



Communication et cryptographie quantiques avec des variables continues

Frédéric Grosshans

► To cite this version:

Frédéric Grosshans. Communication et cryptographie quantiques avec des variables continues. Physique Atomique [physics.atom-ph]. Université Paris Sud - Paris XI, 2002. Français. NNT : . tel-00002343v1

HAL Id: tel-00002343

<https://theses.hal.science/tel-00002343v1>

Submitted on 31 Jan 2003 (v1), last revised 9 Jul 2003 (v2)

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.

N° d'ordre : 7080



INSTITUT D'OPTIQUE
LABORATOIRE CHARLES FABRY

UNIVERSITÉ PARIS XI
UFR SCIENTIFIQUE D'ORSAY

THÈSE

Spécialité : Laser et Matière

présentée pour obtenir
le GRADE de DOCTEUR EN SCIENCES
DE L'UNIVERSITÉ PARIS XI

par
Frédéric GROSSHANS

Sujet :

COMMUNICATION ET CRYPTOGRAPHIE QUANTIQUES AVEC DES VARIABLES CONTINUES

Soutenue le 12 décembre 2002
devant la Commission d'examen :

M. André DUCASSE,	Président
M. Nicolas CERF,	Rapporteur
M. Juan-Ariel LEVENSON,	Rapporteur
M. Thierry DEBUISSCHERT,	Examineur
M. Claude FABRE,	Examineur
M. Gerd LEUCHS,	Membre invité
M. Philippe GRANGIER,	Directeur de thèse

Table des matières

Remerciements	13
Notes diverses	15
Références bibliographiques	16
Publications	16
1 Introduction	17
1.1 Histoire et paradoxes de la mécanique quantique	17
1.1.1 De la vieille mécanique des quanta à la mécanique quantique	17
1.1.2 La mécanique quantique et ses paradoxes	19
1.2 Naissance de l'information quantique	21
1.2.1 Ordinateurs quantiques	21
1.2.2 Communication quantique	22
1.3 Qubits et variables continues	23
1.3.1 Variables discrètes	23
1.3.2 Variables quantiques continues	24
1.3.3 Avantages technologiques des variables continues	24
1.4 Plan de la thèse	25
I Variables Continues	27
2 Variables continues en mécanique quantique	29
2.1 Exemples de variables continues	29
2.1.1 Rôle historique de la position et de l'impulsion	29
2.1.2 Commutateurs	30
2.1.3 Évolution temporelle et quadratures du champ	31
2.1.4 Autres variables continues	32
2.2 États quantiques	33
2.2.1 États propres de $\hat{Q}_\theta$	33
2.2.2 États d'incertitude minimale	35
2.2.3 Propriétés des états gaussiens	39
2.2.4 Superposition linéaire d'états cohérents	40
2.2.5 États de Fock	42

3	La fonction de Wigner	45
3.1	Introduction	45
3.2	Définition	45
3.2.1	Probabilités marginales et tomographie	46
3.2.2	Expression de la fonction de Wigner	47
3.3	Propriétés élémentaires de la fonction de Wigner	49
3.3.1	Linéarité	49
3.3.2	Fonction de Wigner d'opérateurs	49
3.3.3	Fonction de Wigner d'opérateurs hermitiens	50
3.3.4	Traces et valeurs moyennes	50
3.4	Exemples	51
3.4.1	États propres de la position	52
3.4.2	États gaussiens	52
3.4.3	Superposition linéaire de deux états quantiques	53
3.4.4	États de Fock	54
3.5	Fonctions de Wigner à n modes	54
3.5.1	Définition	54
3.5.2	Traces partielles	55
3.5.3	Produits tensoriels et changements de base	55
3.5.4	Changement de base	56
3.5.5	Effet des pertes	57
3.5.6	Pertes et négativité de la fonction de Wigner	58
3.5.7	Effet des pertes sur la fonction de Wigner d'états de Fock	59
4	Formalisme gaussien	61
4.1	Introduction	61
4.2	Notations linéarisées	61
4.2.1	Matrices de covariance	61
4.2.2	Notations simplifiées	62
4.2.3	Bruit quantique	63
4.2.4	Paradoxe Einstein Podolsky Rosen	63
4.3	Canaux quantiques	64
4.3.1	Définition	64
4.3.2	Propriétés	65
4.3.3	Mesure simultanée de l'impulsion et de la position	67
4.3.4	Mesure et reconstruction	67
4.4	Fidélité	68
4.4.1	Définition	68
4.4.2	Fidélité d'un canal	69
4.4.3	Fidélité optimisée	69
4.4.4	Fidélité de la reconstruction	71
II	Détection homodyne impulsionnelle	73
5	Mesure homodyne impulsionnelle du vide	75
5.1	Pourquoi et comment mesurer une quadrature ?	75

5.2	Principes de fonctionnement d'une détection homodyne	76
5.2.1	Étude classique	76
5.2.2	Interprétation ondulatoire et interprétation corpusculaire	77
5.2.3	Étude quantique	78
5.3	Équilibrage d'une détection homodyne	79
5.3.1	Effet d'un déséquilibre	80
5.3.2	De la difficulté expérimentale d'équilibrer une détection homodyne impulsionnelle	81
5.4	Aspects temporels et fréquentiels	82
5.4.1	Détections homodyne résolues en fréquence	82
5.4.2	Détection homodyne résolue en temps	83
5.4.3	Électronique utilisée	84
5.4.4	Notations	85
5.4.5	Réponse percussionnelle	86
5.4.6	Lien avec la fonction de transfert	88
5.5	Résultats expérimentaux	91
6	Mesure homodyne d'un état cohérent	93
6.1	Inefficacité de la détection homodyne	93
6.1.1	Définition de l'efficacité	93
6.1.2	Inefficacités après la lame semi-réfléchissante	94
6.1.3	Efficacité de notre dispositif expérimental	95
6.2	Adaptation des modes temporels	96
6.2.1	Introduction	96
6.2.2	État du champ et opérateurs de mesure	97
6.2.3	Correspondance entre les détecteurs homodynes monomode impar- fait et impulsionnel parfait	98
6.2.4	Généralisations	99
6.3	Bruit de phase	100
6.4	Mesure expérimentale d'un état cohérent	101
III	Communication quantique	103
7	Théorie de l'information	105
7.1	Brefs rappels sur les variables discrètes	105
7.2	Application aux variables continues	106
7.2.1	Définition de l'entropie différentielle	106
7.2.2	Propriétés de l'entropie différentielle	107
7.2.3	Information mutuelle et canaux additifs	108
7.3	Communication avec des états gaussiens	110
7.3.1	États cohérents	111
7.3.2	États comprimés	111
7.3.3	Transmission d'information en présence de pertes	112
7.4	Contributions à l'information mutuelle	113
7.5	Coefficient de transfert d'information	115

8	Clonage quantique	117
8.1	Introduction	117
8.2	Définition et description d'un duplicateur quantique	117
8.3	Clonage optimal	118
8.4	Fidélité	119
8.5	Comment fabriquer un duplicateur ?	120
8.6	Les duplicateurs à gain différent de 1	121
8.6.1	Conditions nécessaires à un duplicateur optimal	121
8.6.2	Domaine d'existence des duplicateurs optimaux	122
8.6.3	Distribution du bruit entre Bob et Charles	123
8.7	Clonage symétrique $1 \rightarrow M$	124
9	Critères de téléportation quantique	127
9.1	Définition de la téléportation	127
9.1.1	Mesure et reconstruction	127
9.1.2	Téléportation quantique	128
9.1.3	Téléportation quantique de variables continues	128
9.2	La téléportation quantique en pratique	130
9.2.1	Intérêts de la téléportation quantique	130
9.2.2	Quand y a-t-il eu téléportation ?	130
9.3	Le seuil classique	131
9.3.1	Définition du seuil	131
9.3.2	Lien entre le seuil classique et l'intrication	132
9.3.3	Rôle des pertes et du facteur de compression	132
9.4	Le seuil $\overline{\mathcal{F}} = \frac{2}{3}$	132
9.4.1	Paradoxe EPR	133
9.4.2	Fonction de Wigner	133
9.4.3	Non-Clonage	133
IV	Cryptographie Quantique	135
10	Généralités sur la cryptographie quantique	137
10.1	Les principes de Kerckhoffs	137
10.1.1	Impossibilité du déchiffrement	138
10.1.2	Système connu de l'ennemi	138
10.1.3	La clef	139
10.1.4	La « correspondance télégraphique »	140
10.1.5	Système portatif	140
10.1.6	Ergonomie	140
10.2	La cryptographie quantique	141
10.2.1	La distribution quantique de clefs	141
10.2.2	Cryptographie quantique avec des variables continues	142

11 Protocoles directs	145
11.1 Introduction	145
11.2 Transfert d'information avec des variables gaussiennes	145
11.3 L'espionnage d'Ève	146
11.4 Émission d'Alice	147
11.5 Mesures de Bob	148
11.6 Information secrète	150
11.7 Rôle des petites et des grandes modulations	150
11.8 Rôle du facteur de compression	151
11.9 Protocoles	152
11.9.1 Cas général	152
11.9.2 États cohérents	154
11.9.3 États intriqués	155
12 Protocoles inverses	157
12.1 Comment franchir la limite des 3 dB	157
12.1.1 La limite des 3 dB	157
12.1.2 Réconciliation inverse	158
12.1.3 Interprétation de BB84 en terme de protocoles inverses	159
12.2 Les cloneuses intriquantes	160
12.2.1 Définition	160
12.2.2 Inégalité de Heisenberg sur les variances conditionnelles	161
12.2.3 Variance conditionnelle d'Alice	162
12.2.4 Variance conditionnelle d'Ève	163
12.2.5 Implémentation	163
12.3 Canaux bruités permettant la cryptographie inverse	166
12.3.1 Condition générale	166
12.3.2 États maximalelement comprimés	167
12.3.3 États cohérents	167
12.3.4 Facteurs de compression intermédiaires	168
12.4 Quantités d'information	169
12.4.1 Informations mutuelles et information secrète	169
12.4.2 États maximalelement comprimés	170
12.4.3 États cohérents	171
12.4.4 Facteurs de compression intermédiaires	171
12.4.5 Fortes pertes	172
13 Expérience de cryptographie quantique	175
13.1 Signaux envoyés par Alice	175
13.1.1 Génération aléatoire du signal	176
13.1.2 Effets du codage discret des nombres réels	178
13.2 Mise en œuvre expérimentale	182
13.2.1 Modulateur d'amplitude d'Alice	182
13.2.2 Modulation de la phase	184
13.2.3 Bob	184
13.2.4 Ève	184
13.3 Hypothèse sur les attaques d'Ève	185

13.4	Encore un peu d'informatique...	186
13.4.1	Estimation	186
13.4.2	Réconciliation par tranches	186
13.4.3	<i>Cascade</i>	187
13.4.4	Amplification de confidentialité	189
13.5	...et Alice et Bob peuvent enfin discuter tranquillement	190
14	Conclusion	193
V	Annexes	195
A	Formules mathématiques	197
A.1	Algèbre d'opérateurs	197
A.1.1	Propriétés élémentaires	197
A.1.2	Commutateurs de fonctions d'opérateurs	197
A.1.3	Formule de Baker–Hausdorff	198
A.2	Gaussiennes et paquets d'ondes	198
A.2.1	Intégrale d'une gaussienne	198
A.2.2	Définition et intégrale d'un paquet d'onde gaussien	199
A.2.3	Paquets d'ondes à variance complexe	200
A.3	Transformées de Fourier	200
A.3.1	Convention	200
A.3.2	Fonctions de Dirac et ondes planes	200
A.3.3	Paquets d'onde gaussiens	200
A.3.4	Produit scalaire	201
B	Propriétés quantiques élémentaires des variables continues	203
B.1	Choix des observables	203
B.2	États propres de $\hat{Q}$ et fonction d'onde	203
B.3	Action de l'opérateur impulsion	205
B.4	Fonction d'onde de $ P\rangle = p$	205
B.5	Relations de Fourier entre impulsion et position	206
B.6	Relation d'incertitude	207
B.7	Opérateur translation et spectre de $\hat{Q}$	208
B.8	Opérateur déplacement	209
B.9	Opérateur amplitude	209
	Bibliographie	211

Table des figures

2.1	Représentation des observables « normales » et « gelées »	32
2.2	Fonction d'onde d'un état minimal	36
2.3	Distribution de probabilité en Q de la superposition linéaire de deux états cohérents	41
2.4	Distribution de probabilité en P comparées d'une superposition linéaire et d'un mélange statistique de deux états cohérents	42
3.1	Représentation de $W_{ q_0\rangle\langle q_0 }(Q, P)$	51
3.2	Représentation schématique de la fonction de Wigner d'un état gaussien . . .	53
3.3	Lame partiellement réfléchissante	57
3.4	Valeur à l'origine de la fonction de Wigner d'états de Fock en fonction des pertes η	60
4.1	Valeurs autorisées de $G\chi$ en fonction de G	66
4.2	Modélisation d'une absorption	66
5.1	Schématisation d'une détection homodyne	76
5.2	Schématisation du système utilisé pour la soustraction des photocourants . .	81
5.3	Schématisation d'un système linéaire	85
5.4	Réponse percutonnelle expérimentale	87
5.5	Réponse du système de détection homodyne finalement utilisé à des impulsions de 120 ns	88
5.6	Transimpédance déduite du bruit et transformée de Fourier de la réponse percutonnelle	89
5.7	Transimpédance complexe de la chaîne amplificatrice	90
5.8	Transformée de Fourier inverse de la transimpédance complexe	91
5.9	Mesures expérimentales et estimation théorique du bruit de photons	92
6.1	Modélisation de l'inefficacité d'une détection homodyne	94
6.2	Modélisation d'une détection homodyne avec des photodiodes d'efficacité η .	95
6.3	Mesure expérimentale d'un état cohérent	101
6.4	Représentation polaire du couplage amplitude-phase du modulateur	101
7.1	Information mutuelle I_{AB} en fonction de la transmission équivalente T_{eq} . . .	113
7.2	Contributions des diverses fluctuations à l'information mutuelle	114
8.1	Schema d'un duplicateur quantique	118
8.2	Implémentation d'un duplicateur quantique	120
8.3	Schéma d'un duplicateur asymétrique	121

8.4	Domaines d'existence des duplicateurs optimaux	122
8.5	Valeurs possibles du coefficient μ	124
11.1	Système de communication quantique entre Alice et Bob	145
11.2	Stratégie possible d'espionnage	146
11.3	Schématisation des différentes variances de la modulation symétrisée d'Alice	147
11.4	Alice	148
11.5	Bob	148
11.6	Évolution des informations mutuelles en fonction des pertes équivalentes T_{eq}	149
11.7	Contributions des différentes amplitudes de modulation aux différentes in- formations mutuelles	150
11.8	Taux d'information secrète que l'on peut atteindre avec différents protocoles directs (états maximalelement comprimés)	152
11.9	Schématisation de la modulation d'Alice dans le cas général	152
11.10	Schématisation de la modulation d'Alice pour un protocole à états cohérents	154
11.11	Schématisation des états maximalelement comprimés envoyés à Bob dans le cas d'un protocole utilisant des faisceaux intriqués	155
12.1	Espionnage par une cloneuse intriquante	160
12.2	Implémentation d'une cloneuse intriquante	164
12.3	Canaux où il est possible de faire de la cryptographie inverse	167
12.4	Les différents régimes utilisant des états comprimés dans des protocoles de réconciliation inverse (pertes seules)	169
12.5	Informations mutuelles en fonction de G (pertes seules)	171
12.6	ΔI en fonction du facteur de compression s	172
13.1	Dispositif expérimental	183
13.2	Taux d'informations mutuelles en fonction des pertes avec notre système de détection homodyne ($V = 40$)	186
A.1	Contour d'intégration dans le plan complexe	199

Liste des tableaux

9.1	Régimes de téléportation en fonction de la fidélité $\overline{\mathcal{F}}$	134
10.1	Les principes de Kerckhoffs	138
13.1	Précision des codages en virgule flottante obéissant à la norme IEEE 754–1985	180
13.2	Taux d’information secrète idéal et pratique	190

Remerciements

Les travaux présentés dans cette thèse ont été effectués dans le Groupe d'Optique Quantique du Laboratoire Charles Fabry de l'Institut d'Optique et je remercie Pierre Chavel et André Ducasse de m'y avoir accueilli.

Parallèlement à ces recherches, j'ai eu la chance d'effectuer mon initiation à l'enseignement à Sup'Optique. Je n'ai pas ici la place pour remercier toute l'équipe pédagogique, au sein de laquelle j'ai vraiment aimé travailler et je me contenterais de remercier, sans être exhaustif, ceux que j'ai le plus côtoyés : Hervé Sauer, François Balembois, Nicolas Dubreuil, Thierry Avignon, Lionel Jacubowicz, Pierre Raybaut, Philippe Delaye, et Gaëlle Lucas-Leclin. Il m'ont tous beaucoup appris, tant sur le plan pédagogique que scientifique.

Je tiens bien entendu à remercier Philippe Grangier pour avoir encadré cette thèse. Il a su m'accorder une grande autonomie tout en étant présent lorsque j'avais besoin de ses conseils : il m'a souvent impressionné en réglant en quelque heures des problèmes, expérimentaux ou théoriques, qui m'avaient fait sécher plusieurs jours, voire plusieurs semaines ! Il a en moyenne une demi-douzaine d'idées intéressantes par jour en ce qui concerne les développements de la manip (et sans doute autant pour chacune des deux autres manip !), ce qui m'a permis de toujours avoir des choses intéressantes à faire. Cette imagination prolifique nous a notamment permis de redéfinir mon sujet de thèse lorsqu'une thèse portant sur le même sujet a été soutenue en Allemagne, huit mois après le début de la mienne.

Ça a été un véritable plaisir de partager le bureau et/ou le labo de Nicolas Schlosser, Georges Reymond et Alexios Beveratos. Les déménagements successifs de la manip et de mon bureau m'ont permis de travailler avec ces trois autres thésards, et de nombreux détails de cette thèse leurs sont dûs, d'un fichier $\text{\TeX}$ pour la mise en forme de ce manuscrit à l'emprunt de miroirs et d'isolateurs optiques pour la manip, sans oublier un (long) cours avec travaux pratiques de couplage dans une fibre. Mais je tiens surtout à les remercier pour l'ambiance joyeuse qu'ils ont fait régner dans le groupe pendant ces trois années.

Gao Jiangrui a subi avec moi les vicissitudes de la détection homodyne impulsienne pendant six mois. Nous avons subi ensemble les difficultés présentées au chapitre 5 de ce manuscrit et bien d'autres. Je tiens à le remercier pour ses nombreuses suggestions qui nous ont permis, malheureusement après son départ, de construire la détection homodyne limitée au bruit de photon qui est au cœur des réalisations expérimentales de cette thèse.

C'est aussi grâce au doigté expérimental de Jérôme Wenger, alors en stage de DEA, que nous sommes parvenu à construire cette détection. Ce stage n'était qu'un début et il a décidé de faire sa thèse sur cette manip. Quasiment tous les résultats expérimentaux présentés dans cette thèse ont été obtenus avec lui. Son perfectionnisme et son sens de l'organisation ne l'empêchent pas d'être quelqu'un de très chaleureux, et ça a été un vrai bonheur de collaborer avec lui. Je lui souhaite bonne chance pour la suite de sa thèse et le remercie encore pour ces presque deux ans de travail en commun.

Je ne saurais oublier Rosa Tualle-Brouiri, qui a encadré cette thèse. Son réalisme et son franc-parler m'ont remonté le moral quand j'avais le « blues du thésard » et m'ont secoué quand j'en avais besoin. Malgré un emploi du temps plus que surchargé, elle a toujours été disponible et à l'écoute. Son encadrement, bien que discret, a été essentiel pour moi dans le déroulement de cette thèse, et je l'en remercie. Je suis impatient de retravailler avec elle à l'avenir, et je sais que ça sera avec plaisir.

Je n'oublie pas non plus les autres membres, temporaires ou non, du Groupe 19, qui en ont fait un endroit où il a été agréable de travailler : Benoît, Sergei, Mohammad, Christophe, Igor, Jean-Philippe, Theresa, Julien, Antoine, Junxiang, Gaëtan, Silvia, l'autre Jérôme et Mathieu. Je tiens aussi à remercier Alain Aide, dont les nombreux conseils m'ont été précieux, ainsi qu'André Villing : sans eux, l'électronique délicate de la manip n'aurait sans doute jamais vu le jour. Je tiens aussi à remercier Thierry Debuisschert, de la société Thales, qui nous a prêté ce qui est sans doute l'un des derniers modulateurs intégrés à 780 nm encore en état de marche, pièce essentielle de l'expérience de cryptographie.

J'ai eu le plaisir, à la fin de cette thèse, de collaborer avec Gilles Van Assche et Nicolas Cerf de l'Université Libre de Bruxelles. Cette collaboration, en plus d'être extrêmement fructueuse, fut des plus agréable.

Je remercie aussi ma petite Cam, qui m'a offert le seul photon unique que j'ai réussi à attraper pendant cette thèse, et mes parents, qui m'ont entourés et soutenus pendant ces 23 ans de scolarité, qui arrivent enfin à leur terme !

Bien entendu, je réserve mes remerciements les plus chaleureux à Anne-Gaël, qui m'a accompagné pendant ces années de thèse. Bien que non physicienne, elle a traqué impitoyablement les fautes d'orthographe et de grammaire à travers ce manuscrit abscons. Elle a aussi subi le bruit nocturne du clavier pendant la rédaction, et a veillé à ce que je dorme quand même un peu. Mais, surtout, elle a fait preuve d'un soutien inébranlable pendant les moments difficiles, et je ne saurais sans doute jamais l'en remercier assez.

Notes diverses

Résumé

Ces dernières années, les variables continues ont émergées en tant qu'alternative aux variables discrètes dans les communications quantiques. Cette thèse s'inscrit dans ce cadre des communications quantiques avec des variables continues.

Les variables continues utilisées ici sont les quadratures d'un mode du champ électromagnétique. Pour les mesurer, nous avons construit une détection homodyne équilibrée, limitée au bruit de photons, impulsionnelle et résolue en temps. Celle-ci peut effectuer 800 000 mesures par seconde.

En se fondant sur la limite de la duplication quantique, nous montrons qu'une valeur de la fidélité supérieure à $2/3$ dans un protocole de téléportation quantique garantit que l'état téléporté est la meilleure copie qui reste de l'état d'entrée.

Nous introduisons de nouveaux protocoles de distribution quantique de clef utilisant des variables quantiques continues, sûrs face à des attaques individuelles pour toute valeur de la transmission de la ligne optique entre Alice et Bob. En particulier, il n'est pas nécessaire que cette transmission soit plus grande que 50 % (moins de 3 dB de pertes). Ni compression des fluctuations quantiques, ni intrication ne sont nécessaires.

Nous avons implémenté expérimentalement ces protocoles, en utilisant la détection homodyne limitée au bruit de photon mentionnée plus haut et des états cohérents. L'extraction complète de la clef secrète est réalisée en utilisant une technique de réconciliation par tranches inversée suivie d'amplification de confidentialité. Notre dispositif expérimental produit un taux net de transmission de clef de 1,7 megabits par seconde pour une ligne sans pertes, et 75 kilobits par seconde pour une ligne avec 3,1 dB de pertes. Les limitations actuelles sont essentiellement techniques et proviennent surtout de l'efficacité limitée du logiciel de réconciliation.

Mots-clefs

Information quantique — Communication quantique — Cryptographie quantique — Téléportation quantique — Variables continues — Détection homodyne impulsionnelle — Clonage quantique

Abstract

In the last years, continuous variables emerged as an alternative to discrete variables in quantum communication. This thesis presents our work in this field of quantum communi-

cations with continuous variables.

The continuous variables we use are the quadratures of a single mode of the electromagnetic field. To measure them, we have built a pulsed time-resolved shotnoise-limited balanced homodyne detection. This detection can measure 800,000 samples per second.

Using an argument based upon 1-to-2 quantum cloning, we show that a fidelity value larger than $2/3$ in a quantum teleportation protocol warrants that the teleported state is the best possible remaining copy of the input state.

We introduce new quantum key distribution protocols using quantum continuous variables, that are secure against individual attacks for any transmission of the optical line between Alice and Bob. In particular, it is not required that this transmission is larger than 50 % (less than 3 dB losses). Neither squeezing nor entanglement are required.

We experimentally demonstrate these protocols, using the shotnoise limited homodyne detection mentioned above and coherent states. Complete secret key extraction is achieved using a reverse sliced reconciliation technique followed by privacy amplification. Our tabletop experiment yields a net key transmission rate of about 1.7 megabits per second for a loss-free line, and 75 kilobits per second for a line with losses of 3.1 dB. The present limitations are essentially technical, arising mainly from the limited efficiency of the reconciliation software.

Keywords

Quantum information — Quantum communication — Quantum cryptography — Quantum teleportation — Continuous variables — Pulsed homodyne detection — Quantum cloning

Références bibliographiques

Les notes bibliographiques sont notée [1, 2] et renvoient à la bibliographie, pages 211 et suivantes. Les différents chapitres des références [3, 4, 5, 6], abondamment citées, bénéficient de références bibliographiques individuelles.

Publications

Les travaux présentés dans cette thèse ont donné lieu à plusieurs publications.

- Notre étude des effets de l'adaptations de modes temporels dans une détection homodyne impulsionnelle (section 6.2) a été publiée dans *The European Physical Journal D* [7].
- L'étude critères de téléportation quantique avec des variables continues (chapitre 9) a fait l'objet de deux notes [8, 9] sur www.arXiv.org et d'une *rapid communication* [10] dans *Physical Review A*.
- Les protocoles de cryptographie directe à variables continues présentés au chapitre 11 ont été publiés dans *Physical Review Letters* [11, 12].
- Les protocoles inverses, présentés au chapitre 12, ont fait l'objet d'une note sur www.arXiv.org et d'une communication invitée à la conférence QCMC'02[13].
- La démonstration expérimentale de cryptographie quantique présentée au chapitre 13 est décrite dans un article publié dans *Nature* [14, 15].

Chapitre 1

Introduction

La mécanique quantique est née au début du xx^e siècle de l'étude de l'interaction lumière-matière. En a émergé une physique très riche, s'appliquant rapidement à une grande variété de systèmes, des particules subatomiques aux étoiles, en passant par la physique nucléaire, la physique atomique et la physique du solide. Malgré ses succès impressionnants, la physique quantique suscite toujours un certain malaise, tant ses bases sont peu intuitives et semblent contraires au bon sens.

