détecter le bon état. L'impact de l'efficacité quantique μ est bien plus crucial que pour un simple détecteur de photons :

le taux de succès varie en effet en μ^n , donnant une décroissance exponentielle avec le nombre de photons n . Celui des coups d'obscurité est quant à lui nouveau par rapport aux expériences réalisées auparavant et vient donc s'ajouter aux sources d'imperfections « habituelles ». À noter qu'ils entraînent tous deux une décroissance du taux de succès d'un bon conditionnement mais augmente celui d'un mauvais conditionnement, dégradant ainsi doublement l'état. Le taux de succès total, et expérimental, peut ainsi augmenter ou diminuer avec ces paramètres suivant leur valeurs et l'état de départ.

Malheureusement améliorer l'un des deux paramètres entraîne la détérioration du second ; comme souvent nous devons donc effectuer un compromis entre les deux. Le premier test à effectuer consiste donc à explorer la tension de polarisation et la température afin de trouver un point de fonctionnement correspondant à un compromis convenable. Ceci permet d'autre part d'étudier l'efficacité du VLPC à 850 nm ainsi que les coups d'obscurité avec notre montage.

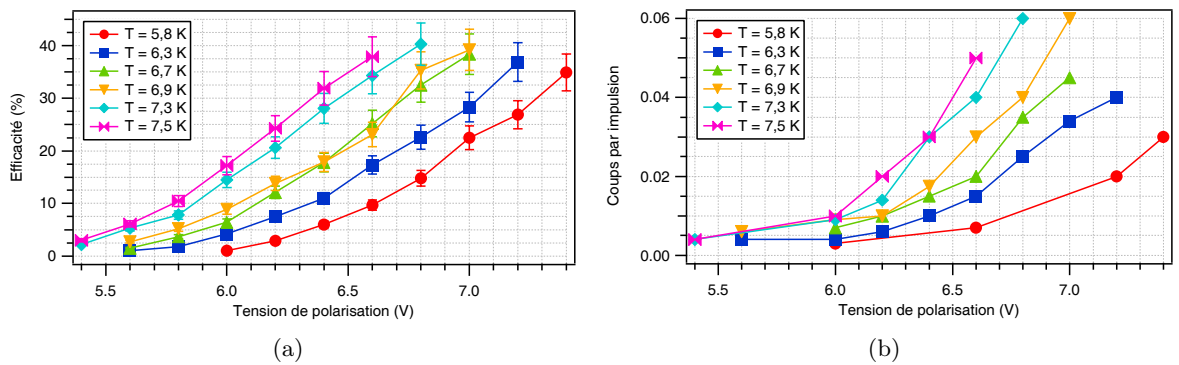


FIGURE 8.9: Étude du point de fonctionnement (a) Efficacité quantique en fonction de la tension pour différentes températures (b) Coups d'obscurités par impulsion en fonction de la tension pour différentes températures

La figure 8.9 montre les résultats de ce test. Nous pouvons remarquer que l'efficacité quantique est environ deux fois plus faible que les précédentes expériences menées à 543 et 694 nm [272, 274]. Nous pouvons cependant noter que le traitement anti-reflet de notre capteur VLPC n'est pas adapté à notre longueur d'onde et est donc en partie responsable de cette baisse d'efficacité. De plus la collection des photons en sortie de la fibre s'est avérée être délicate, et peut donc être elle aussi une cause de la baisse d'efficacité. Le gros point noir de notre dispositif réside dans les coups d'obscurités. Malgré toutes nos précautions ils demeurent en effet extrêmement importants. Il est ainsi difficile d'obtenir moins de 0,01 coup par impulsions, soit de l'ordre de 10^7 cps/s, et de tels efforts font considérablement chuter l'efficacité.

8.3.3 Efficacité en fonction du flux incident

Comme tous les compteurs basés sur du multiplexage, le VLPC voit son efficacité diminuer avec le nombre de photons incidents. Cet effet est traditionnellement restreint au niveau d'une seule impulsion lumineuse ; mais le long temps de relaxation des zones ayant subi une avalanche induit un effet similaire entre les impulsions : non seulement les photons d'une même impulsions ne doivent pas être absorbés dans la même zone mais ils ne doivent pas non plus être absorbés dans une zone ayant servi récemment. La question qui se pose alors est : quelle est l'influence du nombre de photons incidents sur l'efficacité quantique du VLPC ?

Nous pouvons en outre profiter de ce test pour mettre en œuvre une réalisation pratique de

la mesure du nombre de photons sur une impulsion donnée. Le fit utilisé auparavant ne donne en effet que la statistique et non une mesure précise pour chaque impulsion. Pour cela nous devons définir des régions de décision correspondant aux différents nombres de photons : toutes les aires comprises dans une de ces régions seront considérées comme traduisant le même nombre de photons absorbés [234, 275]. Les limites naturelles de ses frontières sont constituées par les points où deux gaussiennes adjacentes du fit se croisent. Seulement ces points dépendent de la hauteur relatives des gaussiennes, c'est-à-dire du poids de chaque état de Fock dans l'état mesuré. Nous avons alors trois possibilités pour définir ces frontières :

- Considérer que l'on a aucune connaissance sur l'état d'entrée, ce qui revient à prendre le même poids pour toutes les gaussiennes.
- Utiliser une connaissance à priori de l'état mesuré, cas où le conditionnement se fait sur un état simple (par exemple un mode d'une paire EPR) ou utilisation d'un modèle le décrivant.
- On effectue d'abord le fit de l'ensemble des données puis on utilise celui-ci pour calculer les frontières et enfin on effectue une post-sélection des données à l'aide de celles-ci.

Pour ce test nous choisirons cette dernière méthode (à la différence près que nous n'effectuerons pas de post-sélection puisqu'il s'agit d'une simple mesure et non d'un conditionnement).

Cette discussion nous amène à une dernière source d'erreur dans la mesure du nombre de photons : de par le bruit des différents pics, les gaussiennes se recoupent, ce qui implique qu'une même mesure peut correspondre à deux nombre de photons différents. Il n'est donc pas toujours possible de discriminer le nombre de photons avec certitude. La solution consiste alors à réduire les régions de décisions en supprimant les zones où le doute est trop important, laissant ainsi des résultats indécis qui ne seront pas utilisés. Cette technique diminue le taux de succès puisque qu'un certain nombre de mesures ne pourront être résolues néanmoins elle n'altère en rien l'efficacité de détection. Notons cependant qu'elle n'est utile que pour un conditionnement ; dans le cadre de ce test, consistant à reconstruire la statistique du nombre de photons d'une impulsion, elle n'est pas utile ; aussi nous ne l'utiliserons pas.

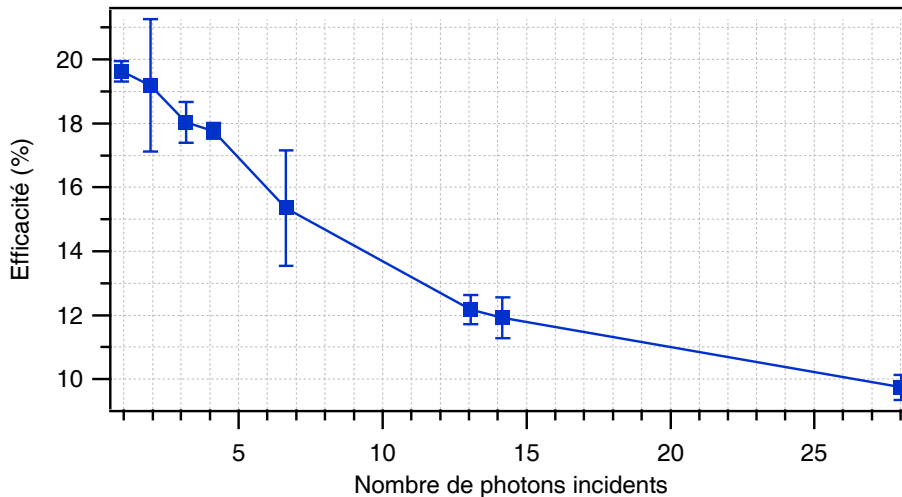


FIGURE 8.10: *Efficacité quantique en fonction du nombre de photons incidents*

La figure 8.10 montre cette évolution de l'efficacité en fonction du nombre de photons moyen par impulsion. En se restreignant à des impulsions de quelques photons (5 au plus) la baisse d'efficacité est peu importante. Mais cette courbe nous apprend autre chose : les effets du flux de photons entre plusieurs impulsions consécutives font qu'augmenter la cadence des impulsions

a un effet très similaire à l'augmentation du nombre de photons moyens d'un même facteur. Actuellement, le taux de photons maximal incidents sur l'APD, venant de l'OPA, est de $20\,000\text{ /s}^2$, ce qui correspond à un flux moyen de 2 photons pour les tests effectués à une cadence de 10,1 kHz, ce qui n'induit qu'une très faible baisse. Il y a cependant plusieurs points à prendre en considération :

- L'un des intérêts de l'amplificateur femtoseconde est d'augmenter le nombre moyen de photons dans les états produits par l'OPA ; c'est d'ailleurs aussi dans cet optique qu'un compteur de photons est intéressant. Ce qui augmente le taux de photons incidents, et diminue l'efficacité du détecteur de quelques pourcents.
- Certaines expériences, comme l'amplificateur sans bruit, nécessitent d'envoyer un faible état cohérent sur l'APD. Le taux de photons incidents peut alors être bien plus important (jusqu'à plus d'un ordre de grandeur).
- Ce taux de photons incidents est obtenu avec un dispositif de filtrage ayant une transmission comprise entre 10 et 15 %, ce qui empêche des utilisations nécessitant une forte efficacité. Celles-ci demanderaient donc de revoir entièrement le système de filtrage, afin d'augmenter fortement sa transmission, et donc le nombre de photons incidents.
- Le fort taux d'obscurité, qui même en l'améliorant restera bien plus important que pour une APD (cf. table 8.1), peut impliquer la nécessiter d'avoir un nombre de photons incidents plus important afin de maintenir une contribution négligeable des coups d'obscurité. Ce dernier point ne vient cependant pas s'ajouter au précédents et pourrait, moyennant une amélioration de notre système, être compatible avec une hausse du taux de photons incidents induite par l'amplificateur femtoseconde ou une augmentation modérée de la transmission du système de filtrage ; ce qui n'entraînerait qu'une faible baisse de l'efficacité.

Au final une utilisation à notre taux de répétition peut s'avérer limitante pour les applications les plus exigeantes, mais ne poserait pas de problème dans les cas les plus courants .

8.4 Conclusion

Ainsi le VLPC ne permet pas de se passer de l'approximation consistant à rendre négligeable la probabilité d'avoir plus de photons que désiré (ce qui est de toutes façons difficilement envisageable avec notre système à cause de la nécessité de filtrer les impulsions en amont du détecteur) ; bien qu'il apporterait tout de même une légère amélioration à ce niveau. Il n'en est cependant pas moins intéressant pour remplacer le multiplexage d'APD, augmentant ainsi le taux de succès. Ceci nécessiterait par contre un travail conséquent pour améliorer le dispositif. Or nous pouvons remarquer que pendant ce travail de thèse sont apparues de nouvelles techniques de comptage qui pourraient s'avérer être bien plus intéressantes. Les TES ont ainsi des performances remarquables, et sont sans aucun doute les meilleurs compteurs de photons actuels ; leur mise en œuvre demande néanmoins un dispositif complexe. À l'opposé les APD fonctionnant en dessous du régime de saturation ont des performances plus modestes mais leur utilisation est très simple. Entre les deux, le CIPD, s'il arrive à dépasser la limitation actuelle concernant sa fréquence de répétition, offre un très bon rapport entre les performances et la complexité.

2. Ce qui, avec l'efficacité d'environ 50% de l'APD, redonne bien les 10 000 cps/s cités dans les chapitres précédents.

Chapitre 9

Conclusion et perspectives

Sommaire

9.1 Conclusion	181
9.2 Perspectives	182
9.2.1 Améliorations du dispositif expérimental	182
9.2.2 Calcul quantique à variables continues	182
9.2.3 De nouveaux outils	183

9.1 Conclusion

Au cours de cette thèse nous avons donc été amenés à aborder sous divers angles l’apport de l’approche discrète au domaine des variables continues. Nous avons ainsi pu dégager de nouveaux intérêts à ce mélange des deux approches et tester de nouveaux outils permettant d’en quantifier l’apport.

Nous avons d’abord utilisé cette technique pour créer de l’intrication. Il était déjà connu que la distillation d’intrication dans le domaine des variables continues requière des opérations non-gaussiennes, et la méthode la plus simple pour les réaliser consiste à utiliser des techniques associées aux variables discrètes. Nous avons montré qu’il était aussi possible de créer de l’intrication à distance et que cette technique est plus efficace que la distillation d’intrication. De plus cette intrication peut être réalisée à travers un canal ayant de fortes pertes, caractéristique qui s’avère très utile pour les communications quantiques à moyenne et longue distance. Cette première expérience illustre les premiers apports de ce mélange, apports qui en sont à l’origine et sont maintenant bien connus.

Nous avons ensuite démontré qu’il était possible d’utiliser les mesures projectives permises par ces techniques afin d’effectuer des opérations normalement interdites par les lois de la physique quantique. Nous avons ainsi pu réaliser l’amplification d’un état cohérent sans en amplifier le bruit. Ceci nous a permis d’obtenir un appareil capable d’augmenter le rapport signal sur bruit. Cette deuxième expérience apporte de nouveaux intérêts au mariage des variables continues et discrètes. Ces intérêts étaient encore méconnus au début de cette thèse et notre démonstration expérimentale a contribué à les populariser. Ils ont fait l’objet d’un assez grand enthousiasme dans la communauté, ouvrant la voie vers la recherche de nouveaux protocoles permettant d’aller plus loin que les limites actuelles.

Enfin nous avons testé et comparé de nouveaux outils concernant l'un des pivots de cette approche : la non-gaussianité des états. Nous avons pu confronter ces mesures à une estimation de la non-classicité, autre propriété importante issue du plein emploi de la dualité onde-corpuscule. Cette étude expérimentale a d'ailleurs permis de soulever l'intérêt d'une analyse conjointe de ces deux caractéristiques.

9.2 Perspectives

Au delà des diverses preuves expérimentales et des développements apportés par ce travail, ce dernier a aussi permis de dégager plusieurs perspectives pour l'avenir. En voici un rapide aperçu.

9.2.1 Améliorations du dispositif expérimental

Le dispositif utilisé au cours de cette thèse fut l'un des premiers à permettre l'exploitation de la combinaison des approches discrète et continue en régime impulsionnel, et a permis de réaliser plusieurs premières mondiales. Cependant nous aurons besoin à l'avenir de produire des états et de réaliser des opérations toujours plus complexes. Il est donc indispensable de l'améliorer, à la fois dans le but de réaliser les futures expériences et dans celui de rester compétitifs dans un domaine où la concurrence c'est affirmée ces dernières années.

L'amélioration la plus simple, effectuée en fin de thèse, concerne le laser lui-même. Nous avons évoqué le fait qu'il avait été poussé au delà de ses spécifications dans le but d'obtenir plus de puissance, conduisant à de longs et difficiles réglages pour le faire fonctionner correctement. Or il existe dorénavant des lasers commerciaux ayant la même puissance de fonctionnement que celle pour laquelle nous nous bâtons avec le *Tiger-CD*. Nous avons donc acquis un nouveau laser, bien plus stable. Il s'agit d'un *Mira 900-F* de la société *Coherent*. Celui-ci permettra de gagner énormément de temps sur les réglages, ce qui est d'autant plus utile que les futures améliorations les complexifieront encore d'avantage.

Les expériences demandent de plus en plus d'états de base, de plus en plus de conditionnements, ainsi que des états plus comprimés ou comprenant plus de photons. L'amplificateur femtoseconde sur lequel nous avons travaillé s'avérera alors capital pour mener à bien ces expériences. Il permettra une meilleure conversion dans les cristaux non-linéaires, augmentant ainsi la compression et le nombre moyen de photons, ce qui augmentera aussi le taux de succès des conditionnements. Il pourra aussi alimenter plusieurs OPA, augmentant ainsi le nombre d'états de base utilisables. Bien que l'intégration de l'amplificateur ne soit pas terminée, nous avons obtenu des résultats encourageants et pu dégager les principaux axes de travail pour la compléter.

Les tests du compteur de photons dont nous disposons ont donné des résultats plus négatifs, cependant la discrimination du nombre de photon est un domaine en plein essor, et de nouvelles techniques très prometteuses sont apparues pendant cette thèse. Ces tests nous ont tout de même permis de soulever un autre point : il ne suffit pas d'avoir un compteur avec une bonne efficacité. Nous avons en effet besoin d'une bonne efficacité globale, ce qui veut dire que nous devons aussi revoir le système de filtrage, afin que ses pertes soient bien plus faibles pour le mode concerné.

9.2.2 Calcul quantique à variables continues

Les expériences que nous avons menées sont plus tournées vers la communication quantique que vers l'ordinateur quantique. Et c'est aussi le cas plus généralement dans l'ensemble du

domaine des variables continues en optique quantique : ce n'est d'ailleurs pas étonnant, si l'on compare au cas classique, nous remarquons déjà que l'optique sert bien plus aux communications qu'aux ordinateurs. Les communications optiques sont omniprésentes, alors que les ordinateurs optiques font encore l'objet de recherches, et sont loin d'une commercialisation.

Il existe néanmoins quelques propositions de protocoles de calcul quantique à variable continues. À la fin de cette thèse, de nouveaux protocoles de portes à un et deux qubits ont été proposés [276]. Ces protocoles sont assez proches des expériences que nous avons déjà réalisées, et notre dispositif est donc bien adapté à leur mise en œuvre. Celle-ci est d'ailleurs déjà entamée, formant ainsi la première expérience réalisée avec le nouveau laser, et complétant la gamme d'applications de l'information quantique abordées.

Cette thèse a permis d'aller plus loin dans le mélange des approches discrètes et continues en démontrant de nouveaux intérêts et en illustrant des mesures qui lui sont liées. Un autre axe intéressant à étudier pour le calcul quantique est celui, déjà évoqué à deux reprises dans ce manuscrit, du calcul hybride. Ce dernier permettrait de renouveler encore une fois l'intérêt de l'approche mixte.