Ses prédictions paradoxales, toujours vérifiées à ce jour, ont d'ailleurs suscité de nombreuses interprétations philosophiques plus ou moins heureuses. Au cours des vingt dernières années, ces paradoxes ont été réinterprétés en termes de complexité algorithmique et de quantité d'information. Cette réinterprétation a posé les fondements d'un domaine de la physique quantique riche et en plein développement : l'information quantique. D'abord limité à l'étude des systèmes à deux niveaux (spins $\frac{1}{2}$, photons uniques ou qubits), ce domaine s'est récemment étendu à l'étude des variables continues, cadre dans lequel s'inscrit cette thèse.

1.1 Histoire et paradoxes de la mécanique quantique

Cette section constitue un exposé rapide de l'histoire de la physique quantique et de ses paradoxes. Pour un exposé plus complet, on pourra lire par exemple les références [17, 18, 19, 20]. En ce qui concerne les paradoxes de la mécanique quantique, je ne peux que recommander les bijoux de vulgarisation que constituent les livres de Gamow [23, 24].

1.1.1 De la vieille mécanique des quanta à la mécanique quantique

La vieille mécanique des quanta est née en 1900 d'une quantification des échanges d'énergie entre lumière et matière. Pour décrire le rayonnement d'un corps chauffé, Planck fut contraint de postuler [27] que l'émission et l'absorption de lumière de fréquence ν par un corps ne peut se faire que par paquets d'énergie $h\nu$, où h est une constante (depuis dénommée *constante de Planck*). On trouve là une caractéristique de la vieille mécanique des quanta, qui repose sur une description discrète de ce que les physiciens croyaient savoir être des quantités continues. Ce postulat, au départ considéré comme un artifice mathématique malheureusement nécessaire, allait se révéler d'un sens profond sur la nature du rayonnement.

Cette théorie s'est développée au cours du premier quart du XX^e siècle, avec un formalisme qui semblait très arbitraire : il s'agissait en général de résoudre un problème avec la mécanique classique, puis de restreindre l'ensemble des solutions au moyen de règles de quantifications plus ou moins *ad hoc* [17], qui consistaient à n'autoriser que les trajectoires dont l'action¹ était un multiple entier de la constante h , qui est parfois appelée *quantum d'action*. Malgré cet arbitraire peu satisfaisant, la vieille mécanique des quanta résolut un certain nombre de problèmes qui résistaient à la mécanique classique. Bohr réussit ainsi en 1913 à construire un modèle planétaire de l'atome d'hydrogène, amélioré en 1916 par Sommerfeld, reproduisant son spectre connu expérimentalement.

De même, la description de la lumière en terme de « grains de lumière » ou *photons* permet à Einstein d'expliquer en 1905 l'effet photo-électrique, déjà à la base de nombreux développements technologiques avant-guerre [28]. Cette vision corpusculaire de la lumière est confirmée en 1923 par l'observation de l'effet Compton, qui peut s'interpréter comme une collision entre un électron et un photon. Elle restait pourtant contradictoire avec les expériences d'interférences qui avaient conduit à l'adoption de la théorie ondulatoire de la lumière au début du XIX^e siècle : suivant l'expérience décrite, il fallait donc décrire la lumière en tant qu'onde ou corpuscule.

Dans sa thèse, soutenue en 1923, de Broglie étend la dualité onde-corpuscule de la lumière aux électrons et aux autres corps matériels en montrant [29] que les conditions de quantification *ad hoc* introduites par Bohr et Sommerfeld pour décrire le mouvement des électrons correspondent à l'existence d'ondes stationnaires, si on associe à chaque corps matériel d'impulsion $p = mv$ une onde de longueur d'onde $\lambda = \frac{h}{p}$. L'existence de ces ondes est assez rapidement vérifiée par Davisson et Germer [30], qui observent la diffraction d'un faisceau d'électrons sur le réseau cristallin du nickel en 1927.

La mécanique quantique proprement dite naît en 1925, lorsque les suppositions plus ou moins arbitraires de la vieille mécanique des quanta sont déduites d'un formalisme cohérent qui résout notamment la contradiction apparente entre les caractères ondulatoires et corpusculaires. Ce formalisme, publié en 1925 par Heisenberg, Born et Jordan, est fondé sur l'algèbre matricielle, alors peu connue des physiciens et suscite un certain rejet par son caractère peu intuitif. Schrödinger publie en 1926 l'équation d'onde vérifiée par les ondes de de Broglie, l'*équation de Schrödinger* et en déduit une description de la mécanique quantique plus intuitive pour les physiciens, habitués à utiliser des équations d'onde depuis le XIX^e siècle. Il démontre la même année que ce formalisme est équivalent à la mécanique des matrices, ce qui n'empêchera pas les partisans de chaque formalisme de s'affronter pour des raisons philosophiques (et parfois personnelles) [20].

Si ces formalismes permettent d'unifier en une vision cohérente les différents aspects apparemment contradictoires de la vieille mécanique des quanta, ils ont un certain nombre de conséquences peu intuitives, et même contradictoires avec la conception qu'on avait jusque là de la physique. La plus célèbre de ces conséquences est le *principe d'incertitude*, découvert par Heisenberg en 1927.

¹L'action d'une trajectoire est définie par l'intégrale curviligne de l'impulsion $\oint p \, dq$ sur cette trajectoire.

1.1.2 La mécanique quantique et ses paradoxes

La plus connue des propriétés paradoxales² de la mécanique quantique le principe d'incertitude de Heisenberg, qui stipule qu'il est impossible de connaître simultanément et parfaitement la position Q et l'impulsion P d'une particule. Plus précisément, si ΔQ (ΔP) représente l'écart type (« l'incertitude ») de la mesure de Q (P), on a

$$\Delta Q \Delta P \geq \frac{\hbar}{2}. \quad (1.1)$$

Cet produit d'incertitudes est une limitation fondamentale et il n'y a aucun moyen d'aller au delà.

Cela est souvent justifié [20, 23, 28], en suivant l'exemple de Heisenberg et de son microscope à rayons gamma³, par le fait que « toute mesure perturbe le système mesuré », et qu'il existe une sorte de perturbation minimale. Heisenberg a introduit le principe d'incertitude au moyen de l'expérience de pensée suivante : on essaye de mesurer la position et l'impulsion d'un électron au moyen d'un microscope à rayons gamma. Pour déterminer la position de l'électron, on l'éclaire avec un photon de longueur d'onde λ . La précision du microscope est limitée par $\Delta q \simeq \lambda$. Malheureusement, le photon interagit avec l'électron et lui transmet une partie de sa quantité de mouvement, de l'ordre de h/λ à l'électron, ce qui induit une incertitude $\Delta p \simeq \frac{h}{\lambda}$ sur toute mesure de l'impulsion de l'électron. On a alors $\Delta q \cdot \Delta p \simeq h$, le choix de privilégier la position ou l'impulsion pouvant se faire par le choix de la longueur d'onde λ .

Un autre point de vue sur le principe d'incertitude est une approche ondulatoire [34, 35]. Si une particule a une impulsion p parfaitement déterminée ($\Delta P = 0$), elle est associée à une onde de de Broglie plane de longueur d'onde $\lambda = \frac{h}{p}$, qu'on peut écrire sous la forme

$$\cos 2\pi \frac{q}{\lambda} = \cos \frac{2\pi}{h} qp. \quad (1.2)$$

Cette onde est définie de $-\infty$ à $+\infty$ et la position de cette particule est parfaitement indéterminée ($\Delta q = \infty$). Pour décrire une particule localisée, il faut la décrire par un paquet d'ondes, qui a une extension finie dans l'espace. Un tel paquet d'ondes s'obtient en superposant des ondes planes de fréquences différentes, de sorte qu'elles ne restent en phase que sur une longueur limitée et interfèrent destructivement. Deux ondes correspondant à deux impulsions décalées de ΔP seront déphasées de π et interféreront destructivement au bout d'une distance $\frac{\Delta Q}{2}$ vérifiant

$$\frac{2\pi}{h} \frac{\Delta Q}{2} (p + \Delta P) = \frac{2\pi}{h} \frac{\Delta Q}{2} p + \pi. \quad (1.3)$$

On a donc

$$\Delta P \Delta Q \simeq h. \quad (1.4)$$

²Ici, le mot *paradoxe* est employé dans son sens premier, *qui heurte le bon sens* [31, 32], notamment le « bon sens » hérité de l'expérience quotidienne et de la physique classique. Il ne suppose notamment pas l'idée de contradiction, la mécanique quantique étant cohérente.

³Une étude plus détaillée du microscope à rayons gamma, qui tient notamment compte de l'ouverture numérique du microscope, se trouve dans la référence [33].

Des calculs d'analyse de Fourier plus approfondis [36] permettent de justifier cette égalité approximative de retrouver le principe d'incertitude de Heisenberg (1.1) dans le cas le plus général.

La différence fondamentale entre cette relation et la relation analogue qui existe pour tout phénomène ondulatoire classique réside dans la nature de l'onde, qui détermine une probabilité : si rien n'empêche, au moins en principe, de mesurer l'amplitude d'une onde classique en tout point et de déterminer parfaitement son étalement en position et en impulsion, l'onde quantique définit la probabilité de trouver une particule à un endroit ou à un autre, et une fois cette position mesurée, on ne peut plus rien mesurer de cette particule. Ce caractère intrinsèquement probabiliste de la mécanique quantique constitue donc les fondements du principe d'incertitude et il serait plus exact de dire que la position et l'impulsion d'un électron ne sont pas définies simultanément, que de sous-entendre qu'elles existent, mais ne sont pas mesurables.

La théorie quantique définit des probabilités de résultats de mesure par des calculs qui n'ont rien à voir avec des calculs de probabilités classiques, mais qui reposent sur ces ondes quantiques de probabilités décrites ci dessus. En examinant le comportement des ondes de probabilités décrivant les systèmes à plusieurs particules, Einstein, Podolsky et Rosen ont soulevé un autre paradoxe de la mécanique quantique, le paradoxe EPR [37]. Ils ont montré que deux particules ayant interagi entre elles ne peuvent plus être décrites par deux fonctions d'ondes distinctes lorsqu'on les éloigne l'une de l'autre, sauf à supposer l'existence d'une « interaction fantomatique à distance » instantanée. Ils en ont déduit que la mécanique quantique était incomplète car elle n'était pas « réaliste locale ». 30 ans plus tard, en 1964, Bell a démontré [38] que si l'on cherchait à construire une théorie classique, déterministe ou probabiliste, reproduisant les résultats de la mécanique quantique, cette théorie devrait permettre des communications instantanées à distance. En d'autres termes, que toute tentative de compléter la mécanique quantique par une théorie réaliste locale (qui aurait satisfait Einstein, Podolsky et Rosen) était vouée à l'échec.

Ces caractéristiques de la mécanique quantique ont soulevé beaucoup d'oppositions, y compris parmi ceux qui ont contribué à sa naissance. Schrödinger était tellement gêné par les développements de la physique quantique qu'il a préféré quitter le domaine [40]. Bien sûr, la figure la plus emblématique de ce refus est Einstein. Son opposition était bien plus profonde et argumentée que ne pourrait le laisser croire sa fameuse réplique « Dieu ne joue pas aux dés ». Il comprenait très bien la mécanique quantique et rejetait certaines de ses implications. La plupart des paradoxes « classiques » de la mécanique quantique ont été soulevés par Einstein, comme le paradoxe EPR [37] que nous avons mentionné plus haut, et celui du chat de Schrödinger.

Ce dernier paradoxe, est lié au problème de la mesure. Ce problème résulte du fait que la théorie quantique décrit la mesure comme un processus distinct de l'évolution du système quantique, sans définir précisément ce qu'est une mesure. Faut-il considérer le chat comme effectuant la mesure ou faisant partie du système ? Certains, parmi lesquels Wigner [18], ont vu dans l'acte de mesurer une action de l'esprit sur la matière, ce qui a ouvert la porte à de nombreuses interprétations fumeuses de la mécanique quantique⁴. Ce problème est aujourd'hui compris en terme d'interaction avec l'environnement [43], plutôt qu'avec un esprit conscient.

⁴ On m'a ainsi affirmé avec assurance que « maintenant, avec la mécanique quantique, on a démontré la réincarnation » !

Si la mécanique quantique s'est imposée malgré ces difficultés apparentes, c'est qu'elle a parallèlement résolu un nombre impressionnant de problèmes, tant fondamentaux qu'appliqués. Dès les années 1920, son domaine d'application allait des électrons aux étoiles, et il s'étend aujourd'hui de la physique des particules à la cosmologie. Les applications technologiques de la physique quantique n'ont pas tardé non plus ; ainsi Leprince-Ringuet cite en 1940 [28] l'importance de l'effet photoélectrique, notamment dans les développements en cours de la télévision. Depuis, le nombre de technologies reposant sur la physique quantique s'est tellement multiplié qu'il serait fastidieux d'en dresser une liste. Citons seulement certains développements marquants comme l'énergie nucléaire, l'électronique moderne et le laser.

De plus, les situations paradoxales évoquées plus haut n'apparaissaient que dans des expériences de pensée qui semblaient difficilement réalisables, voire impossibles. Certaines de ces expériences ont pourtant été construites à partir des années 1970, et la mécanique quantique a jusqu'à présent passé tous les tests expérimentaux avec succès, contrairement aux théories alternatives, en général réalistes locales, supposées plus « raisonnables ».

1.2 Naissance de l'information quantique

La mécanique quantique, et plus particulièrement le problème de la mesure, se prête extrêmement bien à une description en termes informationnels [44, 45, 46, 47]. Les errements « spiritualistes » évoqués plus hauts sont sans doute liés à l'association abusive du concept d'information à celui d'esprit.

L'information est rigoureusement définie et quantifiée autour de la seconde guerre mondiale [48, 49, 50], en s'affranchissant de toute subjectivité. S'en est suivi un effet de mode cherchant à appliquer la « cybernétique » à peu près partout, des sciences exactes aux sciences humaines, avec plus ou moins de bonheur [48]. Il est rétrospectivement curieux de constater que la physique quantique, regorgeant de paradoxes liés à l'information, s'est tenue relativement à l'écart de la théorie de l'information jusqu'à récemment.

L'information quantique proprement dite ne naît en effet qu'au début des années 1980, époque où les physiciens utilisent quotidiennement des ordinateurs et considèrent l'information comme une quantité concrète, mesurable en bits et en octets. L'époque a sans doute un rôle déterminant, car les deux publications fondant les deux domaines habituellement regroupés sous le nom d'information quantique paraissent quasiment simultanément, en 1982 et 1983.

1.2.1 Ordinateurs quantiques

Le premier de ces domaines, l'informatique quantique [44, 45, 46, 47], ne sera qu'évoqué ici. Feynmann découvre en 1982 [51] que certains systèmes quantiques sont exponentiellement difficiles à simuler sur un ordinateur classique, mais qu'il est néanmoins possible de les simuler en un temps raisonnable avec d'autres systèmes quantiques. La difficulté pour un système classique de simuler un système quantique vient de la taille de l'espace de Hilbert de ce dernier : en effet, un système de n spins $\frac{1}{2}$ (ou qubits) est décrit par un vecteur dans un espace de Hilbert de dimension 2^n , alors qu'un registre de n bits classiques évolue dans un espace de dimension n seulement. Cette différence induit une sorte de parallélisme

massif du système quantique, celui-ci étant en quelque sorte dans tous les états à la fois au cours de son évolution.

Un ordinateur quantique est un système qui exploite ce parallélisme [52] pour résoudre un problème, que ce soit la simulation d'un autre système quantique ou un problème classique. Ce gain exponentiel par rapport à l'ordinateur classique n'existe malheureusement pas pour tous les algorithmes, et les premiers algorithmes utilisant des ordinateurs quantiques étaient assez artificiels, mais Shor a trouvé en 1994 un algorithme quantique rapide de factorisation d'un nombre en ses facteurs premiers [53] qui a relancé l'intérêt du domaine. Un autre algorithme où un ordinateur quantique est plus rapide qu'un homologue classique est la recherche dans une base de données [54, 55], mais le gain n'est que quadratique et non plus exponentiel.

La construction d'un ordinateur quantique se heurte à des difficultés expérimentales impressionnantes, mais elle fournit un cadre conceptuel intéressant pour construire des expériences mettant en évidence des effets quantiques. Réussir à construire un ordinateur quantique reviendrait en quelque sorte à construire un système capable de mettre en oeuvre toutes les expériences de pensée imaginables. L'ordinateur quantique le plus puissant construit à ce jour est une molécule dont 7 atomes adressés par résonance magnétique nucléaire (RMN) forment les 7 qubits [56]. Cet ordinateur est capable de factoriser le nombre 15 en ses facteurs premiers, alors que le record avec des ordinateurs classiques est un nombre de 158 chiffres [57]. D'autres candidats pour fabriquer à (long) terme des ordinateurs quantiques sont des atomes individuels dans des pièges dipolaires [59, 60, 61, 62] ou des ions individuels dans des pièges électrostatiques [63], pour n'en citer que quelques uns.

1.2.2 Communication quantique

Si l'informatique quantique est le pendant quantique de l'informatique classique, la communication quantique est le pendant de la théorie de l'information. On peut qualifier de dispositif de communication quantique tout système qui exploite le principe d'incertitude [64, 46, 47] ou l'intrication quantique.⁵

Le premier à considérer le principe d'incertitude comme une ressource plutôt que comme une limitation semble être Wiesner [65], qui imagine vers 1969 des billets de banque infalsifiables à base de spins $\frac{1}{2}$. Cette idée ne sera malheureusement pas publiée avant 1983. Chaque billet a un numéro de série et une série de spins « stockés » sur le billet. L'orientation de chaque spin est connue de la banque seule, qui reconnaît le billet à son numéro de série. Le principe d'incertitude empêche les faux-monnayeurs de mesurer l'orientation des spins sans erreur et de les copier. Par contre, il n'empêche pas le banquier de vérifier que les photons sont bien dans l'état voulu, ce qui lui permet de vérifier l'authenticité des billets. Ce procédé intéressant permet en théorie de faire des billets de banques incopiables, mais semble condamné à rester une statue d'expérience de pensée. Le stockage de spins $\frac{1}{2}$ (ou d'autres formes de qubits) pendant des durées assez grandes (de l'ordre de l'heure) pour présenter un intérêt semble en effet un objectif hors de portée, même pour les rêveurs les plus optimistes.

⁵Les systèmes de tests d'inégalité de Bell sont en général exclus de cette définition, car l'intrication est plus mise en évidence qu'exploitée. Les systèmes de mesure quantique non-destructive aussi, car le principe d'incertitude n'y est pas exploité, mais contourné. Ils sont cependant apparentés à la communication quantique, qui leur emprunte nombre d'outils théoriques et expérimentaux.

Ce codage en polarisation a été repris par Bennett et Brassard en 1984 [66, 64, 67] pour un protocole de distribution quantique de clés, qui semblait, lui, plus utile et plus réaliste. Ce protocole, appelé BB84, permet à deux partenaires, Alice et Bob, de se communiquer une clé cryptographique secrète. Au lieu de stocker les photons sur un billet, Alice les envoie à Bob, qui les mesure immédiatement. Comme le faux-monnayeur mentionné plus haut, Bob se trompera de base la moitié du temps, mais ces erreurs seront simplement éliminées par une conversation publique entre Alice et Bob. Par contre, les erreurs induites par un espion éventuel seront détectées, ce qui garantit la confidentialité de la communication.

Ce protocole a été assez rapidement mis en œuvre, d’abord sous la forme de démonstration de principe, puis de dispositifs de plus en plus opérationnels [68, 69]. La distribution quantique de clés, souvent simplement appelée cryptographie quantique, est à ce jour le seul domaine de l’information quantique où des systèmes commerciaux sont disponibles. Certaines rumeurs difficilement vérifiables font également état de dispositifs déployés entre les bâtiments d’agences gouvernementales à Washington. La mise au point de ces systèmes en présence de bruit a également stimulé des recherches d’algorithmes classiques de traitement de l’information utiles en cryptographie quantique, comme les algorithmes de réconciliation et d’amplification de confidentialité [70, 71, 72, 73, 74].

D’autres protocoles de distribution de clé ont été proposés, que ce soit avec des qubits [68], ou plus récemment avec des variables continues (voir le chapitre 10 pour une brève revue). D’autres applications cryptographiques ont été étudiées, comme le partage de secret ou le problème de l’accord byzantin. L’exploitation plus spécifique de l’intrication quantique a également conduit au codage dense [75], à la téléportation quantique [76] et à d’autres applications.

1.3 Qubits et variables continues

1.3.1 Variables discrètes

Les systèmes quantiques à deux états, comme le spin $\frac{1}{2}$, l’atome à deux niveaux ou la polarisation d’un photon unique, ont été utilisés très tôt comme systèmes modèles en mécanique quantique. Ce sont en effet les plus simples des systèmes quantiques : leur état est représenté par un vecteur dans un espace de Hilbert de dimension 2 seulement, ce qui ne les empêche pas de manifester tous les comportements paradoxaux de la mécanique quantique. L’information quantique a rebaptisé ces systèmes quantiques *qubits*, ou bits quantiques, par analogie au bit, quantité élémentaire d’information introduite par Shannon en 1948 [50]. Les registres des ordinateurs quantiques sont constitués de qubits, et les protocoles de cryptographie quantique étaient initialement décrits avec la polarisation de photons uniques.

On a souvent dit et pensé que l’unicité de ces photons était essentielle pour la communication quantique, car elle est essentielle pour que ces photons soient de vrais qubits. Si c’est vrai pour les protocoles généralement considérés, notamment BB84, ça n’est pas une généralité : il existe notamment des protocoles fondés sur les variables continues, adaptés aux faisceaux (relativement) intenses, notamment ceux que nous avons développés au cours de cette thèse.

Les variables continues sont arrivées assez vite en cryptographie quantique, mais sur des systèmes simulant des variables discrètes. Le codage en polarisation canonique de BB84 n’est pas adapté aux expériences de cryptographie quantique par fibres optiques, en rai-

son de la biréfringence des fibres. Ces expériences utilisent en général un codage en temps ou en fréquence, beaucoup plus robuste [68]. Même si le temps et la fréquence sont des variables intrinsèquement continues, il s'agit d'un codage discret, qui exige toujours des photons unique.

1.3.2 Variables quantiques continues

Si le temps et la fréquence (ou l'énergie) sont parfois commodes à utiliser sur un plan expérimental, ces variables ont un statut un peu particulier qui rend leur usage un peu délicat. On préfère en général formuler les problèmes impliquant des variables continues en se référant aux variables continues « canoniques » que sont la position (notée Q) et l'impulsion (notée P) d'une particule dans un espace à une dimension. Ces variables continues, étudiées dès l'origine de la mécanique quantique, ont des analogues optiques plus fréquemment utilisés lors de réalisations expérimentales. Ainsi, les expériences réalisées au cours de cette thèse se font sur les quadratures du champ électromagnétique, qui peuvent être l'intensité et la phase, par exemple. D'autres expériences utilisent la polarisation de faisceaux lumineux [77, 78] ou de nuages atomiques polarisés par pompe optique.

La richesse des variables continues réside dans leur capacité à prendre un nombre infini de valeurs sur $\mathbb{R}$. Elles définissent alors un espace de Hilbert de dimension infinie, beaucoup plus riche que l'espace de dimension 2 (ou 2^n) considéré avec les qubits, même si on se restreint souvent au sous-ensemble des états gaussiens.

Ces variables sont étudiées depuis longtemps en optique quantique, tant sur un plan expérimental que théorique. Elles sont au cœur de la problématique de nombreux domaines d'études, comme les mesures quantiques non destructives [79, 80, 81, 82, 83, 84, 85], les communications cohérentes, l'étude du bruit des lasers [86, 87, 88], etc. Dès 1994, Vaidman propose [89] une adaptation aux variables continues du protocole de téléportation quantique de Bennett et ses collaborateurs [76]. Un protocole expérimentalement plus réaliste a été proposé [90] et réalisé [91] en 1998 par l'équipe de Kimble. Ces articles sont en général considérés comme le début de la communication quantique avec des variables continues et ils ont été suivis de nombreux autres travaux théoriques et expérimentaux.

1.3.3 Avantages technologiques des variables continues

L'étude des systèmes à variables continues, un peu plus délicate sur le plan théorique, présente de nombreux avantages expérimentaux. Les qubits ont en effet deux inconvénients : ils sont difficiles à fabriquer et difficiles à mesurer.

Les protocoles à base de qubits exigent des sources de photons uniques, parfois limitées par Fourier, délicates à réaliser. Si des sources satisfaisantes ont été développées ces dernières années [92, 93, 94, 95, 96], la plupart des systèmes étudiés actuellement simulent les photons uniques par des impulsions cohérentes fortement atténuées. Les protocoles théoriques à base de variables continues utilisent en général des états comprimés, plus simples à fabriquer que les photons uniques, voire des états cohérents, comme nous l'avons montré au cours de cette thèse [11, 13], ce qui nous a permis d'effectuer la première expérience complète de cryptographie quantique avec des variables continues [14].

Mais l'avantage essentiel des variables continues se place au niveau de la détection. Les systèmes à base de photons uniques doivent en effet utiliser des détecteurs sensibles au quantum individuel, en général des photodiodes à avalanche. En général, ces détecteurs ont une

efficacité quantique limitée, sont chers et assez lents, notamment dans le domaine des longueurs d'ondes télécoms, où ils ne sont pas disponibles commercialement. La mesure des variables continues peut utiliser des photodiodes usuelles, ce qui permet de construire des systèmes de détection interférométrique rapides, efficaces et bon marché.

1.4 Plan de la thèse

Cette thèse est divisée en quatre parties.

Nous commencerons cette thèse en introduisant, dans la partie I, un certain nombre d'outils utiles pour décrire les variables continues en mécanique quantique. Le chapitre 2 est un rappel des propriétés des variables continues en mécanique quantique, qui nous permettra notamment d'introduire les notations que nous utiliserons tout au long de cette thèse. Nous définirons ensuite la fonction de Wigner au chapitre 3, qui est une distribution dans l'espace des phases. Pour des états gaussiens, cette distribution peut être assimilée à une distribution de probabilité, ce qui nous permettra d'introduire un formalisme adapté aux états gaussiens, au chapitre 4 et de simplifier nettement les calculs pour ces états. Nous nous restreindrons dans toute la suite à l'étude de ces états.

Pour mesurer les quadratures du champ électromagnétique, variables continues que nous utiliserons dans cette thèse, nous avons construit une détection homodyne impulsionnelle que nous décrirons dans la partie II. Nous présenterons dans le chapitre 5 les principes de fonctionnement d'un tel système de détection et les difficultés que posent sa réalisation expérimentale. Nous y présentons également les résultats de mesures homodynes du vide quantique qui montrent que notre détection est équilibrée et limitée au bruit de photons. Le chapitre 6 est consacré à l'efficacité de cette détection et à sa stabilité en phase. Il se termine par des mesures expérimentales d'états cohérents.

Nous aborderons la communication quantique proprement dite dans la partie III, qui commence par le chapitre 7, où l'on rappelle les limites imposées par la théorie de l'information à la transmission d'information avec des variables gaussiennes continues. Nous étudierons ensuite, au chapitre 8, les limites imposées par la mécanique quantique au clonage de variables continues, et nous appliquerons ces résultats au chapitre 9 pour définir un critère de téléportation quantique avec des variables continues.

Enfin, la cryptographie quantique sera l'objet de la partie IV. Nous commencerons, au chapitre 10, par définir ce qu'on attend d'un protocole de cryptographie quantique et faire un état des lieux des différents protocoles de cryptographie quantique avec des variables continues qui étaient publiés au début de cette étude. Nous exposerons au chapitre 11, la famille de protocoles directs que nous avons développée, qui fonctionne notamment avec des états cohérents. Nous présenterons ensuite, au chapitre 12 les protocoles inverses, qui en sont une variante robuste aux pertes, et qui fonctionne pour des pertes arbitrairement grandes. Nous terminerons cette étude par une description, au chapitre 13, d'une implémentation expérimentale de ces protocoles.

Première partie

Variables Continues

Chapitre 2

Variables continues en mécanique quantique

Les variables utilisées dans le contexte de la théorie de l'information (quantique ou classique) sont généralement discrètes, c'est-à-dire que les valeurs qu'elles peuvent prendre sont définies par des entiers, ce qui présente un certain nombre d'avantages théoriques et pratiques. Cette thèse s'inscrit dans une démarche différente : elle repose sur l'utilisations de variables continues, qui peuvent prendre toutes les valeurs d'un intervalle de $\mathbb{R}$.

Nous définirons dans la section 2.1 les variables quantiques continues utilisées dans cette thèse. La section 2.2 décrits les états usuels de l'optique quantique et leurs propriétés élémentaires¹. Les propriétés de bases des variables continues sont redémontrées rapidement dans l'annexe B.

2.1 Exemples de variables continues

Les variables continues auxquelles on se référera dans l'ensemble de cette thèse seront appelées *position* et *impulsion* et notées respectivement P et Q . Elles peuvent être effectivement vues comme la position et l'impulsion d'une particule dans une direction de l'espace, mais ces dénominations sont essentiellement historiques et les réalisations expérimentales utilisent en général d'autres variables, comme on le verra ici.

2.1.1 Rôle historique de la position et de l'impulsion

La position et l'impulsion sont en effet les variables continues « canoniques » en mécanique quantique. Ces deux variables ont un rôle analogue en mécanique lagrangienne [98, 99, 100], où l'on parle souvent d'*impulsion généralisée*. Ces similitudes ne sont pas fortuites, vu les liens très étroits entre la vieille théorie des quanta et la mécanique lagrangienne. Ces théories décrivant à leur origine les mouvements de particules ponctuelles, les dénominations de position et d'impulsion sont souvent restées pour décrire des variables qui n'ont parfois qu'un rapport mathématique avec la position ou l'impulsion d'une particule.

¹Toutes les notations introduites ici ainsi que les formules mathématiques utilisées sont rappelées dans l'annexe A

L'enseignement de la mécanique quantique commence très souvent par l'étude de ces variables [3, 4, 5]. L'analogie formelle employée ici a pour but essentiel de clarifier les choses par l'emploi d'une notation familière au lecteur.

2.1.2 Commutateurs

En mécanique quantique, les grandeurs mesurables, appelées *observables*, sont décrites par des opérateurs hermitiens. On notera $\hat{Q}$ et $\hat{P}$ les opérateurs associés aux observables Q et P . Toutes leurs propriétés sont décrites par les propriétés algébriques de ces opérateurs, qui sont elles-mêmes entièrement déterminées par leur commutateur

$$[\hat{Q}, \hat{P}] \equiv \hat{Q} \hat{P} - \hat{P} \hat{Q}. \quad (2.1)$$

Dans le cas de la position et de l'impulsion d'une particule matérielle, ce commutateur est scalaire et non nul :

$$[\hat{Q}, \hat{P}] = i\hbar. \quad (2.2)$$

La relation d'incertitude (1.1) et la plupart des propriétés de la position et de l'impulsion découlent de cette relation de commutation, qui fixe leurs propriétés algébriques, comme nous le rappelons dans l'annexe B.

Les variables que nous étudierons devront donc avoir une relation de commutation similaire, c'est à dire scalaire :

$$[\hat{Q}, \hat{P}] \equiv 2iN_0 \quad \text{avec } N_0 \in \mathbb{R}_+^*. \quad (2.3)$$

N_0 représente un facteur d'échelle, qui dépend des variables considérées et du système d'unités choisi. Le commutateur est nécessairement imaginaire pur, puisque les opérateurs considérés sont hermitiens. L'inégalité (1.1) devient alors

$$\Delta P \Delta Q \geq N_0 \quad (2.4)$$

(une démonstration de cette inégalité est donnée dans la section B.6).

Si P et Q sont exprimés dans les mêmes unités et ont la même variance, ce qui est souvent le cas expérimentalement, N_0 représente la variance de P et de Q :

$$N_0 = \Delta P^2 = \Delta Q^2. \quad (2.5)$$

N_0 est souvent appelé le *bruit quantique standard* ou, suivant le contexte expérimental, *bruit de photons* ou *bruit de grenaille* (*shotnoise* en anglais).

Dans la littérature, certains expérimentateurs ont tendance à poser $N_0 = 1$, car la mesure du bruit quantique standard sert d'étalon, alors que la plupart des théoriciens ont tendance à poser $N_0 = \frac{1}{2}$, car cela correspond au système d'unités où $[\hat{Q}, \hat{P}] = i$ et, souvent, où $\hbar = 1$. Cette double convention étant souvent source de confusions, nous avons préféré conserver ce paramètre explicite par souci de clarté, même si les notations s'en trouvent légèrement alourdies.

2.1.3 Évolution temporelle et quadratures du champ

L'information quantique repose sur les propriétés de l'espace de Hilbert, c'est-à-dire sur les propriétés relatives des états quantiques. Leur évolution temporelle a peu d'importance, et est en général perçue comme une gêne. En pratique, on n'utilisera des hamiltoniens que pour préparer des états particuliers, et on considérera en général que les états quantiques et les observables ne varient pas pendant la durée de l'expérience. On peut définir en pratique des observables « gelées », qui ne varient pas au cours du temps, par une approche similaire à celle qui permet de passer de la représentation de Schrödinger à la représentation de Heisenberg [101].

La quasi-totalité des expériences de communication quantique utilisent le champ électromagnétique plutôt que des particules ponctuelles. L'électrodynamique quantique nous dit que chaque mode du champ électromagnétique est l'analogue strict d'un oscillateur harmonique à une dimension [102, 103, 104, 105]. Dans ce cas, Q et P peuvent désigner les champs électrique et magnétique dans un mode du champ, par exemple. Les observables « gelées » $Q^G(t)$ et $P^G(t)$ définies à partir de Q et P sont les *quadratures* du champ électromagnétique.

Dans la suite, nous supposons que Q et P s'expriment dans les mêmes unités. Pour un oscillateur harmonique de fréquence ω , on peut écrire le hamiltonien sous la forme

$$\hat{H} = \frac{\hbar\omega}{2N_0} \left(\frac{1}{2}\hat{P}^2 + \frac{1}{2}\hat{Q}^2 \right), \quad (2.6)$$

où $\hat{P}$, $\hat{Q}$ et $\hat{H}$ sont indépendants du temps dans la représentation de Schrödinger.

Le changement de variable par rapport au cas usuel de la particule ponctuelle de masse m est $Q \rightarrow \frac{1}{\sqrt{m\omega}} Q$ et $P \rightarrow \sqrt{m\omega} P$. Ce simple changement d'échelle conserve la relation de commutation (2.3) et permettra de réduire l'évolution de l'oscillateur harmonique à une simple rotation dans le plan (Q, P) .

On peut montrer que les observables des quadratures vérifient l'équation différentielle

$$\frac{d}{dt} \begin{bmatrix} \hat{Q}^G(t) \\ \hat{P}^G(t) \end{bmatrix} = \begin{bmatrix} 0 & -\omega \\ \omega & 0 \end{bmatrix} \begin{bmatrix} \hat{Q}^G(t) \\ \hat{P}^G(t) \end{bmatrix}. \quad (2.7)$$

Cette équation est la même que l'équation classique, à ceci près qu'elle concerne des opérateurs. La solution est donc la même :

$$\begin{bmatrix} \hat{Q}^G(t) \\ \hat{P}^G(t) \end{bmatrix} = \begin{bmatrix} \cos \omega(t - t_0) & -\sin \omega(t - t_0) \\ \sin \omega(t - t_0) & \cos \omega(t - t_0) \end{bmatrix} \begin{bmatrix} \hat{Q} \\ \hat{P} \end{bmatrix} \quad (2.8)$$

Les observables $Q^G(t)$ et $P^G(t)$ se déduisent donc de Q et P par une simple rotation d'angle $-\omega(t - t_0)$. Elles définissent un repère tournant dans l'espace des phases dans lequel l'état quantique est figé, car il « tourne » à la même vitesse que le repère (voir la FIG. 2.1). Ces observables, les quadratures, sont mesurables expérimentalement avec une détection homodyne, comme on le verra au chapitre 5.

Dans la suite, sauf précision contraire, on se référera à ces observables « tournantes » en omettant leur dépendance temporelle et en les désignant directement par Q et P . Comme elles suivent l'évolution des états, nous omettrons également la dépendance temporelle de ces derniers.

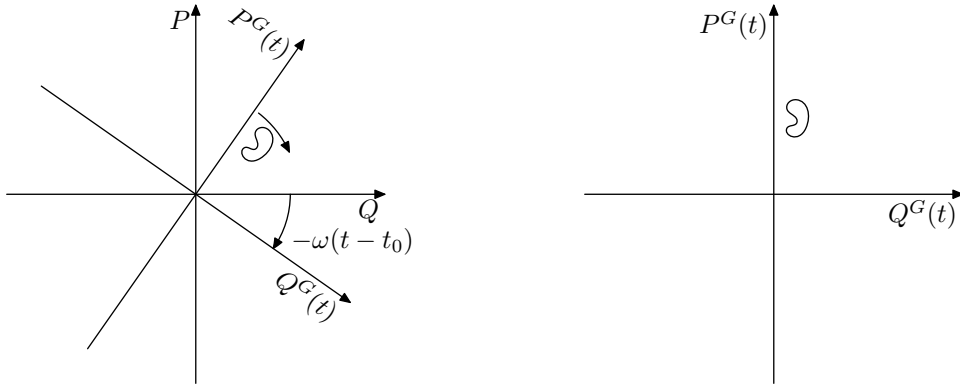


FIG. 2.1 – Représentation des observables « normales » et « gelées »

2.1.4 Autres variables continues

D'autres variables continues exploitables expérimentalement sont également analogues à Q et P . Le champ électromagnétique en présente essentiellement deux, qui peuvent être utilisées dans des expériences de cryptographie : la polarisation d'un faisceau brillant et le temps d'arrivée d'un photon unique.

2.1.4.1 Polarisation d'un faisceau brillant

Les composantes du vecteur de Stokes décrivant la polarisation d'un faisceau brillant [78, 77] sont des variables continues intéressantes expérimentalement : elles se propagent bien dans l'air qui n'est pas biréfringent et se mesurent facilement. Ce sont des analogues des composantes du moment cinétique, leurs relations de commutation ne sont pas sous la forme (2.3) mais

$$[\hat{I}_1, \hat{I}_2] = i\alpha \hat{I}_3 \quad (2.9a)$$

$$[\hat{I}_2, \hat{I}_3] = i\alpha \hat{I}_1 \quad (2.9b)$$

$$[\hat{I}_3, \hat{I}_1] = i\alpha \hat{I}_2, \quad (2.9c)$$

où $\alpha \in \mathbb{R}$. Le fait que le commutateur soit lui-même un opérateur et non un scalaire limite l'analogie avec la position et l'impulsion étudiées dans cette thèse et oblige à s'intéresser aux trois composantes du vecteur de Stokes.

Cependant, dans le cas de faisceaux brillants fortement polarisés, on peut choisir la base dans laquelle on décrit le vecteur de Stokes, de sorte que

$$\langle \hat{I}_3 \rangle \simeq I_3 \gg \hat{I}_1, \hat{I}_2. \quad (2.10)$$

Dans ce cas, l'opérateur $\hat{I}_3$ se comporte presque comme une variable classique, et commute presque avec les deux autres. Les relations de commutation s'écrivent alors sous la forme :

$$[\hat{I}_1, \hat{I}_2] \simeq i\alpha I_3 \equiv 2iN_0 \quad (2.11a)$$

$$[\hat{I}_2, \hat{I}_3] \simeq 0 \quad (2.11b)$$

$$[\hat{I}_3, \hat{I}_1] \simeq 0. \quad (2.11c)$$

L'équation (2.11a) est alors sous la forme (2.3), mais l'analogie avec la position et l'impulsion n'est qu'approximative.

Il est donc possible d'adapter à des faisceaux polarisés tous les protocoles de communication quantique décrits pour des quadratures du champ électromagnétique. On peut également définir des observables analogues pour des nuages macroscopiques d'atomes froids, qui pourraient être utilisés comme mémoires quantiques continues.

2.1.4.2 Variables temps-fréquence

Les relations entre le temps d'arrivée et la fréquence d'un photon unique sont très similaires à celles qui existent entre la position et l'impulsion. Si certaines similitudes existent sur le plan théorique, il est difficile de mesurer ces deux variables avec le même type d'appareil de mesure contrairement aux quadratures du champ électromagnétique et aux variables de polarisation.

En pratique, les protocoles qui utilisent la dualité temps-fréquence sont donc très différents des autres protocoles utilisant des variables continues. Comme ils nécessitent l'utilisation de photons uniques et se limitent en général à l'utilisation d'un sous-espace de Hilbert de dimension finie pour définir des qubits, ces protocoles sont en fait des protocoles à variables discrètes, qui n'utilisent le caractère continu du temps que par commodité expérimentale.

2.2 États quantiques

En plus des états propres de l'impulsion et de la position introduits section B.2 et utilisés pour définir les fonctions d'ondes, il existe un certain nombre d'autres états quantiques utiles, tant sur le plan expérimental que théorique. Le but de cette section est d'exposer la liste des états les plus couramment utilisés et d'en exposer certaines propriétés.

Les états présentés section 2.2.1 sont une généralisation des états propres de la position et de l'impulsion, qui nous permettront de définir des rotations dans l'espace des phases. Nous définirons ensuite (section 2.2.2) les états d'incertitude minimales, ou *états gaussiens*, omniprésents en optique quantique avec des variables continues, avant de présenter certaines de leurs propriétés (2.2.3). Nous étudierons ensuite (section 2.2.4) une superposition quantique de deux états gaussien, qui est le modèle généralement utilisé pour illustrer le paradoxe du chat de Schrödinger. Nous clorons cette liste (section 2.2.5) par les états de Fock, qui correspondent à un nombre de photons bien défini.

2.2.1 États propres de $\hat{Q}_\theta$

Comme on l'a vu section B.1, les observables $\hat{Q}$ et $\hat{P}$ ne sont que deux cas particuliers d'une famille plus générale d'observables $\hat{Q}_\theta$, correspondant aux cas $\theta = 0$ et $\theta = \frac{\pi}{2}$. L'étude des états propres $|Q\rangle_\theta = q_\theta$ de $\hat{Q}_\theta$ nous permettra aussi d'exprimer l'action du hamiltonien (2.6) de l'oscillateur harmonique, qui correspond à une rotation dans l'espace des phases comme on l'a vu section 2.1.3. Ces états propres ont les mêmes propriétés que ceux de l'opérateur position, notamment l'orthogonalité et la complétude.

Le calcul de leur fonction d'onde $\langle q | q_\theta \rangle$ est analogue à celui qu'on a utilisé pour calculer la fonction d'onde $\langle q | p \rangle$.

$$\langle q | \hat{Q}_\theta | q_\theta \rangle = q_\theta \langle q | q_\theta \rangle \quad (2.12)$$

$$\langle q | \hat{Q}_\theta | q_\theta \rangle = \langle q | (\cos \theta \hat{Q} + \sin \theta \hat{P}) | q_\theta \rangle \quad (2.13a)$$

$$= q \cos \theta \langle q | q_\theta \rangle - 2iN_0 \sin \theta \frac{d}{dq} \langle q | q_\theta \rangle \quad (2.13b)$$

La fonction d'onde $\langle q | q_\theta \rangle$ est donc solution de l'équation différentielle

$$2iN_0 \sin \theta \frac{d}{dq} \langle q | q_\theta \rangle = (q \cos \theta - q_\theta) \langle q | q_\theta \rangle \quad (2.14)$$

Si $\theta = k\pi$ ($k \in \mathbb{Z}$), $\sin \theta$ est nul et cette équation se réduit à $|Q_{k\pi} = q\rangle = |Q = (-1)^k q\rangle$. Sinon, les solutions de l'équation (2.14) sont de la forme

$$\langle q | q_\theta \rangle = \mathcal{Z} e^{-i \frac{\cos \theta q^2 - 2q q_\theta}{4N_0 \sin \theta}} \quad (2.15)$$

Ces fonctions d'ondes ne sont visiblement pas plus normalisables que celles des états $|q\rangle$ ou $|p\rangle$, puisque leur module est constant de $-\infty$ à $+\infty$. Comme dans le cas de $|p\rangle$, le calcul du coefficient $\mathcal{Z}$ se fera donc en se reposant sur la propriété de complétude de ces états propres

$$\int dq_\theta |q_\theta\rangle \langle q_\theta| = \hat{\mathbb{I}}, \quad (2.16)$$

d'où on déduit

$$\langle q | q' \rangle = \int dq_\theta \langle q | q_\theta \rangle \langle q_\theta | q' \rangle \quad (2.17)$$

$$= |\mathcal{Z}|^2 e^{\frac{i \cos \theta}{4N_0 \sin \theta} (q'^2 - q^2)} \int dq_\theta e^{i \frac{q_\theta}{2N_0 \sin \theta} (q - q')}. \quad (2.18)$$

Le terme de gauche est donné par la relation d'orthogonalité (B.8), alors que l'intégrale se calcule à l'aide de la formule (A.20). On a donc

$$\delta(q - q') = |\mathcal{Z}|^2 e^{\frac{i \cos \theta}{4N_0 \sin \theta} (q'^2 - q^2)} 2\pi 2N_0 |\sin \theta| \delta(q - q'), \quad (2.19)$$

la valeur absolue $|\sin \theta|$ apparaissant lorsqu'on effectue le changement de variable pour se ramener à l'équation (A.20). De plus, comme le terme de droite est nul quand $q \neq q'$, on peut remplacer q' par q dans l'argument de l'exponentielle complexe, ce qui simplifie l'équation précédente en

$$\delta(q - q') = |\mathcal{Z}|^2 2\pi 2N_0 |\sin \theta| \delta(q - q'), \quad (2.20)$$

d'où on déduit, en omettant le terme de phase arbitraire

$$\mathcal{Z} = \frac{1}{\sqrt{4\pi N_0 |\sin \theta|}} \quad (2.21)$$

La fonction d'onde de q_θ peut donc s'écrire

$$\langle q | q_\theta \rangle = \frac{1}{\sqrt{4\pi N_0 |\sin \theta|}} e^{-i \frac{\cos \theta q^2 - 2q q_\theta}{4N_0 \sin \theta}} \quad (2.22)$$

La variation de phase est quadratique (sauf dans le cas $\theta = \frac{\pi}{2} + k\pi$, qui annule le cosinus, ce qui correspond à $|Q_\theta = p\rangle = (-1)^k |P = p\rangle$, où la fonction d'onde est une onde plane), et la phase oscille de plus en plus vite à mesure que q s'éloigne de $\frac{q_\theta}{\cos \theta}$. Si on regarde la limite $\theta \rightarrow 0$, on voit que les oscillations autour d'un point $q \neq q_\theta$ seront de plus en plus rapides, et seront moyennées à 0 dès qu'on regarde cette fonction d'onde avec une résolution finie. La valeur en q_θ , par contre tend vers $+\infty$ comme le facteur de normalisation $\mathcal{Z}$. Cette fonction d'onde ressemble donc de plus en plus à la fonction de Dirac qui est celle de $|Q = q_\theta\rangle = |Q_0 = q_\theta\rangle$.

Ces fonctions d'ondes peuvent nous servir à étudier l'effet d'une rotation d'angle θ dans l'espace des phases sur la fonction d'onde d'un état quelconque $|\psi\rangle$. L'effet du hamiltonien de l'oscillateur harmonique (2.6) peut ainsi se calculer en effectuant la substitution $\theta = -\omega(t - t_0)$.

$$\hat{U}_\theta |\psi\rangle = \int dq |Q_\theta = q\rangle \langle Q = q | \psi \rangle \quad (2.23)$$

$$\psi_\theta(q) = \langle q | \hat{U}_\theta | \psi \rangle = \int dq' \frac{1}{\sqrt{4\pi N_0 |\sin \theta|}} e^{i \frac{\cos \theta q'^2 - 2q' q}{4N_0 \sin \theta}} \psi(q') \quad (2.24)$$

2.2.2 États d'incertitude minimale

2.2.2.1 Définition et caractérisation

Les états d'incertitude minimale, ou *états minimaux*, sont ceux pour lesquels l'inégalité de Heisenberg (2.4) est saturée, c'est à dire ceux pour lesquels l'égalité

$$\Delta Q \Delta P = N_0 \quad (2.25)$$

est vérifiée.

Pour caractériser ces états, nous reprendrons la démonstration de la section B.6 de l'inégalité de Heisenberg. Pour que l'inégalité soit saturée, il faut que le discriminant Δ défini en (B.34) soit nul. Le trinôme qui constitue le membre de droite de (B.33d) a pour racine double

$$\lambda_0 = \frac{i 2i N_0}{2 \Delta Q^2} = -\frac{N_0}{\Delta Q^2}. \quad (2.26)$$

L'équation (B.33d) nous dit alors que le vecteur $\hat{A}' |\psi\rangle$ associé au nombre réel λ_0 a une norme nulle et est donc nul. On a alors

$$0 = \hat{A}' |\psi\rangle = (\hat{P} - \langle P \rangle + i\lambda_0 \hat{Q} - i\lambda_0 \langle Q \rangle) |\psi\rangle \quad (2.27a)$$

$$= \left(\hat{P} - i \frac{N_0}{\Delta Q^2} \hat{Q} - \langle P \rangle + i \frac{N_0}{\Delta Q^2} \langle Q \rangle \right) |\psi\rangle, \quad (2.27b)$$

où $\langle Q \rangle$ et $\langle P \rangle$ désignent les valeurs moyennes des observable P et Q .

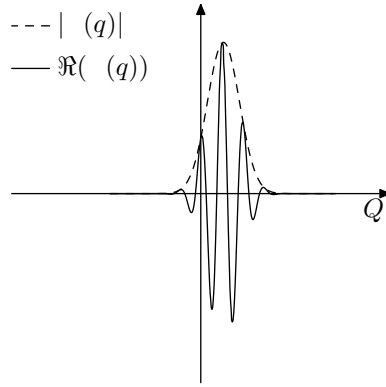


FIG. 2.2 – Fonction d'onde d'un état minimal

En multipliant cette équation à gauche par $\langle q|$ et en tenant compte de (B.20), on obtient l'expression suivante, vérifiée par la fonction d'onde

$$0 = -2iN_0 \frac{d\psi}{dq} - i \frac{N_0}{\Delta Q^2} q \psi(q) - \left(\langle P \rangle - i \frac{N_0}{\Delta Q^2} \langle Q \rangle \right) \psi(q). \quad (2.28)$$

La fonction d'onde ψ vérifie donc l'équation différentielle

$$\frac{d\psi}{dq} = -\frac{1}{2\Delta Q^2} q \psi(q) + \left(\frac{\langle Q \rangle}{2\Delta Q^2} + i \frac{\langle P \rangle}{2N_0} \right) \psi(q), \quad (2.29)$$

qu'on peut récrire sous la forme

$$\frac{d\psi}{\psi} = -\frac{1}{2\Delta Q^2} q dq + \left(\frac{\langle Q \rangle}{2\Delta Q^2} + i \frac{\langle P \rangle}{2N_0} \right) dq \quad (2.30)$$

et intégrer pour obtenir l'ensemble des solutions

$$\psi(q) = \mathcal{Z} e^{-\frac{(q-\langle Q \rangle)^2}{4\Delta Q^2} + \frac{i\langle P \rangle q}{2N_0}}. \quad (2.31)$$

Le coefficient de normalisation $\mathcal{Z}$ est déterminé par la condition de normalisation (B.13), qui s'exprime ici sous la forme

$$1 = \int dq |\psi(q)|^2 = |\mathcal{Z}|^2 \int dq e^{-\frac{(q-\langle Q \rangle)^2}{2\Delta Q^2}} = |\mathcal{Z}|^2 \sqrt{2\pi \Delta Q^2}, \quad (2.32)$$

d'où on déduit, en omettant un facteur de phase $e^{i\phi}$ arbitraire,

$$\mathcal{Z} = \left(2\pi \Delta Q^2 \right)^{-\frac{1}{4}}. \quad (2.33)$$

Les fonctions d'ondes des états d'incertitude minimale sont donc des paquets d'ondes gaussiens :

$$\psi(q) = \left(2\pi \Delta Q^2 \right)^{-\frac{1}{4}} e^{-\frac{(q-\langle Q \rangle)^2}{4\Delta Q^2} + \frac{i\langle P \rangle q}{2N_0}}. \quad (2.34)$$

Comme ces états vérifient par définition l'égalité (2.25), cette équation peut également s'écrire sous la forme

$$\psi(q) = \left(2\pi \Delta Q^2\right)^{-\frac{1}{4}} e^{-\frac{1}{4}\left(\frac{q-\langle Q \rangle}{\Delta Q}\right)^2 + \frac{i}{2} \frac{\langle P \rangle}{\Delta P} \frac{q}{\Delta Q}}. \quad (2.35)$$

Chaque état d'incertitude minimale est donc défini de manière unique par sa variance en position ΔQ^2 (ou en impulsion ΔP^2) et ses valeurs moyennes en position $\langle Q \rangle$ et en impulsion $\langle P \rangle$. Ces états, également appelés *états gaussiens*² en raison de la forme de leur fonction d'onde, sont avec les états de Fock, les états les plus employés en optique quantique. Les paragraphes suivants décrivent les états minimaux les plus courants.