9.2.3 De nouveaux outils

Enfin un dernier axe abordé dans cette thèse, et qui aura certainement des développements intéressants par la suite, est la conception et le test de nouveaux outils. Les tests que nous avons été amenés à faire ont été riches en enseignements et ont aussi mis en évidence l'intérêt de disposer de plus d'outils que nous n'en avons actuellement.

Nos expériences vont aussi de plus en plus de la production d'état à la manipulation de ceux-ci. Si la tomographie d'un état produit suffit pour le décrire, une transformation nécessite plus que de tomographier l'état de sortie pour la caractériser entièrement. L'équivalent de la tomographie d'un état pour une opération est la tomographie de processus, qui permet de reconstruire le superopérateur décrivant cette opération. Déjà utilisée dans d'autres domaines de l'information quantique, elle commence à se développer aussi en optique quantique. Cependant il nous manque encore de bons outils pour la réaliser en variables continues et avec des processus conditionnés.

Outre d'offrir une description complète d'un processus quantique réel, cette tomographie possède un second avantage non négligeable : pour la réaliser nous devons appliquer le processus à un ensemble d'état formant une base de l'espace de Hilbert ; ainsi, il suffit d'envoyer en entrée des états cohérents [277], y compris si le processus est prévu pour fonctionner avec des états plus compliqués¹. Les états cohérents étant très faciles à produire, cela simplifie la réalisation des expériences ; et c'est surtout un autre moyen de faire face aux processus quantiques qui nécessitent plus d'états de bases normalement produits par l'OPA.

1. Notons que, pour un processus conditionné, il ne faut pas uniquement utiliser les états en sortie mais aussi la probabilité de succès. Ainsi, les états cohérents peuvent aussi être utilisés pour tomographier des processus qui les laissent inchangés, comme la dégaussification ou la porte de phase pour des qubits encodés sur la phase d'un état cohérent.

Quatrième partie

Annexes

Annexe A

Formulaire et considérations mathématiques

A.1 Intégrales gaussiennes

Les fonctions de Wigner des états que nous manipulons sont toujours sous forme de gaussiennes ou de gaussiennes multipliées par des polynômes. Les calculs théoriques sont donc essentiellement des calculs d'intégrales gaussiennes. Celles-ci sont fort heureusement relativement simples, sans quoi nous ne pourrions avoir de formules mathématiques aussi pratiques et précises pour la description de nos expériences.

Cette section présente les différentes formules utiles pour calculer ces intégrales, de la plus simple à la plus complète. Nous nous contenterons de les énoncer, les démonstrations pouvant être trouvées dans de nombreux livres de mathématiques.

Commençons par la formule la plus simple, à savoir l'intégrale d'une gaussienne seule :

$$\int_{-\infty}^{\infty} e^{-ax^2+bx} dx = \sqrt{\frac{\pi}{a}} e^{\frac{b^2}{4a}} \quad (\text{A.1})$$

où a est de partie réelle strictement positive et b peut prendre n'importe quelle valeur complexe.

Ajoutons maintenant un polynôme. Grâce à la linéarité de l'intégration nous pouvons décomposer le calcul en intégrale d'une gaussienne multiplié par un monôme. Nous ne verrons donc que ce cas de figure. Commençons par le cas d'une gaussienne centrée sur l'origine, le résultat diffère alors suivant que la puissance de ce monôme est paire ou impaire :

$$\int_{-\infty}^{\infty} x^{2n} e^{-x^2/\sigma^2} dx = \sigma^{2n} \frac{(2n)!}{2^{2n} n!} \sqrt{\pi \sigma^2} \quad (\text{A.2})$$

$$\int_{-\infty}^{\infty} x^{2n+1} e^{-x^2/\sigma^2} dx = 0 \quad (\text{A.3})$$

avec $\sigma \in \mathbb{R}$ et $n \in \mathbb{N}$. Passons enfin à la formule complète :

$$\int_{-\infty}^{\infty} x^n e^{-ax^2+bx} dx = \sqrt{\frac{\pi}{a}} e^{\frac{b^2}{4a}} \sum_{k=0}^{\lfloor n/2 \rfloor} C_n^{2k} (2k-1)!! (2a)^{k-n} b^{n-2k} \quad (\text{A.4})$$

avec la convention $(-1)!! = 1$ et b non nul, les autres conditions étant les mêmes que précédemment.

A.2 Pertes et amplification parasite

D'un point de vue mathématique, l'ajout de perte homodynes et d'une amplification parasite sur une fonction de Wigner monomode sont similaires. Les deux reviennent à appliquer une matrice de changement de variables sur le mode de cette dernière et un mode vide, puis à tracer sur celui-ci. Dans les deux cas la matrice peut être décomposée en deux sous-matrice identiques, l'une agissant sur les quadratures x et l'autre sur les quadratures p . Ces deux sous-matrices sont de la forme

$$M = \begin{pmatrix} \mu & \nu \\ \epsilon\nu & \mu \end{pmatrix} \quad (\text{A.5})$$

avec μ et ν réels, $\epsilon = \pm 1$ et $\det(M) = \mu^2 - \epsilon\nu^2 = 1$.

La transformation sur un état $W(x, p) = f(x)f(p)$ avec

$$f(x) = \frac{1}{\sqrt{\pi}\sigma} e^{-\frac{(x-\sqrt{2}\alpha_x)^2}{\sigma_x^2}} \quad (\text{A.6})$$

donne donc :

$$\int_{-\infty}^{\infty} \frac{1}{\sqrt{\pi}\sigma_x} e^{-\frac{(\mu x - \nu y - \sqrt{2}\alpha_x)^2}{\sigma^2} - (\mu y - \epsilon\nu x)^2} dy = \frac{1}{\sqrt{\pi(\mu^2\sigma_x^2 + \nu^2)}} e^{-\frac{(x-\sqrt{2}\mu\alpha_x)^2}{\mu^2\sigma_x^2 + \nu^2}} \quad (\text{A.7})$$

où y correspond au mode vide. Ceci nous donne donc la transformation suivante :

$$\alpha \mapsto \mu\alpha \quad (\text{A.8})$$

$$\sigma^2 \mapsto \mu^2\sigma^2 + \nu^2 \quad (\text{A.9})$$

Nous pouvons effectuer le même calcul pour une fonction de Wigner bimode $W(x_1, p_1, x_2, p_2) = f(x_1, x_2)f(p_1, p_2)$ telle que

$$f(x_1, x_2) = \frac{1}{\pi\sigma_{1,x}\sigma_{2,x}} e^{-\frac{x_1+x_2-\sqrt{2}\alpha_{1,x}}{2\sigma_{1,x}} - \frac{x_1-x_2-\sqrt{2}\alpha_{2,x}}{2\sigma_{2,x}}} \quad (\text{A.10})$$

À condition d'effectuer le même changement de variable M sur les deux modes, nous obtenons les mêmes transformations.

A.3 Bruit électronique

Appel *et al.* ont démontré que le bruit électronique de la détection homodyne pouvait être modélisé comme une efficacité, pour peu que l'on ne corrige pas de celui-ci le bruit du vide utilisé pour normaliser les données prises [90]. Sans refaire la démonstration, voici un petit calcul simpliste qui s'intéresse à la variance et à la moyenne, et est donc plus proche des cas pratiques que l'on rencontre.

La variance du vide V_{SNL} , utilisée pour la normalisation, vaut :

$$V_{\text{SNL}} = N_0 + V_e \quad (\text{A.11})$$

où N_0 est la valeur du vide en l'absence de bruit électronique (c'est le « vrai » facteur de normalisation) et V_e le bruit électronique (tel que mesuré en l'absence d'oscillateur local).

Si nous utilisons le vrai facteur de normalisation nous devrions convoluer la fonction de Wigner mesurée par la détection homodyne avec une gaussienne de largeur égale à la variance normalisée du bruit électronique e [82]

$$e = \frac{V_e}{V_{\text{SNL}}} = \frac{V_e}{N_0 + V_e} \quad (\text{A.12})$$

Ici, nous allons introduire le paramètre η_{elec} , tel que

$$\eta_{\text{elec}} = 1 - e = \frac{N_0}{N_0 + V_e} \quad (\text{A.13})$$

nous pouvons noter que nous avons bien $0 < \eta_{\text{elec}} \leq 1$.

Avec cette notation nous pouvons en déduire que le facteur de normalisation utilisé (V_{SNL}) est relié au facteur de normalisation réel (N_0) par

$$N_0 = \eta_{\text{elec}} V_{\text{SNL}} \quad (\text{A.14})$$

Nous pouvons aussi réécrire le bruit électronique sous la forme

$$V_e = (1 - \eta_{\text{elec}}) V_{\text{SNL}} \quad (\text{A.15})$$

Considérons un état de valeur moyenne α et de variance σ^2 , la valeur moyenne mesurée m vaut alors :

$$m = \alpha \sqrt{N_0} \quad (\text{A.16})$$

$$= \sqrt{\eta_{\text{elec}}} \alpha \sqrt{V_{\text{SNL}}} \quad (\text{A.17})$$

De même la variance mesurée V vaut :

$$V = \sigma^2 N_0 + V_e \quad (\text{A.18})$$

$$= (\eta_{\text{elec}} \sigma^2 + 1 - \eta_{\text{elec}}) V_{\text{SNL}} \quad (\text{A.19})$$

Après normalisation des mesures nous obtenons bien les mêmes relations que pour les pertes optiques. Ce calcul nous a aussi permis, au passage, d'exprimer la valeur de ces pertes en fonction des données mesurées expérimentalement.

Annexe B

Le laser femtoseconde

B.1 Critères importants

Maintenant que nous savons quel type de laser utiliser nous devons définir quelles sont les critères importants à respecter pour effectuer les expériences dans de bonnes conditions. Passons-les en revue :

- Le premier critère est l'énergie par impulsion, elle doit être suffisamment importante pour permettre d'obtenir des effets non-linéaires en simple passage. Nous avons besoin d'au minimum quelques dizaines de nanojoules par impulsions, sachant que plus elle sera importante plus nos possibilités d'expériences seront grandes.
- Le deuxième est la cadence de répétition, elle doit être suffisamment élevée pour acquérir rapidement les données afin de pouvoir réaliser des statistiques sur celles-ci en évitant les dérives expérimentales. Par contre une cadence trop élevée donne des puissances trop faibles et pourrait dépasser les capacités de l'électronique des systèmes de détections. Une cadence raisonnable serait comprise entre quelques centaines de kilohertz et la dizaine de mégahertz.
- Le suivant est la durée des impulsions. Ce critère n'est pas totalement indépendant des précédents. Il est en effet difficile d'obtenir des impulsions longues ayant suffisamment d'énergie avec un taux de répétition convenable. Il nous faut donc des impulsions femtosecondes. Elles ne devront cependant pas être trop courtes afin de minimiser l'étalement des impulsions lors du passage dans les cristaux non-linéaires. Une durée comprise entre 100 fs et 300 fs est un bon compromis.
- La longueur d'onde est aussi un critère important ; afin de limiter les pertes lors de la détection, elle doit être dans une gamme où ceux-ci possèdent une forte efficacité quantique. Pour les photodiodes à silicium, qui associent une bonne efficacité à un faible courant d'obscurité, il s'agit de la gamme 800 – 900 nm correspondant à l'infrarouge très proche. De plus après doublage de fréquence les faisceaux restent dans le domaine visible et peuvent donc être manipulés avec des composants optiques standards.
- Les enveloppes temporelles et spectrales doivent être aussi proches que possible de la limite de Fourier afin de pouvoir considérer les impulsions comme monomodes. Le non respect de cette limite peut être dû à des impulsions monomodes temporellement, induisant une dégradation de la pureté des états produits, ou à un glissement de fréquence (ou « chirp » en anglais) qui diminuerait la puissance crête. Le milieu amplificateur devra aussi avoir une bande de gain suffisamment large pour couvrir la largeur spectrale des impulsions qui est de l'ordre de la dizaine de nanomètres.

- Enfin le mode spatial devra être de la meilleure qualité possible et dans le mode gaussien TEM_{00} . Ceci afin d'obtenir d'une part des états monomodes spatialement et d'autre part une bonne adaptation des modes lorsqu'il faudra faire interférer plusieurs faisceaux.

Ces différents critères demandent clairement de faire des compromis, le principal étant entre la cadence de répétition et l'énergie par impulsion. Les lasers standards, qui ont un taux de répétition de 80 MHz génèrent des impulsions dont l'énergie est encore un peu faible malgré les progrès techniques et le succès d'expériences récentes les utilisant [278, 279, 280]. Nous utilisons donc un système en cavité fermée : les impulsions font plusieurs fois le tour de la cavité avant d'en être extraites, ce qui donne une meilleure énergie par impulsion mais en diminuant la cadence de répétition.

B.2 Principe d'un laser femtoseconde

B.2.1 Amplification

Dans notre gamme de longueur d'onde le milieu amplificateur le plus commun est le cristal de saphir dopé au titane. Il possède l'avantage d'avoir une bande très large, entre 700 nm et 1 100 nm, une forte densité de stockage d'énergie ainsi qu'une excellente conductivité thermique.

B.2.2 Verrouillage de mode

La génération d'impulsion nécessite non seulement la présence simultanée de nombreux modes longitudinaux, c'est-à-dire avec des fréquences différentes, mais aussi que tous ces modes soient en phase (dans le cas contraire nous aurions un laser chaotique). Pour les durées d'impulsion que l'on s'est fixé la largeur spectrale est d'environ 2 THz, de plus les modes sont typiquement séparés par un intervalle spectral libre de 78 MHz (correspondant à une longueur cavité d'environ 1 m) ce qui nous donne dans les 30 000 modes à synchroniser.

Il existe plusieurs moyens de verrouiller les modes pour que leurs phases soient accordées, nous nous intéresserons uniquement aux méthodes dites *passives*. Elles ont pour point commun d'introduire un élément dont le rôle est de favoriser le fonctionnement impulsionnel en introduisant des pertes qui sont bien plus importantes pour les autres modes. Il existe principalement deux types de verrouillage passifs.

Verrouillage par absorbant saturable Il s'agit d'un dispositif dont l'absorption diminue quand la puissance incidente augmente, privilégiant ainsi les pics les plus intenses. Mais l'absorbant saturable n'est pas le seul à agir, une fois qu'il a présélectionné l'impulsion la plus intense c'est l'action combinée de celui-ci et du milieu amplificateur qui permet de stabiliser l'impulsion. Lorsque l'impulsion arrive sur l'absorbant saturable, le front subit une absorption tandis que le reste de l'impulsion bénéficie de la transparence induite par la saturation de l'absorption. Au niveau du milieu c'est la partie avant qui est amplifiée alors que la queue est peu modifiée à cause de la saturation du gain. Au final l'impulsion s'est donc fait tronqué à l'avant par l'absorbant saturable et à l'arrière par le milieu amplificateur, entraînant un raccourcissement jusqu'à ce qu'un régime stationnaire soit atteint.

Verrouillage par effet Kerr Lorsque cette méthode fut découverte elle fut d'abord appelée « magic mode-locking » car elle favorisait le fonctionnement impulsionnel sans ajout d'un quelconque dispositif spécialisé. C'est en effet le milieu amplificateur lui même qui, associé à un diaphragme, effectue la sélection. Le principe repose sur un effet non-linéaire du troisième

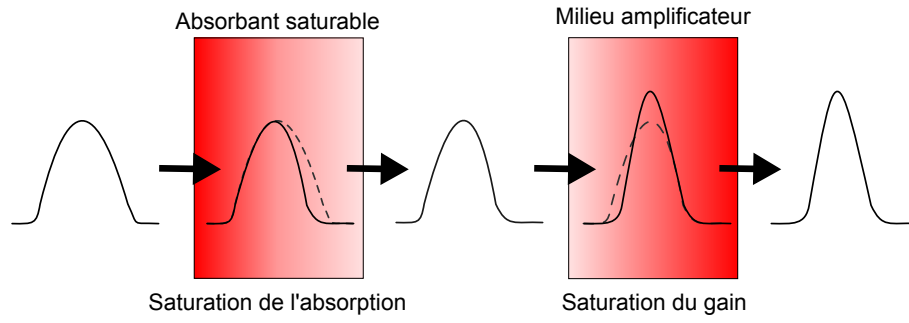


FIGURE B.1: Action de l'absorbant saturable et du milieu amplificateur sur une impulsion

ordre, l'effet Kerr, qui induit une augmentation de l'indice optique proportionnelle à l'intensité : $n = n_0 + n_2 I$. Le faisceau ayant un profil spatial gaussien, l'intensité est plus forte en son centre, le milieu amplificateur agit donc comme une lentille qui focalise d'autant plus le faisceau que celui-ci est intense. En plaçant correctement un diaphragme on peut donc privilégier les modes plus focalisés. Par contre cette méthode nécessite un processus de démarrage chargé de créer l'instabilité de départ.

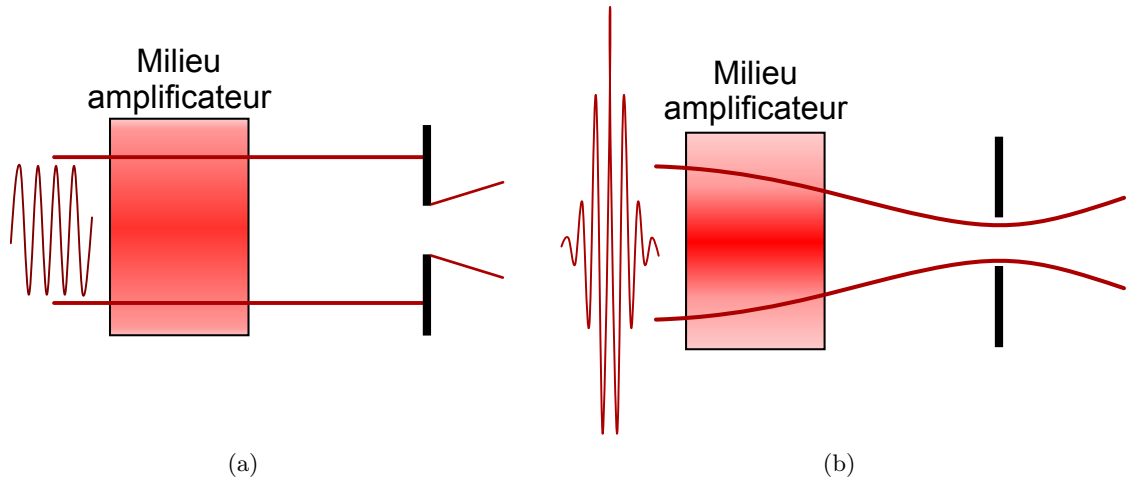


FIGURE B.2: Principe du verrouillage par effet Kerr (a) Mode continu (b) Mode impulsif

B.2.3 Dispersion de la vitesse de groupe

Le milieu amplificateur étant dispersif l'indice optique diminue avec la longueur d'onde ce qui induit un étirement de l'impulsion : les composantes décalées vers le bleu se déplacent moins vite que celles décalées vers le rouge. Nous pouvons néanmoins compenser cet effet à l'aide d'un ensemble de quatre prismes disposés de manière à allonger le chemin optique en fonction de la longueur d'onde. Pour les cavités linéaires il est possible de réduire le nombre de prismes à deux en plaçant le miroir de fin de cavité au milieu de ce dispositif.