2.2.2.2 État fondamental de l'oscillateur harmonique

L'état fondamental $|0\rangle$ de l'oscillateur harmonique est celui pour lequel la valeur moyenne du hamiltonien (2.6) est minimale. Cette valeur peut s'écrire

$$\langle 0 | \hat{H} | 0 \rangle = \frac{\hbar\omega}{2N_0} \left(\frac{1}{2} \Delta Q^2 + \frac{1}{2} \langle Q \rangle^2 + \frac{1}{2} \Delta P^2 + \frac{1}{2} \langle P \rangle^2 \right), \quad (2.36)$$

où les valeurs moyennes et les écarts-types sont pris sur l'état $|0\rangle$. Pour que cette moyenne soit minimale, il faut nécessairement que $\langle Q \rangle$ et $\langle P \rangle$ soient nulles. La relation d'incertitude (2.4) fournit alors une borne inférieure pour $\langle \hat{H} \rangle$ qui est atteinte pour $|0\rangle$, état minimal centré en 0 tel que $\Delta Q^2 = \Delta P^2 = N_0$:

$$\langle 0 | \hat{H} | 0 \rangle = \frac{\hbar\omega}{2} \equiv E_0. \quad (2.37)$$

La fonction d'onde de cet état minimal est donc de la forme (2.35)

$$\langle q | 0 \rangle \equiv \psi_0(q) = (2\pi N_0)^{-\frac{1}{4}} e^{-\frac{1}{4} \frac{q^2}{N_0}}. \quad (2.38)$$

On vérifie aisément en utilisant la relation (B.21) que $|0\rangle$ est un état propre de l'énergie $\hat{H}$:

$$\hat{H} \psi_0(q) = \frac{\hbar\omega}{4N_0} \left(q^2 - 4N_0^2 \frac{d^2}{dq^2} \right) (2\pi N_0)^{-\frac{1}{4}} e^{-\frac{1}{4} \frac{q^2}{N_0}} \quad (2.39a)$$

$$= \frac{\hbar\omega}{4N_0} \left(q^2 - 2N_0 - q^2 \right) (2\pi N_0)^{-\frac{1}{4}} e^{-\frac{1}{4} \frac{q^2}{N_0}} \quad (2.39b)$$

$$= \frac{\hbar\omega}{2} \psi_0(q) \equiv E_0 \psi_0(q). \quad (2.39c)$$

L'énergie $E_0 = \frac{1}{2} \hbar\omega$ du fondamental n'est pas nulle et est appelée *énergie de point zéro*. Cette énergie est due aux fluctuations de Q et P autour de 0, ces fluctuations étant elles-mêmes nécessaires pour ne pas violer la relation d'incertitude (2.4). Dans le cas où P et Q représentent les quadratures d'un mode du champ électromagnétique, l'état fondamental $|0\rangle$ est appelé *vide*, car il correspond à un mode vide de tout photon.

²Nous nous contenterons ici de considérer les *états purs gaussiens*. Plus généralement, lorsqu'un état quantique n'est pas pur, il n'est plus décrit par une fonction d'onde mais par une matrice densité. Dans ce cas, on continue à parler d'*états gaussiens* lorsque leur fonction de Wigner (voir chapitre 3) est gaussienne.

2.2.2.3 États cohérents

En physique classique, l'amplitude $Q + iP$ d'un état de l'oscillateur harmonique est parfaitement définie, ce que la mécanique quantique interdit. On s'attend à ce que le pendant quantique de ces états soient des états minimaux centrés autour de l'amplitude de leur contrepartie classique. Les états dont l'amplitude est « la mieux définie » sont les états propres $|\alpha\rangle$ de l'opérateur $\hat{a}$ associé à l'amplitude, défini dans la section B.9. Un état d'amplitude $\alpha \in \mathbb{C}$ vérifiera donc l'équation

$$\hat{a} |\alpha\rangle = \alpha |\alpha\rangle \quad (2.40)$$

qui se développe sous la forme

$$\left(\frac{\hat{Q} + i\hat{P}}{\sqrt{4N_0}} - \alpha \right) |\alpha\rangle = 0. \quad (2.41)$$

En multipliant cette équation par $-i\sqrt{4N_0}$, on obtient

$$(\hat{P} - i\hat{Q} + i\sqrt{4N_0}\alpha) |\alpha\rangle = 0. \quad (2.42)$$

On retrouve l'équation (2.27b) définissant un état minimal, avec les conditions

$$\Delta Q^2 = N_0 \quad \alpha = \frac{\langle Q \rangle + i\langle P \rangle}{\sqrt{4N_0}}. \quad (2.43)$$

Les $|\alpha\rangle$, états propres de $\hat{a}$, sont donc des états minimaux, centrés autour de l'amplitude qui correspond à leur valeur propre α . De plus, leurs variances en Q et en P sont égales au bruit quantique standard : $\Delta Q^2 = \Delta P^2 = N_0$. Ces états, appelés *états cohérents*, *états quasi-classiques* ou *états de Glauber*, sont ceux obtenus à la sortie d'un laser de bonne qualité.

Par exemple, le vide, état minimal symétrique centré en 0 est l'état cohérent $|\alpha = 0\rangle$. C'est donc l'état propre de $\hat{a}$ de valeur propre 0 et on retrouve l'expression (2.36) de l'énergie de point zéro en utilisant l'expression (B.57) du hamiltonien de l'oscillateur harmonique.

L'état obtenu en appliquant l'opérateur déplacement $\hat{D}_\alpha$, défini dans la section B.8, à l'état cohérent $|\beta\rangle$ est l'état cohérent $|\alpha + \beta\rangle$. En effet, si on calcule l'effet de l'opérateur $\hat{a}$ sur cet état, on obtient

$$\hat{a} \hat{D}_\alpha |\beta\rangle = \hat{D}_\alpha \hat{D}_\alpha^\dagger \hat{a} \hat{D}_\alpha |\beta\rangle = \hat{D}_\alpha (\alpha + \hat{a}) |\beta\rangle = (\alpha + \beta) \hat{D}_\alpha |\beta\rangle. \quad (2.44)$$

L'état $\hat{D}_\alpha |\beta\rangle$ est donc l'état propre de $\hat{a}$ associé à la valeur propre $\alpha + \beta$, c'est à dire l'état cohérent $|\alpha + \beta\rangle$. Tous les états cohérents peuvent donc être obtenus par déplacement du vide :

$$|\alpha\rangle = \hat{D}_\alpha |0\rangle. \quad (2.45)$$

Cette relation fixe le terme de phase globale arbitraire omis dans l'équation (2.35). Les équations (B.48) et (2.38) nous permettent de calculer la fonction d'onde de l'état cohérent $|\alpha\rangle$ avec $\alpha = \frac{\langle Q \rangle + i\langle P \rangle}{\sqrt{4N_0}}$.

$$\langle q | \alpha \rangle = \frac{e^{-i\frac{\langle Q \rangle \langle P \rangle}{4N_0}}}{(2\pi N_0)^{\frac{1}{4}}} e^{-\frac{(q - \langle Q \rangle)^2}{4N_0} + i\frac{\langle P \rangle q}{2N_0}} \quad (2.46)$$

2.2.2.4 États comprimés

Les autres états minimaux, pour lesquels $\Delta Q \neq \Delta P$, sont appelés *états comprimés* (*squeezed states* en anglais). Ils sont caractérisés par le *facteur de compression* s défini par

$$\Delta Q^2 = s N_0 \quad \Delta P^2 = \frac{1}{s} N_0. \quad (2.47)$$

Le *paramètre de compression* r souvent utilisé dans la littérature est relié à s par la relation $s = e^{-2r}$. La fonction d'onde (2.35) en fonction de ce paramètre s s'écrit

$$\psi(q) = (2\pi s N_0)^{-\frac{1}{4}} e^{-\frac{1}{4} \frac{(q - \langle Q \rangle)^2}{s N_0} + i \frac{\langle P \rangle \hat{Q}}{2 N_0}} \quad (2.48)$$

Lorsque $s \rightarrow 0$, la gaussienne devient de plus en plus étroite, et tend bien vers une fonction de Dirac, au facteur de renormalisation près. L'état propre $|Q = \langle Q \rangle\rangle$ correspond à la limite d'un état comprimé centré en $\hat{Q}$ pour $s \rightarrow 0$. Si, au contraire, on fait tendre s vers $+\infty$, la gaussienne devient plate et la fonction d'onde devient une onde plane. On retrouve l'état propre $|P = \langle P \rangle\rangle$, au facteur de normalisation près. Les états propres de $\hat{Q}$ et $\hat{P}$ correspondent donc à la limite d'états infiniment comprimés.

2.2.3 Propriétés des états gaussiens

2.2.3.1 Fonction d'onde en P

Ces fonctions d'ondes sont des paquets d'ondes gaussiens, définis section A.2.2, page 199. Pour calculer la transformée de Fourier, on peut récrire l'équation (2.35) sous une forme similaire à (A.12) :

$$\psi(q) = \frac{e^{-\frac{\langle Q \rangle^2}{4 \Delta Q^2}}}{(2\pi)^{\frac{1}{4}} \sqrt{\Delta Q}} e^{-\frac{q^2}{2(\sqrt{2} \Delta Q)^2} + \frac{1}{2} \left(\frac{\langle Q \rangle}{\Delta Q} + i \frac{\langle P \rangle}{\Delta P} \right) \frac{q}{\Delta Q}}. \quad (2.49)$$

L'équation (A.21b) devient donc, aux changements d'échelle (B.30) près,

$$\psi[p] = \frac{1}{\sqrt{2 N_0}} \frac{e^{-\frac{\langle Q \rangle^2}{4 \Delta Q^2}}}{(2\pi)^{\frac{1}{4}} \sqrt{\Delta Q}} \sqrt{2 \Delta Q} e^{\frac{2 \Delta Q^2}{2} \frac{1}{4} \left(\frac{\langle Q \rangle}{\Delta Q} + i \frac{\langle P \rangle}{\Delta P} \right)^2 \frac{1}{\Delta Q^2} + \frac{2 \Delta Q^2}{2} \frac{p^2}{4 N_0^2} - i 2 \Delta Q^2 \frac{1}{2} \left(\frac{\langle Q \rangle}{\Delta Q} + i \frac{\langle P \rangle}{\Delta P} \right) \frac{1}{\Delta Q} \frac{p}{2 N_0}} \quad (2.50a)$$

Cette expression se simplifie, en tenant compte notamment de l'égalité (2.25), pour obtenir

$$\psi[p] = \frac{1}{(2\pi \Delta P^2)^{\frac{1}{4}}} e^{-\frac{\langle P \rangle^2}{4 \Delta P^2} + \frac{i}{2} \frac{\langle Q \rangle \langle P \rangle}{\Delta Q \Delta P} - \frac{p^2}{4 \Delta P^2} + \frac{1}{2} \frac{p}{\Delta P} \left(\frac{\langle P \rangle}{\Delta P} - i \frac{\langle Q \rangle}{\Delta Q} \right)} \quad (2.50b)$$

$$= \frac{1}{(2\pi \Delta P^2)^{\frac{1}{4}}} e^{-\frac{1}{4} \left(\frac{p - \langle P \rangle}{\Delta P} \right)^2 + \frac{i}{2} \frac{\langle Q \rangle}{\Delta Q} \frac{p - \langle P \rangle}{\Delta P}} \quad (2.50c)$$

2.2.3.2 Produit scalaire de deux paquets d'ondes

Supposons que l'on ait deux paquets d'ondes gaussiens $\psi(Q)$ et $\varphi(Q)$. Il est aisé de se convaincre que leur produit scalaire

$$\langle \varphi | \psi \rangle = \int dq \varphi^*(q) \psi(q) \quad (2.51)$$

est une intégrale gaussienne et n'est jamais nul. Deux paquets d'ondes gaussiens ne sont donc jamais orthogonaux. Dans le cas où ils ont le même écart-type ΔQ , l'équation précédente devient

$$\langle \varphi | \psi \rangle = \frac{e^{-\frac{\langle Q \rangle_\varphi^2 + \langle Q \rangle_\psi^2}{4\Delta Q^2}}}{\sqrt{2\pi\Delta Q^2}} \int dq e^{-\frac{q^2}{2\Delta Q^2} + \frac{1}{2} \left(\frac{\langle Q \rangle_\varphi + \langle Q \rangle_\psi}{\Delta Q} + i \frac{\langle P \rangle_\psi - \langle P \rangle_\varphi}{\Delta P} \right) \frac{q}{\Delta Q}} \quad (2.52a)$$

$$= \frac{e^{-\frac{\langle Q \rangle_\varphi^2 + \langle Q \rangle_\psi^2}{4\Delta Q^2}}}{\sqrt{2\pi\Delta Q^2}} e^{\Delta Q^2 \frac{1}{4} \left(\frac{\langle Q \rangle_\varphi + \langle Q \rangle_\psi}{\Delta Q} + i \frac{\langle P \rangle_\psi - \langle P \rangle_\varphi}{\Delta P} \right)^2 \frac{1}{\Delta Q^2}} \sqrt{2\pi\Delta Q^2} \quad (2.52b)$$

$$= e^{-\frac{1}{4} \left(\frac{\langle Q \rangle_\psi - \langle Q \rangle_\varphi}{\Delta Q} \right)^2 - \frac{1}{4} \left(\frac{\langle P \rangle_\psi - \langle P \rangle_\varphi}{\Delta P} \right)^2 + \frac{i}{2} \frac{\langle Q \rangle_\varphi + \langle Q \rangle_\psi}{\Delta Q} \frac{\langle P \rangle_\psi - \langle P \rangle_\varphi}{\Delta P}}. \quad (2.52c)$$

Ce produit scalaire est donc, à une phase près, une fonction gaussienne de la « distance » des deux paquets d'ondes dans l'espace des phases à une échelle définie par ΔQ et ΔP . On voit donc que, même s'ils ne sont jamais tout à fait orthogonaux, $|\varphi\rangle$ et $|\psi\rangle$ sont presque orthogonaux lorsque $\langle Q \rangle_\psi$ et $\langle Q \rangle_\varphi$ diffèrent de quelques ΔQ (ou lorsque $\langle P \rangle_\psi$ et $\langle P \rangle_\varphi$ diffèrent de quelques ΔP).

2.2.4 Superposition linéaire d'états cohérents

La superposition linéaire de deux états cohérents est souvent désignée par le nom de *chat de Schrödinger*. Ces états illustrent en effet le fameux paradoxe du chat de Schrödinger. Le paradoxe, dans sa description initiale, repose sur un dispositif expérimental plaçant un chat dans une superposition d'états mort et vivant. Si chacun sait ce qu'est un chat mort et un chat vivant, personne ne sait vraiment ce qu'est un chat $\frac{1}{\sqrt{2}}(\text{mort} + \text{vivant})$, état pourtant autorisé par la mécanique quantique.

Des états cohérents d'amplitude différentes correspondent à des états classiques bien définis, et leur superposition quantique est donc analogue fameux au chat qui se retrouve dans la superposition de deux états classiques bien définis. Par souci de simplicité, nous nous limiterons à la superposition des états $|\alpha\rangle$ et $|\alpha\rangle$ avec $\alpha \in \mathbb{R}$, les autres cas s'y ramenant par déplacement et rotation dans l'espace des phases.

L'état $|\psi\rangle$ défini par la superposition

$$|\psi\rangle = \frac{|\alpha\rangle + |-\alpha\rangle}{\sqrt{2 + \langle \alpha | -\alpha \rangle + \langle -\alpha | \alpha \rangle}} = \frac{1}{\sqrt{2(1+e^{-4\alpha^2})}} (|\alpha\rangle + |-\alpha\rangle) \quad (2.53)$$

ne peut être qualifié de *chat de Schrödinger* que lorsque les états $|\alpha\rangle$ et $|\alpha\rangle$ sont vraiment différents, c'est à dire lorsque $|\langle \alpha | -\alpha \rangle|^2 \ll 1$. D'après l'équation (2.52c), cette approximation est déjà valable pour $|\alpha| > 2$, en raison de la décroissance rapide de la fonction gaussienne. Dans ce cas, l'égalité précédente se récrit

$$|\psi\rangle \simeq \frac{1}{\sqrt{2}} (|\alpha\rangle + |-\alpha\rangle). \quad (2.54)$$

La fonction d'onde en Q de cette superposition est la somme pondérée de la fonction d'onde de chacun des termes

$$\langle q | \psi \rangle = \frac{1}{\sqrt{2(1+e^{-4\alpha^2})}} (\langle q | \alpha \rangle + \langle q | -\alpha \rangle) \quad (2.55)$$

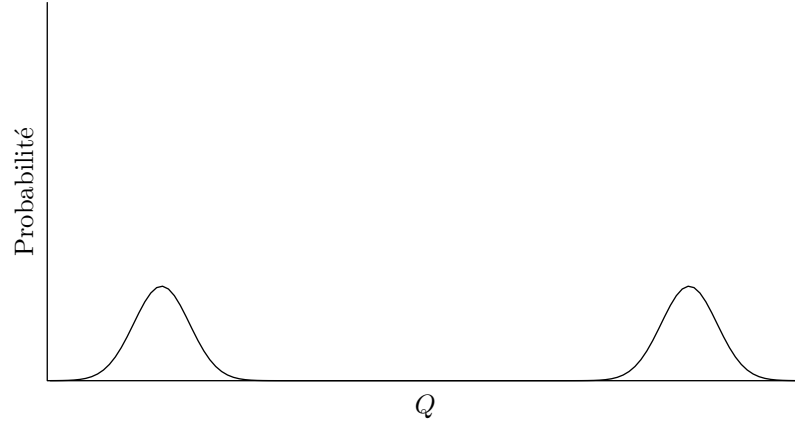


FIG. 2.3 – Distribution de probabilité en Q de la superposition linéaire de deux états cohérents

Dans le cas où $\alpha \in \mathbb{R}$, cette équation devient

$$\psi(q) = \frac{1}{(2\pi N_0)^{\frac{1}{4}} \sqrt{2(1+e^{-4\alpha^2})}} \left(e^{-\left(\frac{q}{\sqrt{4N_0}} - \alpha\right)^2} + e^{-\left(\frac{q}{\sqrt{4N_0}} + \alpha\right)^2} \right) \quad (2.56)$$

La distribution de probabilité en Q associée $|\psi(Q)|^2$ correspond presque à la moyenne des distributions de probabilité des deux états si les gaussiennes se recouvrent peu :

$$|\psi(q)|^2 = \frac{1}{2(1+e^{-4\alpha^2})} \left(|\langle q|\alpha \rangle|^2 + |\langle q|-\alpha \rangle|^2 + 2\mathcal{R}(\langle \alpha|q \rangle \langle q|-\alpha \rangle) \right) \quad (2.57)$$

$$\simeq \frac{1}{2} \left(|\langle q|\alpha \rangle|^2 + |\langle q|-\alpha \rangle|^2 \right) \quad (2.58)$$

La distribution de probabilités en position de l'état quantique superposé est simplement la moyenne des distributions de probabilités de chaque état. Elle ne se distingue donc pas d'un mélange statistique où on a une fois sur deux l'état $|\alpha\rangle$ et une fois sur deux l'état $|\alpha\rangle$.

La différence apparaît lorsqu'on regarde ce qui se passe en P . En effet, les fonctions d'ondes en P des deux états cohérents ne diffèrent que par le facteur de phase et leurs distributions de probabilité en P sont identiques.

$$\langle p|\pm\alpha \rangle = (2\pi N_0)^{-\frac{1}{4}} e^{-\frac{p^2}{4N_0} \pm i\frac{\alpha p}{\sqrt{N_0}}} \quad (2.59)$$

Comme les distributions de probabilités en P sont identiques, un mélange statistique ne change pas cette distribution, qui reste gaussienne. Par contre, les termes de phase vont interférer dans la superposition quantique et la fonction d'onde en P de la superposition peut s'écrire :

$$\psi[p] = \frac{1}{(2\pi N_0)^{\frac{1}{4}} \sqrt{2(1+e^{-4\alpha^2})}} e^{-\frac{p^2}{4N_0}} 2 \cos \frac{\alpha p}{\sqrt{N_0}}. \quad (2.60)$$

La distribution de probabilité en P est alors une gaussienne modulée par un cosinus, très

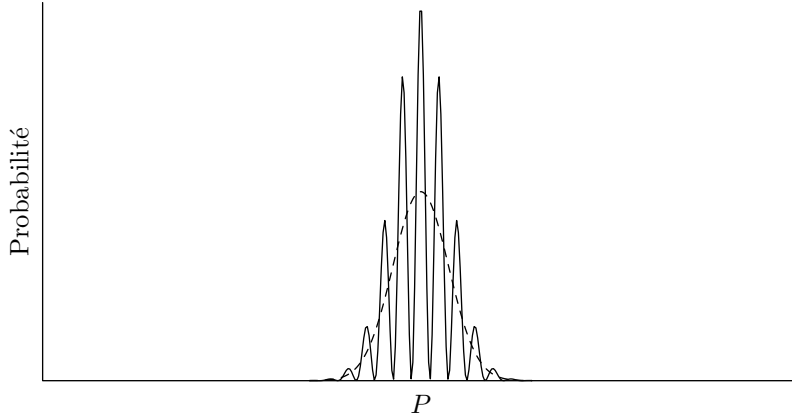


FIG. 2.4 – Distribution de probabilité en P comparées d’une superposition linéaire (trait plein) et d’un mélange statistique (pointillés) de deux états cohérents

différente de la gaussien du simple mélange statistique.

$$|\psi[p]|^2 = \frac{2}{\sqrt{2\pi N_0} (1 + e^{-4\alpha^2})} e^{-\frac{p^2}{2N_0}} \cos^2 \frac{\alpha p}{\sqrt{N_0}}. \quad (2.61)$$

Par contre cette différence s’effacera en présence d’un bruit assez important pour brouiller les franges. La superposition quantique est fragile.

2.2.5 États de Fock

Les *états de Fock* ou *états nombres de photons* sont les états propres de l’énergie. Comme nous le verrons ici, il s’agit d’une base discrète de l’espace de Hilbert. Elle est très utilisée en optique quantique, même si les états de Fock ont peu de réalité expérimentale au delà de quelques photons.

Par définition, un état de Fock $|n\rangle$ est un état propre de l’hamiltonien de l’oscillateur harmonique défini par l’équation (2.6) et vérifie

$$\hat{H} |n\rangle = E_n |n\rangle, \quad (2.62)$$

avec $E_n \in \mathbb{R}$. Comme on l’a vu dans le paragraphe précédent l’équation (2.36), E_n est toujours supérieur ou égal à $E_0 = \frac{\hbar\omega}{2}$. L’état fondamental de l’oscillateur $|0\rangle$ défini par l’équation (2.38) est donc l’état de Fock d’énergie minimale.

Nous utiliserons ici l’expression (B.57) de l’hamiltonien $\hat{H}$ de l’oscillateur harmonique :

$$\hat{H} = \hbar\omega(\hat{a}^\dagger \hat{a} + \frac{1}{2}), \quad (2.63)$$

Appelons $\hat{n}$ l’opérateur $\hat{a}^\dagger \hat{a}$. Les états de Fock sont des états propres de $\hat{n}$ de valeur propres $n \in \mathbb{R}_+$. En particulier, le vide $|0\rangle$ est le seul état qui annule l’opérateur $\hat{n}$:

$$\hat{n} |0\rangle = 0. \quad (2.64)$$

Il sera intéressant d’étudier les effets de l’opérateur amplitude $\hat{a}$ sur les états propres $|n\rangle$. On a la relation, pour tout $|\psi\rangle$,

$$\langle \psi | \hat{n} | \psi \rangle = \|\hat{a} | \psi \rangle\|^2, \quad (2.65a)$$

qui devient, en substituant $|n\rangle$ à $|\psi\rangle$,

$$n = \|\hat{a}|n\rangle\|^2. \quad (2.65b)$$

Le vide est donc le seul état qui annule l'opérateur $\hat{a}$.

De plus, si $n > 0$,

$$\hat{n}\hat{a}|n\rangle = \hat{a}(\hat{n} - 1)|n\rangle = (n - 1)\hat{a}|n\rangle. \quad (2.66)$$

$\hat{a}|n\rangle$ est donc un vecteur propre de $\hat{n}$ valeur propre $n - 1$. On peut donc écrire, si $n > 0$, en tenant compte de l'équation (2.65b),

$$\hat{a}|n\rangle = \sqrt{n}|n - 1\rangle \quad (2.67)$$

Si $n \in \mathbb{N}$, on peut ainsi reconstruire à partir de $|n\rangle$ les états de Fock $|n - 1\rangle$, $|n - 2\rangle$, et ainsi de suite jusqu'à $|0\rangle$ en appliquant successivement l'opérateur $\hat{a}$. On ne peut plus continuer à descendre au dessous de $|0\rangle$ car $\hat{a}|0\rangle = 0$.

Par contre, si n n'était pas entier, rien ne nous empêcherait de continuer à appliquer l'opérateur jusqu'à obtenir un état propre de $\hat{n}$ de valeur propre négative, qui aurait une énergie inférieure à E_0 . Comme E_0 est la valeur minimale de l'énergie, c'est impossible. Les valeurs propres de $\hat{n}$ sont donc toutes entières positives.

On a donc

$$n \in \mathbb{N} \quad \text{et} \quad E_n = \hbar\omega \left(n + \frac{1}{2}\right) = (2n + 1)E_0. \quad (2.68)$$

Les valeurs propres de l'énergie sont donc (à E_0 près) des multiples entiers d'une excitation élémentaire $\hbar\omega$, appelée *photon*. L'opérateur $\hat{n}$ est appelé opérateur nombre de photons, et l'état $|n\rangle$ est un état à n photons, chacun de ces photons apportant une énergie $\hbar\omega$. L'équation (2.67) justifie le nom d'*opérateur annihilation* souvent donné à $\hat{a}$, car son action détruit effectivement un photon.

De même, il n'est pas difficile de se convaincre que

$$\hat{a}^\dagger|n\rangle = \sqrt{n + 1}|n + 1\rangle, \quad (2.69)$$

pour tout n . L'opérateur $\hat{a}^\dagger$ est donc appelé *opérateur création*, car il crée un photon.

Nous pouvons donc définir l'état de Fock $|n\rangle$ pour tout entier positif n par application successive de $\hat{a}^\dagger$ sur le vide $|0\rangle$, l'unicité du fondamental $|0\rangle$ garantissant la non dégénérescence des états de Fock. On a donc

$$\boxed{|n\rangle = \frac{\hat{a}^{\dagger n}}{\sqrt{n!}}|0\rangle} \quad (2.70)$$

Ces états sont orthogonaux entre eux. En particulier, ils sont tous orthogonaux au vide, qui est un état gaussien. Leurs fonctions d'ondes sont donc toutes non gaussiennes (sauf pour $|0\rangle$, bien sûr). Pour trouver l'expression de la fonction d'onde d'un état de Fock quelconque, il suffit de remplacer $\hat{a}^\dagger$ par son expression (B.49) dans (2.70). On a alors

$$\langle q|n\rangle = \left\langle q \left| \frac{(\hat{Q} - i\hat{P})^n}{\sqrt{n!}(4N_0)^{\frac{n}{2}}} \right| 0 \right\rangle \quad (2.71a)$$

$$= \frac{1}{\sqrt{n!}} \left(\frac{q}{2\sqrt{N_0}} - \sqrt{N_0} \frac{d}{dq} \right)^n \langle q|0\rangle \quad (2.71b)$$

$$= \frac{1}{(2\pi N_0)^{\frac{1}{4}} \sqrt{n!}} \left(\frac{q}{2\sqrt{N_0}} - \sqrt{N_0} \frac{d}{dq} \right)^n e^{-\frac{q^2}{4N_0}}. \quad (2.71c)$$

Nous pouvons par exemple exprimer ainsi la fonction d'onde de l'état de Fock à 1 photon :

$$\langle q|1\rangle = \frac{1}{(2\pi N_0)^{\frac{1}{4}}} \left(\frac{q}{2\sqrt{N_0}} - \sqrt{N_0} \frac{d}{dq} \right) e^{-\frac{q^2}{4N_0}} \quad (2.72a)$$

$$= \frac{q}{(2\pi)^{\frac{1}{4}} N_0^{\frac{3}{4}}} e^{-\frac{q^2}{4N_0}}. \quad (2.72b)$$

Pour un nombre quelconque de photons , l'expression, plus complexe, fait intervenir les polynômes de Hermite [106].

Chapitre 3

La fonction de Wigner : une distribution dans l'espace des phases

3.1 Introduction

La mécanique quantique étant fondamentalement probabiliste, il est tentant de chercher une distribution de probabilité $\Pi(P, Q)$ dans l'espace des phases, correspondant à un état quantique $|\psi\rangle$ ou $\hat{\rho}$ donné, ce qui nous permettra de justifier avec plus de rigueur les schémas de la figure FIG. 2.1. Cette démarche sous-entend l'existence simultanée des deux variables P et Q , ce qui est contradictoire avec la mécanique quantique. Cette probabilité est donc mal définie, et il existe de fait une infinité de fonctions pouvant presque prétendre au titre de distribution de probabilités dans l'espace des phases. Aucune n'a toutes les propriétés requises pour être une vraie distribution de probabilité Π . Elles en ont cependant un certain nombre de propriétés intéressantes et fournissent des représentations plus intuitives des états quantiques que les fonctions d'onde. Parmi cette infinité de fonctions, trois distributions sont fréquemment employées, selon les propriétés sur lesquelles on veut mettre l'accent : la fonction- P , la fonction- Q et la fonction de Wigner. On se contentera ici d'étudier cette dernière, reliée à la tomographie quantique. Pour une étude plus générale de l'espace des phases en mécanique quantique, on pourra se référer à [5, 97].

Nous définirons donc la fonction de Wigner dans la section 3.2 et étudierons quelques-unes de ses nombreuses propriétés dans la section 3.3. La section 3.4 décrit les fonctions de Wigner associées aux états les plus couramment utilisés en optique quantique. Nous élargirons ensuite la définition de la fonction de Wigner aux états à plusieurs modes (section 3.5), ce qui nous permettra notamment d'étudier les effets des pertes sur la fonction de Wigner.

3.2 Définition

Nous introduirons dans cette section la fonction de Wigner par les probabilités marginales suivant l'approche de Bertrand et Bertrand [107, 97]. Nous définirons, section 3.2.1, les propriétés que devraient avoir une « vraie » distribution de probabilité $\Pi(Q, P)$ dans l'espace des phases indépendamment de toute considération quantique. Nous en déduirons l'existence d'une distribution $W(Q, P)$, la fonction de Wigner, qui a presque toutes les propriétés de $\Pi(Q, P)$, dont nous calculerons l'expression dans la section 3.2.2.

3.2.1 Probabilités marginales et tomographie

Une distribution de probabilité est toujours positive et normée à un :

$$\forall p, q, \quad \Pi(P = p, Q = q) \geq 0 \quad (3.1a)$$

$$\iint dq dp \Pi(P = q, Q = p) = 1. \quad (3.1b)$$

Ses distributions de probabilités marginales devraient aussi correspondre aux distributions de probabilité mesurables. Si on se limite aux rotations dans l'espace des phases, définies dans les sections 2.2.1 et B.1, on peut mesurer les distributions de probabilités

$$\mathcal{P}_\theta(q_\theta) \equiv \mathcal{P}_{Q_\theta=q_\theta}. \quad (3.2)$$

Une « bonne » distribution de probabilité vérifiera donc

$$\begin{aligned} \forall \theta \quad \mathcal{P}_\theta(q_\theta) &= \int dp_\theta \Pi(Q_\theta = q_\theta, P_\theta = p_\theta) \\ &= \int dp_\theta \Pi(Q = q_\theta \cos \theta - p_\theta \sin \theta, P = q_\theta \sin \theta + p_\theta \cos \theta). \end{aligned} \quad (3.3)$$

Le fait que $\mathcal{P}_\theta(Q_\theta)$ soit une probabilité, d'intégrale unité, garantit la normalisation (3.1b).

Cet ensemble d'équations intégrales définit la transformation de Radon, qui est le fondement de la méthode d'imagerie médicale appelée *tomographie*¹. Ces systèmes, couramment appelés *scanners*, reconstituent la répartition de matière à l'intérieur du corps depuis l'extérieur, et produisent des images analogues à des coupes anatomiques sans qu'il soit nécessaire de découper le patient. Pour ce faire, ils illuminent la « tranche » du corps considérée avec des rayons X selon plusieurs directions θ . Pour chaque angle, l'ombre portée par cette tranche représente alors l'équivalent de la fonction $\mathcal{P}_\theta(Q_\theta)$. Elle est en effet obtenue par une intégrale similaire à (3.3), le rôle de la distribution de probabilité $\Pi(Q, P)$ étant ici tenu par la répartition de l'opacité aux rayons X dans la tranche.

La reconstitution de la distribution Π à partir des probabilités marginales $\mathcal{P}_\theta$ sera appelée tomographie par analogie.

Il nous reste donc à trouver la distribution de probabilité vérifiant l'équation (3.3) sous la condition (3.1a). On peut démontrer que, si elle existe, cette distribution est unique, en comparant sa fonction caractéristique² [108]

$$\Pi[\mu, \nu] \equiv \frac{1}{2\pi} \iint dq dp e^{-i\mu q - i\nu p} \Pi(q, p) \quad (3.4)$$

à celle de la distribution de probabilité marginale

$$\mathcal{P}_\theta[\xi] = \frac{1}{\sqrt{2\pi}} \int dq_\theta e^{-i\xi q_\theta} \mathcal{P}_\theta(q_\theta). \quad (3.5)$$

Le remplacement, dans l'équation précédente, de $\mathcal{P}_\theta(q_\theta)$ par son expression (3.3) nous conduit à

$$\mathcal{P}_\theta[\xi] = \frac{1}{\sqrt{2\pi}} \iint dq_\theta dp_\theta e^{-i\xi q_\theta} \Pi(q_\theta \cos \theta - p_\theta \sin \theta, q_\theta \sin \theta + p_\theta \cos \theta). \quad (3.6)$$

¹Du grec *τομός* : *morceau coupé, part* et *γραφίς* : *dessin* [32]. Littéralement : dessin d'une tranche

²On abrégera fréquemment dans la suite la notation $\Pi(Q = q, P = p)$ en $\Pi(q, p)$. Les conventions utilisées dans les transformations de Fourier sont explicitées dans le paragraphe A.3.1.

Le changement de variables (B.1) nous conduit alors à

$$\mathcal{P}_\theta[\xi] = \frac{1}{\sqrt{2\pi}} \iint dq dp e^{-i\xi q \cos \theta - i\xi p \sin \theta} \Pi(q, p) = \sqrt{2\pi} \Pi[\xi \cos \theta, \xi \sin \theta], \quad (3.7)$$

en utilisant la définition (3.4) de la fonction caractéristique $\Pi[\mu, \nu]$. Celle-ci est donc définie, si elle existe, de manière unique par l'ensemble des fonctions caractéristiques $\mathcal{P}_\theta[\xi]$, qui ne sont rien d'autre que son expression en coordonnées polaires.

On peut donc reconstituer $\Pi(Q, P)$ à partir des distributions de probabilité marginales mesurées. Il suffit de calculer les transformés de Fourier $\mathcal{P}_\theta[\xi]$ à partir des distributions de probabilités $\mathcal{P}_\theta(Q_\theta)$. Ensuite, l'équation (3.7) nous permet d'en déduire la fonction caractéristique $\Pi[M, V]$, dont la transformée de Fourier inverse³ nous donne la distribution $\Pi(Q, P)$.

3.2.2 Expression de la fonction de Wigner

L'équation (3.7) nous donne donc une recette pour reconstruire une distribution de probabilité $\Pi(Q, P)$ à partir de ses distributions de probabilité marginales. Or la mécanique quantique, à défaut de définir la distribution Π , nous permet de calculer les probabilités marginales $\mathcal{P}_\theta$, qui sont mesurables. En appliquant notre « recette », nous trouvons la *fonction de Wigner* $W(P, Q)$, qui n'est malheureusement pas toujours positive. On ne peut donc pas parler de probabilité au sens strict ; cependant cette distribution présente cependant de nombreuses propriétés des distributions de probabilités, ce qui nous permettra de parler de *quasiprobabilité*.

Nous allons exprimer W en fonction de la matrice densité $\hat{\rho}$ du champ. Les distributions de probabilités $\mathcal{P}_\theta$ valent

$$\forall \theta, \quad \mathcal{P}_\theta(q_\theta) = \langle Q_\theta = q_\theta | \hat{\rho} | Q_\theta = q_\theta \rangle. \quad (3.8)$$

Si on injecte ce résultat dans (3.5), on a

$$\mathcal{P}_\theta[\xi] = \frac{1}{\sqrt{2\pi}} \int dq_\theta e^{-i\xi q_\theta} \langle Q_\theta = q_\theta | \hat{\rho} | Q_\theta = q_\theta \rangle \quad (3.9a)$$

$$= \frac{1}{\sqrt{2\pi}} \int dq_\theta \langle Q_\theta = q_\theta | \hat{\rho} e^{-i\xi \hat{Q}_\theta} | Q_\theta = q_\theta \rangle \quad (3.9b)$$

$$= \frac{1}{\sqrt{2\pi}} \text{Tr} \left\{ \hat{\rho} e^{-i\xi \hat{Q}_\theta} \right\} = \frac{1}{\sqrt{2\pi}} \text{Tr} \left\{ \hat{\rho} e^{-i\xi \hat{Q} \cos \theta - i\xi \hat{P} \sin \theta} \right\}. \quad (3.9c)$$

L'équation (3.7) nous permet de réexprimer le dernier terme sous la forme

$$W[\xi \cos \theta, \xi \sin \theta] = \frac{1}{2\pi} \text{Tr} \left\{ \hat{\rho} e^{-i\xi \hat{Q} \cos \theta - i\xi \hat{P} \sin \theta} \right\} \quad (3.10)$$

$$\boxed{W[\mu, \nu] = \frac{1}{2\pi} \text{Tr} \left\{ \hat{\rho} e^{-i\mu \hat{Q} - i\nu \hat{P}} \right\}} \quad (3.11)$$

³Si on cherche effectivement à reconstituer expérimentalement cette distribution de probabilité, les fluctuations statistiques compliqueront la tâche, et on ne peut malheureusement pas effectuer directement la transformée de Fourier. On pourra se référer à [97] pour plus de détails.

et la formule de Baker–Hausdorff (A.8) nous permet de calculer l'exponentielle

$$e^{-i\mu\hat{Q}-i\nu\hat{P}} = e^{-i\mu\nu N_0} e^{-i\nu\hat{P}} e^{-i\mu\hat{Q}}. \quad (3.12)$$

On a alors

$$W[\mu, \nu] = \frac{e^{-i\mu\nu N_0}}{2\pi} \int dq \langle Q = q | \hat{\rho} e^{-i\nu\hat{P}} e^{-i\mu\hat{Q}} | Q = q \rangle \quad (3.13a)$$

$$= \frac{e^{-i\mu\nu N_0}}{2\pi} \int dq e^{-i\mu q} \langle q | \hat{\rho} \hat{T}_{2N_0\nu} | q \rangle, \quad (3.13b)$$

où l'opérateur $\hat{T}_{2N_0\nu}$ désigne l'opérateur translation, qui est défini par l'équation (B.37) avec $x = 2N_0\nu$.

On a donc

$$W[\mu, \nu] = \frac{1}{2\pi} \int dq e^{-i\mu(q+N_0\nu)} \langle q | \hat{\rho} | q + 2N_0\nu \rangle. \quad (3.13c)$$

Le changement de variable $q' = q + N_0\nu$ permet de rendre l'expression plus symétrique :

$$W[\mu, \nu] = \frac{1}{2\pi} \int dq' e^{-i\mu q'} \langle q' - N_0\nu | \hat{\rho} | q' + N_0\nu \rangle. \quad (3.13d)$$

La fonction $W[\mu, \nu]$ étant par définition la transformée de Fourier de la fonction de Wigner $W(q, p)$ dont on cherche l'équation, il suffit d'inverser l'équation (3.4).

$$W(p, q) = \frac{1}{2\pi} \iint d\mu d\nu W[\mu, \nu] e^{i\mu q + i\nu p} \quad (3.14a)$$

$$W(p, q) = \frac{1}{(2\pi)^2} \iiint d\mu d\nu dq' e^{i\mu q + i\nu p - i\mu q'} \langle q' - N_0\nu | \hat{\rho} | q' + N_0\nu \rangle \quad (3.14b)$$

On peut alors intégrer sur μ , en appliquant l'équation (A.20) :

$$W(p, q) = \frac{1}{2\pi} \iint d\nu dq' e^{i\nu p} \delta(q - q') \langle q' - N_0\nu | \hat{\rho} | q' + N_0\nu \rangle \quad (3.14c)$$

$$= \frac{1}{2\pi} \int d\nu e^{i\nu p} \langle q - N_0\nu | \hat{\rho} | q + N_0\nu \rangle \quad (3.14d)$$

Le changement de variable $\frac{x}{2} = N_0\nu$ nous permet d'exprimer la fonction de Wigner sous la forme

$$W(p, q) = \frac{1}{2\pi 2N_0} \int dx e^{\frac{ixp}{2N_0}} \langle q - \frac{x}{2} | \hat{\rho} | q + \frac{x}{2} \rangle \quad (3.15a)$$

Dans la démonstration ci-dessus, nous avons choisi de développer la trace (3.11) sur la base des états propres $|q\rangle$ de la position. Nous aurions tout aussi bien pu choisir de développer cette trace sur la base des états propres $|p\rangle$ de la position, et par un raisonnement strictement analogue, nous aurions alors obtenu l'expression

$$W(p, q) = \frac{1}{2\pi 2N_0} \int dy e^{\frac{iyq}{2N_0}} \langle p + \frac{y}{2} | \hat{\rho} | p - \frac{y}{2} \rangle, \quad (3.15b)$$

qui nous sera parfois utile.

3.3 Propriétés élémentaires de la fonction de Wigner

3.3.1 Linéarité

L'expression (3.15a) est manifestement linéaire en $\hat{\rho}$. On en déduit donc que la fonction de Wigner d'une superposition incohérente d'états quantiques est la moyenne pondérée des superpositions.

$$\forall \lambda \in [0, 1], \quad W_{\lambda \hat{\rho}_1 + (1-\lambda) \hat{\rho}_2} = \lambda W_{\hat{\rho}_1} + (1 - \lambda) W_{\hat{\rho}_2}. \quad (3.16)$$

Cela correspond à ce qu'on attend d'une « vraie » distribution de probabilité dans le cas d'un mélange statistique.

3.3.2 Fonction de Wigner d'opérateurs

On peut définir la fonction de Wigner d'un opérateur quelconque par l'équation (3.15a), que cet opérateur soit hermitien ou non.

La fonction de Wigner d'un opérateur exprimé sous la forme $f(\hat{Q}) + g(\hat{P})$ est particulièrement simple et peut s'écrire

$$W_{f(\hat{Q})+g(\hat{P})}(q, p) = \frac{1}{2\pi 2N_0} (f(q) + g(p)), \quad (3.17)$$

comme nous allons le montrer ci-dessous. La fonction de Wigner en chaque point de l'espace des phases semble donc bien correspondre, à un facteur près, à la « valeur » de cet opérateur en ce point. On pourrait démontrer des propriétés similaires pour des fonctions plus générales de $\hat{Q}$ et $\hat{P}$, mais il faudrait alors tenir compte de l'ordre des termes qui ne commutent pas [5, 97].

La linéarité de la fonction de Wigner nous permet de traiter séparément la position et l'impulsion. Elle nous permet même de traiter séparément chaque terme du développement en série entière de f et de g , de sorte qu'il suffit d'étudier $W_{\hat{Q}^n}$ et $W_{\hat{P}^n}$ pour démontrer l'équation (3.17). Le calcul de $W_{\hat{Q}^n}$ est assez direct à partir de l'équation (3.15a) :

$$W_{\hat{Q}^n}(q, p) \equiv \frac{1}{2\pi 2N_0} \int dx e^{i \frac{px}{2N_0}} \left\langle q - \frac{x}{2} \left| \hat{Q}^n \right| q + \frac{x}{2} \right\rangle \quad (3.18a)$$

$$= \frac{1}{2\pi 2N_0} \int dx e^{i \frac{px}{2N_0}} \left(q + \frac{x}{2} \right)^n \delta(x) \quad (3.18b)$$

$$= \frac{q^n}{2\pi 2N_0}. \quad (3.18c)$$

Le calcul de $W_{\hat{P}}$ est analogue, mais part de l'expression (3.15b). On obtient alors

$$W_{\hat{P}^n}(q, p) = \frac{p^n}{2\pi 2N_0}, \quad (3.19)$$

ce qui nous permet de démontrer la propriété (3.17) par linéarité.

3.3.3 Fonction de Wigner d'opérateurs hermitiens

Un changement de variable $x \rightarrow -x$ dans l'équation (3.15a) permet de vérifier aisément que $W_{\hat{A}}^*(q, p) = W_{\hat{A}^\dagger}(q, p)$. Si l'opérateur $\hat{A}$ est hermitien, on a alors $W_{\hat{A}}^*(q, p) = W_{\hat{A}}(q, p)$. La fonction de Wigner des opérateurs hermitiens est donc réelle. En particulier, la distribution de quasiprobabilité W associée à une matrice densité peut être négative, mais n'est jamais complexe. Il en est de même pour les fonctions de Wigner associées à des observables.

3.3.4 Traces et valeurs moyennes

La moyenne du produit de deux fonctions de Wigner se calcule par l'intégrale suivante :

$$\begin{aligned} \iint dq dp W_{\hat{A}}(q, p) W_{\hat{B}}(q, p) \\ = \frac{1}{2\pi 2N_0} \iiint dq dx dy \langle q - \frac{x}{2} | \hat{A} | q + \frac{x}{2} \rangle \langle q - \frac{y}{2} | \hat{B} | q + \frac{y}{2} \rangle \\ \times \frac{1}{2\pi} \int \frac{dp}{2N_0} e^{i \frac{p}{2N_0} (x+y)}. \end{aligned} \quad (3.20a)$$

L'intégrale (A.20) sur p nous donne la fonction de Dirac $\delta(x + y)$, qui nous permet d'intégrrer sur y et d'obtenir

$$\iint dq dp W_{\hat{A}}(q, p) W_{\hat{B}}(q, p) = \frac{1}{2\pi 2N_0} \iint dq dx \langle q - \frac{x}{2} | \hat{A} | q + \frac{x}{2} \rangle \langle q + \frac{x}{2} | \hat{B} | q - \frac{x}{2} \rangle \quad (3.20b)$$

Le changement de variable

$$q' = q - \frac{x}{2} \qquad q'' = q + \frac{x}{2} \quad (3.21)$$

de jacobien 1 permet de terminer le calcul et de montrer que

$$\boxed{\iint dq dp W_{\hat{A}}(q, p) W_{\hat{B}}(q, p) = \frac{1}{2\pi 2N_0} \text{Tr} \{ \hat{A} \hat{B} \}} \quad (3.22)$$

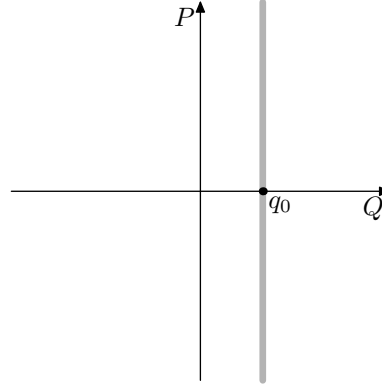
Cette expression est utile et montre bien à quel point la fonction de Wigner d'une matrice densité ressemble à une distribution de probabilité. En effet, de nombreux calculs de mécanique quantique se traduisent par des traces de produits d'opérateurs, et les quantités analogues dans le cas stochastique classique se calculent par des intégrales.

Par exemple, on a vu que la fonction de Wigner d'une observable pouvait se décrire comme la « valeur » de cette observable en tout point de l'espace des phases, cette expression étant parfois très intuitive (3.17) et identique à sa valeur classique. La valeur moyenne de cette observable se calcule alors, d'après (3.22), par la même intégrale que dans le cas classique si on remplace la distribution de probabilité par la fonction de Wigner. On a alors :

$$\langle \hat{A} \rangle_{\hat{\rho}} = \text{Tr} \{ \hat{\rho} \hat{A} \} = 2\pi 2N_0 \iint dq dp W_{\hat{\rho}}(q, p) W_{\hat{A}}(q, p). \quad (3.23)$$

Elle nous permet aussi de calculer la probabilité de transition entre deux états

$$|\langle \psi_1 | \psi_2 \rangle|^2 = \text{Tr} \{ |\psi_1\rangle \langle \psi_1| \psi_2 \rangle \langle \psi_2| \}. \quad (3.24)$$

FIG. 3.1 – Représentation de $W_{|q_0\rangle\langle q_0|}(Q, P)$

Cette probabilité de transition correspond à la fidélité $\mathcal{F}$ avec laquelle un état pur est reproduit par un autre état, pur ou non. Il s'agit en fait de la probabilité qu'on a de « confondre » ces deux états avec la meilleure mesure possible. Ce calcul de la fidélité s'étend aisément au cas où l'un des états est décrit par une matrice densité et est donc donnée par l'intégrale de recouvrement entre leurs deux fonctions de Wigner :

$$\mathcal{F} = \text{Tr} \{ |\psi_1\rangle \langle \psi_1| \hat{\rho}_2 \} = 2\pi 2N_0 \iint dq dp W_1(q, p) W_2(q, p). \quad (3.25)$$

Cette fidélité correspond à la fidélité d'intrication. Il serait tentant d'étendre cette définition à des matrices densités, mais cette extension présente un certain nombre de subtilités et n'est pas si directe qu'on pourrait le penser [109].

L'équation (3.25) impose l'existence de fonctions de Wigner négatives. En effet, la fidélité de deux états orthogonaux vaut 0, ce qui impose l'existence de régions où leurs fonctions de Wigner sont de signes opposés pour annuler l'intégrale. Elle limite également la taille de ces zones négatives. En effet, nous verrons section 3.4.2 que la fonction de Wigner d'un état gaussien est une gaussienne qui couvre une surface de l'ordre de N_0 dans l'espace des phases. Comme une fidélité est toujours positive, les régions négatives d'une fonction de Wigner couvrent toujours une surface plus petite que N_0 .

3.4 Exemples

Nous énumérerons ici les fonctions de Wigner de divers états quantiques, ce qui nous permettra parfois de nous faire une idée plus intuitive de leur nature.

3.4.1 États propres de la position

La relation d'orthogonalité (B.8) nous permet de calculer explicitement l'équation (3.15a) pour un état propre $|q_0\rangle$ de la position :

$$W_{|q_0\rangle\langle q_0|}(q, p) = \frac{1}{2\pi 2N_0} \int dx e^{i\frac{px}{2N_0}} \delta\left(q - q_0 - \frac{x}{2}\right) \delta\left(q - q_0 + \frac{x}{2}\right) \quad (3.26a)$$

$$= \frac{1}{2\pi 2N_0} \delta(q - q_0). \quad (3.26b)$$

Cette distribution dans l'espace des phases, représentée FIG. 3.1, est bien infiniment étroite en position et infiniment large en impulsion, comme on s'y attendait.

Bien entendu, les fonctions de Wigner des états propres de l'impulsion sont analogues, à une rotation dans l'espace des phases près, de même que celles des états propres de $\hat{Q}_\theta$.

3.4.2 États gaussiens

Les fonctions de Wigner des états gaussiens présentent de nombreuses propriétés intéressantes, comme on le verra dans la suite. Un état gaussien $|\psi\rangle$ est caractérisé par les valeurs moyennes de sa position $\langle Q \rangle$, de son impulsion $\langle P \rangle$, et par son écart-type en position ΔQ . Comme ses états sont minimaux, l'écart-type en impulsion vaut $\Delta P = \frac{N_0}{\Delta Q}$. La fonction d'onde d'un tel état est donnée par l'équation (2.35), qui, injectée dans la relation (3.15a), nous permet de calculer sa fonction de Wigner :

$$W(q, p) = \frac{1}{(2\pi)^{\frac{3}{2}} 2N_0 \Delta Q} \int dx e^{i\frac{xp}{2N_0}} e^{\left(\frac{q - \frac{x}{2} - \langle Q \rangle}{2\Delta}\right)^2 + i\frac{\langle P \rangle \left(q - \frac{x}{2}\right)}{2N_0}} e^{\left(\frac{q + \frac{x}{2} - \langle Q \rangle}{2\Delta}\right)^2 - i\frac{\langle P \rangle \left(q + \frac{x}{2}\right)}{2N_0}} \quad (3.27a)$$

$$= \frac{1}{(2\pi)^{\frac{3}{2}} 2N_0 \Delta Q} e^{-\frac{(q - \langle Q \rangle)^2}{2\Delta Q^2}} \int dx e^{-\frac{x^2}{8\Delta Q^2} - i\frac{(p - \langle P \rangle)x}{2N_0}} \quad (3.27b)$$

$$= \frac{1}{(2\pi)^{\frac{3}{2}} 2N_0 \Delta Q} e^{-\frac{(q - \langle Q \rangle)^2}{2\Delta Q^2}} \sqrt{2\pi 4\Delta Q^2} e^{-\frac{4\Delta Q^2 (p - \langle P \rangle)^2}{24N_0^2}} \quad (3.27c)$$

$$= \frac{1}{2\pi N_0} e^{-\frac{(q - \langle Q \rangle)^2}{2\Delta Q^2} - \frac{(p - \langle P \rangle)^2}{2\Delta P^2}} \quad (3.27d)$$

La fonction de Wigner d'un état gaussien est donc identique à ce que serait une « vraie » distribution de probabilité gaussienne dans l'espace des phases, de mêmes moyennes et variances. Ces états peuvent donc être décrits par une telle distribution de probabilité, positive partout, et leur seule propriété quantique semble être la présence incompressible du bruit quantique standard. Cette description probabiliste classique permet une description assez intuitive de ces états, qui sera détaillée au chapitre 4.