B.2.4 Extraction des impulsions (cavity dumper)

Le laser fonctionnant en cavité fermée nous avons besoin d'un déflecteur optique qui viendra extraire l'impulsion de la cavité [281, 282]. Celui-ci peut être réalisé soit à l'aide d'un modulateur

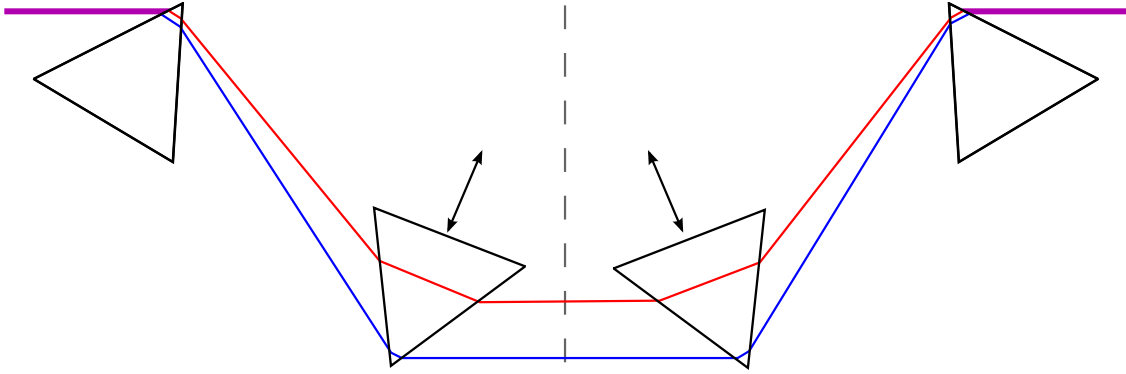


FIGURE B.3: Compensation de la dispersion de vitesse de groupe

électro-optique, ou cellule de Pockels, soit à l'aide d'un modulateur acousto-optique, ou cellule de Bragg. Si le premier a une meilleure efficacité, le deuxième est plus rapide et surtout moins dispersif. De plus une configuration en double passage permet d'obtenir un taux d'extraction convenable, c'est donc celui-ci qui est généralement utilisé. La synchronisation de l'extraction se fait à l'aide d'une photodiode rapide positionnée sur une fuite d'un miroir ou sur une partie prélevée afin de surveiller le fonctionnement de la cavité laser.

La cellule de Bragg est placée au centre de courbure du miroir de fin de cavité, orienté suivant l'angle de Bragg pour maximiser l'ordre 1 et en incidence de Brewster pour minimiser les pertes. Afin de décrire son fonctionnement, représentons le champ incident par $E(t) = E_0 \cos(\omega t)$ et

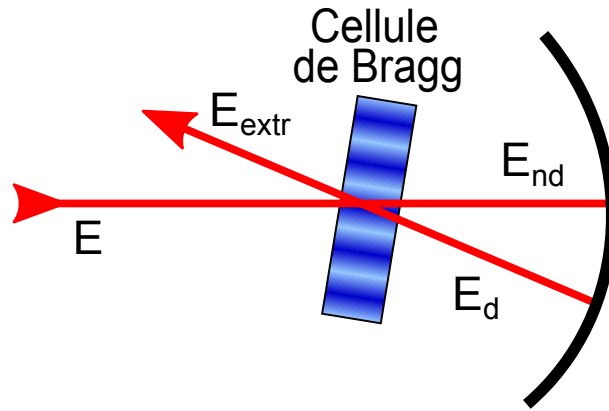


FIGURE B.4: Fonctionnement de la cellule de Bragg en double passage

l'onde RF par sa pulsation Ω et sa phase ϕ . Lors du premier passage le faisceau est séparé en une partie déviée E_d et une partie non déviée E_{nd} suivant une efficacité η

$$E_d = \sqrt{\eta} E_0 \cos((\omega + \Omega)t + \phi) \quad (\text{B.1})$$

$$E_{nd} = \sqrt{1 - \eta} E_0 \cos(\omega t) \quad (\text{B.2})$$

Après réflexion sur le miroir sphérique de fin de cavité ces deux faisceaux sont à nouveau focalisés dans la cellule de Bragg. Une seconde diffraction a lieu donnant le faisceau extrait qui sera finalement issue de l'interférence de deux contributions : la partie déviée uniquement au premier passage est celle déviée uniquement au second.

$$E_{\text{extr}} = \sqrt{1-\eta}\sqrt{\eta}E_0 \cos((\omega + \Omega)t + \phi) + \sqrt{\eta}\sqrt{1-\eta}E_0 \cos((\omega - \Omega)t - \phi) \quad (\text{B.3a})$$

$$= 2\sqrt{\eta}\sqrt{1-\eta}E_0 \cos(\omega t) \cos(\Omega t + \phi) \quad (\text{B.3b})$$

Cette configuration donne une intensité extraite proportionnelle à

$$I_{\text{extr}} = \eta(1 - \eta) |E_0|^2 \cos^2(\Omega t + \phi) \quad (\text{B.4})$$

En ajustant correctement la phase de l'onde RF l'extraction est maximale pour une efficacité de diffraction de $\eta = 50\%$. Il faut cependant veiller à ne pas trop extraire afin de laisser une « graine » dans la cavité pour la formation de l'impulsion suivante. Dans le cas contraire cette dernière se reformerait à un moment aléatoire conduisant à un fonctionnement chaotique du laser.

La configuration en double passage possède un autre avantage important : elle permet de supprimer les impulsions entourant celle devant être extraite. Le temps de réponse de la cellule de Bragg est en effet trop long (le temps de montée et le temps de descente sont de l'ordre de 7 ns) par rapport au temps d'aller retour dans la cavité (de l'ordre de la dizaine de nanosecondes) pour que les impulsions précédentes et suivantes ne soient pas aussi déviées. L'astuce réside dans le fait que si ce temps d'aller-retour vaut $\tau_{\text{cav}} = (2n + 1)/4\tau_{\text{RF}}$ où τ_{RF} est la période de l'onde RF, alors l'interférence devient destructive annulant ainsi la partie extraite.

Annexe C

Calculs pour la comparaison des protocoles de communication

Cette annexe présente les modèles des divers protocoles de communications comparés dans la section 4.4.1 avec la production d'intrication à distance. Nous nous contenterons du calcul des fonctions de Wigner ; les matrices densités, nécessaires au calcul de l'intrication, peuvent en être déduites analytiquement à l'aide de l'équation 2.49, certaines expressions peuvent d'ailleurs être trouvées dans [82], ou obtenues numériquement à partir de ces même fonctions de Wigner.

C.1 Envoi direct : état EPR et état de Bell

La fonction de Wigner de l'état EPR avec les imperfections à déjà été vu plusieurs fois, à commencer par le calcul du photon unique. Nous nous contenterons donc d'en rappeler la formule

$$W_{\eta}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 ab} e^{-\frac{(x_1-x_2)^2}{2a} - \frac{(p_1-p_2)^2}{2b} - \frac{(x_1+x_2)^2}{2b} - \frac{(p_1+p_2)^2}{2a}} \quad (\text{C.1})$$

avec

$$a = \eta(hs + h - 1) + 1 - \eta \quad (\text{C.2})$$

$$b = \eta\left(\frac{h}{s} + h - 1\right) + 1 - \eta \quad (\text{C.3})$$

Les pertes des fibres optique utilisées pour la transmission pourront être intégrées à l'efficacité homodyne η .

La fonction de Wigner d'un état de Bell à déjà été calculé dans la section 4.3 (équation 4.34), nous disposons même de sa matrice densité (équation 4.34). Comme pour l'état EPR, la transmission des canaux de communication peut être intégré à l'efficacité homodyne.

C.2 Distillation d'états EPR

La méthode de distillation retenue consiste à soustraire localement un photon dans chacun des deux modes de l'état EPR (figure C.3). La transmission des canaux de communication devra cette fois être intégrée à la fois à l'efficacité homodyne η et à l'efficacité des voies APD μ .

Comme d'habitude nous allons commencer par déplacer les pertes homodynes en sortie de l'OPA. Nous pouvons ainsi réutiliser l'état EPR vu précédemment (équation C.1). Nous prélevons

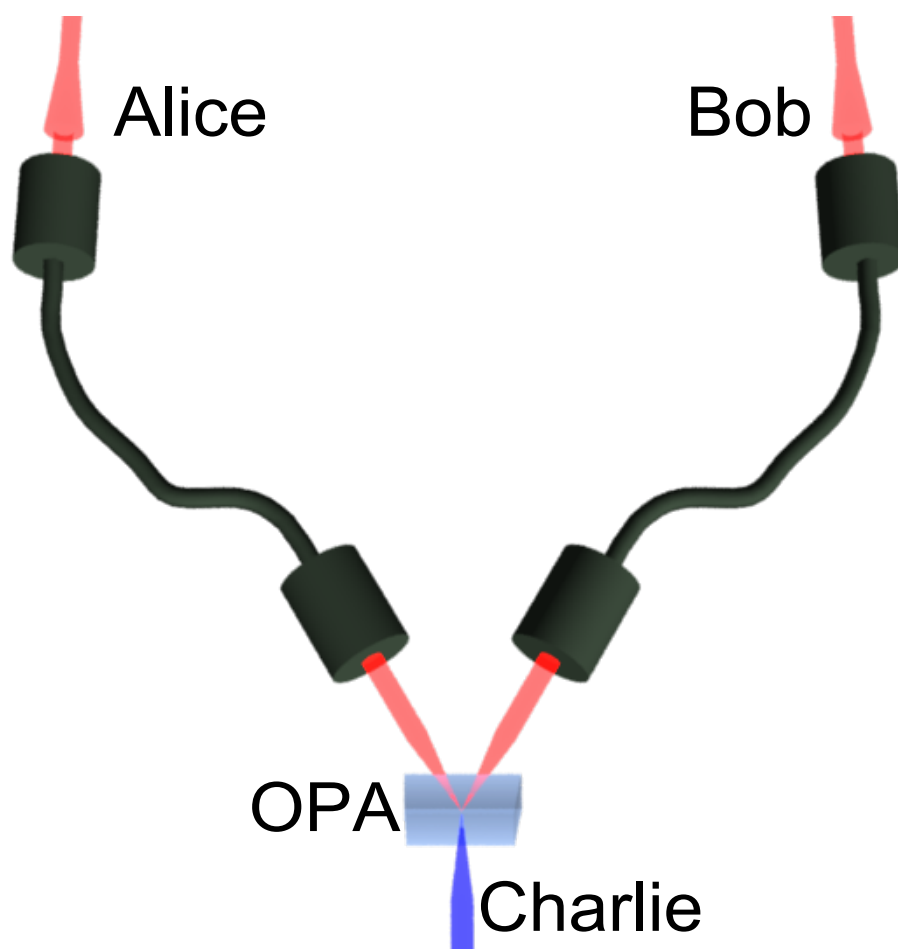


FIGURE C.1: *Protocole de communication par envoie d'un état EPR*

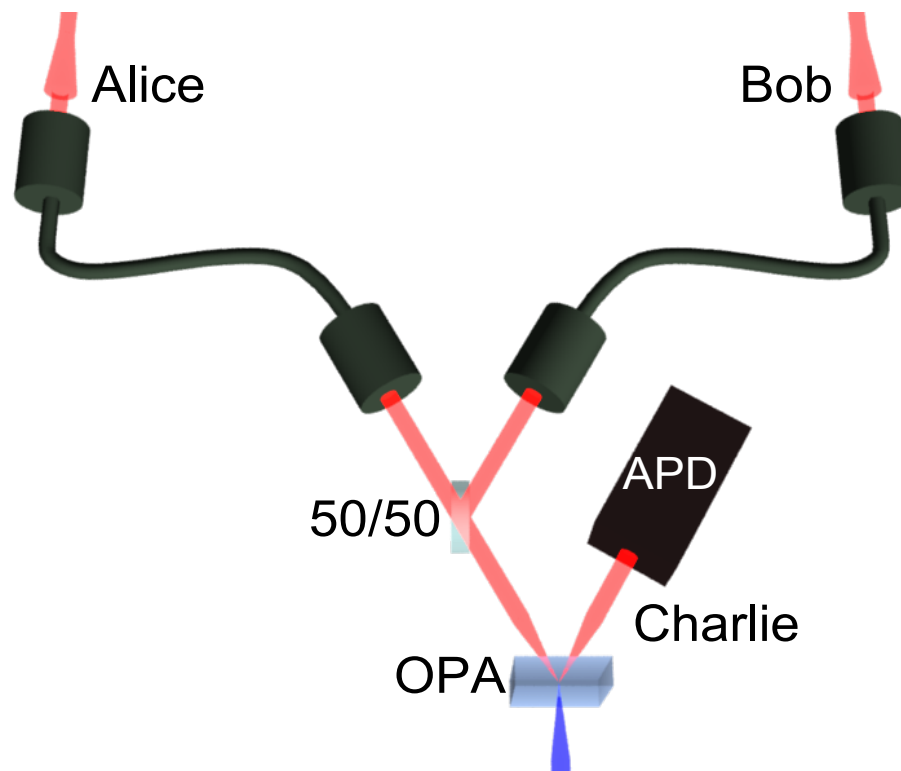


FIGURE C.2: Protocole de communication par envoi d'un état de Bell (discret)

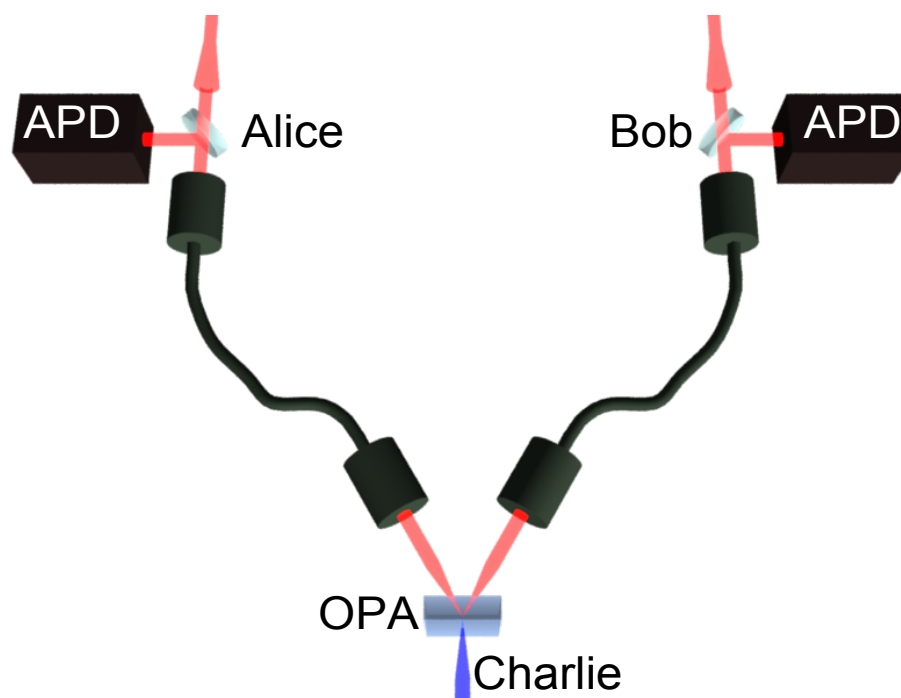


FIGURE C.3: Protocole de communication par distillation d'un état EPR

alors une fraction $R = r^2$ (associé à une transmission $T = t^2$) de chacun des deux modes, conduisant ainsi à un mélange avec deux vides $W_0(x, p)$:

$$W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4) = W_\eta \left(\frac{x_1 + x_3}{\sqrt{2}}, \frac{p_1 + p_3}{\sqrt{2}}, \frac{x_2 + x_4}{\sqrt{2}}, \frac{p_2 + p_4}{\sqrt{2}} \right) \times W_0 \left(\frac{x_1 - x_3}{\sqrt{2}}, \frac{p_1 - p_3}{\sqrt{2}} \right) W_0 \left(\frac{x_2 - x_4}{\sqrt{2}}, \frac{p_2 - p_4}{\sqrt{2}} \right) \quad (\text{C.4})$$

Un bon conditionnement consiste à compter un photon dans chacun des modes réfléchis 3 et 4. Nous garderons l'approximation d'une faible transmission de la voie APD la fonction de Wigner de l'état bien conditionné vaut donc :

$$W_{\text{cond}}(x_1, p_1, x_2, p_2) = \frac{\text{Tr}_{3,4}(W_n(x_3, p_3)W_n(x_4, p_4)W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4))}{\text{Tr}(W_n(x_3, p_3)W_n(x_4, p_4)W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4))} \quad (\text{C.5})$$