Les seuls états purs dont la fonction de Wigner est partout positive et peut tenir le rôle d'une distribution de probabilité sont les états gaussiens, ainsi que ceux obtenus par rotation dans l'espace des phases à partir d'états gaussiens [110, 111]. L'ensemble des états quantiques dont les distributions de probabilité marginales ne peuvent pas découler d'une vraie distribution de probabilité dans l'espace des phases est donc constitué de tous les états non gaussiens. Les états quantiques usuels présentant ces propriétés sont les chats de Schrödinger et les états de Fock.

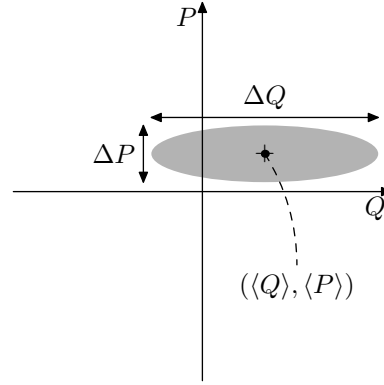


FIG. 3.2 – Représentation schématique de la fonction de Wigner d'un état gaussien

3.4.3 Superposition linéaire de deux états quantiques

La fonction de Wigner de $\alpha |\psi\rangle + \beta |\varphi\rangle$, superposition linéaire de deux états purs peut s'écrire sous la forme

$$W_{(\alpha|\psi\rangle+\beta|\varphi\rangle)(\alpha^*\langle\psi|+\beta^*\langle\varphi|)} = |\alpha|^2 W_{|\psi\rangle\langle\psi|} + |\beta|^2 W_{|\varphi\rangle\langle\varphi|} + 2\mathcal{R}(\alpha\beta^* W_{|\psi\rangle\langle\varphi|}) \quad (3.28)$$

Les deux premiers termes correspondent à ce que donnerait une superposition incohérente. Dans le cas de la superposition de deux états cohérent que nous avons défini section 2.2.4, ces deux termes correspondent à deux gaussiennes de même variance centrées en $\pm\alpha$.

L'intégrale du dernier terme, le terme d'interférence, est nulle en raison de la normalisation à 1 de l'intégrale de $W_{(\alpha|\psi\rangle+\beta|\varphi\rangle)(\alpha^*\langle\psi|+\beta^*\langle\varphi|)}$. Ce terme présente donc nécessairement des parties négatives. Si la fonction de Wigner de chaque état ($|\psi\rangle$ et $|\varphi\rangle$) est localisée dans l'espace des phases, le terme d'interférence est localisé et à mi-chemin des deux fonctions de Wigner. Si les états $|\psi\rangle$ et $|\varphi\rangle$ sont suffisamment éloignés dans l'espace des phases, la partie négative du terme d'interférence ne peut pas être masquée par d'éventuels termes positifs des autres termes. La fonction de Wigner totale présente alors des parties négatives et ne peut plus être interprétée comme une probabilité au sens strict.

Pour la superposition définie par l'équation (2.53), on a

$$W_{|-\alpha\rangle\langle\alpha|} = \frac{1}{2(2\pi N_0)^{\frac{3}{2}}} \int dx e^{i\frac{px}{2N_0} - \frac{(q - \frac{x}{2} + \sqrt{4N_0}\alpha)^2}{4N_0} - \frac{(q + \frac{x}{2} - \sqrt{4N_0}\alpha)^2}{4N_0}} \quad (3.29a)$$

$$= \frac{1}{2(2\pi N_0)^{\frac{3}{2}}} e^{-\frac{q^2}{2N_0} - 2\alpha^2} \int dx e^{-\frac{x^2}{8N_0} + \frac{(\sqrt{4N_0}\alpha + ip)x}{2N_0}} \quad (3.29b)$$

$$= \frac{1}{2(2\pi N_0)^{\frac{3}{2}}} e^{-\frac{q^2}{2N_0} + 2\alpha^2} \sqrt{2\pi 4N_0} e^{\frac{4N_0(\sqrt{4N_0}\alpha + ip)^2}{2(2N_0)^2}} \quad (3.29c)$$

$$= \frac{1}{2\pi N_0} e^{-\frac{q^2 + p^2}{2N_0} + 2i\frac{\alpha p}{\sqrt{N_0}}}. \quad (3.29d)$$

C'est donc une gaussienne centrée en 0, avec un terme phase oscillant en p .

Le terme d'interférence vaut donc

$$\frac{1}{1 + e^{-4\alpha^2}} \mathcal{R}(W_{|-\alpha\rangle\langle\alpha|}(q, p)) = \frac{1}{2\pi N_0 (1 + e^{-4\alpha^2})} e^{-\frac{q^2 + p^2}{2N_0}} \cos \frac{2\alpha p}{\sqrt{N_0}}. \quad (3.30)$$

Il présente des oscillations positives et négatives en P , de période $\pi \frac{\sqrt{N_0}}{\alpha}$. Comme nous le verrons section 3.5.5, l'effet des pertes sur la fonction de Wigner est équivalent à un bruit gaussien, qui brouillera ces franges d'autant plus vite qu'elles sont resserrées, c'est à dire qu' α est grand. Une fois le terme d'interférence brouillé, la fonction de Wigner est identique à celle d'une superposition incohérente.

Comme α doit être assez grand pour que le terme d'interférence ne soit pas masqué par les deux termes classiques, un chat est très fragile et disparaît, ou plutôt se transforme en un mélange statistique classique, à la moindre perte. Ces états sont donc très difficiles à réaliser expérimentalement.

3.4.4 États de Fock

Les états de Fock $|n\rangle$ sont d'autres états ayant une fonction de Wigner localement négative. En effet, quand $n \neq 0$, le produit scalaire $\langle 0 | n \rangle = 0$ et l'équation (3.25) impose que la fonction de Wigner ait des zones négatives.

Ces états présentent deux avantages essentiels par rapport aux chats de Schrödinger mentionnés ci-dessus : il est actuellement possible de les produire expérimentalement, au moins pour $n = 1$, et ils sont plus robustes aux pertes, comme on le verra dans la section 3.5.7. Ces avantages ont récemment permis la reconstruction expérimentale d'états de Fock à un photon [112, 113, 114, 115].

La fonction de Wigner $W_{|n\rangle}$ de l'état de Fock $|n\rangle$ vaut [104] :

$$W_{|n\rangle}(q, p) = \frac{(-1)^n}{2\pi N_0} e^{-\frac{p^2+q^2}{2N_0}} L_n\left(\frac{p^2+q^2}{N_0}\right) \quad (3.31)$$

$$= \frac{(-1)^n}{2\pi N_0} e^{-\frac{r^2}{2N_0}} L_n\left(\frac{r^2}{N_0}\right), \quad (3.32)$$

où $r^2 = q^2 + p^2$ et L_n désigne le n -ième polynôme de Laguerre :

$$L_n(x) \equiv \sum_{k=0}^n C_n^k \frac{(-x)^k}{k!}. \quad (3.33)$$

Ces fonctions présentent toutes un pic d'amplitude maximale à l'origine, ce pic étant négatif pour n impair. En effet

$$W_{|n\rangle}(0, 0) = \frac{(-1)^n}{2\pi N_0}. \quad (3.34)$$

Dans le cas de l'état de Fock à 1 photon, cette fonction de Wigner devient

$$W_{|1\rangle}(q, p) = -\frac{1}{2\pi N_0} e^{-\frac{r^2}{2N_0}} \left(1 - \frac{r^2}{N_0}\right), \quad (3.35)$$

et est négative pour $r < \sqrt{N_0}$.

3.5 Fonctions de Wigner à n modes

3.5.1 Définition

Dans tout ce qui précède, nous avons travaillé avec un seul mode du champ électromagnétique ; or, une caractéristique de la mécanique quantique est l'existence d'états collectifs,

qui ne peuvent pas se décrire comme combinaison d'états individuels. On doit donc pouvoir définir une fonction de Wigner pour des ensembles de modes. La fonction de Wigner d'un état défini par la matrice densité $\hat{\rho}$ vaut

$$W(q_1, p_1, \dots, q_n, p_n) = \frac{1}{(2\pi 2N_0)^n} \int \dots \int dx_1 \dots dx_n e^{\frac{ix_1 p_1}{2N_0} + \dots + \frac{ix_n p_n}{2N_0}} \times \left\langle q_1 - \frac{x_1}{2}, \dots, q_n - \frac{x_n}{2} \left| \hat{\rho} \right| q_1 + \frac{x_1}{2}, \dots, q_n + \frac{x_n}{2} \right\rangle \quad (3.36)$$

Cette fonction de Wigner dans un espace des phases à $2n$ dimensions a toutes les propriétés de la fonction de Wigner usuelle, dans l'espace des phases à 2 dimensions. Elle en a aussi quelques autres, explicitées dans les sections ci-après.

3.5.2 Traces partielles

Les fonctions de Wigner correspondant aux traces partielles des matrices densité se déduisent par intégration sur les modes « tracés ». En effet, on a

$$W(q_1, p_1, \dots, q_j, p_j) \equiv \int \dots \int dq_{j+1} dp_{j+1} \dots dq_n dp_n \times W(q_1, p_1, \dots, q_j, p_j, q_{j+1}, p_{j+1}, \dots, q_n, p_n). \quad (3.37a)$$

Cette intégrale peut s'exprimer sous la forme

$$\frac{1}{(2\pi 2N_0)^j} \int \dots \int dx_1 \dots dx_j e^{\frac{ix_1 p_1}{2N_0} + \dots + \frac{ix_j p_j}{2N_0}} \times \left\langle q_1 - \frac{x_1}{2}, \dots, q_j - \frac{x_j}{2} \left| \hat{\rho}' \right| q_1 + \frac{x_1}{2}, \dots, q_j + \frac{x_j}{2} \right\rangle. \quad (3.37b)$$

C'est donc la fonction de Wigner associée à l'opérateur $\hat{\rho}'$, défini par l'intégrale

$$\hat{\rho}' \equiv \int \dots \int dx_{j+1} \dots dx_n dq_{j+1} \dots dq_n \frac{dp_{j+1}}{2\pi 2N_0} \dots \frac{dp_n}{2\pi 2N_0} e^{\frac{ix_{j+1} p_{j+1}}{2N_0} + \dots + \frac{ix_n p_n}{2N_0}} \times \left\langle q_{j+1} - \frac{x_{j+1}}{2}, \dots, q_n - \frac{x_n}{2} \left| \hat{\rho}' \right| q_{j+1} + \frac{x_{j+1}}{2}, \dots, q_n + \frac{x_n}{2} \right\rangle. \quad (3.38a)$$

L'intégration sur les variables de position p_k donne des fonctions de Dirac $\delta(x_k)$, qui permet d'effectuer simplement l'intégration sur les x_k .

$$\hat{\rho}' = \int \dots \int dq_{j+1} \dots dq_n \left\langle q_{j+1}, \dots, q_n \left| \hat{\rho}' \right| q_{j+1}, \dots, q_n \right\rangle \equiv \text{Tr}_{j+1, \dots, n} \{ \hat{\rho} \}. \quad (3.38b)$$

L'opérateur $\hat{\rho}'$ est bien la trace de la matrice densité sur les modes intégrés de la fonction de Wigner.

3.5.3 Produits tensoriels et changements de base

La matrice densité de certains états peut s'écrire sous la forme du produit tensoriel $\hat{\rho} = \hat{\rho}_1 \otimes \hat{\rho}_2$. Il est aisé de se convaincre que la fonction de Wigner de tels états peut s'écrire sous une forme de produit :

$$W_{\hat{\rho}_1 \otimes \hat{\rho}_2} \left(\begin{pmatrix} Q_1 \\ P_1 \\ Q_2 \\ P_2 \end{pmatrix} \right) = W_{\hat{\rho}_1} \left(\begin{pmatrix} Q_1 \\ P_1 \end{pmatrix} \right) W_{\hat{\rho}_2} \left(\begin{pmatrix} Q_2 \\ P_2 \end{pmatrix} \right). \quad (3.39)$$

On retrouve encore une fois une propriété des distributions de probabilité. En effet, l'expression de $\hat{\rho}$ sous la forme du produit tensoriel $\hat{\rho}_1 \otimes \hat{\rho}_2$ est équivalent à l'indépendance statistique des résultats de mesures dans les modes définissant les bases dans lesquelles sont exprimées $\hat{\rho}_1$ et $\hat{\rho}_2$.

Cette propriété est souvent utilisée pour écrire les fonctions de Wigner d'états à plusieurs modes. Par exemple, si on envoie deux états quantiques non corrélés de fonctions de Wigner respectives W_1 et W_2 sur une lame semi-réfléchissante, la fonction de Wigner décrivant l'état des modes de sortie sera

$$W \begin{pmatrix} q_+ \\ p_+ \\ q_- \\ p_- \end{pmatrix} = W_1 \left(\frac{1}{\sqrt{2}} \begin{bmatrix} q_+ - q_- \\ p_+ - p_- \end{bmatrix} \right) W_2 \left(\frac{1}{\sqrt{2}} \begin{bmatrix} q_+ + q_- \\ p_+ + p_- \end{bmatrix} \right). \quad (3.40)$$

Dans le cas particulier où l'état 1 (respectivement 2) est le vide comprimé en Q (en P) avec un facteur de compression s , on a

$$W \begin{pmatrix} q_+ \\ p_+ \\ q_- \\ p_- \end{pmatrix} = \frac{1}{(2\pi N_0)^2} e^{-\frac{(q_+ - q_-)^2}{4sN_0} - \frac{s(p_+ - p_-)^2}{4N_0} - \frac{s(q_+ + q_-)^2}{4N_0} - \frac{(p_+ + p_-)^2}{4sN_0}}. \quad (3.41)$$

Cet état est en général appelé état EPR, car dans la limite des forts facteurs de compression $s \rightarrow 0$, la fonction de Wigner devient

$$W \begin{pmatrix} q_+ \\ p_+ \\ q_- \\ p_- \end{pmatrix} \rightarrow \delta(q_+ - q_-) \delta(p_+ + p_-), \quad (3.42)$$

et cet état est celui introduit par Einstein, Podolsky et Rosen dans leur fameux article [37] qui introduit le paradoxe qui porte leur nom.

Plus généralement, les états gaussiens multimodes peuvent s'écrire dans une certaine base orthonormée directe comme un produit tensoriels d'états gaussiens monomodes. Un simple changement de variable permet alors d'écrire la fonction de Wigner sous la forme souhaitée. On notera au passage que les fonctions de Wigner d'états gaussiens multimodes sont des gaussiennes à plusieurs dimensions et qu'elles restent positives.

Un tel changement de base correspond à une transformation canonique dans l'espace des phases à $2n$ dimensions. Elle est réalisable expérimentalement avec une série de lames partiellement réfléchissantes, qui mélangent les modes, de « compresseurs », qui compriment Q et dilatent P , et de déphasages.

3.5.4 Changement de base

Une lame partiellement réfléchissante lie les opérateurs de position et d'impulsion entre ses modes d'entrée et de sortie par la relation linéaire

$$\begin{bmatrix} \hat{Q}_{o1} \\ \hat{P}_{o1} \\ \hat{Q}_{o2} \\ \hat{P}_{o2} \end{bmatrix} = \begin{bmatrix} t & 0 & -r & 0 \\ 0 & t & 0 & -r \\ r & 0 & t & 0 \\ 0 & r & 0 & t \end{bmatrix} \begin{bmatrix} \hat{Q}_{i1} \\ \hat{P}_{i1} \\ \hat{Q}_{i2} \\ \hat{P}_{i2} \end{bmatrix}, \quad (3.43)$$

avec les notations définies FIG. 3.3. Il est simple de voir que la fonction de Wigner en sortie de cette lame partiellement réfléchissante se déduit de la fonction de Wigner en entrée par la même rotation.

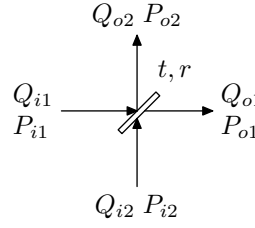


FIG. 3.3 – lame partiellement réfléchissante

Plus généralement, si on a un système de lames séparatrices, de déphaseurs et de « compresseurs »⁴, il réalise une transformation canonique dans l'espace des phases à $2n$ dimensions caractérisée par une matrice orthogonale M . Si on appelle $\hat{A}_i$ ($\hat{A}_o$) le vecteur constitué par les $2n$ opérateurs position et impulsion en entrée (en sortie) du dispositif, on a :

$$\hat{A}_o = M \hat{A}_i. \quad (3.44)$$

Si $W_i(\vec{A}_i)$ et $W_o(\vec{A}_o)$ sont les fonctions de Wigner des états respectivement à l'entrée et à la sortie du système, on a donc

$$W_o(\vec{A}_o) = W_i(M^{-1} \vec{A}_o). \quad (3.45)$$

3.5.5 Effet des pertes

Si on envoie un état décrit par la fonction de Wigner W_1 dans un canal de transmission η , l'approche précédente nous permet de calculer la fonction de Wigner en sortie du canal. En effet un canal avec pertes peut être modélisé par une lame partiellement réfléchissante de transmission η , avec du vide, de fonction de Wigner W_0 dans l'autre mode d'entrée de la lame partiellement réfléchissante.

La fonction de Wigner totale à l'entrée peut alors s'écrire.

$$W_i \begin{pmatrix} q_0 \\ p_0 \\ q_1 \\ p_1 \end{pmatrix} = W_0 \begin{pmatrix} q_0 \\ p_0 \end{pmatrix} W_1 \begin{pmatrix} q_1 \\ p_1 \end{pmatrix} \quad (3.46)$$

La matrice M^{-1} correspondant à la lame partiellement réfléchissante s'écrit

$$M^{-1} = \begin{bmatrix} t & 0 & r & 0 \\ 0 & t & 0 & r \\ -r & 0 & t & 0 \\ 0 & -r & 0 & t \end{bmatrix} \quad \text{avec } t = \sqrt{\eta} \text{ et } r = \sqrt{1 - \eta}. \quad (3.47)$$

La fonction de Wigner de l'état de sortie s'écrit alors

$$W_o \begin{pmatrix} q_{o0} \\ p_{o0} \\ q_{o1} \\ p_{o1} \end{pmatrix} = W_i \left(\begin{bmatrix} t & 0 & r & 0 \\ 0 & t & 0 & r \\ -r & 0 & t & 0 \\ 0 & -r & 0 & t \end{bmatrix} \begin{bmatrix} q_{o0} \\ p_{o0} \\ q_{o1} \\ p_{o1} \end{bmatrix} \right) \quad (3.48a)$$

$$= W_0 \begin{pmatrix} tq_{o0} + rq_{o1} \\ tp_{o0} + rp_{o1} \end{pmatrix} W_1 \begin{pmatrix} -rq_{o0} + tq_{o1} \\ -rp_{o0} + tp_{o1} \end{pmatrix} \quad (3.48b)$$

⁴C'est à dire quelque chose (en général un cristal non linéaire) qui transforme un état cohérent en état comprimé. Le terme anglais est *squeezer*.

Seul le mode $o1$ est accessible, et il faut faire la trace sur le mode $o0$ pour avoir la matrice densité du mode de sortie. On a donc

$$W_{o1} \begin{pmatrix} q_{o1} \\ p_{o1} \end{pmatrix} = \iint dq_{o0} dp_{o0} W_o \begin{pmatrix} q_{o0} \\ p_{o0} \\ q_{o1} \\ p_{o1} \end{pmatrix} \quad (3.49)$$

$$= \iint dq_{o0} dp_{o0} W_0 \begin{pmatrix} tq_{o0} + rq_{o1} \\ tp_{o0} + rp_{o1} \end{pmatrix} W_1 \begin{pmatrix} -rq_{o0} + tq_{o1} \\ -rp_{o0} + tp_{o1} \end{pmatrix}. \quad (3.50)$$

Cette intégrale est essentiellement une convolution. Cela apparaît lorsque l'on fait le changement de variable

$$\begin{bmatrix} q \\ p \end{bmatrix} \equiv rt \begin{bmatrix} q_{o0} \\ p_{o0} \end{bmatrix} + r^2 \begin{bmatrix} q_{o1} \\ p_{o1} \end{bmatrix}, \quad (3.51)$$

à (q_{o1}, p_{o1}) fixé. On a alors

$$W_{o1} \begin{pmatrix} q_{o1} \\ p_{o1} \end{pmatrix} = \iint dq dp \frac{1}{r^2 t^2} W_0 \left(\frac{1}{r} \begin{bmatrix} q \\ p \end{bmatrix} \right) W_1 \left(\frac{1}{t} \begin{bmatrix} q_{o1} - q \\ p_{o1} - t \end{bmatrix} \right) \quad (3.52)$$

$$\boxed{W_{o1} \begin{pmatrix} q_{o1} \\ p_{o1} \end{pmatrix} = \frac{1}{\eta(1-\eta)} \iint dq dp W_0 \left(\frac{1}{\sqrt{1-\eta}} \begin{bmatrix} q \\ p \end{bmatrix} \right) W_1 \left(\frac{1}{\sqrt{\eta}} \begin{bmatrix} q_{o1} - q \\ p_{o1} - p \end{bmatrix} \right)}. \quad (3.53)$$

Comme la fonction de Wigner W_0 du vide est une gaussienne centrée de variance N_0 , cette fonction de Wigner est la fonction de Wigner initiale, convoluée avec une gaussienne bidimensionnelle de variance $\chi N_0 \equiv \frac{1-\eta}{\eta} N_0$ et réduite d'un facteur d'échelle η . La variance χN_0 correspond à ce qui est couramment appelé en électronique *bruit équivalent ramené à l'entrée*.

3.5.6 Pertes et négativité de la fonction de Wigner

Les détails occupant une surface inférieure à χN_0 sont donc brouillés. On en déduit donc que la fonction de Wigner ne peut plus être localement négative pour $\chi \geq 1$, c'est à dire pour $\eta \leq \frac{1}{2}$, puisque les régions négatives occupent une surface inférieure à N_0 . Le calcul est assez simple si on regarde à l'origine, l'ensemble de l'espace des phases s'en déduisant par application d'opérateurs déplacement. On a en effet

$$W_{o1} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = \frac{1}{\eta(1-\eta)} \iint dq dp W_0 \left(\frac{1}{\sqrt{1-\eta}} \begin{bmatrix} q \\ p \end{bmatrix} \right) W_1 \left(\frac{1}{\sqrt{\eta}} \begin{bmatrix} -q \\ -p \end{bmatrix} \right) \quad (3.54a)$$

$$= \frac{1}{\eta(1-\eta)} \iint dq dp W_0 \left(\frac{1}{\sqrt{1-\eta}} \begin{bmatrix} q \\ p \end{bmatrix} \right) W_1 \left(\frac{1}{\sqrt{\eta}} \begin{bmatrix} q \\ p \end{bmatrix} \right), \quad (3.54b)$$

par symétrie de la fonction de Wigner W_0 du vide autour de l'origine. On a donc

$$W_{o1} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = \frac{1+\chi}{\chi} \iint dq dp W_0 \left(\frac{1}{\sqrt{\chi}} \begin{bmatrix} q \\ p \end{bmatrix} \right) W_1 \begin{pmatrix} q \\ p \end{pmatrix}. \quad (3.54c)$$

La fonction

$$W' \begin{pmatrix} Q \\ P \end{pmatrix} \equiv W_0 \left(\frac{1}{\sqrt{\chi}} \begin{bmatrix} Q \\ P \end{bmatrix} \right) \quad (3.55)$$

est une gaussienne de variance χN_0 . C'est, à un facteur d'échelle près, la fonction de Wigner d'un état quantique obéissant aux relations d'Heisenberg, défini par la matrice densité $\hat{\rho}_\chi$, si et seulement si $\chi \geq 1$. Dans ce cas, en appelant ρ_1 la matrice densité de l'état initial, on a

$$W_{o1} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = (1 + \chi) \iint dq dp W_{\hat{\rho}_\chi} \begin{pmatrix} q \\ p \end{pmatrix} W_1 \begin{pmatrix} q \\ p \end{pmatrix} \quad (3.56a)$$

$$= \text{Tr} \{ \hat{\rho}_\chi \hat{\rho}_1 \} \geq 0. \quad (3.56b)$$

La fonction de Wigner à l'origine d'un état qui a subi des pertes supérieures à 3 dB ($\eta \leq \frac{1}{2}$, ce qui correspond à $\chi \geq 1$) est donc nécessairement positive ou nulle. Cette propriété s'étend aisément à l'ensemble de l'espace des phases par l'application d'opérateurs déplacement.

3.5.7 Effet des pertes sur la fonction de Wigner d'états de Fock

Pour mettre en évidence la négativité de la fonction de Wigner, il faudra donc utiliser un système de mesure d'efficacité quantique globale supérieure à 50 %. Il faudra aussi, bien entendu, générer un état quantique qui a une fonction de Wigner négative, c'est-à-dire un état non-gaussien, comme les états de Fock. Nous étudierons ici l'évolution de la valeur à l'origine de la fonction de Wigner d'états de Fock en fonction des pertes.

La valeur à l'origine de la fonction de Wigner de l'état de Fock $|n\rangle$ vaut

$$W_{|n\rangle} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = \frac{(-1)^n}{2\pi N_0}. \quad (3.57)$$

Le calcul de la valeur à l'origine de la fonction de Wigner d'un état de Fock qui a subi des pertes η est plus simple en regardant directement sa matrice densité $\hat{\rho}_{|n\rangle, \eta}$ qu'en essayant de calculer la convolution mentionnée plus haut. Comme

$$\hat{\rho}_{|n\rangle, \eta} = \sum_{k=0}^n C_n^k \eta^k (1 - \eta)^{n-k} |k\rangle \langle k|, \quad (3.58)$$

on a

$$W_{|n\rangle, \eta} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = \sum_{k=0}^n C_n^k \eta^k (1 - \eta)^{n-k} W_{|k\rangle} \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad (3.59a)$$

$$= \frac{1}{2\pi N_0} \sum_{k=0}^n C_n^k (-\eta)^k (1 - \eta)^{n-k} \quad (3.59b)$$

$$= \frac{(1 - 2\eta)^n}{2\pi N_0} \quad (3.59c)$$

Cette valeur est bien nulle pour $\eta = \frac{1}{2}$ (sauf si $n = 0$, bien sûr) et toujours positive pour $\eta < \frac{1}{2}$. Si n est impair, la fonction de Wigner est négative à l'origine dès que les pertes sont inférieures à 3 dB, c'est à dire dès que $\eta > \frac{1}{2}$. Cette négativité est d'autant plus forte que n est petit. L'état de Fock dont la négativité de la fonction de Wigner⁵ est la plus robuste aux pertes est donc l'état de Fock à un photon. Par chance, c'est également, avec le vide, le

⁵Pour être exhaustif, il faudrait examiner également les points différents de l'origine, notamment pour les états à nombre pairs de photons.