Le mauvais conditionnement correspond soit à un prélèvement du photon sur un seul des modes (et des pertes sur l'autre) :

$$W_{\text{n.c.1}}(x_1, p_1, x_2, p_2) = \frac{\text{Tr}_{3,4}(W_n(x_3, p_3)W_{\text{mel}})}{\text{Tr}(W_n(x_3, p_3)W_{\text{mel}})} + \frac{\text{Tr}_{3,4}(W_n(x_4, p_4)W_{\text{mel}})}{\text{Tr}(W_n(x_4, p_4)W_{\text{mel}})} \quad (\text{C.6})$$

soit à un état EPR ayant subi des pertes, correspondant à la transmission T , ce qui donne la fonction de Wigner suivante (cf. section A.2) :

$$W_{\text{n.c.2}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 a' b'} e^{-\frac{(x_1 - x_2)^2}{2a'} - \frac{(p_1 - p_2)^2}{2b'} - \frac{(x_1 + x_2)^2}{2b'} - \frac{(p_1 + p_2)^2}{2a'}} \quad (\text{C.7})$$

avec

$$a' = \eta T(hs + h - 2) + 1 \quad (\text{C.8})$$

$$b' = \eta T\left(\frac{h}{s} + h - 2\right) + 1 \quad (\text{C.9})$$

Nous pouvons alors calculer la fonction de Wigner finale

$$W_{\text{distil}}(x_1, p_1, x_2, p_2) = \xi W_{\text{cond}}(x_1, p_1, x_2, p_2) + \xi(1 - \xi)W_{\text{n.c.1}}(x_1, p_1, x_2, p_2) + (1 - \xi)W_{\text{n.c.2}}(x_1, p_1, x_2, p_2) \quad (\text{C.10})$$

avec plus précisément

$$W_{\text{n.c.1}} = \frac{W_{\text{n.c.2}}}{a' + b' - 2} \left(\left(2 - \frac{1}{a'} - \frac{1}{b'} \right) + \left(1 - \frac{1}{a'} \right)^2 \frac{(x_1 - x_2)^2 + (p_1 + p_2)^2}{2} + \left(1 - \frac{1}{b'} \right)^2 \frac{(x_1 + x_2)^2 + (p_1 - p_2)^2}{2} \right) \quad (\text{C.11})$$

$$W_{\text{cond}} = \frac{W_{\text{n.c.2}}}{(a' - 1)^2 + (b' - 1)^2} \left(\left(1 - \frac{1}{a'} \right)^2 + \left(1 - \frac{1}{b'} \right)^2 + 2 \left(\left(1 - \frac{1}{a'} \right)^3 \frac{(x_1 - x_2)^2 + (p_1 + p_2)^2}{2} + \left(1 - \frac{1}{b'} \right)^3 \frac{(x_1 + x_2)^2 + (p_1 - p_2)^2}{2} \right) + \frac{1}{2} \left(\left(1 - \frac{1}{a'} \right)^2 \frac{(x_1 - x_2)^2 + (p_1 + p_2)^2}{2} + \left(1 - \frac{1}{b'} \right)^2 \frac{(x_1 + x_2)^2 + (p_1 - p_2)^2}{2} \right)^2 + 2 \left(1 - \frac{1}{a'} \right)^2 \left(1 - \frac{1}{b'} \right)^2 \left(\frac{(x_1 - x_2)(x_1 + x_2)}{2} + \frac{(p_1 - p_2)(p_1 + p_2)}{2} \right)^2 \right) \quad (\text{C.12})$$

C.3 Transfert d'intrication

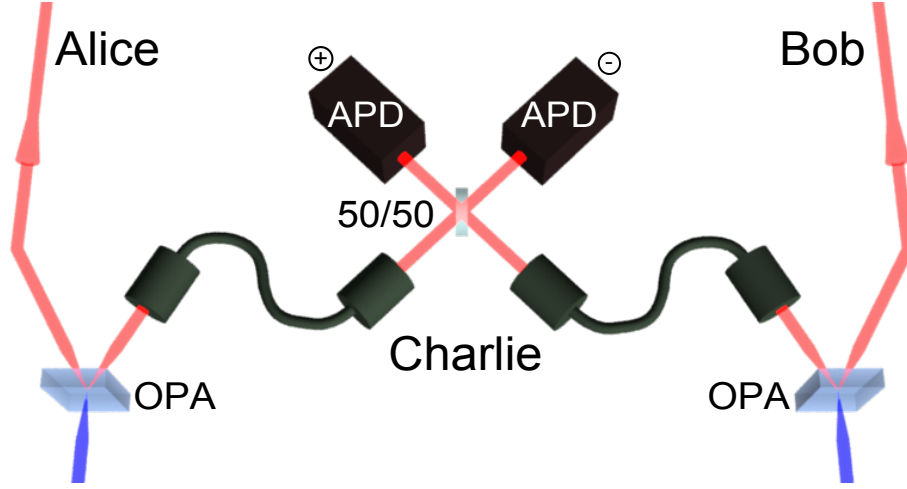


FIGURE C.4: Protocole de communication par transfert d'intrication entre deux états EPR

Pour cette comparaison nous nous intéressons à un protocole de transfert d'intrication entre deux états EPR comprenant peu de photons. Le transfert s'effectue alors en effectuant une mesure conjointe à l'aide d'une APD entre deux modes venant chacun d'une paire EPR différente (figure C.4). La transmission des canaux de communication doit être intégrée à l'efficacité des voies APD μ , elle ne nuit donc qu'au taux de succès, comme pour notre protocole d'intrication à distance.

À nouveau, nous commençons par déplacer les pertes homodynes et partons de la fonction de Wigner de l'état EPR. Cette fois nous aurons cependant deux paires EPR, modélisée chacune par l'équation C.1, dont nous mélangeons deux modes :

$$W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4) = W_{\eta} \left(x_1, p_1, \frac{x_3 - x_4}{\sqrt{2}}, \frac{p_3 - p_4}{\sqrt{2}} \right) W_{\eta} \left(x_2, p_2, \frac{x_3 + x_4}{\sqrt{2}}, \frac{p_3 + p_4}{\sqrt{2}} \right) \quad (\text{C.13})$$

Il ne nous reste plus qu'à faire agir Charlie, c'est-à-dire à modéliser le conditionnement sur le mode 3, suivi d'une trace partielle sur les modes 3 et 4. Nous utilisons toujours l'approximation d'une faible transmission de la voie APD, ce qui nous donne

$$W_{\text{cond}}(x_1, p_1, x_2, p_2) = \frac{\iiint W_n(x_3, p_3) W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4) dx_3 dp_3 dx_4 dp_4}{\text{Tr}(W_n(x_3, p_3) W_{\text{mel}}(x_1, p_1, x_2, p_2, x_3, p_3, x_4, p_4))} \quad (\text{C.14})$$

auquel nous devons ajouter les mauvais conditionnement, qui conduit à un produit de deux états thermique avec une probabilité $1 - \xi$:

$$W_{\text{transf}}(x_1, p_1, x_2, p_2) = \xi W_{\text{cond}}(x_1, p_1, x_2, p_2) + (1 - \xi) W_{\text{th}}(x_1, p_1) W_{\text{th}}(x_2, p_2) \quad (\text{C.15})$$

Au final nous obtenons une fonction de Wigner qui s'exprime en fonction des paramètres δ et σ du photon unique (cf. section 3.5) :

$$W_{\text{transf}}(x_1, p_1, x_2, p_2) = \frac{1}{\pi^2 \sigma^4} \left(1 - \delta + \delta \frac{(x_1 - x_2)^2 + (p_1 - p_2)^2}{2\sigma^2} \right) e^{-\frac{x_1^2 + p_1^2 + x_2^2 + p_2^2}{\sigma^2}} \quad (\text{C.16})$$

Annexe D

Influence de la cadence sur l'amplification femtoseconde

L'équation de Frantz-Nodvick (équation 7.1) qui donne le gain de l'amplificateur fait intervenir le paramètre $G_0 = e^{J_{sto}/J_{sat}}$. Ce dernier peut aussi s'écrire en fonction de la densité de population d'état excité en régime stationnaire N_0 [209] :

$$G_0 = e^{\sigma N_0 L} \quad (D.1)$$

où σ est la section efficace d'émission du cristal amplificateur et L sa longueur. Chaque amplification d'une impulsion va évidemment réduire la population d'états excités ; si la cadence est trop élevée elle n'aura alors pas le temps de retrouver le régime stationnaire avant le passage de l'impulsion suivante. Il convient alors de remplacer N_0 dans l'équation D.1 par une nouvelle densité de population que nous appellerons N_{av} . Nous nommerons de même N_{ap} la densité de population après le passage de l'impulsion.

Rappelons que la densité de population d'état excité est régie par l'équation

$$\frac{dN(t)}{dt} = \frac{N_0 - N(t)}{\tau} \quad (D.2)$$

où τ est le temps de fluorescence de la transition laser et N_0 la densité de population de l'état excité en régime stationnaire. En appelant t_{rep} le temps entre deux impulsions, nous obtenons donc :

$$N_{av} = N_{ap} e^{-\frac{t_{rep}}{\tau}} + N_0 \left(1 - e^{-\frac{t_{rep}}{\tau}}\right) \quad (D.3)$$

En première approximation, nous pouvons considérer que le nombre d'états désexcités est égal au nombre de photons gagnés par l'impulsion, ce qui donne

$$\frac{G J_{in} - J_{in}}{\hbar \omega} = (N_{av} - N_{ap}) L \quad (D.4)$$

On en déduit que

$$N_{av} = N_{ap} + \frac{(G - 1) J_{in}}{\hbar \omega L} \quad (D.5)$$

À partir des équations D.3 et D.5 nous pouvons calculer la densité de population d'états excités juste avant le passage de l'impulsion :

$$N_{av} = \frac{(G - 1) J_{in}}{\hbar \omega L} \frac{1}{1 - e^{-\frac{t_{rep}}{\tau}}} + N_0 \quad (D.6)$$

En posant $\kappa = \frac{1}{1 - e^{t_{\text{rep}}/\tau}}$ nous obtenons donc la nouvelle valeur du paramètre G_0 :

$$G_0 = e^{\sigma N_0 L} e^{\kappa(G-1) \frac{J_{\text{in}}}{J_{\text{sat}}}} \quad (\text{D.7})$$

Ce qui revient à effectuer la transformation

$$G_0 \mapsto G_0 e^{\kappa(G-1) \frac{J_{\text{in}}}{J_{\text{sat}}}} \quad (\text{D.8})$$

Nous avons dit que $J_{\text{in}} \ll J_{\text{sat}}$ (cf. section 7.1.3), ce qui donne $G \approx G_0$. Nous pouvons pousser plus loin l'approximation avec la nouvelle valeur de G_0 :

$$G_0 \approx e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}} \left(1 - \kappa(G-1) \frac{J_{\text{in}}}{J_{\text{sat}}} \right) \quad (\text{D.9})$$

Ce qui donne

$$\frac{J_{\text{sat}} + \kappa e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}} J_{\text{in}}}{J_{\text{sat}}} G \approx \frac{J_{\text{sat}} + \kappa e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}} J_{\text{in}}}{J_{\text{sat}}} e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}} \quad (\text{D.10})$$

ou

$$G \approx \eta_G e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}} \quad (\text{D.11})$$

avec

$$\eta_G = \frac{J_{\text{sat}} - \kappa J_{\text{in}}}{J_{\text{sat}} - \kappa J_{\text{in}} e^{\frac{J_{\text{sto}}}{J_{\text{sat}}}}} \quad (\text{D.12})$$

Bibliographie

- [1] Claude Shannon : A Mathematical Theory of Communication, *Bell System Technical Journal*, **27**, pp. 379–423pp. 623–656 (1948).
- [2] Peter W. Shor : Polynomial-Time Algorithms for Prime Factorization and Discrete Logarithms on a Quantum Computer, *SIAM J. Comput.*, **26**, pp. 1484–1509 (1997). DOI: 10.1137/S0097539795293172.
- [3] Lov K. Grover : A fast quantum mechanical algorithm for database search, *Proceedings of the twenty-eighth annual ACM symposium on Theory of computing*, pp. 212–219 (1996). DOI: <http://doi.acm.org/10.1145/237814.237866>.
- [4] Gilbert Grynberg, Alain Aspect et Claude Fabre : *Introduction aux lasers et a l'optique quantique*, Ellipses (1997).
- [5] Ulf Leonhardt : *Measuring the Quantum State of Light*, Cambridge University Press (1997).
- [6] Hans-A. Bachor et Timothy C. Ralph : *A Guide To Experiments In Quantum Optics*, Wiley (2004).
- [7] D. F. Walls et Gerard J. Milburn : *Quantum Optics*, Springer (1995).
- [8] Wolfgang Schleich : *Quantum optics in phase space*, Wiley (2001).
- [9] Claude Cohen-Tannoudji, Jacques Dupont-Roc et Gilbert Grynberg : *Processus d'interaction entre photons et atomes*, EDP Sciences (1996).
- [10] Gilbert Grynberg, Alain Aspect et Claude Fabre : *Introduction aux lasers et a l'optique quantique*, Ellipses (1997).
- [11] Max Planck : Über das Gezetz des Energieverteilung im Normalspektrum, *Annalen der Physik*, **4**, p. 553 (1901).
- [12] A. B. Arons et M. B. Peppard : Einstein's Proposal of the Photon Concept—a Translation of the Annalen der Physik Paper of 1905, *American Journal of Physics*, **33**(5), pp. 367–374 (1965). DOI: 10.1119/1.1971542.
- [13] D. F. Phillips, A. Fleischhauer, A. Mair, R. L. Walsworth et M. D. Lukin : Storage of Light in Atomic Vapor, *Phys. Rev. Lett.*, **86**(5), pp. 783–786 (2001). DOI: 10.1103/PhysRevLett.86.783.
- [14] M. D. Lukin et A. Imamoglu : Nonlinear Optics and Quantum Entanglement of Ultraslow Single Photons, *Phys. Rev. Lett.*, **84**(7), pp. 1419–1422 (2000). DOI: 10.1103/PhysRevLett.84.1419.
- [15] A. Rauschenbeutel, P. Bertet, S. Osnaghi, G. Nogues, M. Brune, J. M. Raimond et S. Haroche : Controlled entanglement of two field modes in a cavity quantum electrodynamics experiment, *Phys. Rev. A*, **64**(5), p. 050301 (2001). DOI: 10.1103/PhysRevA.64.050301.

- [16] Q. A. Turchette, C. J. Hood, W. Lange, H. Mabuchi et H. J. Kimble : Measurement of Conditional Phase Shifts for Quantum Logic, *Phys. Rev. Lett.*, **75**(25), pp. 4710–4713 (1995). DOI: 10.1103/PhysRevLett.75.4710.
- [17] E. Knill, R. Laflamme et G. J. Milburn : A scheme for efficient quantum computation with linear optics, *Nature Physics*, **409**, pp. 46–52 (2001). DOI: 10.1038/35051009.
- [18] Jeremy L. O’Brien : Optical Quantum Computing, *Science*, **318**, pp. 1567–1570 (2007). DOI: 10.1126/science.1142892.
- [19] J. L. O’Brien, G. J. Pryde, A. G. White, T. C. Ralph et D. Branning : Demonstration of an all-optical quantum controlled-NOT gate, *Nature Physics*, **426**, pp. 264–267 (2003). DOI: 10.1038/nature02054.
- [20] C. H. Bennett et G. Brassard : Quantum cryptography : public-key distribution and coin tossing, *Proceedings of the IEEE International Conference on Computers, Systems and Signal Processing, Bangalore, India*, p. 175 (1984).
- [21] B. P. Lanyon, T. J. Weinhold, N. K. Langford, J. L. O’Brien, K. J. Resch, A. Gilchrist et A. G. White : Manipulating Biphotonic Qutrits, *Phys. Rev. Lett.*, **100**(6), p. 060504 (2008). DOI: 10.1103/PhysRevLett.100.060504.
- [22] N. K. Langford, R. B. Dalton, M. D. Harvey, J. L. O’Brien, G. J. Pryde, A. Gilchrist, S. D. Bartlett et A. G. White : Measuring Entangled Qutrits and Their Use for Quantum Bit Commitment, *Phys. Rev. Lett.*, **93**(5), p. 053601 (2004). DOI: 10.1103/PhysRevLett.93.053601.
- [23] Thomas Durt, Nicolas J. Cerf, Nicolas Gisin et Marek Żukowski : Security of quantum key distribution with entangled qutrits, *Phys. Rev. A*, **67**(1), p. 012311 (2003). DOI: 10.1103/PhysRevA.67.012311.
- [24] T. C. Ralph, K. J. Resch et A. Gilchrist : Efficient Toffoli gates using qudits, *Phys. Rev. A*, **75**(2), p. 022313 (2007). DOI: 10.1103/PhysRevA.75.022313.
- [25] Nicolas J. Cerf, Mohamed Bourennane, Anders Karlsson et Nicolas Gisin : Security of Quantum Key Distribution Using d -Level Systems, *Phys. Rev. Lett.*, **88**(12), p. 127902 (2002). DOI: 10.1103/PhysRevLett.88.127902.
- [26] J. McKeever, A. Boca, A. D. and Miller, R. Boozer, J. R. Buck, A. Kuzmich et H. J. Kimble : Deterministic Generation of Single Photons from One Atom Trapped in a Cavity, *Science*, **303**, pp. 1992–1994 (2004). DOI: 10.1126/science.1095232.
- [27] B. Darquié, M. P. A. Jones, J. Dingjan, J. Beugnon, S. Bergamini, Y. Sortais, G. Messin, A. Browaeys et P. Grangier : Controlled Single-Photon Emission from a Single Trapped Two-Level Atom, *Science*, **309**, pp. 454–456 (2005). DOI: 10.1126/science.1113394.
- [28] Markus Hilkema, Bernhard Weber, Specht Holger P., Simon C. Webster, Axel Kuhn et Gerhard Rempe : A single-photon server with just one atom, *Nature Physics*, **3**, pp. 253–255 (2007). DOI: 10.1038/nphys569.
- [29] Rosa Brouri, Alexios Beveratos, Jean-Philippe Poizat et Philippe Grangier : Photon antibunching in the fluorescence of individual color centers in diamond, *Opt. Lett.*, **25**(17), pp. 1294–1296 (2000). DOI: 10.1364/OL.25.001294.
- [30] Christian Kurtsiefer, Sonja Mayer, Patrick Zarda et Harald Weinfurter : Stable Solid-State Source of Single Photons, *Phys. Rev. Lett.*, **85**(2), pp. 290–293 (2000). DOI: 10.1103/PhysRevLett.85.290.