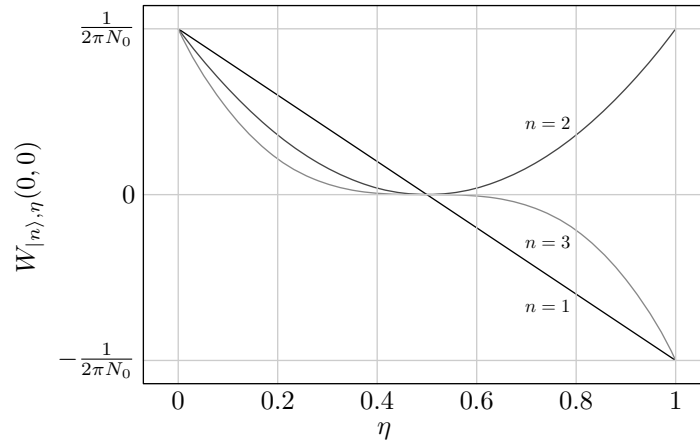


FIG. 3.4 – Valeur à l’origine de la fonction de Wigner d’états de Fock en fonction des pertes η

seul état de Fock réalisable aujourd’hui. La fonction de Wigner d’états de Fock à un photon a été reconstituée expérimentalement par Hansen et ses collaborateurs, à l’université de Constance [113, 114, 115] avec une détection homodyne pour des photons générés par fluorescence paramétrique, et par Bertet et ses collaborateurs [112] avec des atomes de Rydberg dans une cavité micro-ondes.

Chapitre 4

Formalisme gaussien

4.1 Introduction

Les états de Fock sont très utilisés pour les calculs d'optique quantique, car ils forment une base discrète de l'espace de Hilbert. Ils sont cependant difficiles à fabriquer et ont peu de réalité expérimentale au delà de quelques photons. Les états gaussiens, comprimés ou non, sont couramment utilisés et produits en laboratoire. C'est l'une des raisons pour lesquelles l'essentiel de l'optique quantique avec des variables continues se limite à l'étude des états gaussiens, et l'information quantique ne fait pas exception.

L'autre raison est de nature théorique et repose sur la positivité de leurs fonctions de Wigner. On peut donc faire « comme si » il s'agissait d'une vraie distribution de probabilité et exploiter tous les théorèmes statistiques. De plus, ces distributions de probabilité sont gaussiennes, et ont donc de nombreuses propriétés mathématiques qui en facilitent l'étude.

Nous présenterons dans ce chapitre des outils fréquemment utilisés dans l'étude des états gaussiens. Nous commencerons (section 4.2) par définir un système de notations linéarisées, fréquemment utilisées dans le cadre de l'optique quantique avec des états gaussiens. Nous utiliserons ensuite ces notations pour caractériser les canaux quantiques (section 4.3) et calculer des formules simplifiant les calculs de fidélité dans le cas d'états gaussiens (section 4.4).

4.2 Notations linéarisées

4.2.1 Matrices de covariance

Une distribution gaussienne est entièrement caractérisée par son vecteur moyen $\vec{\mu}$ et sa matrice de covariance K [108]. Dans le cas d'un état gaussien à n modes, ce vecteur est de dimension $2n$, et la matrice est $2n \times 2n$. En général, ils sont exprimés dans la base $\{Q_1, P_1, Q_2, P_2 \dots Q_n, P_n\}$ et on a

$$\vec{\mu} \equiv \langle \vec{A} \rangle = \begin{bmatrix} \langle Q_1 \rangle \\ \langle P_1 \rangle \\ \langle Q_2 \rangle \\ \langle P_2 \rangle \\ \vdots \\ \langle Q_n \rangle \\ \langle P_n \rangle \end{bmatrix} \quad (4.1)$$

$$K \equiv \left\langle (\vec{A} - \vec{\mu}) (\vec{A} - \vec{\mu})^T \right\rangle \equiv \left\langle \delta \vec{A} \delta \vec{A}^T \right\rangle \quad (4.2a)$$

$$= \begin{bmatrix} \langle \delta Q_1^2 \rangle & \langle \delta Q_1 \delta P_1 \rangle & \langle \delta Q_1 \delta Q_2 \rangle & \langle \delta Q_1 \delta P_2 \rangle & \dots & \langle \delta Q_1 \delta Q_n \rangle & \langle \delta Q_1 \delta P_n \rangle \\ \langle \delta P_1 \delta Q_1 \rangle & \langle \delta P_1^2 \rangle & \langle \delta P_1 \delta Q_2 \rangle & \langle \delta P_1 \delta P_2 \rangle & \dots & \langle \delta P_1 \delta Q_n \rangle & \langle \delta P_1 \delta P_n \rangle \\ \langle \delta Q_2 \delta Q_1 \rangle & \langle \delta Q_2 \delta P_1 \rangle & \langle \delta Q_2^2 \rangle & \langle \delta Q_2 \delta P_2 \rangle & \dots & \langle \delta Q_2 \delta Q_n \rangle & \langle \delta Q_2 \delta P_n \rangle \\ \langle \delta P_2 \delta Q_1 \rangle & \langle \delta P_2 \delta P_1 \rangle & \langle \delta P_2 \delta Q_2 \rangle & \langle \delta P_2^2 \rangle & \dots & \langle \delta P_2 \delta Q_n \rangle & \langle \delta P_2 \delta P_n \rangle \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \langle \delta Q_n \delta Q_1 \rangle & \langle \delta Q_n \delta P_1 \rangle & \langle \delta Q_n \delta Q_2 \rangle & \langle \delta Q_n \delta P_2 \rangle & \dots & \langle \delta Q_n^2 \rangle & \langle \delta Q_n \delta P_n \rangle \\ \langle \delta P_n \delta Q_1 \rangle & \langle \delta P_n \delta P_1 \rangle & \langle \delta P_n \delta Q_2 \rangle & \langle \delta P_n \delta P_2 \rangle & \dots & \langle \delta P_n \delta Q_n \rangle & \langle \delta P_n^2 \rangle \end{bmatrix} \quad (4.2b)$$

La densité de probabilité du vecteur $\vec{A}$ peut alors s'exprimer sous la forme

$$\mathcal{P}_{\vec{A}} = \frac{1}{(2\pi)^n \sqrt{|K|}} e^{\frac{1}{2} (\vec{A} - \vec{\mu})^T K^{-1} (\vec{A} - \vec{\mu})} \quad (4.3)$$

De plus, on peut montrer [116, 117] que des opérations locales (déphasage et compression) dans chaque mode permettent de redéfinir la base des $\{Q_i, P_i\}_{1 \leq i \leq n}$ de façon à découpler les variables de position et les variables d'impulsion. Tous les coefficients de la matrice de covariance qui s'expriment sous la forme $\langle Q_i P_i \rangle$ sont alors nuls. La matrice K devient alors

$$K = \begin{bmatrix} \langle \delta Q_1^2 \rangle & 0 & \langle \delta Q_1 \delta Q_2 \rangle & 0 & \dots & \langle \delta Q_1 \delta Q_n \rangle & 0 \\ 0 & \langle \delta P_1^2 \rangle & 0 & \langle \delta P_1 \delta P_2 \rangle & \dots & 0 & \langle \delta P_1 \delta P_n \rangle \\ \langle \delta Q_2 \delta Q_1 \rangle & 0 & \langle \delta Q_2^2 \rangle & 0 & \dots & \langle \delta Q_2 \delta Q_n \rangle & 0 \\ 0 & \langle \delta P_2 \delta P_1 \rangle & 0 & \langle \delta P_2^2 \rangle & \dots & 0 & \langle \delta P_2 \delta P_n \rangle \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \langle \delta Q_n \delta Q_1 \rangle & 0 & \langle \delta Q_n \delta Q_2 \rangle & 0 & \dots & \langle \delta Q_n^2 \rangle & 0 \\ 0 & \langle \delta P_n \delta P_1 \rangle & 0 & \langle \delta P_n \delta P_2 \rangle & \dots & 0 & \langle \delta P_n^2 \rangle \end{bmatrix} \quad (4.4)$$

4.2.2 Notations simplifiées

Pour limiter les recours à l'algèbre linéaire, nous utiliserons des notations linéarisées, plus intuitives dans le cas d'un faible nombre de modes, introduites dans les références [79, 118, 119, 84, 86, 85] pour l'étude des mesures quantiques non destructives. Dans cette approche le fait que la fonction de Wigner se comporte comme une distribution de probabilité nous permet de supposer que les variables Q_i et P_i sont simultanément définies, même si on ne peut en mesurer qu'une à la fois. On effectue tous les calculs en supposant que ces valeurs sont définies, sans se poser de question.

Le rôle de la mécanique quantique se limite ici à imposer des valeurs minimales aux variances. Si cette approche rend beaucoup de calculs plus intuitifs, il ne faut pas oublier qu'il ne s'agit que d'une fiction commode.

Sur un plan mathématique, cette approche est valable pour tous les états, gaussiens ou non, si on considère qu'on manipule des opérateurs en représentations de Heisenberg. Cependant, si l'état n'est pas gaussien, il n'est pas entièrement déterminé par $\vec{\mu}$ et K , et il faudrait calculer les moments d'ordres supérieurs en Q et P , s'ils existent. C'est pourquoi nous nous contenterons en général d'étudier des états gaussiens.

4.2.3 Bruit quantique

Ce formalisme permet de justifier l'appellation de *bruit quantique* par des notations similaires à celles utilisées en traitement du signal. Si Alice¹ prépare un état quantique donné, elle peut connaître (et choisir) simultanément ses valeurs moyennes en position $\langle Q \rangle \equiv A_Q$ et en impulsion $\langle P \rangle \equiv A_P$. Comme ces valeurs moyennes sont définies simultanément, elles commutent : $[A_Q, A_P] = 0$. On peut alors définir le bruit quantique par

$$B_Q \equiv Q - A_Q \quad B_P \equiv P - A_P. \quad (4.5)$$

Comme A_Q et A_P sont des valeurs moyennes, définies classiquement, elles commutent avec tous les opérateurs. Le commutateur des bruits vérifie donc

$$[B_Q, B_P] = [Q, P] = 2iN_0, \quad (4.6)$$

d'où on déduit, en utilisant l'équation (B.35), la relation d'incertitude de Heisenberg

$$\langle B_Q^2 \rangle \langle B_P^2 \rangle \geq N_0^2. \quad (4.7)$$

Comme on l'a vu section 2.2.2, cette inégalité est saturée pour les états gaussiens, et seulement pour les états gaussiens. Les états cohérents ont un bruit indépendant de la phase et vérifient

$$\langle B_Q^2 \rangle = \langle B_P^2 \rangle = N_0. \quad (4.8)$$

Plus généralement, on peut définir un facteur de compression s tel que

$$\langle B_Q^2 \rangle = sN_0 \quad \langle B_P^2 \rangle = \frac{1}{s}N_0. \quad (4.9)$$

Si $s < 1$ l'état sera dit comprimé en Q , si $s > 1$ il sera dit comprimé en P .

4.2.4 Paradoxe Einstein Podolsky Rosen

Les calculs précédents ne font que reproduire quasiment à l'identique les calculs du chapitre 2. Ce formalisme est plus intéressant lorsqu'il s'agit d'étudier les corrélations entre plusieurs modes. Les corrélations du paradoxe Einstein Podolsky Rosen (EPR) [37] apparaissent assez naturellement dans ce formalisme. Si on a deux modes du champ $+$ et $-$, les relations de commutation suivantes sont vérifiées :

$$[Q_+, Q_-] = [P_+, P_-] = [Q_+, P_-] = [P_+, Q_-] = 0 \quad (4.10)$$

$$[Q_+, P_+] = [Q_-, P_-] = 2iN_0. \quad (4.11)$$

On en déduit alors aisément les relations de commutation

$$[Q_+ - Q_-, P_+ + P_-] = [Q_+ + Q_-, P_+ - P_-] = 0. \quad (4.12)$$

¹Conformément à la tradition du domaine, les différents protagonistes des protocoles de communications quantiques seront appelés *Alice*, *Bob*, *Charles*, *Yolande* et *Zaccharie*, en raison de l'ordre alphabétique de leurs initiales. Un(e) espion(ne) intervenant dans les protocoles de cryptographie sera appelée *Ève*, de l'anglais *eavesdrop*, écouter aux portes. Un vérificateur chargé de vérifier le bon déroulement d'un protocole sera appelé *Victor*.

Celles-ci nous autorisent à avoir simultanément des variables de position parfaitement anti-corrélées et des variables d'impulsions parfaitement corrélées dans les deux modes (ou *vice versa*).

Les corrélations sont limitées par la relation

$$[Q_+ - Q_-, P_+] = 2iN_0 \quad \text{d'où} \quad \langle (Q_+ - Q_-)^2 \rangle \langle P_+^2 \rangle \geq N_0^2 \quad (4.13)$$

et des relations similaires. Cette relation montre que la limite de corrélation parfaite en position ($\langle (Q_+ - Q_-)^2 \rangle \rightarrow 0$) correspond à une variance infinie en impulsion ($\langle P_+^2 \rangle \rightarrow \infty$), c'est à dire à un coût énergétique infini.

Ces inégalités sont saturées pour les états gaussiens définis par l'équation (3.41). Les modes + et - sont en général désignés sous le nom de *faisceaux jumeaux* ou *faisceaux EPR*. Comme ces états sont gaussiens, leur fonction de Wigner est positive et peut s'interpréter directement comme une distribution de probabilité dans l'espace des phases. Les positions et impulsions de ces faisceaux dans l'espace des phases peuvent donc être décrites par des bruits gaussiens mieux corrélés (ou anticorrélés, selon la quadrature) que le bruit quantique standard.

Supposons qu'Alice prépare une paire de faisceaux jumeaux et mesure $Q_- = q_-$ de l'un des faisceaux. L'équation (4.13), saturée, nous dit que l'autre faisceau (+) peut être décrit par un état maximalelement comprimé en Q , de facteur de compression $\frac{N_0}{\langle P_+^2 \rangle}$, centré en $\begin{bmatrix} q_- \\ 0 \end{bmatrix}$. Si Alice envoie ce faisceau à Bob, et lui révèle *a posteriori* q_- , Bob n'aura aucun moyen de savoir comment Alice a préparé ce faisceau. Le scénario que nous venons de décrire, en particulier, est indiscernable d'un scénario où Alice aurait choisi aléatoirement q_- avec la loi gaussienne appropriée et fabriqué un état comprimé centré en $\begin{bmatrix} q_- \\ 0 \end{bmatrix}$.

On peut bien sûr tenir le même raisonnement en P , à un signe près, dans la mesure où les impulsions sont anticorrélées si les positions sont corrélées.

4.3 Canaux quantiques

4.3.1 Définition

Ce formalisme linéarisé est particulièrement adapté à l'étude de canaux bruités. Si Alice dispose d'un état décrit par les quadratures Q et P , et qu'elle le transmet à Bob par un canal bruité linéaire, les quadratures de Bob peuvent s'écrire

$$Q_B = g_Q (Q + B_Q) + h_Q P \quad (4.14a)$$

$$P_B = g_P (P + B_P) + h_P Q, \quad (4.14b)$$

sauf dans le cas très particulier où $g_Q = 0$ ou $g_P = 0$. Pour simplifier les calculs, on se limitera au cas où il n'y a pas de contamination entre les variables de position et d'impulsion, c'est à dire qu'on supposera $h_P = h_Q = 0$. Dans le cas contraire, les calculs deviennent un peu plus complexes, mais le principe reste le même.

On a donc

$$Q_B = g_Q (Q + B_Q) \quad (4.15a)$$

$$P_B = g_P (P + B_P). \quad (4.15b)$$

On a gardé les gains $g_{Q,P}$ en facteurs des bruits $B_{Q,P}$, de sorte que ceux-ci représentent les bruits équivalents ramenés à l'entrée, ce qui est le paramètre pertinent pour étudier la perte d'information dans le canal. Ces bruits ne sont, par définition, pas corrélés (ni classiquement, ni quantiquement) au « signal » P, Q . Nous supposons en plus que leur valeur moyenne est nulle et qu'ils ne sont pas corrélés entre eux.² Leurs variances respectives seront notées χ_Q et χ_P .

Lorsque l'état passe par plusieurs canaux successifs, le canal composé ainsi constitué peut être décrit par le même formalisme. On a ainsi, pour deux canaux successifs,

$$Q_C = g'_Q (Q_B + C_Q) = g'_Q g_Q (Q + D_Q) \quad \text{avec} \quad D_Q = B_Q + \frac{1}{g_Q} C_Q \quad (4.16a)$$

$$P_C = g'_P (P_B + C_P) = g'_P g_P (Q + D_P) \quad \text{avec} \quad D_P = B_P + \frac{1}{g_P} C_P. \quad (4.16b)$$

4.3.2 Propriétés

Cette description des canaux quantiques nous permettra d'examiner les contraintes que la mécanique quantique impose à de tels canaux.

Nous pouvons calculer le commutateur

$$[Q_B, P_B] = g_Q g_P ([Q, P] + [B_Q, B_P]). \quad (4.17)$$

Comme Q_B et P_B d'une part et Q et P d'autre part décrivent des modes du champ, leurs commutateurs respectifs valent $2iN_0$. On en déduit la valeur du commutateur

$$[B_Q, B_P] = \left(\frac{1}{g_Q g_P} - 1 \right) 2iN_0 = \left(\frac{1}{G} - 1 \right) 2iN_0, \quad (4.18)$$

avec $G \equiv g_Q g_P$.

Si $G < 1$, le préfacteur $\frac{1}{G} - 1$ est positif, et le bruit $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ peut être modélisé par un mode du champ qui est mélangé au signal et on peut écrire

$$\begin{bmatrix} B_Q \\ B_P \end{bmatrix} = \sqrt{\frac{1}{G} - 1} \begin{bmatrix} Q_0 \\ P_0 \end{bmatrix}. \quad (4.19a)$$

Par contre, si $G > 1$, ce préfacteur est négatif, et le bruit est proportionnel au complexe conjugué d'un mode du champ, c'est à dire qu'on a

$$\begin{bmatrix} B_Q \\ B_P \end{bmatrix} = \sqrt{1 - \frac{1}{G}} \begin{bmatrix} Q_0 \\ -P_0 \end{bmatrix}. \quad (4.19b)$$

La relation de commutation (4.18) et la relation (B.35) nous conduisent à l'inégalité suivante :

$$\chi_Q \chi_P \geq \left(\frac{1}{G} - 1 \right)^2, \quad (4.20)$$

où $\langle B_{Q,P}^2 \rangle = \chi_{Q,P} N_0$. Cette équation ressemble à une inégalité de Heisenberg pondérée. On peut toujours, au moins en théorie, avoir un bruit arbitrairement faible sur une quadrature en ayant beaucoup de bruit sur l'autre. Pour « oublier » l'équilibre entre position et impulsion, on peut utiliser le paramètre

$$\chi \equiv \sqrt{\chi_Q \chi_P}, \quad (4.21)$$

²Cette dernière supposition n'est pas fondamentale et permet de simplifier les calculs. Le cas plus général où les bruits B_Q et P_Q sont corrélés est plus complexe mais peut être traité avec le même formalisme [84].

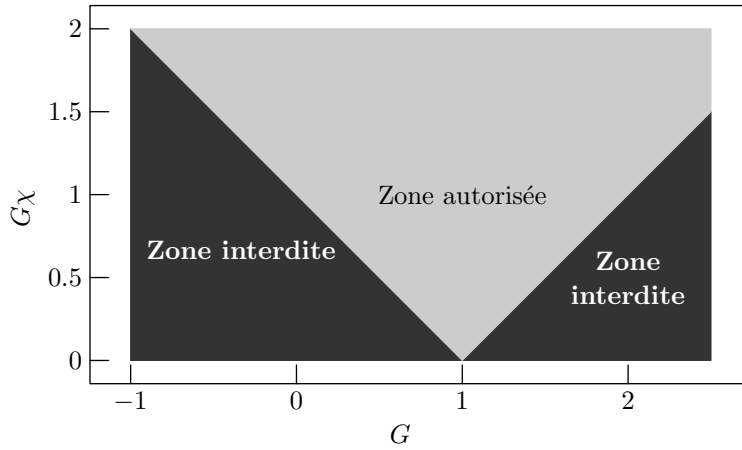
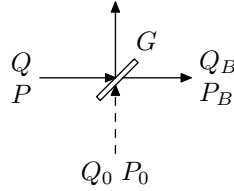
FIG. 4.1 – Valeurs autorisées de $G\chi$ en fonction de G 

FIG. 4.2 – Modélisation d'une absorption

et le bruit ramené à l'entrée minimum ajouté pour un canal de gain G vaut

$$\chi \geq \left| 1 - \frac{1}{G} \right|. \quad (4.22)$$

Cette équation peut aussi s'écrire

$$|G\chi| \geq |1 - G|, \quad (4.23)$$

où $|G\chi|$ représente l'amplitude du bruit *en sortie*.

Tout canal quantique ajoute donc du bruit, sauf si $G = 1$ [120]. En plus du cas trivial d'une transmission parfaite, cela autorise l'amplification sans bruit d'une quadrature si l'autre est déamplifiée (sans bruit également), c'est à dire si $g_Q = \frac{1}{g_P}$. Si $G > 1$, cette équation nous donne le bruit minimum ajouté par un amplificateur indépendant de la phase [120]; si $0 < G < 1$ elle nous donne le bruit minimum ajouté lors d'une atténuation; enfin, si $G < 0$ elle donne le bruit minimum ajouté par une conjugaison de phase.

Dans le cas d'une absorption, $0 < g_Q = g_P = \sqrt{G} < 1$, $\chi_Q = \chi_P = \chi$ et l'inégalité est saturée :

$$\chi = \frac{1}{G} - 1. \quad (4.24)$$

Une telle absorption est en général modélisée par une lame partiellement réfléchissante de coefficient de transmission G et de coefficient de réflexion $1 - G$. Dans ce modèle, le mode

vide Q_0, P_0 se couple au signal et on a, avec les notations de la FIG. 4.2,

$$Q_B = \sqrt{G} Q + \sqrt{1-G} Q_0 = \sqrt{G} \left(Q + \sqrt{\frac{1}{G} - 1} Q_0 \right) \quad (4.25a)$$

$$P_B = \sqrt{G} P + \sqrt{1-G} P_0 = \sqrt{G} \left(P + \sqrt{\frac{1}{G} - 1} P_0 \right). \quad (4.25b)$$

Dans le cas où le mode 0 est vide, il est aisé de retrouver la valeur de χ . Si le mode 0 contient un état thermique, alors on peut modéliser un canal de pertes avec excès de bruit. Si on veut construire un canal atténué où le bruit n'est pas également réparti entre les quadratures, il nous suffit d'utiliser cette lame semi-réfléchissante en remplaçant le mode 0 par du vide comprimé.

Cet ajout d'une variable gaussienne a le même effet sur la distribution de probabilité que la convolution de l'équation (3.53).

4.3.3 Mesure simultanée de l'impulsion et de la position

L'approche de la section précédente permet de traiter également le problème de la mesure, de deux manières différentes.

On peut partir des équations (4.15) en considérant que les mesures de Bob Q_B et P_B sont classiques et commutent. L'équation (4.17) devient alors

$$[Q_B, P_B] = 0 = g_Q g_P ([Q, P] + [B_Q, B_P]) \quad (4.26)$$

$$[B_Q, B_P] = -[Q, P] = 2iN_0 \quad (4.27)$$

d'où on déduit la relation d'incertitude de Heisenberg, qui s'écrit avec les notations introduites dans la section précédente

$$\sqrt{\chi_Q \chi_P} \equiv \chi \geq 1. \quad (4.28)$$

Une autre approche, strictement équivalente, consiste à dire que la première étape d'une mesure est une amplification du microscopique, quantique, au macroscopique, classique. Il suffit alors de prendre la limite $G \rightarrow \infty$ dans l'équation (4.22) pour retrouver la relation (4.28). Cette limite revient effectivement à négliger le commutateur $[Q_B, P_B]$ dans l'équation (4.17).