- [31] P. Michler, A. Kiraz, C. Becher, W. V. Schoenfeld, P. M. Petroff, Lidong Zhang, E. Hu et A. Imamoglu : A Quantum Dot Single-Photon Turnstile Device, *Science*, **290**(5500), pp. 2282–2285 (2000). DOI: 10.1126/science.290.5500.2282.
- [32] Charles Santori, Matthew Pelton, Glenn Solomon, Yseulte Dale et Yoshihisa Yamamoto : Triggered Single Photons from a Quantum Dot, *Phys. Rev. Lett.*, **86**(8), pp. 1502–1505 (2001). DOI: 10.1103/PhysRevLett.86.1502.
- [33] E. Arthurs et M. S. Goodman : Quantum Correlations : A Generalized Heisenberg Uncertainty Relation, *Phys. Rev. Lett.*, **60**(24), pp. 2447–2449 (1988). DOI: 10.1103/PhysRevLett.60.2447.
- [34] Serge Reynaud, Astrid Lambrecht, Cyriaque Genet et Marc-Thierry Jaekel : Quantum vacuum fluctuations, *C. R. Acad. Sci.*, **2**(IV), pp. 1287–1298 (2001). <http://arxiv.org/abs/quant-ph/0105053>.
- [35] Frédéric Grosshans, Gilles Van Assche, Jérôme Wenger, Rosa Brouri, Nicolas J. Cerf et Philippe Grangier : Quantum key distribution using gaussian-modulated coherent states, *Nature*, **421**, pp. 238–241 (2003). DOI: 10.1038/nature01289.
- [36] Jérôme Lodewyck, Matthieu Bloch, Raúl García-Patrón, Simon Fossier, Evgueni Karpov, Eleni Diamanti, Thierry Debuisschert, Nicolas J. Cerf, Rosa Tualle-Brouri, Steven W. McLaughlin et Philippe Grangier : Quantum key distribution over 25km with an all-fiber continuous-variable system, *Phys. Rev. A*, **76**(4), p. 042305 (2007). DOI: 10.1103/PhysRevA.76.042305.
- [37] Seth Lloyd et Samuel L. Braunstein : Quantum Computation over Continuous Variables, *Phys. Rev. Lett.*, **82**(8), pp. 1784–1787 (1999). DOI: 10.1103/PhysRevLett.82.1784.
- [38] Samuel L. Braunstein et Peter van Loock : Quantum information with continuous variables, *Rev. Mod. Phys.*, **77**(2), pp. 513–577 (2005). DOI: 10.1103/RevModPhys.77.513.
- [39] T. C. Ralph, A. Gilchrist, G. J. Milburn, W. J. Munro et S. Glancy : Quantum computation with optical coherent states, *Phys. Rev. A*, **68**(4), p. 042319 (2003). DOI: 10.1103/PhysRevA.68.042319.
- [40] E. Wigner : On the Quantum Correction For Thermodynamic Equilibrium, *Phys. Rev.*, **40**(5), pp. 749–759 (1932). DOI: 10.1103/PhysRev.40.749.
- [41] Ulf Leonhardt : *Measuring the Quantum State of Light*, Cambridge University Press (1997).
- [42] Claude Cohen-Tannoudji : Cours VII - Opérateurs densité d’une particule quantique - Représentation de Wigner, *Cours au Collège de France* (1983-1984). <http://www.phys.ens.fr/cours/college-de-france/1983-84/cours7/cours7.pdf>.
- [43] J.E. Moyal : Quantum mechanics as a statistical theory., *Proc. Camb. Philos. Soc.*, **45**, pp. 99–124 (1949).
- [44] A. I. Lvovsky et M. G. Raymer : Continuous-variable optical quantum-state tomography, *Rev. Mod. Phys.*, **81**(1), pp. 299–332 (2009). DOI: 10.1103/RevModPhys.81.299.
- [45] Jaromír Fiurášek et Zdeněk Hradil : Maximum-likelihood estimation of quantum processes, *Phys. Rev. A*, **63**(2), p. 020101 (2001). DOI: 10.1103/PhysRevA.63.020101.
- [46] J. Řeháček, Z. Hradil et M. Ježek : Iterative algorithm for reconstruction of entangled states, *Phys. Rev. A*, **63**(4), p. 040303 (2001). DOI: 10.1103/PhysRevA.63.040303.
- [47] A. I. Lvovsky : Iterative maximum-likelihood reconstruction in quantum homodyne tomography, *J. Opt. B*, **6**(6), p. S556 (2004). DOI: 10.1088/1464-4266/6/6/014.

- [48] A. Einstein, B. Podolsky et N. Rosen : Can Quantum-Mechanical Description of Physical Reality Be Considered Complete?, *Phys. Rev.*, **47**(10), pp. 777–780 (1935). DOI: 10.1103/PhysRev.47.777.
- [49] W. P. Bowen, R. Schnabel, P. K. Lam et T. C. Ralph : Experimental Investigation of Criteria for Continuous Variable Entanglement, *Phys. Rev. Lett.*, **90**(4), p. 043601 (2003). DOI: 10.1103/PhysRevLett.90.043601.
- [50] W. P. Bowen, R. Schnabel, P. K. Lam et T. C. Ralph : Experimental characterization of continuous-variable entanglement, *Phys. Rev. A*, **69**(1), p. 012304 (2004). DOI: 10.1103/PhysRevA.69.012304.
- [51] Jun Mizuno, Kentaro Wakui, Akira Furusawa et Masahide Sasaki : Experimental demonstration of entanglement-assisted coding using a two-mode squeezed vacuum state, *Phys. Rev. A*, **71**(1), p. 012304 (2005). DOI: 10.1103/PhysRevA.71.012304.
- [52] A. Furusawa, J. L. Sørensen, S. L. Braunstein, C. A. and, Kimble, H. J. Fuchs et E. S. Polzik : Unconditional Quantum Teleportation, *Science*, **282**, pp. 706–709 (1998). DOI: 10.1126/science.282.5389.706.
- [53] Z. Y. Ou, S. F. Pereira, H. J. Kimble et K. C. Peng : Realization of the Einstein-Podolsky-Rosen paradox for continuous variables, *Phys. Rev. Lett.*, **68**(25), pp. 3663–3666 (1992). DOI: 10.1103/PhysRevLett.68.3663.
- [54] Z. Y. Ou, S. F. Pereira et H. J. Kimble : Realization of the Einstein-Podolsky-Rosen paradox for continuous variables in nondegenerate parametric amplification, *Appl. Phys. B*, **55**(3), pp. 265–278 (1992). DOI: 10.1007/BF00325015.
- [55] Yun Zhang, Hai Wang, Xiaoying Li, Jietai Jing, Changde Xie et Kunchi Peng : Experimental generation of bright two-mode quadrature squeezed light from a narrow-band nondegenerate optical parametric amplifier, *Phys. Rev. A*, **62**(2), p. 023813 (2000). DOI: 10.1103/PhysRevA.62.023813.
- [56] Christian Schori, Jens L. Sørensen et Eugene S. Polzik : Narrow-band frequency tunable light source of continuous quadrature entanglement, *Phys. Rev. A*, **66**(3), p. 033802 (2002). DOI: 10.1103/PhysRevA.66.033802.
- [57] J. Laurat, T. Coudreau, G. Keller, N. Treps et C. Fabre : Compact source of Einstein-Podolsky-Rosen entanglement and squeezing at very low noise frequencies, *Phys. Rev. A*, **70**(4), p. 042315 (2004). DOI: 10.1103/PhysRevA.70.042315.
- [58] R. L. Hudson : When is the Wigner quasi-probability density non-negative?, *Rep. Math. Phys.*, **6**, p. 249 (1974).
- [59] J. Eisert, S. Scheel et M. B. Plenio : Distilling Gaussian States with Gaussian Operations is Impossible, *Phys. Rev. Lett.*, **89**(13), p. 137903 (2002). DOI: 10.1103/PhysRevLett.89.137903.
- [60] Jaromír Fiurášek : Gaussian Transformations and Distillation of Entangled Gaussian States, *Phys. Rev. Lett.*, **89**(13), p. 137904 (2002). DOI: 10.1103/PhysRevLett.89.137904.
- [61] Géza Giedke et J. Ignacio Cirac : Characterization of Gaussian operations and distillation of Gaussian states, *Phys. Rev. A*, **66**(3), p. 032316 (2002). DOI: 10.1103/PhysRevA.66.032316.
- [62] Lu-Ming Duan, G. Giedke, J. I. Cirac et P. Zoller : Entanglement Purification of Gaussian Continuous Variable Quantum States, *Phys. Rev. Lett.*, **84**(17), pp. 4002–4005 (2000). DOI: 10.1103/PhysRevLett.84.4002.

- [63] Stephen D. Bartlett, Barry C. Sanders, Samuel L. Braunstein et Kae Nemoto : Efficient Classical Simulation of Continuous Variable Quantum Information Processes, *Phys. Rev. Lett.*, **88**(9), p. 097904 (2002). DOI: 10.1103/PhysRevLett.88.097904.
- [64] A Gilchrist, Kae Nemoto, W J Munro, T C Ralph, S Glancy, Samuel L Braunstein et G J pp Milburn : Schrödinger cats and their power for quantum information processing, *J. Opt. B*, **6**(8), pp. S828–S833 (2004). DOI: 10.1088/1464-4266/6/8/032.
- [65] Julien Niset, Jaromír Fiurášek et Nicolas J. Cerf : No-Go Theorem for Gaussian Quantum Error Correction, *Phys. Rev. Lett.*, **102**(12), p. 120501 (2009). DOI: 10.1103/PhysRevLett.102.120501.
- [66] J.A. Wheeler et W.H. Zurek : *Quantum Theory and Measurement*, Princeton university Press (1983).
- [67] Z. Y. Ou, S. F. Pereira et H. J. Kimble : Realization of the Einstein-Podolsky-Rosen paradox for continuous variables in nondegenerate parametric amplification, *Appl. Phys. B*, **55**(3), pp. 265–278 (1992). DOI: 10.1007/BF00325015.
- [68] Yuishi Takeno, Mitsuyoshi Yukawa, Hidehiro Yonezawa et Akira Furusawa : Observation of -9 dB quadrature squeezing with improvement of phase stability in homodyne measurement, *Optics Express*, **15**(7), pp. 4321–4327 (2007).
- [69] Tobias Eberle, Sebastian Steinlechner, Jöran Bauchrowitz, Vitus Händchen, Henning Vahlbruch, Moritz Mehmet, Helge Müller-Ebhardt et Roman Schnabel : Quantum Enhancement of the Zero-Area Sagnac Interferometer Topology for Gravitational Wave Detection, *Phys. Rev. Lett.*, **104**(25), p. 251102 (2010). DOI: 10.1103/PhysRevLett.104.251102.
- [70] Kentaro Wakui, Hiroki Takahashi, Akira Furusawa et Masahide Sasaki : Photon subtracted squeezed states generated with periodically poled KTiOPO₄, *Optics Express*, **15**(6), pp. 3568–3574 (2007).
- [71] J. S. Neergaard-Nielsen, B. Melholt Nielsen, C. Hettich, K. Mølmer et E. S. Polzik : Generation of a Superposition of Odd Photon Number States for Quantum Information Networks, *Phys. Rev. Lett.*, **97**(8), p. 083604 (2006). DOI: 10.1103/PhysRevLett.97.083604.
- [72] B. Yurke et D. Stoler : Generating quantum mechanical superpositions of macroscopically distinguishable states via amplitude dispersion, *Phys. Rev. Lett.*, **57**(1), pp. 13–16 (1986). DOI: 10.1103/PhysRevLett.57.13.
- [73] Matteo G. A. Paris, Mary Cola et Rodolfo Bonifacio : Quantum-state engineering assisted by entanglement, *Phys. Rev. A*, **67**(4), p. 042104 (2003). DOI: 10.1103/PhysRevA.67.042104.
- [74] A. P. Lund, H. Jeong, T. C. Ralph et M. S. Kim : Conditional production of superpositions of coherent states with inefficient photon detection, *Phys. Rev. A*, **70**(2), p. 020101 (2004). DOI: 10.1103/PhysRevA.70.020101.
- [75] Siegfried Schneider, A. Stockmann et Wolfgang Schuesslbauer : Self-starting mode-locked cavity-dumped femtosecond Ti :sapphire laser employing a semiconductor saturable absorber mirror, *Opt. Express*, **6**(11), pp. 220–226 (2000). <http://www.opticsexpress.org/abstract.cfm?URI=oe-6-11-220>.
- [76] Beat Zysset, Ivan Biaggio et Peter Günter : Refractive indices of orthorhombic KNbO₃. I. Dispersion and temperature dependence, *J. Opt. Soc. Am. B*, **9**(3), pp. 380–386 (1992). <http://josab.osa.org/abstract.cfm?URI=josab-9-3-380>.
- [77] Y.R. Shen : *Principles of nonlinear optics*, Wiley Classics Library (1983).

- [78] Richard L. Sutherland : *Handbook of nonlinear optics*, Marcel Dekker Inc. (1995).
- [79] Roger J. Reeves, Mahendra G. Jani, Bahaeddin Jassemnejad, Richard C. Powell, Greg J. Mizell et William Fay : Photorefractive properties of $KNbO_3$, *Phys. Rev. B*, **43**(1), pp. 71–82 (1991). DOI: 10.1103/PhysRevB.43.71.
- [80] H. Mabuchi, E. S. Polzik et H. J. Kimble : Blue-light-induced infrared absorption in $KNbO_3$, *J. Opt. Soc. Am. B*, **11**(10), pp. 2023–2029 (1994). <http://josab.osa.org/abstract.cfm?URI=josab-11-10-2023>.
- [81] Jérôme Wenger : *Dispositifs impulsionsnels pour la communication quantique à variables continues*. Thèse de doctorat, Université Paris XI Orsay (2004). <http://tel.archives-ouvertes.fr/docs/00/04/71/18/PDF/tel-00006926.pdf>.
- [82] Alexei Ourjountsev : *Étude théorique et expérimentale de superpositions quantiques cohérentes et d'états intriqués non-gaussiens de la lumière*. Thèse de doctorat, Université Paris XI Orsay (2007). http://tel.archives-ouvertes.fr/docs/00/29/22/34/PDF/TheseA0_Main.pdf.
- [83] Kahraman G. Köprülü et Orhan Aytür : Analysis of Gaussian-beam degenerate optical parametric amplifiers for the generation of quadrature-squeezed states, *Phys. Rev. A*, **60**(5), pp. 4122–4134 (1999). DOI: 10.1103/PhysRevA.60.4122.
- [84] Jérôme Wenger, Rosa Tualle-Brouri et Philippe Grangier : Pulsed homodyne measurements of femtosecond squeezed pulses generated by single-pass parametric deamplification, *Opt. Lett.*, **29**(11), pp. 1267–1269 (2004). <http://ol.osa.org/abstract.cfm?URI=ol-29-11-1267>.
- [85] Arthur La Porta et Richard E. Slusher : Squeezing limits at high parametric gains, *Phys. Rev. A*, **44**(3), pp. 2013–2022 (1991). DOI: 10.1103/PhysRevA.44.2013.
- [86] Rosa Tualle-Brouri, Alexei Ourjountsev, Aurelien Dantan, Philippe Grangier, Martijn Wubs et Anders S. Sørensen : Multimode model for projective photon-counting measurements, *Phys. Rev. A*, **80**(1), p. 013806 (2009). DOI: 10.1103/PhysRevA.80.013806.
- [87] Horace P. Yuen et Vincent W. S. Chan : Noise in homodyne and heterodyne detection, *Opt. Lett.*, **8**(3), pp. 177–179 (1983). <http://ol.osa.org/abstract.cfm?URI=ol-8-3-177>.
- [88] R. E. Slusher, P. Grangier, A. LaPorta, B. Yurke et M. J. Potasek : Pulsed Squeezed Light, *Phys. Rev. Lett.*, **59**(22), pp. 2566–2569 (1987). DOI: 10.1103/PhysRevLett.59.2566.
- [89] F. Grosshans et P. Grangier : Effective quantum efficiency in the pulsed homodyne detection of a n-photon state, *Eur. Phys. J. D*, **14**(1), pp. 119–125 (2001). DOI: 10.1007/s100530170243.
- [90] Jürgen Appel, Dallas Hoffman, Eden Figueroa et A. I. Lvovsky : Electronic noise in optical homodyne tomography, *Phys. Rev. A*, **75**(3), p. 035802 (2007). DOI: 10.1103/PhysRevA.75.035802.
- [91] Alexei Ourjountsev, Rose Tualle-Brouri et Philippe Grangier : Quantum Homodyne Tomography of a Two-Photon Fock State, *Phys. Rev. Lett.*, **96**(21), p. 213601 (2006). <http://arxiv.org/abs/quant-ph/0603284>.
- [92] Alexei Ourjountsev, Hyunseok Jeong, Rose Tualle-Brouri et Philippe Grangier : Generation of optical "Schrödinger cats" from photon number states, *Nature*, **448**, pp. 784–786 (2007). DOI: doi:10.1038/nature06054.
- [93] Jérôme Wenger, Rose Tualle-Brouri et Philippe Grangier : Non-Gaussian Statistics from Individual Pulses of Squeezed Light, *Phys. Rev. Lett.*, **Apr**(15), p. 153601 (2004). <http://arxiv.org/abs/quant-ph/0403234>.

- [94] Alexei Ourjoumtsev, Rose Tualle-Brouri, Julien Laurat et Philippe Grangier : Generating Optical Schrödinger Kittens for Quantum Information Processing, *Science*, **312**, pp. 83–86 (2006).
- [95] Thomas Gerrits, Scott Glancy, Tracy S. Clement, Brice Calkins, Adriana E. Lita, Aaron J. Miller, Alan L. Migdall, Sae Woo Nam, Richard P. Mirin et Emanuel Knill : Generation of optical coherent-state superpositions by number-resolved photon subtraction from the squeezed vacuum, *Phys. Rev. A*, **82**(3), p. 031802 (2010). DOI: 10.1103/PhysRevA.82.031802.
- [96] A. I. Lvovsky, H. Hansen, T. Aichele, O. Benson, J. Mlynek et S. Schiller : Quantum State Reconstruction of the Single-Photon Fock State, *Phys. Rev. Lett.*, **87**(5), p. 050402 (2001). DOI: 10.1103/PhysRevLett.87.050402.
- [97] Alessandro Zavatta, Silvia Viciani et Marco Bellini : Tomographic reconstruction of the single-photon Fock state by high-frequency homodyne detection, *Phys. Rev. A*, **70**(5), p. 053821 (2004). DOI: 10.1103/PhysRevA.70.053821.
- [98] S. R. Huisman, Nitin Jain, S. A. Babichev, Frank Vewinger, A. N. Zhang, S. H. Youn et A. I. Lvovsky : Instant single-photon Fock state tomography, *Opt. Lett.*, **34**(18), pp. 2739–2741 (2009). DOI: 10.1364/OL.34.002739.
- [99] Jérôme Wenger, Rose Tualle-Brouri et Philippe Grangier : Pulsed homodyne measurements of femtosecond squeezed pulses generated by single-pass parametric deamplification, *Optics Letters*, **29**(11), (2004). <http://arxiv.org/abs/quant-ph/0402193>.
- [100] Jérôme Wenger, Ourjoumtsev Alexei, Rose Tualle-Brouri et Philippe Grangier : Time-resolved homodyne characterization of individual quadrature-entangled pulses, *Eur. Phys. J. D*, **32**, pp. 391–396 (2005). <http://arxiv.org/abs/quant-ph/0409211>.
- [101] Jérôme Wenger, Rose Tualle-Brouri et Philippe Grangier : Non-Gaussian Statistics from Individual Pulses of Squeezed Light, *Phys. Rev. Lett.*, **Apr**(15), p. 153601 (2004). <http://arxiv.org/abs/quant-ph/0403234>.
- [102] Alexei Ourjoumtsev, Aurelien Dantan, Rose Tualle-Brouri et Philippe Grangier : Increasing Entanglement between Gaussian States by Coherent Photon Subtraction, *Phys. Rev. Lett.*, **98**(03), p. 030502 (2007). <http://arxiv.org/abs/quant-ph/0608230>.
- [103] Anthony Leverrier : *Etude théorique de la distribution quantique de clés à variables continues*. Thèse de doctorat, École Nationale Supérieure des Télécommunications (2009). <http://tel.archives-ouvertes.fr/tel-00451021/fr/>.
- [104] Charles H. Bennett, Gilles Brassard, Claude Crépeau, Richard Jozsa, Asher Peres et William K. Wootters : Teleporting an unknown quantum state via dual classical and Einstein-Podolsky-Rosen channels, *Phys. Rev. Lett.*, **70**(13), pp. 1895–1899 (1993). DOI: 10.1103/PhysRevLett.70.1895.
- [105] Dik Bouwmeester, Jian-Wei Pan, Klaus Mattle, Manfred Eibl, Harald Weinfurter et Anton Zeilinger : Experimental quantum teleportation, *Nature*, **390**, p. 575 (1997). DOI: 10.1038/37539.
- [106] D. Boschi, S. Branca, F. De Martini, L. Hardy et S. Popescu : Experimental Realization of Teleporting an Unknown Pure Quantum State via Dual Classical and Einstein-Podolsky-Rosen Channels, *Phys. Rev. Lett.*, **80**(6), pp. 1121–1125 (1998). DOI: 10.1103/PhysRevLett.80.1121.
- [107] M. Żukowski, A. Zeilinger, M. A. Horne et A. K. Ekert : “Event-ready-detectors” Bell experiment via entanglement swapping, *Phys. Rev. Lett.*, **71**(26), pp. 4287–4290 (1993). DOI: 10.1103/PhysRevLett.71.4287.

- [108] Jian-Wei Pan, Dik Bouwmeester, Harald Weinfurter et Anton Zeilinger : Experimental Entanglement Swapping : Entangling Photons That Never Interacted, *Phys. Rev. Lett.*, **80**(18), pp. 3891–3894 (1998). DOI: 10.1103/PhysRevLett.80.3891.
- [109] H.-J. Briegel, W. Dür, J. I. Cirac et P. Zoller : Quantum Repeaters : The Role of Imperfect Local Operations in Quantum Communication, *Phys. Rev. Lett.*, **81**(26), pp. 5932–5935 (1998). DOI: 10.1103/PhysRevLett.81.5932.
- [110] Lu-Ming Duan, Mikhail Lukin, Ignacio Cirac et Peter Zoller : Long-distance quantum communication with atomic ensembles and linear optics, *Nature*, **414**(6862), pp. 413–418 (2001). DOI: 10.1038/35106500.
- [111] Paul G. Kwiat, Salvador Barraza-Lopez, Andre Stefanov et Nicolas Gisin : Experimental entanglement distillation and "hidden" non-locality, *Nature*, **409**(6823), pp. 1014–1017 (2001). DOI: 10.1038/35059017.
- [112] Jian-Wei ; Gasparoni, Sara ; Ursin, Rupert ; Weihs, Gregor ; Zeilinger, Anton Pan : Experimental entanglement purification of arbitrary unknown states, *Nature*, **423**(6938), pp. 417–422 (2003). DOI: 10.1038/nature01623.
- [113] P. Walther, K. J. Resch, Č. Brukner, A. M. Steinberg, J.-W. Pan et A. Zeilinger : Quantum Nonlocality Obtained from Local States by Entanglement Purification, *Phys. Rev. Lett.*, **94**(4), p. 040504 (2005). DOI: 10.1103/PhysRevLett.94.040504.
- [114] Charles H. Bennett, Gilles Brassard, Sandu Popescu, Benjamin Schumacher, John A. Smolin et William K. Wootters : Purification of Noisy Entanglement and Faithful Teleportation via Noisy Channels, *Phys. Rev. Lett.*, **76**(5), pp. 722–725 (1996). DOI: 10.1103/PhysRevLett.76.722.
- [115] Jaromír Fiurášek : Distillation and purification of symmetric entangled Gaussian states, *Phys. Rev. A*, **82**(4), p. 042331 (2010). DOI: 10.1103/PhysRevA.82.042331.
- [116] C. W. Chou, H. de Riedmatten, D. Felinto, S. V. Polyakov, S. J. van Enk et H. J. Kimble : Measurement-induced entanglement for excitation stored in remote atomic ensembles, *Nature*, **438**(7069), pp. 828–832 (2005). DOI: 10.1038/nature04353.
- [117] T. Chaneliere, D. N. Matsukevich, S. D. Jenkins, S.-Y. Lan, T. A. B. Kennedy et A. Kuzmich : Storage and retrieval of single photons transmitted between remote quantum memories, *Nature*, **438**(7069), pp. 833–836 (2005). DOI: 10.1038/nature04315.
- [118] M. D. Eisaman, A. Andre, F. Massou, M. Fleischhauer, A. S. Zibrov et M. D. Lukin : Electromagnetically induced transparency with tunable single-photon pulses, *Nature*, **438**(7069), pp. 837–841 (2005). DOI: 10.1038/nature04327.
- [119] Chin-Wen Chou, Julien Laurat, Hui Deng, Kyung Soo Choi, Hugues de Riedmatten, Daniel Felinto et H. Jeff Kimble : Functional Quantum Nodes for Entanglement Distribution over Scalable Quantum Networks, *Science*, **316**(5829), pp. 1316–1320 (2007). DOI: 10.1126/science.1140300.
- [120] Robert Raussendorf et Hans J. Briegel : A One-Way Quantum Computer, *Phys. Rev. Lett.*, **86**(22), pp. 5188–5191 (2001). DOI: 10.1103/PhysRevLett.86.5188.
- [121] M. D. Reid : Demonstration of the Einstein-Podolsky-Rosen paradox using nondegenerate parametric amplification, *Phys. Rev. A*, **40**(2), pp. 913–923 (1989). DOI: 10.1103/PhysRevA.40.913.
- [122] Z. Y. Ou, S. F. Pereira, H. J. Kimble et K. C. Peng : Realization of the Einstein-Podolsky-Rosen paradox for continuous variables, *Phys. Rev. Lett.*, **68**(25), pp. 3663–3666 (1992). DOI: 10.1103/PhysRevLett.68.3663.

- [123] Frédéric Grosshans et Philippe Grangier : Quantum cloning and teleportation criteria for continuous quantum variables, *Phys. Rev. A*, **64**(1), p. 010301 (2001). DOI: 10.1103/PhysRevA.64.010301.
- [124] W. P. Bowen, N. Treps, B. C. Buchler, R. Schnabel, T. C. Ralph, T. Symul et P. K. Lam : Unity and non-unity gain quantum teleportation, *IEEE J. of Quant. Elec.*, **9**(6), pp. 1519–1532 (2003).
- [125] G. Vidal et R. F. Werner : Computable measure of entanglement, *Phys. Rev. A*, **65**(3), p. 032314 (2002). DOI: 10.1103/PhysRevA.65.032314.
- [126] Kentaro Wakui, Hiroki Takahashi, Akira Furusawa et Masahide Sasaki : Photon subtracted squeezed states generated with periodically poled KTiOPO₄, *Optics Express*, **15**(6), pp. 3568–3574 (2007).
- [127] Alexei Ourjoumtsev, Franck Ferreyrol, Rosa Tualle-Brouri et Philippe Grangier : Preparing non-local superpositions of quasi-classical states of the light, *Nature Physics*, **5**, pp. 189–192 (2009). DOI: 10.1038/nphys1199.
- [128] P. van Loock : Optical hybrid approaches to quantum information, *Laser & Photonics Reviews*, **n/a**(n/a), pp. 1–34 (2010). DOI: 10.1002/lpor.201000005.
- [129] Jonatan B. Brask, Ioannes Rigas, Eugene S. Polzik, Ulrik L. Andersen et Anders S. Sørensen : A Hybrid Long-Distance Entanglement Distribution Protocol, (2010). <http://arxiv.org/abs/1004.0083>.
- [130] W. K. Wootters et W. H. Zurek : A single quantum cannot be cloned, *Nature*, **299**, pp. 802–803 (1982). DOI: 10.1038/299802a0.
- [131] Valerio Scarani, Sofyan Iblisdir, Nicolas Gisin et Antonio Acín : Quantum cloning, *Rev. Mod. Phys.*, **77**(4), pp. 1225–1256 (2005). DOI: 10.1103/RevModPhys.77.1225.
- [132] G. Y. Xiang, Ralph T. C., Lund A. P., N. Walk et G. J. Pryde : Heralded Noiseless Linear Amplification and Distillation of Entanglement, *Nature Physics*, **4**(5), pp. 316–319 (2010). DOI: 10.1038/nphoton.2010.35.
- [133] Carlton M. Caves : Quantum limits on noise in linear amplifiers, *Phys. Rev. D*, **26**(8), pp. 1817–1839 (1982). DOI: 10.1103/PhysRevD.26.1817.
- [134] E. Arthurs et M. S. Goodman : Quantum Correlations : A Generalized Heisenberg Uncertainty Relation, *Phys. Rev. Lett.*, **60**(24), pp. 2447–2449 (1988). DOI: 10.1103/PhysRevLett.60.2447.
- [135] J. A. Levenson, I. Abram, T. Rivera, P. Fayolle, J. C. Garreau et P. Grangier : Quantum optical cloning amplifier, *Phys. Rev. Lett.*, **70**(3), pp. 267–270 (1993). DOI: 10.1103/PhysRevLett.70.267.
- [136] J. A. Levenson, I. Abram, T. Rivera et P. Grangier : Reduction of quantum noise in optical parametric amplification, *JOSA B*, **10**(11), pp. 2233–2238 (1993). DOI: 10.1364/JOSAB.10.002233.
- [137] Jaromír Fiurášek : Optimal probabilistic cloning and purification of quantum states, *Phys. Rev. A*, **70**(3), p. 032308 (2004). DOI: 10.1103/PhysRevA.70.032308.
- [138] Lu-Ming Duan et Guang-Can Guo : Probabilistic Cloning and Identification of Linearly Independent Quantum States, *Phys. Rev. Lett.*, **80**(22), pp. 4999–5002 (1998). DOI: 10.1103/PhysRevLett.80.4999.
- [139] Charles H. Bennett, Gilles Brassard, Sandu Popescu, Benjamin Schumacher, John A. Smolin et William K. Wootters : Purification of Noisy Entanglement and Faithful Teleportation

- via Noisy Channels, *Phys. Rev. Lett.*, **76**(5), pp. 722–725 (1996). DOI: 10.1103/PhysRevLett.76.722.
- [140] David Menzies et Sarah Croke : Noiseless linear amplification via weak measurements, (2009). <http://arxiv.org/abs/0903.4181>.
- [141] Michel Le Bellac : *Physique Quantique*, EDP SCIENCES (2007).
- [142] Michael A. Nielsen et Isaac L. Chuang : *Quantum Computation and Quantum Information*, Cambridge University Press (2000).
- [143] T. C. Ralph et A. P. Lund : Nondeterministic Noiseless Linear Amplification of Quantum Systems, *QUANTUM COMMUNICATION, MEASUREMENT AND COMPUTING (QCMC)*, pp. 155–160 (2009).
- [144] David T. Pegg, Lee S. Phillips et Stephen M. Barnett : Optical State Truncation by Projection Synthesis, *Phys. Rev. Lett.*, **81**(8), pp. 1604–1606 (1998). DOI: 10.1103/PhysRevLett.81.1604.
- [145] S. A. Babichev, J. Ries et A. I. Lvovsky : Quantum scissors : Teleportation of single-mode optical states by means of a nonlocal single photon, *EPL*, **64**(1), p. 1 (2003). DOI: 10.1209/epl/i2003-00504-y.
- [146] F. Laloë et Mullin W.J. : Quantum properties of a single beam splitter, *sera publié dans Physical Review Letters* (2010). <http://arxiv.org/abs/1004.1731>.
- [147] Franck Ferreyrol, Marco Barbieri, Rémi Blandino, Simon Fossier, Rosa Tualle-Brouri et Philippe Grangier : Implementation of a non-deterministic optical noiseless amplifier, *Phys. Rev. Lett.*, **104**, p. 123603 (2010). DOI: 10.1103/PhysRevLett.104.123603.
- [148] Sze M Tan : A computational toolbox for quantum and atomic optics, *Journal of Optics B : Quantum and Semiclassical Optics*, **1**(4), p. 424 (1999). <http://stacks.iop.org/1464-4266/1/i=4/a=312>.
- [149] Frédéric Grosshans : *Communication et cryptographie quantiques avec des variables continues*. Thèse de doctorat, Université Paris XI Orsay (2002). <http://tel.archives-ouvertes.fr/docs/00/04/54/62/PDF/tel-00003104.pdf>.
- [150] Petr Marek et Radim Filip : Coherent-state phase concentration by quantum probabilistic amplification, *Phys. Rev. A*, **81**(2), p. 022302 (2010). DOI: 10.1103/PhysRevA.81.022302.
- [151] Jaromír Fiurášek : Engineering quantum operations on traveling light beams by multiple photon addition and subtraction, *Phys. Rev. A*, **80**(5), p. 053822 (2009). DOI: 10.1103/PhysRevA.80.053822.
- [152] A. Zavatta, J. Fiurášek et M. Bellini, *Nature Physics* A high-fidelity noiseless amplifier for quantum light states, **5**(1), (2011). DOI: 10.1038/nphoton.2010.260.
- [153] Mario A. Usuga, Christian R. Müller, Wittmann Christoffer, Marek Petr, Radim Filip, Christoph Marquardt, Gerd Leuchs et Ulrik L. Andersen : Noise-powered probabilistic concentration of phase information, *Nature Physics*, **6**(10), pp. 767–771 (2010). DOI: 10.1038/nphys1743.
- [154] Simon Fossier : *Mise en oeuvre et évaluation de dispositifs de cryptographie quantique à longueur d'onde télécom*. Thèse de doctorat, Université Paris XI Orsay (2009). <http://tel.archives-ouvertes.fr/tel-00429450/fr/>.
- [155] Simon Fossier, Eleni Diamanti, Thierry Debuisschert, Rosa Tualle-Brouri et Philippe Grangier : Improvement of continuous-variable quantum key distribution systems by

- using optical preamplifiers, *J. Phys. B*, **42**(11), p. 114014 (2009). DOI: 10.1088/0953-4075/42/11/114014.
- [156] G. Y. Xiang, Ralph T. C., Lund A. P., N. Walk et G. J. Pryde : Heralded Noiseless Linear Amplification and Distillation of Entanglement, *Nature Physics*, **4**(5), pp. 316–319 (2010). DOI: 10.1038/nphoton.2010.35.
- [157] Nicolas Gisin, Stefano Pironio et Nicolas Sangouard : Proposal for Implementing Device-Independent Quantum Key Distribution Based on a Heralded Qubit Amplifier, *Phys. Rev. Lett.*, **105**(7), p. 070501 (2010). DOI: 10.1103/PhysRevLett.105.070501.
- [158] A. S. Holevo, M. Sohma et O. Hirota : Capacity of quantum Gaussian channels, *Phys. Rev. A*, **59**(3), pp. 1820–1828 (1999). DOI: 10.1103/PhysRevA.59.1820.
- [159] Frédéric Grosshans et Nicolas J. Cerf : Continuous-Variable Quantum Cryptography is Secure against Non-Gaussian Attacks, *Phys. Rev. Lett.*, **92**(4), p. 047905 (2004). DOI: 10.1103/PhysRevLett.92.047905.
- [160] J. Eisert et Wolf M.M. : Gaussian quantum channels Quantum Information with Continuous Variables of Atoms and Light (2007). <http://arxiv.org/abs/quant-ph/0505151>.
- [161] Michael M. Wolf, David Pérez-García et Geza Giedke : Quantum Capacities of Bosonic Channels, *Phys. Rev. Lett.*, **98**(13), p. 130501 (2007). DOI: 10.1103/PhysRevLett.98.130501.
- [162] A. S. Holevo et R. F. Werner : Evaluating capacities of bosonic Gaussian channels, *Phys. Rev. A*, **63**(3), p. 032312 (2001). DOI: 10.1103/PhysRevA.63.032312.
- [163] Michael M. Wolf, Geza Giedke et J. Ignacio Cirac : Extremality of Gaussian Quantum States, *Phys. Rev. Lett.*, **96**(8), p. 080502 (2006). DOI: 10.1103/PhysRevLett.96.080502.
- [164] T. Opatrny, G. Kurizki et D.-G. Welsch : Improvement on teleportation of continuous variables by photon subtraction via conditional measurement, *Phys. Rev. A*, **61**(3), p. 032302 (2000). DOI: 10.1103/PhysRevA.61.032302.
- [165] P. T. Cochrane, T. C. Ralph et G. J. Milburn : Teleportation improvement by conditional measurements on the two-mode squeezed vacuum, *Phys. Rev. A*, **65**(6), p. 062306 (2002). DOI: 10.1103/PhysRevA.65.062306.
- [166] Stefano Olivares, Matteo G. A. Paris et Rodolfo Bonifacio : Teleportation improvement by inconclusive photon subtraction, *Phys. Rev. A*, **67**(3), p. 032314 (2003). DOI: 10.1103/PhysRevA.67.032314.
- [167] N. J. Cerf, O. Krüger, P. Navez, R. F. Werner et M. M. Wolf : Non-Gaussian Cloning of Quantum Coherent States is Optimal, *Phys. Rev. Lett.*, **95**(7), p. 070501 (2005). DOI: 10.1103/PhysRevLett.95.070501.
- [168] Federico Casagrande, Alfredo Lulli et Matteo G. A. Paris : Improving the entanglement transfer from continuous-variable systems to localized qubits using non-Gaussian states, *Phys. Rev. A*, **75**(3), p. 032336 (2007). DOI: 10.1103/PhysRevA.75.032336.
- [169] Marco G. Genoni, Carmen Invernizzi et Matteo G. A. Paris : Enhancement of parameter estimation by Kerr interaction, *Phys. Rev. A*, **80**(3), p. 033842 (2009). DOI: 10.1103/PhysRevA.80.033842.
- [170] G. Adesso, F. Dell’Anno, S. De Siena, F. Illuminati et L. A. M. Souza : Optimal estimation of losses at the ultimate quantum limit with non-Gaussian states, *Phys. Rev. A*, **79**(4), p. 040305 (2009). DOI: 10.1103/PhysRevA.79.040305.

- [171] Alessandro Zavatta, Valentina Parigi et Marco Bellini : Experimental nonclassicality of single-photon-added thermal light states, *Phys. Rev. A*, **75**(5), p. 052106 (2007). DOI: 10.1103/PhysRevA.75.052106.
- [172] Stefano Olivares et Matteo Paris : Squeezed Fock state by inconclusive photon subtraction, *J. Opt. B*, **7**(12), pp. S616–S621 (2005). DOI: 10.1088/1464-4266/7/12/025.
- [173] M S Kim : Recent developments in photon-level operations on travelling light fields, *J. Phys. B*, **41**(13), p. 133001 (2008). DOI: 10.1088/0953-4075/41/13/133001.
- [174] Daniel E. Browne, Jens Eisert, Stefan Scheel et Martin B. Plenio : Driving non-Gaussian to Gaussian states with linear optics, *Phys. Rev. A*, **67**(6), p. 062320 (2003). DOI: 10.1103/PhysRevA.67.062320.
- [175] Tomas Tyc et Natalia Korolkova : Highly non-Gaussian states created via cross-Kerr non-linearity, *New Journal of Physics*, **10**(2), p. 023041 (20072008). <http://arxiv.org/abs/0709.2011>.
- [176] V. D’Auria, C. de Lisio, A. Porzio, S. Solimeno, Javaid Anwar et M. G. A. Paris : Non-Gaussian states produced by close-to-threshold optical parametric oscillators : Role of classical and quantum fluctuations, *Phys. Rev. A*, **81**(3), p. 033846 (2010). DOI: 10.1103/PhysRevA.81.033846.
- [177] Virginia D’Auria, Antonino Chiummo, Martina De Laurentis, Alberto Porzio, Salvatore Solimeno et Matteo Paris : Tomographic characterization of OPO sources close to threshold, *Opt. Express*, **13**(3), pp. 948–956 (2005). <http://www.opticsexpress.org/abstract.cfm?URI=oe-13-3-948>.
- [178] G. M. D’Ariano, L. Maccone, M. G. A. Paris et M. F. Sacchi : Optical Fock-state synthesizer, *Phys. Rev. A*, **61**(5), p. 053817 (2000). DOI: 10.1103/PhysRevA.61.053817.
- [179] B. Hladký, G. Drobný et V. Bužek : Quantum synthesis of arbitrary unitary operators, *Phys. Rev. A*, **61**(2), p. 022102 (2000). DOI: 10.1103/PhysRevA.61.022102.
- [180] Matteo G. A. Paris : Optical qubit by conditional interferometry, *Phys. Rev. A*, **62**(3), p. 033813 (2000). DOI: 10.1103/PhysRevA.62.033813.
- [181] Samuel Deleglise, Igor Dotsenko, Clement Sayrin, Julien Bernu, Michel Brune, Jean-Michel Raimond et Serge Haroche : Reconstruction of non-classical cavity field states with snapshots of their decoherence, *Nature*, **455**(7212), pp. 510–514 (2008). DOI: 10.1038/nature07288.
- [182] Max Hofheinz, E. M. Weig, M. Ansmann, Radoslaw C. Bialczak, Erik Lucero, M. Neeley, A. D. O’Connell, H. Wang, John M. Martinis et A. N. Cleland : Generation of Fock states in a superconducting quantum circuit, *Nature*, **454**(7202), pp. 310–314 (2008). DOI: 10.1038/nature07136.
- [183] F. Dell’Anno, S. De Siena et F. Illuminati : Realistic continuous-variable quantum teleportation with non-Gaussian resources, *Phys. Rev. A*, **81**(1), p. 012333 (2010). DOI: 10.1103/PhysRevA.81.012333.
- [184] Gerardo Adesso : Experimentally friendly bounds on non-Gaussian entanglement from second moments, *Phys. Rev. A*, **79**(2), p. 022315 (2009). DOI: 10.1103/PhysRevA.79.022315.
- [185] Milton Abramowitz et Irene A. Stegun : *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*, Dover (1964).
- [186] A. Mansour et C. Jutten : What should we say about the kurtosis?, *Signal Processing Letters, IEEE*, **6**(12), pp. 321–322 (1999). DOI: 10.1109/97.803435.

- [187] Jean-Louis Lacoume, Pierre-Olivier Amblard et Pierre Como : Statistiques d'ordre supérieur pour le traitement du signal (1997).
- [188] Marco G. Genoni, Matteo G. A. Paris et Konrad Banaszek : Measure of the non-Gaussian character of a quantum state, *Phys. Rev. A*, **76**(4), p. 042327 (2007). DOI: 10.1103/PhysRevA.76.042327.
- [189] Marco G. Genoni et Matteo G. A. Paris : Quantifying non-Gaussianity for quantum information, *Phys. Rev. A*, **82**(5), p. 052341 (2010). DOI: 10.1103/PhysRevA.82.052341.
- [190] M. G. Genoni et M. G. A. Paris : Non-Gaussianity and purity in finite dimension, *Int. J. Quant. Inf.*, **7**, pp. 97–103 (2009). DOI: 10.1142/S0219749909004712.
- [191] A. Hyvärinen, J. Karhunen et E. Oja : *Independent Component Analysis*, John Wiley and Sons (2001).
- [192] Marco G. Genoni, Matteo G. A. Paris et Konrad Banaszek : Quantifying the non-Gaussian character of a quantum state by quantum relative entropy, *Phys. Rev. A*, **78**(6), p. 060303 (2008). DOI: 10.1103/PhysRevA.78.060303.
- [193] V. Vedral : The role of relative entropy in quantum information theory, *Rev. Mod. Phys.*, **74**(1), pp. 197–234 (2002). DOI: 10.1103/RevModPhys.74.197.
- [194] J. Solomon Ivan, M. Sanjay Kumar et R. Simon : A measure of non-Gaussianity for quantum states, , (2008). <http://arxiv.org/abs/0812.2800>.
- [195] V. V. Dodonov, O. V. Man'ko, V. I. Man'ko et A. Wünsche : Hilbert-Schmidt distance and non-classicality of states in quantum optics, *J. Mod. Opt.*, **47**(4), pp. 633–654 (2000).
- [196] V. V. Dodonov et M. B. Renó : Classicality and anticlassicality measures of pure and mixed quantum states, *Physics Letters A*, **308**(4), pp. 249–255 (2003). DOI: 10.1016/S0375-9601(03)00066-5.
- [197] Paulina Marian, Tudor A. Marian et Horia Scutaru : Distinguishability and nonclassicality of one-mode Gaussian states, *Phys. Rev. A*, **69**(2), p. 022104 (2004). DOI: 10.1103/PhysRevA.69.022104.
- [198] A. Mari, K. Kieling, B. Melholt Nielsen, E. S. Polzik et J. Eisert : Directly estimating non-classicality, , (2010). <http://arxiv.org/abs/1005.1665>.
- [199] Marco Barbieri, Nicolò Spagnolo, Marco G. Genoni, Franck Ferreyrol, Rémi Blandino, Matteo G.A. Paris, Philippe Grangier et Rosa Tualle-Brouri : Non-Gaussianity of quantum states : an experimental test on single-photon added coherent states, *Physical Review A*, **82**(6), p. 063833 (2010). DOI: 10.1103/PhysRevA.82.063833.
- [200] J. Solomon Ivan, M. Sanjay Kumar et R. Simon : A measure of non-Gaussianity for quantum states, , (2008). <http://arxiv.org/abs/0812.2800>.
- [201] Michal Horodecki, Jonathan Oppenheim et Andreas Winter : Partial quantum information, *Nature*, **436**(7051), pp. 673–676 (2005). DOI: doi:10.1038/nature03909.
- [202] P. T. Cochrane, T. C. Ralph et A. Dolińska : Optimal cloning for finite distributions of coherent states, *Phys. Rev. A*, **69**(4), p. 042313 (2004). DOI: 10.1103/PhysRevA.69.042313.
- [203] R. García-Patrón, J. Fiurášek, N. J. Cerf, J. Wenger, R. Tualle-Brouri et Ph. Grangier : Proposal for a Loophole-Free Bell Test Using Homodyne Detection, *Phys. Rev. Lett.*, **93**(13), p. 130409 (2004). DOI: 10.1103/PhysRevLett.93.130409.
- [204] Donna Strickland et Gerard Mourou : Compression of amplified chirped optical pulses, *Optics Communications*, **56**(3), pp. 219–221 (1985). DOI: 10.1016/0030-4018(85)90120-8.

- [205] A. Dubietis, G. Jonusauskas et A. Piskarskas : Powerful femtosecond pulse generation by chirped and stretched pulse parametric amplification in BBO crystal, *Optics Communications*, **88**(4-6), pp. 437–440 (1992). DOI: 10.1016/0030-4018(92)90070-8.
- [206] Igor Jovanovic, Brian J. Comaskey, Christopher A. Ebbers, Randal A. Bonner, Deanna M. Pennington et Edward C. Morse : Optical Parametric Chirped-Pulse Amplifier as an Alternative to Ti :Sapphire Regenerative Amplifiers, *Appl. Opt.*, **41**(15), pp. 2923–2929 (2002). <http://ao.osa.org/abstract.cfm?URI=ao-41-15-2923>.
- [207] Giulio Cerullo et Sandro De Silvestri : Ultrafast optical parametric amplifiers, *Review of Scientific Instruments*, **74**(1), pp. 1–18 (2003). DOI: 10.1063/1.1523642.
- [208] R. Butkus, R. Danielius, A. Dubietis, A. Piskarskas et A. Stabinis : Progress in chirped pulse optical parametric amplifiers, *Applied Physics B : Lasers and Optics*, **79**, pp. 693–700 (2004). DOI: 10.1007/s00340-004-1614-3.
- [209] Sébastien Forget : *Source laser picoseconde à haute cadence dans l'ultraviolet*. Thèse de doctorat, Université Paris XI Orsay (2003). <http://tel.archives-ouvertes.fr/tel-00004272/fr/>.
- [210] Jeff Squier, François Salin, Gerard Mourou et Donald Harter : 100-fs pulse generation and amplification in Ti :Al₂O₃, *Opt. Lett.*, **16**(5), pp. 324–326 (1991). <http://ol.osa.org/abstract.cfm?URI=ol-16-5-324>.
- [211] C. P. J. Barty, T. Guo, C. Le Blanc, F. Raksi, C. Rose-Petruck, J. Squier, K. R. Wilson, V. V. Yakovlev et K. Yamakawa : Generation of 18-fs, multiterawatt pulses by regenerative pulse shaping and chirped-pulse amplification, *Opt. Lett.*, **21**(9), pp. 668–670 (1996). <http://ol.osa.org/abstract.cfm?URI=ol-21-9-668>.
- [212] Y. Nabekawa, Y. Kuramoto, T. Togashi, T. Sekikawa et S. Watanabe : Generation of 0.66-TW pulses at 1 kHz by a Ti :sapphire laser, *Opt. Lett.*, **23**(17), pp. 1384–1386 (1998). <http://ol.osa.org/abstract.cfm?URI=ol-23-17-1384>.
- [213] Joe Z. H. Yang et Barry C. Walker : 0.09-terawatt pulses with a 31% efficient, kilohertz repetition-rate Ti :sapphire regenerative amplifier, *Opt. Lett.*, **26**(7), pp. 453–455 (2001). <http://ol.osa.org/abstract.cfm?URI=ol-26-7-453>.
- [214] Isao Matsushima, Hidehiko Yashiro et Toshihisa Tomie : 10 kHz 40 W Ti :sapphire regenerative ring amplifier, *Opt. Lett.*, **31**(13), pp. 2066–2068 (2006). <http://ol.osa.org/abstract.cfm?URI=ol-31-13-2066>.
- [215] Kyung-Han Hong, Sergei Kostritsa, Tae Jun Yu, Jae Hee Sung, Il Woo Choi, Young-Chul Noh, Do-Kyeong Ko et Jongmin Lee : 100-kHz high-power femtosecond Ti :sapphire laser based on downchirped regenerative amplification, *Opt. Express*, **14**(2), pp. 970–978 (2006). <http://www.opticsexpress.org/abstract.cfm?URI=oe-14-2-970>.
- [216] T. B. Norris : Femtosecond pulse amplification at 250 kHz with a Ti :sapphire regenerative amplifier and application to continuum generation, *Opt. Lett.*, **17**(14), pp. 1009–1011 (1992). <http://ol.osa.org/abstract.cfm?URI=ol-17-14-1009>.
- [217] Murray K. Reed, Michael K. Steiner-Shepard et Daniel K. Negus : Widely tunable femtosecond optical parametric amplifier at 250 kHz with a Ti :sapphire regenerative amplifier, *Opt. Lett.*, **19**(22), pp. 1855–1857 (1994). <http://ol.osa.org/abstract.cfm?URI=ol-19-22-1855>.
- [218] Zhenlin Liu, Hidetoshi Murakami, Toshimasa Kozeki, Hideyuki Ohtake et Nobuhiko Sarukura : High-gain, reflection-double pass, Ti :sapphire continuous-wave amplifier delivering

- 5.77 W average power, 82 MHz repetition rate, femtosecond pulses, *Applied Physics Letters*, **76**(22), pp. 3182–3184 (2000). DOI: 10.1063/1.126639.
- [219] R. Huber, F. Adler, A. Leitenstorfer, M. Beutler, P. Baum et E. Riedle : 12-fs pulses from a continuous-wave-pumped 200-nJ Ti:sapphire amplifier at a variable repetition rate as high as 4 MHz, *Opt. Lett.*, **28**(21), pp. 2118–2120 (2003). <http://ol.osa.org/abstract.cfm?URI=ol-28-21-2118>.
- [220] Lee M. Frantz et John S. Nodvik : Theory of Pulse Propagation in a Laser Amplifier, *Journal of Applied Physics*, **34**(8), pp. 2346–2349 (1963). DOI: 10.1063/1.1702744.
- [221] Mark Csele : *Fundamentals of Light Sources and Lasers*, Wiley (2004).
- [222] R. Szipocs et A. Koházi-Kis : Theory and design of chirped dielectric laser mirrors, *Applied Physics B : Lasers and Optics*, **65**, pp. 115–135 (1997). DOI: 10.1007/s003400050258.
- [223] Aurélien Dantan, Julien Laurat, Alexei Ourjoumtsev, Rosa Tualle-Brouiri et Philippe Grangier : Femtosecond Ti:sapphire cryogenic amplifier with high gain and MHz repetition rate, *Opt. Express*, **15**(14), pp. 8864–8870 (2007). <http://www.opticsexpress.org/abstract.cfm?URI=oe-15-14-8864>.
- [224] Francois Salin, Catherine Le Blanc, Jeff Squier et Chris Barty : Thermal eigenmode amplifiers for diffraction-limited amplification of ultrashort pulses, *Opt. Lett.*, **23**(9), pp. 718–720 (1998). <http://ol.osa.org/abstract.cfm?URI=ol-23-9-718>.
- [225] M. Zavelani-Rossi, F. Lindner, C. Le Blanc, G. Chériaux et J.P. Chambaret : Control of thermal effects for high-intensity Ti:sapphire laser chains, *Applied Physics B : Lasers and Optics*, **70**, pp. S193–S196 (2000). DOI: 10.1007/s003400000364.
- [226] M. Pittman, S. Ferré, J.P. Rousseau, L. Notebaert, J.P. Chambaret et G. Chériaux : Design and characterization of a near-diffraction-limited femtosecond 100-TW 10-Hz high-intensity laser system, *Applied Physics B : Lasers and Optics*, **74**, pp. 529–535 (2002). DOI: 10.1007/s003400200838.
- [227] P. F. Moulton : Spectroscopic and laser characteristics of Ti:Al₂O₃, *J. Opt. Soc. Am. B*, **3**(1), pp. 125–133 (1986). <http://josab.osa.org/abstract.cfm?URI=josab-3-1-125>.
- [228] Martin Delaigue : *Étude et réalisation de sources femtosecondes haute puissance moyenne* Étude et réalisation de sources femtosecondes haute puissance moyenne. Thèse de doctorat, Université de Bordeaux 1 (2006).
- [229] R. Huber, F. Adler, A. Leitenstorfer, M. Beutler, P. Baum et E. Riedle : 12-fs pulses from a continuous-wave-pumped 200-nJ Ti:sapphire amplifier at a variable repetition rate as high as 4 MHz, *Opt. Lett.*, **28**(21), pp. 2118–2120 (2003). <http://ol.osa.org/abstract.cfm?URI=ol-28-21-2118>.
- [230] Thomas Planchon : *Modélisation des processus liés à l'amplification et à la propagation d'impulsions étirées dans des chaînes laser de très haute intensité* Modélisation des processus liés à l'amplification et à la propagation d'impulsions étirées dans des chaînes laser de très haute intensité. Thèse de doctorat, École Polytechnique (2004). <http://tel.archives-ouvertes.fr/tel-00005388/fr/>.
- [231] A. M. Weiner : Femtosecond pulse shaping using spatial light modulators, *Review of Scientific Instruments*, **71**(5), pp. 1929–1960 (2000). DOI: 10.1063/1.1150614.
- [232] <http://www.amplificationtechnologies.com>.
- [233] Edo Waks, Eleni Diamanti, Barry C. Sanders, Stephen D. Bartlett et Yoshihisa Yamamoto : Direct Observation of Nonclassical Photon Statistics in Parametric Down-Conversion, *Phys. Rev. Lett.*, **92**(11), p. 113602 (2004). DOI: 10.1103/PhysRevLett.92.113602.

- [234] Edo Waks, Eleni Diamanti et Yoshihisa Yamamoto : Generation of photon number states, *New Journal of Physics*, **8**(1), p. 4 (2006). <http://stacks.iop.org/1367-2630/8/i=1/a=004>.
- [235] A. Imamoglu : High Efficiency Photon Counting Using Stored Light, *Phys. Rev. Lett.*, **89**(16), p. 163602 (2002). DOI: 10.1103/PhysRevLett.89.163602.
- [236] Daniel F. V. James et Paul G. Kwiat : Atomic-Vapor-Based High Efficiency Optical Detectors with Photon Number Resolution, *Phys. Rev. Lett.*, **89**(18), p. 183601 (2002). DOI: 10.1103/PhysRevLett.89.183601.
- [237] Pieter Kok et Samuel L. Braunstein : Detection devices in entanglement-based optical state preparation, *Phys. Rev. A*, **63**(3), p. 033812 (2001). DOI: 10.1103/PhysRevA.63.033812.
- [238] H. Paul, P. Törmä, T. Kiss et I. Jex : Photon Chopping : New Way to Measure the Quantum State of Light, *Phys. Rev. Lett.*, **76**(14), pp. 2464–2467 (1996). DOI: 10.1103/PhysRevLett.76.2464.
- [239] Leaf A. Jiang, Eric A. Dauler et Joshua T. Chang : Photon-number-resolving detector with 10bits of resolution, *Phys. Rev. A*, **75**(6), p. 062325 (2007). DOI: 10.1103/PhysRevA.75.062325.
- [240] K. Yamamoto, K. Yamamura, K. Sato, T. Ota, H. Suzuki et S. Ohsuka : Development of Multi-Pixel Photon Counter (MPPC), *IEEE Nucl. Sci. Symp. Conf. Record 2006 2*, pp. 1094–1097 (2006).
- [241] Daryl Achilles, Christine Silberhorn, Cezary Śliwa, Konrad Banaszek, Ian A. Walmsley, Michael J. Fitch, Bryan C. Jacobs, Todd B. Pittman et James D. Franson : Photon-number-resolving detection using time-multiplexing, *Journal of Modern Optics*, **51**(9), pp. 1499–1515 (2004). <http://www.informaworld.com/10.1080/09500340408235288>.
- [242] M. J. Fitch, B. C. Jacobs, T. B. Pittman et J. D. Franson : Photon-number resolution using time-multiplexed single-photon detectors, *Phys. Rev. A*, **68**(4), p. 043814 (2003). DOI: 10.1103/PhysRevA.68.043814.
- [243] O. Haderka, M. Hamar et J. Perina Jr : Experimental multi-photon-resolving detector using a single avalanche photodiode, *Eur. Phys. J. D*, **28**(1), pp. 149–154 (2004). DOI: 10.1140/epjd/e2003-00287-1.
- [244] Konrad Banaszek et Ian A. Walmsley : Photon counting with a loop detector, *Opt. Lett.*, **28**(1), pp. 52–54 (2003). DOI: 10.1364/OL.28.000052.
- [245] J. Řeháček, Z. Hradil, O. Haderka, J. Peřina et M. Hamar : Multiple-photon resolving fiber-loop detector, *Phys. Rev. A*, **67**(6), p. 061801 (2003). DOI: 10.1103/PhysRevA.67.061801.
- [246] Peter P Rohde, James G Webb, Elanor H Huntington et Timothy C Ralph : Photon number projection using non-number-resolving detectors, *New Journal of Physics*, **9**(7), p. 233 (2007). <http://stacks.iop.org/1367-2630/9/i=7/a=233>.
- [247] Roy S. Bondurant, Prem Kumar, Jeffrey H. Shapiro et Michael M. Salour : Photon-counting statistics of pulsed light sources, *Opt. Lett.*, **7**(11), pp. 529–531 (1982). <http://ol.osa.org/abstract.cfm?URI=ol-7-11-529>.
- [248] Guido Zambra, Maria Bondani, Alessandro S. Spinelli, Fabio Paleari et Alessandra Andreoni : Counting photoelectrons in the response of a photomultiplier tube to single picosecond light pulses, *Review of Scientific Instruments*, **75**(8), pp. 2762–2765 (2004). DOI: 10.1063/1.1777407.

- [249] Makoto Akiba, Mikio Fujiwara et Masahide Sasaki : Ultrahigh-sensitivity high-linearity photodetection system using a low-gain avalanche photodiode with an ultralow-noise readout circuit, *Opt. Lett.*, **30**(2), pp. 123–125 (2005). <http://ol.osa.org/abstract.cfm?URI=ol-30-2-123>.
- [250] B. E. Kardynal, Z. L. Yuan et A. J. Shields : An avalanche-photodiode-based photon-number-resolving detector, *Nat. Photon.*, **2**(7), pp. 425–428 (2008). DOI: 10.1038/nphoton.2008.101.
- [251] Guang Wu, Yi Jian, E. Wu et Heping Zeng : Photon-number-resolving detection based on InGaAs/InP avalanche photodiode in the sub-saturated mode, *Opt. Express*, **17**(21), pp. 18782–18787 (2009). <http://www.opticsexpress.org/abstract.cfm?URI=oe-17-21-18782>.
- [252] Shigeki Takeuchi, Jungsang Kim, Yoshihisa Yamamoto et Henry H. Hogue : Development of a high-quantum-efficiency single-photon counting system, *Applied Physics Letters*, **74**(8), pp. 1063–1065 (1999). DOI: 10.1063/1.123482.
- [253] J. Kim, S. Somani et Y. Yamamoto : *Nonclassical Light from Semiconductor Lasers and LEDs*, Springer (2001).
- [254] Stephen D. Bartlett, E. Diamanti, Barry C. Sanders et Yoshihisa Yamamoto : Photon counting schemes and performance of non-deterministic nonlinear gates in linear optics, *Free-Space Laser Communication and Laser Imaging II*, pp. 427–435 (2002). DOI: 10.1117/12.451332.
- [255] B. E. Kardynał, S. S. Hees, A. J. Shields, C. Nicoll, I. Farrer et D. A. Ritchie : Photon number resolving detector based on a quantum dot field effect transistor, *Applied Physics Letters*, **90**(18), p. 181114 (2007). DOI: 10.1063/1.2735281.
- [256] E. J. Gansen, M. A. Rowe, M. B. Greene, D. Rosenberg, T. E. Harvey, M. Y. Su, R. H. Hadfield, S. W. Nam et R. P. Mirin : Photon-number-discriminating detection using a quantum-dot, optically gated, field-effect transistor, *Nat. Photon.*, **1**(10), pp. 585–588 (2007). DOI: 10.1038/nphoton.2007.173.
- [257] Mikio Fujiwara et Masahide Sasaki : Multiphoton discrimination at telecom wavelength with charge integration photon detector, *Applied Physics Letters*, **86**(11), p. 111119 (2005). DOI: 10.1063/1.1886903.
- [258] Mikio Fujiwara et Masahide Sasaki : Photon-number-resolving detection at a telecommunications wavelength with a charge-integration photon detector, *Opt. Lett.*, **31**(6), pp. 691–693 (2006). <http://ol.osa.org/abstract.cfm?URI=ol-31-6-691>.
- [259] J. H. J. de Bruijne, A. P. Reynolds, M. A. C. Perryman, F. Favata et A. Peacock : Analysis of astronomical data from optical superconducting tunnel junctions, *Optical Engineering*, **41**(6), pp. 1158–1169 (2002). DOI: 10.1117/1.1475334.
- [260] Veronica Savu, Luigi Frunzio et Daniel E. Prober : Enhancing the Energy Resolution of a Single Photon STJ Spectrometer Using Diffusion Engineering, *IEEE trans. appl. Supercond.*, **17**(2), pp. 324–327 (2003).
- [261] Robert H. Hadfield, Martin J. Stevens, Steven S. Gruber, Aaron J. Miller, Robert E. Schwall, Richard P. Mirin et Sae Woo Nam : Single photon source characterization with a superconducting single photon detector, *Opt. Express*, **13**(26), pp. 10846–10853 (2005). <http://www.opticsexpress.org/abstract.cfm?URI=oe-13-26-10846>.
- [262] R. Sobolewski, A. Verevkin, G.N. Gol'tsman, A. Lipatov et K. Wilsher : Ultrafast superconducting single-photon optical detectors and their applications, *IEEE trans. appl. Supercond.*, **13**(2), pp. 1151–1157 (2003). DOI: 10.1109/TASC.2003.814178.

- [263] A. J. Annunziata, A. Frydman, M. O. Reese, L. Frunzio, M. Rooks et D. E. Prober : Superconducting niobium nanowire single photon detectors, *Advanced Photon Counting Techniques*, p. 63720V (2006). DOI: 10.1117/12.686301.
- [264] Eric A. Dauler, Andrew J. Kerman, Bryan S. Robinson, Joel K. W. Yang, Boris Voronov, Gregory Goltsman, Scott A. Hamilton et Karl K. Berggren : Photon-number-resolution with sub-30-ps timing using multi-element superconducting nanowire single photon detectors, *Journal of Modern Optics*, **56**(2), pp. 364–373 (2009). DOI: 10.1080/09500340802411989.
- [265] Aleksander Divochiy, Francesco Marsili, David Bitauld, Alessandro Gaggero, Roberto Leoni, Francesco Mattioli, Alexander Korneev, Vitaliy Seleznev, Nataliya Kaurova, Olga Minaeva, Gregory Gol'tsman, Konstantinos G. Lagoudakis, Moushab Benkhaoul, Francis Levy et Andrea Fiore : Superconducting nanowire photon-number-resolving detector at telecommunication wavelengths, *Nat. Photon.*, **2**(5), pp. 302–306 (2008). DOI: 10.1038/nphoton.2008.51.
- [266] M Ejrnaes, A Casaburi, O Quaranta, S Marchetti, A Gaggero, F Mattioli, R Leoni, S Pagano et R Cristiano : Characterization of parallel superconducting nanowire single photon detectors, *Superconductor Science and Technology*, **22**(5), p. 055006 (2009). <http://stacks.iop.org/0953-2048/22/i=5/a=055006>.
- [267] Aaron J. Miller, Sae Woo Nam, John M. Martinis et Alexander V. Sergienko : Demonstration of a low-noise near-infrared photon counter with multiphoton discrimination, *Applied Physics Letters*, **83**(4), pp. 791–793 (2003). DOI: 10.1063/1.1596723.
- [268] E. Waks, K. Inoue, W.D. Oliver, E. Diamanti et Y. Yamamoto : High-efficiency photon-number detection for quantum information processing, *IEEE Journal of Selected Topics in Quantum Electronics*, **9**(6), pp. 1502–1511 (2003). DOI: 10.1109/JSTQE.2003.820917.
- [269] Radhika Rangarajan : *Photonic Sources and Detectors for Quantum Information Processing : A Trilogy in Eight Parts*. Thèse de doctorat, University of Illinois (2010). <http://research.physics.illinois.edu/QI/Photonics/theses/rangarajan-thesis.pdf>.
- [270] Randall A. LaViolette et M. G. Stapelbroek : A non-Markovian model of avalanche gain statistics for a solid-state photomultiplier, *Journal of Applied Physics*, **65**(2), pp. 830–836 (1989). DOI: 10.1063/1.343073.
- [271] Edo Waks : *Quantum Information Processing with Non-Classical Light*. Thèse de doctorat, Stanford University (2003). <http://www.stanford.edu/group/yamamotogroup/Thesis/EWthesis.pdf>.
- [272] E. Waks, K. Inoue, W.D. Oliver, E. Diamanti et Y. Yamamoto : High-efficiency photon-number detection for quantum information processing, *IEEE Journal of Selected Topics in Quantum Electronics*, **9**(6), pp. 1502–1511 (2003). DOI: 10.1109/JSTQE.2003.820917.
- [273] A. Bross, V. Büscher, J. Estrada, G. Ginther et J. Molina : Gain dispersion in visible light photon counters as a function of counting rate, *Applied Physics Letters*, **87**(21), p. 214102 (2005). DOI: 10.1063/1.2133921.
- [274] Shigeki Takeuchi, Jungsang Kim, Yoshihisa Yamamoto et Henry H. Hogue : Development of a high-quantum-efficiency single-photon counting system, *Applied Physics Letters*, **74**(8), pp. 1063–1065 (1999). DOI: 10.1063/1.123482.
- [275] Edo Waks : *Quantum Information Processing with Non-Classical Light*. Thèse de doctorat, Stanford University (2003). <http://www.stanford.edu/group/yamamotogroup/Thesis/EWthesis.pdf>.

- [276] Petr Marek et Jaromír Fiurášek : Elementary gates for quantum information with superposed coherent states, *Phys. Rev. A*, **82**(1), p. 014304 (2010). DOI: 10.1103/PhysRevA.82.014304.
- [277] Saleh Rahimi-Keshari, Artur Scherer, Ady Mann, Ali T. Rezakhani, A. I. Lvovsky et Barry C. Sanders : Quantum process tomography with coherent states, *New Journal of Physics*, **13**(1), p. 013006 (2010). DOI: 10.1088/1367-2630/13/1/013006.
- [278] Alessandro Zavatta, Silvia Viciani et Marco Bellini : Quantum-to-Classical Transition with Single-Photon-Added Coherent States of Light, *Science*, **306**(5696), pp. 660–662 (2004). DOI: 10.1126/science.1103190.
- [279] Valentina Parigi, Alessandro Zavatta, Myungshik Kim et Marco Bellini : Probing Quantum Commutation Rules by Addition and Subtraction of Single Photons to/from a Light Field, *Science*, **317**(5846), pp. 1890–1893 (2007). DOI: 10.1126/science.1146204.
- [280] S. R. Huisman, Nitin Jain, S. A. Babichev, Frank Vewinger, A. N. Zhang, S. H. Youn et A. I. Lvovsky : Instant single-photon Fock state tomography, *Opt. Lett.*, **34**(18), pp. 2739–2741 (2009). DOI: 10.1364/OL.34.002739.
- [281] Maxim S. Pshenichnikov, Wim P. de Boeij et Douwe A. Wiersma : Generation of 13-fs, 5-MW pulses from a cavity-dumped Ti :sapphire laser, *Opt. Lett.*, **19**(8), pp. 572–574 (1994). <http://ol.osa.org/abstract.cfm?URI=ol-19-8-572>.
- [282] M. Ramaswamy, M. Ulman, J. Paye et J. G. Fujimoto : Cavity-dumped femtosecond Kerr-lens mode-locked Ti :A12O3laser, *Opt. Lett.*, **18**(21), pp. 1822–1824 (1993). <http://ol.osa.org/abstract.cfm?URI=ol-18-21-1822>.

Note : pour accéder à un article sur internet à partir de son DOI il faut entrer l'adresse "http://dx.doi.org/" suivi du DOI.

Résumé

Résumé L'objectif de cette thèse concerne la manipulation à l'échelle quantique du champ électromagnétique dans le cadre de l'information quantique à variables continues. Pour ce faire nous mélangeons les outils de l'optique quantique à variables discrètes, où la lumière est décrite en termes de photons, avec l'approche continue, traitant des quadratures du champ. Cette technique permet de produire des états non-classiques décrits par des fonctions de Wigner prenant des valeurs négatives. Nous avons pu générer des états intriqués à partir d'impulsions lumineuses initialement indépendantes et pouvant être séparées par une longue distance, l'intrication s'effectuant au travers d'un canal acceptant de fortes pertes. Nous avons ensuite démontré et caractérisé expérimentalement un protocole non-déterministe permettant d'amplifier de faibles signaux sans en amplifier le bruit quantique, augmentant ainsi le rapport signal sur bruit. Puis nous avons mis en œuvre et comparé expérimentalement différentes mesures de non-gaussianité d'un état quantique : ce caractère propre à une description continue de la lumière est d'un intérêt capital pour l'information quantique. Enfin nous avons développé et testé deux améliorations pour notre dispositif. La première est un amplificateur femtoseconde pour notre laser impulsif, qui permettra d'obtenir de meilleurs états de départ pour nos expériences. La deuxième est un appareil capable de discriminer le nombre de photon, donnant ainsi des résultats plus précis que ceux des détecteurs dont nous disposons actuellement qui sont uniquement capable de détecter la présence de photons.

Mots-clés

Optique quantique – Optique non linéaire – Information quantique – Variables continues – Détection homodyne impulsionnelle – Tomographie quantique – États non-gaussiens – Intrication quantique

Abstract

Abstract This thesis aims at handling the electromagnetic field at a quantum scale in the area of quantum information processing. For this purpose we mixed tools of discrete variable quantum optics, where light is described in terms of photons, with the continuous approach, which uses the quadratures of the field. This technique enables the production of non-classical states which should be described by Wigner functions that takes negative values. We have generated entangled states from ultra-short light pulses initially independent and which can be separated by a long distance : the entanglement is indeed performed through a low-transmission channel. Then we have experimentally demonstrated and characterized a protocol that non-deterministically amplifies low signals without amplifying the quantum noise, increasing the signal to noise ratio. Furthermore we experimentally implement and compared several measures of the non-gaussianity of a quantum state : this characteristic, which belongs to continuous description of light, is of essential interest for quantum information processing. Finally we develop and test two improvements for our setup. The first one is a femtosecond amplifier for our pulsed laser. It will enable us to obtain better primitive states for our experiments. The second one is an apparatus that can discriminate the number of photon in a pulse, giving more accurate results than the detectors we used up to now that are only able to detect the presence of photons.

Keywords

Quantum optics – Non-linear optics – Quantum information – Continuous variables – Pulsed homodyne detection – Quantum tomography – Non-Gaussian states – Quantum entanglement