4.3.4 Mesure et reconstruction

Nous pouvons utiliser les résultats précédents pour étudier le canal constitué par une mesure suivie d'une reconstruction, ce qui est parfois appelé *téléportation classique*. On a alors, en notant $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ les résultats de la mesure et $\begin{bmatrix} Q_C \\ P_C \end{bmatrix}$ les quadratures de l'état préparé,

$$Q_C = g_Q (Q_B + C_Q) = g_Q (Q + B_Q + C_Q) \quad (4.29a)$$

$$P_C = g_P (P_B + C_P) = g_P (P + B_P + C_P). \quad (4.29b)$$

Comme on l'a vu plus haut

$$\langle B_Q^2 \rangle \langle B_P^2 \rangle \geq N_0 \quad \text{et} \quad \langle (g_Q C_Q)^2 \rangle \langle (g_P C_P)^2 \rangle \geq N_0. \quad (4.30)$$

Si la mesure est symétrique et que l'état reconstruit est un état cohérent, et que $g_Q = g_P = \sqrt{G}$, on a

$$\chi = 1 + \frac{1}{G}. \quad (4.31)$$

4.4 Fidélité

4.4.1 Définition

La mesure usuelle pour comparer deux états quantiques est la fidélité $\mathcal{F}$ [109], qui correspond à la probabilité que l'on a de confondre ces deux états en effectuant la mesure optimale. D'après l'équation (3.25), la fidélité se calcule aisément à partir de la fonction de Wigner. Dans le cas de deux états gaussiens, elle s'écrit donc

$$\mathcal{F} = \frac{N_0^n}{\pi^n \sqrt{|K_1 K_2|}} \int \dots \int d\vec{A} e^{-\frac{1}{2} [(\vec{A} - \vec{\mu}_1)^T K_1^{-1} (\vec{A} - \vec{\mu}_1) + (\vec{A} - \vec{\mu}_2)^T K_2^{-1} (\vec{A} - \vec{\mu}_2)]}. \quad (4.32)$$

Si on se limite aux états à 1 mode, avec des matrices de covariance diagonales dans la même base, cette égalité peut se calculer plus simplement. La condition de diagonalité des matrices de covariance exclut de ces calculs les états comprimés dans des directions inclinées l'une par rapport à l'autre. Les états 1 et 2 que l'on considérera seront définis par leurs variances $\Delta Q_{1,2}$, $\Delta P_{1,2}$, et leurs valeurs moyennes $q_{1,2}$, $p_{1,2}$.

$$\mathcal{F} = \frac{N_0}{\pi \Delta Q_1 \Delta P_1 \Delta Q_2 \Delta P_2} \int dq e^{-\frac{(q-q_1)^2}{2\Delta Q_1^2} - \frac{(q-q_2)^2}{2\Delta Q_2^2}} \int dp e^{-\frac{(p-p_1)^2}{2\Delta P_1^2} - \frac{(p-p_2)^2}{2\Delta P_2^2}}, \quad (4.33)$$

Les intégrales en q et p peuvent se calculer séparément. L'intégrande de l'intégrale en q vaut

$$e^{-\frac{(q-q_1)^2}{2\Delta Q_1^2} - \frac{(q-q_2)^2}{2\Delta Q_2^2}} = e^{-\frac{1}{2} \left(\frac{1}{\Delta Q_1^2} + \frac{1}{\Delta Q_2^2} \right) q^2 - \left(\frac{q_1}{\Delta Q_1^2} + \frac{q_2}{\Delta Q_2^2} \right) q - \frac{q_1^2}{2\Delta Q_1^2} - \frac{q_2^2}{2\Delta Q_2^2}} \quad (4.34)$$

et peut s'intégrer en utilisant l'équation (A.16). Cette intégrale se simplifie ensuite avec des calculs un peu longs mais ne présentant aucune difficulté en

$$\sqrt{\frac{2\pi}{\Delta Q_1^2 + \Delta Q_2^2}} \Delta Q_1 \Delta Q_2 e^{-\frac{(q_1 - q_2)^2}{2(\Delta Q_1^2 + \Delta Q_2^2)}}. \quad (4.35)$$

Cette valeur, ainsi que l'intégrale équivalente en p , peut être injectée dans (4.33), qui devient

$$\mathcal{F} = \frac{2N_0}{\sqrt{(\Delta Q_1^2 + \Delta Q_2^2)(\Delta P_1^2 + \Delta P_2^2)}} e^{-\frac{(q_1 - q_2)^2}{2(\Delta Q_1^2 + \Delta Q_2^2)} - \frac{(p_1 - p_2)^2}{2(\Delta P_1^2 + \Delta P_2^2)}} \quad (4.36)$$

A variances fixées, cette fidélité est maximale pour $q_1 = q_2$ et $p_1 = p_2$: deux gaussiennes se recouvrent d'autant mieux qu'elles sont centrées au même endroit. De plus, on a les inégalités de Schwarz

$$\Delta Q_1^2 + \Delta Q_2^2 \geq 2 \Delta Q_1 \Delta Q_2 \quad (4.37a)$$

$$\Delta P_1^2 + \Delta P_2^2 \geq 2 \Delta P_1 \Delta P_2, \quad (4.37b)$$

qui sont saturées si et seulement si les variances sont égales entre 1 et 2. On a donc

$$\mathcal{F} \leq \frac{N_0}{\sqrt{\Delta Q_1 \Delta Q_2 \Delta P_1 \Delta P_2}} = 1, \quad (4.38)$$

la fidélité étant maximale pour des variances égales.

Cette valeur est toujours plus petite que 1 lorsque les états 1 et 2 sont des mélanges statistiques, même s'ils sont indiscernables et ont la même matrice densité. En effet, la fidélité tient compte du fait qu'ils risquent alors de ne pas être dans le même état pur.

4.4.2 Fidélité d'un canal

La formule (4.36) est surtout utilisée pour comparer les états à l'entrée et à la sortie d'un canal décrit par les équations (4.15). La fidélité dépendra de l'état à l'entrée du canal. On choisira ici un état comprimé en Q , de facteur de compression s et centré en $\left[\frac{Q_i}{P_i}\right]$, pour rester le plus général possible. Si on note $\beta \equiv \frac{\chi_Q}{\chi}$ et $\gamma \equiv \frac{g_Q^2}{G}$, l'équation (4.36) devient

$$\mathcal{F} = \frac{2}{\sqrt{(s + G\gamma s + G\gamma\beta\chi) \left(\frac{1}{s} + \frac{G}{\gamma s} + \frac{G\chi}{\beta\gamma}\right)}} e^{-\frac{(1-g_Q)^2 Q_i^2}{2(s+G\gamma s+G\gamma\beta\chi)N_0} - \frac{(1-g_P)^2 P_i^2}{2\left(\frac{1}{s} + \frac{G}{\gamma s} + \frac{G\chi}{\beta\gamma}\right)N_0}}. \quad (4.39)$$

Si les gains g_Q et g_P ne sont pas égaux à 1, cette fidélité dépendra de la valeur moyenne $\left[\frac{Q_i}{P_i}\right]$ de l'état considéré. Un canal n'est intéressant que parce qu'on peut y envoyer plusieurs états différents. Nous nous intéresserons donc à la fidélité moyenne $\overline{\mathcal{F}}$ d'un tel canal. Nous supposerons que, les états envoyés ne diffèrent que par leurs valeurs moyennes. Si celles-ci sont distribuées suivant des lois gaussiennes indépendantes de variances $\sigma_Q^2 N_0$ et $\sigma_P^2 N_0$, il est aisé de calculer

$$\overline{\mathcal{F}} = \frac{2}{\sqrt{((1 - \sqrt{\gamma G})^2 \alpha \sigma^2 + s + G\gamma s + G\gamma\beta\chi) \left(\left(1 - \sqrt{\frac{G}{\gamma}}\right)^2 \frac{\sigma^2}{\alpha} + \frac{1}{s} + \frac{G}{\gamma s} + \frac{G\chi}{\beta\gamma} \right)}}, \quad (4.40)$$

où on a noté $\sigma^2 \equiv \sigma_Q \sigma_P$ et $\alpha \equiv \frac{\sigma_Q^2}{\sigma^2}$. Cette équation se simplifie dans le cas symétrique, où tout est indépendant de la phase, c'est-à-dire où $s = \gamma = \beta = \alpha = 1$. On a alors

$$\boxed{\overline{\mathcal{F}} = \frac{2}{\sigma^2 + 1 - 2\sqrt{G}\sigma^2 + G(\sigma^2 + 1 + \chi)}}. \quad (4.41)$$

Dans le cas où $G = 1$, cette fidélité est indépendante de la variance σ^2 . Cela est dû au fait que la fidélité devient indépendante de la valeur moyenne de l'état envoyé.

4.4.3 Fidélité optimisée

La perte d'information dans le canal est fonction du bruit équivalent, caractérisé par χ et β , mais pas du gain, caractérisé par G et γ . Dans beaucoup de problèmes (téléportation et clonage quantique, notamment), les valeurs des paramètres décrivant les bruits sont bornées inférieurement par une relation de Heisenberg, qui s'écrit souvent sous la forme $\chi \geq \chi_{\min}$. Il est souvent intéressant de calculer la fidélité optimale autorisée par de tels canaux, qui s'obtient par optimisation sur les paramètres G et γ en fixant $\chi = \chi_{\min}$.

Pour simplifier l'analyse, nous étudierons dans un premier temps le cas où tout est indépendant de la phase. Le cas plus général sera traité plus bas.

$\overline{\mathcal{F}}$, donné par (4.41), est maximal lorsque le trinôme en $\sqrt{G}$ au dénominateur est minimal, ce qui est obtenu pour le gain optimal

$$\sqrt{G_{\text{opt}}} = \frac{\sigma^2}{1 + \chi_{\min} + \sigma^2} < 1. \quad (4.42)$$

La valeur de G_{opt} donnée par (4.42) est inférieure à 1 pour les variances σ^2 finies, en raison de la régression vers la moyenne, bien connue des statisticiens [108]. En effet, cette valeur du gain sert en quelque sorte à compenser l'élargissement dû au bruit³. L'écart à 1 devient négligeable pour les grandes variances de modulation.

Cette valeur n'est malheureusement pas toujours physique. Par exemple, lorsque $\chi_{\min} = 0$, on trouve $\sqrt{G_{\text{opt}}} = \frac{\sigma^2}{1+\sigma^2} \neq 1$. Or pour cette valeur du gain, l'équation (4.22) interdit d'avoir $\chi = \chi_{\min} = 0$, hypothèse pourtant nécessaire pour établir l'équation (4.42).

Plus généralement, pour les autres valeurs de χ , l'inégalité (4.22) limite les valeurs de G accessibles à $\chi = \chi_{\min}$ constant. La valeur optimale donnée par (4.42) reste accessible pour

$$\sigma^2 \geq \frac{1+\chi_{\min}}{\chi_{\min}} (\sqrt{1 + \chi_{\min}} + 1). \quad (4.43)$$

Si cette inégalité n'est pas vérifiée, la valeur de χ n'est plus limitée aux faibles gains par $\chi_{\min}$ mais par les pertes, à $\chi = \frac{1}{G} - 1 > \chi_{\min}$. On a alors

$$\overline{\mathcal{F}} = \frac{2}{\sigma^2 + 2 - 2\sqrt{G}\sigma^2 + G\sigma^2}, \quad (4.44)$$

qui croît avec G et est maximal pour $G = 1$ ($\chi = 0$). Si (4.43) n'est pas vérifiée, $\overline{\mathcal{F}}$ est maximal au moment où $G = \frac{1}{1+\chi_{\min}}$ et on a

$$\overline{\mathcal{F}}_{\text{opt}} = \frac{2(1 + \chi_{\min})}{2(1 + \chi_{\min}) + \sigma^2(\sqrt{1 + \chi_{\min}} - 1)^2} \quad (4.45)$$

Nous nous intéresserons dans la suite au cas où σ est assez grand, c'est-à-dire où l'inégalité (4.43) est vérifiée. La fidélité optimale est alors calculée lorsque G est donné par (4.42), et vaut

$$\boxed{\overline{\mathcal{F}}_{\text{opt}} = \frac{2(1 + \chi + \sigma^2)}{1 + \chi + \sigma^2(2 + \chi)} \xrightarrow{\sigma^2 \rightarrow \infty} \frac{2}{2 + \chi}}. \quad (4.46)$$

Des calculs similaires nous permettent de traiter le cas asymétrique en Q et P , et d'obtenir, pour des variances σ^2 assez grandes,

$$\overline{\mathcal{F}}_{\text{opt}} = \sqrt{\frac{4(s + \beta\chi_{\min} + \alpha\sigma^2) \left(\frac{1}{s} + \frac{1}{\beta}\chi_{\min} + \frac{1}{\alpha}\sigma^2 \right)}{(s^2 + s\beta\chi_{\min} + \alpha\sigma^2(2s + \beta\chi_{\min})) \left(\frac{1}{s^2} + \frac{1}{\beta s}\chi_{\min} + \frac{1}{\alpha}\sigma^2 \left(\frac{2}{s} + \frac{1}{\beta}\chi_{\min} \right) \right)}} \quad (4.47)$$

³Plus quantitativement, si la variable aléatoire X_B est définie par $X_B = g(X + B)$, où X et B sont deux variables aléatoires indépendantes de variance connues, et que l'on cherche à minimiser l'écart quadratique moyen $\langle (X_B - X)^2 \rangle = (g - 1)^2 \langle X^2 \rangle + g^2 \langle B^2 \rangle$, il est aisé de se convaincre que la valeur optimale du gain est $g = \frac{\langle X^2 \rangle}{\langle X^2 \rangle + \langle B^2 \rangle} < 1$. Cette valeur correspond bien à celle donnée par (4.42), si on tient compte du bruit de photons et du bruit ajouté par le canal.

À la limite des fortes modulations, on a

$$\overline{\mathcal{F}}_{\text{opt}} \xrightarrow{\sigma^2 \rightarrow \infty} \frac{2}{\sqrt{(2s + \beta\chi_{\min}) \left(\frac{2}{s} + \frac{1}{\beta}\chi_{\min} \right)}} = \frac{2}{\sqrt{4 + 2\chi_{\min} \left(\frac{s}{\beta} + \frac{\beta}{s} \right) + \chi_{\min}^2}}. \quad (4.48)$$

On voit alors que $\overline{\mathcal{F}}_{\text{opt}}$ est maximum pour $\beta = s$, c'est à dire que les meilleurs canaux pour transmettre des états comprimés sont ceux où le bruit à la même structure que le bruit quantique de ces états.

4.4.4 Fidélité de la reconstruction

Les calculs précédents sont assez généraux, mais ne s'appliquent pas au canal constitué par une mesure et une reconstruction, introduit section 4.3.4. En effet, dans ce cas χ est donné par l'équation (4.31) et dépend de G . Dans ce cas l'équation (4.41) devient alors

$$\overline{\mathcal{F}} = \frac{2}{\sigma^2 + 2 - 2\sqrt{G}\sigma^2 + G(\sigma^2 + 2)}. \quad (4.49)$$

Le gain optimal vaut alors

$$\sqrt{G} = \frac{\sigma^2}{\sigma^2 + 2} \quad (4.50)$$

et la fidélité optimisée de ce canal vaut

$$\boxed{\overline{\mathcal{F}}_{\text{opt}} = \frac{\sigma^2 + 2}{2(\sigma^2 + 1)} \xrightarrow{\sigma^2 \rightarrow \infty} \frac{1}{2}.} \quad (4.51)$$

Ces calculs de fidélité nous seront utiles au chapitre 9 pour caractériser les systèmes de téléportation quantique.

Deuxième partie

Détection homodyne impulsionnelle

Chapitre 5

Mesure homodyne impulsienne du vide

5.1 Pourquoi et comment mesurer une quadrature ?

Le but initial de cette thèse était de construire un système de détection homodyne rapide, afin de mettre en évidence les propriétés statistiques non classiques d'états non gaussiens [7, 114, 121, 122]. Un tel système permet en effet de reconstituer la fonction de Wigner d'un état, en accumulant une quantité significative d'échantillons statistiques en suffisamment peu de temps pour pouvoir négliger les dérives lentes de la phase présentes dans le dispositif expérimental.

L'équilibrage de la détection homodyne s'est révélé beaucoup plus délicat que prévu, et nous avons été devancés par Hansen et ses collaborateurs, à l'université de Constance [113, 123], qui ont reconstitué la fonction de Wigner d'états de Fock à un photon [114, 115] et mis en évidence le caractère non classique de ces états [114, 121].

Ce type de système est également un composant fondamental pour la réalisation expérimentale de protocoles de communication quantique avec des variables continues [14, 91, 124]. Ainsi, le système de cryptographie quantique étudié dans la partie IV de cette thèse est construit autour d'une détection homodyne. Le taux de répétition élevé se traduit alors directement par un débit important [14, 15].

Nous envisageons également, à plus long terme, de l'utiliser pour des expériences d'inégalités de Bell avec des variables continues. [16, 125].

Nous étudierons (section 5.2) les principes de fonctionnement d'une détection homodyne, et le rôle crucial que joue l'équilibrage dans le cas d'une détection impulsienne (section 5.3). Nous examinerons ensuite les particularités des systèmes impulsions (section 5.4). Nous terminerons ce chapitre par les résultats expérimentaux d'une mesure homodyne du vide (section 5.5), qui montre que notre système de détection est limité au bruit de photons.¹

¹L'étude expérimentale de toutes les difficultés posées par le caractère impulsien de notre détection homodyne a été faite avec l'aide de Gao Jiangrui, en stage post-doctoral dans notre laboratoire, et l'équilibrage de la détection homodyne a été obtenu en juillet 2001 [126], en collaboration avec Jérôme Wenger, alors en stage de DEA.

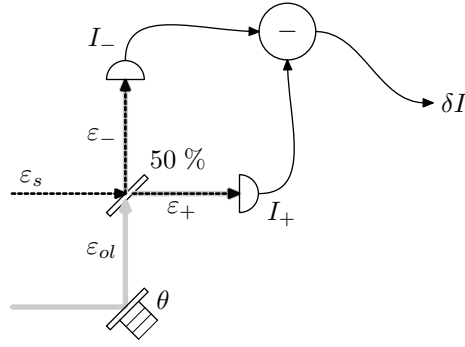


FIG. 5.1 – Schématisation d'une détection homodyne

5.2 Principes de fonctionnement d'une détection homodyne

Nous étudierons ici le fonctionnement d'une détection homodyne. La section 5.2.1 en est une étude classique, qui nous permet d'en comprendre les principes de fonctionnement. L'étude classique privilégiait un point de vue ondulatoire, et nous mentionnerons brièvement (section 5.2.2) une explication corpusculaire, avant d'en faire une véritable description quantique (section 5.2.3) qui justifiera le point de vue classique précédent.

5.2.1 Étude classique

Nous étudierons dans cette section le fonctionnement d'une détection homodyne en des termes purement classiques, qui correspondent à une description ondulatoire de l'état du champ électromagnétique. Comme une détection homodyne équilibrée mesure une quadrature du champ, quantité fondamentalement ondulatoire, cette approche permet de bien comprendre les principes de fonctionnement d'un tel détecteur. Les effets quantiques pourront être introduits par la statistique *ad hoc* de ces quadratures.

Dans une détection homodyne équilibrée, un champ signal $\mathcal{E}_s$ faible interfère avec un oscillateur local beaucoup plus intense $\mathcal{E}_{ol}$ sur une lame semi-réfléchissante. L'intensité des deux ports de sortie + et - est mesurée et le signal de la détection est la différence de ces deux intensités.

Le champ de l'oscillateur local est décrit par le nombre complexe $\mathcal{E}_{ol} \equiv r_{ol} e^{i\theta} e^{i\omega t}$. L'état du champ signal de même fréquence est alors décrit par le nombre complexe

$$\mathcal{E}_s = (q_s + ip_s) e^{i\omega t} = (q_\theta + ip_\theta) e^{i\theta} e^{i\omega t}, \quad (5.1)$$

avec $|q_\theta|, |p_\theta| \ll r_{ol}$. Le champ dans les deux ports de sortie de la lame semi-réfléchissante s'écrit alors

$$\mathcal{E}_\pm = \frac{1}{\sqrt{2}} (q_\theta + ip_\theta \pm r_{ol}) e^{i\theta + i\omega t} \quad (5.2)$$

La seule différence entre les deux modes est le changement de signe indiqué $\pm$ dans l'équa-

tion ci-dessus. Comme l'intensité mesurée est le module carré du champ², on a

$$I_{\pm} = \frac{1}{2}[(q_{\theta} \pm r_{\text{ol}})^2 + p_{\theta}^2] \quad (5.3a)$$

$$= \frac{1}{2}[r_{\text{ol}}^2 \pm 2r_{\text{ol}} q_{\theta} + p_{\theta}^2 + q_{\theta}^2] \quad (5.3b)$$

$$= \frac{I_{\text{ol}}}{2} \pm \sqrt{I_{\text{ol}}} q_{\theta} + \frac{I_s}{2}. \quad (5.3c)$$

Le terme dominant, $\frac{I_{\text{ol}}}{2}$, est identique pour les intensités I_+ et I_- . Le terme suivant, environ $\sqrt{\frac{I_s}{I_{\text{ol}}}}$ fois plus petit, est le seul qui subsiste dans la différence δI :

$$\delta I \equiv I_+ - I_- = 2\sqrt{I_{\text{ol}}} q_{\theta}. \quad (5.4)$$

Le signal de sortie de la détection homodyne, proportionnel à δI , est donc proportionnel à la quadrature Q_{θ} . L'angle θ est déterminé par la phase de l'oscillateur local, qui peut être facilement ajustée expérimentalement, par une cale piézoélectrique par exemple.

Comme δI est proportionnel à l'amplitude $\sqrt{I_{\text{ol}}}$ de l'oscillateur local, ce dernier agit comme un amplificateur optique, ce qui permet de mesurer des champs d'amplitude q_{θ} très faible avec des photodiodes ordinaires, du moment que l'oscillateur local est suffisamment intense.

En pratique, l'intensité de l'oscillateur local présentera des fluctuations, petites devant I_{ol} , mais grandes devant le signal I_s . Ces fluctuations masqueront le signal dans I_+ et I_- , mais elles seront identiques sur les deux voies. Les fluctuations de l'oscillateur local se traduiront donc par des corrélations entre les deux voies, qui disparaîtront de la différence δI si la détection est bien équilibrée. Cette soustraction du bruit de l'oscillateur local et l'amplification optique du signal permettent de mesurer des signaux extrêmement faibles et de mettre en évidence des effets quantiques.

5.2.2 Interprétation ondulatoire et interprétation corpusculaire

S'il n'y a pas de champ signal, il serait tentant de supposer que $q_{\theta} = 0$, et que δI doit être nul, mais cela ne correspond pas aux observations expérimentales. Cependant, l'absence de champ signal signifie qu'on est dans l'état fondamental de l'oscillateur harmonique décrit par l'équation (2.38) : q_{θ} présente des petites fluctuations gaussiennes d'origine quantique, même en l'absence de signal. On peut donc écrire

$$q_{\theta} = \sqrt{N_0} \mathcal{G}, \quad (5.5)$$

où $\mathcal{G}$ représente une variable aléatoire réelle gaussienne centrée de variance 1. Comme δI reproduit cette distribution, mais amplifiée d'un facteur $2\sqrt{I_{\text{ol}}}$, on en déduit

$$\delta I = 2\sqrt{I_{\text{ol}}} \sqrt{N_0} \mathcal{G}, \quad (5.6)$$

ce qui correspond aux mesures expérimentales.

²L'intensité est souvent définie comme *la moitié* du module carré du champ et non comme le module carré du champ.

Cette interprétation ondulatoire du bruit observé à la sortie d'une détection homodyne en l'absence de signal n'est pas la seule : une approche corpusculaire décrit l'oscillateur local comme un ensemble constitué de $n_{\text{ol}} = \frac{I_{\text{ol}}}{4N_0}$ photons. Chaque photon, en l'absence de photon signal, a une chance sur deux d'aller dans la voie + et une chance sur deux d'aller dans la voie -. La probabilité d'avoir n_+ photons dans la voie + sera donc donnée par la loi binomiale de paramètres n_{ol} et $\frac{1}{2}$ [108]. Comme l'oscillateur local est brillant par hypothèse, $n_{\text{ol}} \gg 1$. Cette loi binomiale est donc pratiquement identique à une loi gaussienne de mêmes moyenne ($\frac{n_{\text{ol}}}{2}$) et variance ($\frac{n_{\text{ol}}}{4}$). On peut alors écrire

$$n_+ \simeq \frac{n_{\text{ol}}}{2} + \sqrt{\frac{n_{\text{ol}}}{4}} \mathcal{G}. \quad (5.7)$$

La différence du nombre de photons dans les deux voies s'écrit alors

$$\delta n \equiv n_+ - n_- = 2n_+ - n_{\text{ol}} = \sqrt{n_{\text{ol}}} \mathcal{G}. \quad (5.8)$$

Comme $\delta I = 4N_0 \delta n$, on retrouve bien le même résultat que l'équation (5.6) :

$$\delta I = 2\sqrt{N_0} \underbrace{\sqrt{4N_0 n_{\text{ol}}}}_{\sqrt{I_{\text{ol}}}} \mathcal{G}. \quad (5.9)$$

Cette interprétation peut sembler plus naturelle que de postuler des fluctuations intrinsèques au vide quantique, et c'est pourquoi on parle plus souvent de *bruit de photons* que de mesure des fluctuations du vide quantique pour expliquer ce type de mesures. Cependant, les deux explications sont strictement équivalentes, et l'interprétation corpusculaire devient plus complexe en présence d'un signal différent du vide.

Dans ce cas, l'explication corpusculaire doit tenir compte de la statistique bosonique des photons alors que l'interprétation ondulatoire ne fait que reproduire la statistique de la quadrature Q_θ , au facteur d'échelle $2\sqrt{I_{\text{ol}}}$ près. Dans la référence [127], Ou étudie en termes corpusculaires ce qui se passe lorsque le champ signal est dans un état de Fock à un photon. Il faut alors expliquer comment l'unique photon du mode signal parvient à influencer le bruit de partition des nombreux (n_{ol}) photons de l'oscillateur local, en « déplaçant » en quelque sorte un grand nombre $\sim \sqrt{n_{\text{ol}}}$ de photons d'une voie à l'autre. Les calculs sont nettement plus compliqués que pour l'approche ondulatoire, qui se réduit à supposer que δI reproduit la statistique de la quadrature Q_θ du signal à l'entrée, amplifiée du facteur $2\sqrt{I_{\text{ol}}}$ [122].

5.2.3 Étude quantique

Les calculs de la section 5.2.1 étaient purement classiques ; les effet quantiques sont apparus lorsqu'on a introduit la statistique *ad hoc* de la quadrature $\hat{Q}_\theta$. Le caractère quantique des champs peut être inclus dès le début de l'analyse, ce qui donne un fondement plus rigoureux aux résultats de la section précédente.

Les observables décrivant l'intensité des modes + et - peuvent s'écrire au moyen des opérateurs amplitude définis dans la section B.9 :

$$\hat{I}_\pm = 4N_0 \hat{a}_\pm^\dagger \hat{a}_\pm = 2N_0 (\hat{a}_{\text{ol}}^\dagger \pm \hat{a}_s^\dagger) (\hat{a}_{\text{ol}} \pm \hat{a}_s) \quad (5.10a)$$

$$= 2N_0 (\hat{a}_{\text{ol}}^\dagger \hat{a}_{\text{ol}} + \hat{a}_s^\dagger \hat{a}_s \pm \hat{a}_{\text{ol}}^\dagger \hat{a}_s \pm \hat{a}_s^\dagger \hat{a}_{\text{ol}}) \quad (5.10b)$$

$$= \frac{1}{2} (\hat{I}_{\text{ol}} + \hat{I}_s \pm \delta \hat{I}) \quad (5.10c)$$

et la différence d'intensité $\delta\hat{I}$ s'écrit

$$\delta\hat{I} \equiv \hat{I}_+ - \hat{I}_- = 4N_0(\hat{a}_{\text{ol}}^\dagger \hat{a}_s + \hat{a}_s^\dagger \hat{a}_{\text{ol}}). \quad (5.11)$$

Certaines expériences d'inégalités de Bell utilisent ce type de détection homodyne avec des oscillateurs locaux peu brillants en régime de comptage de photons. Dans ce cas, il faut utiliser directement l'opérateur $\delta\hat{I}$ tel qu'il est défini ci dessus.

Cependant, dans les expériences de variables continues, l'oscillateur local est un état cohérent intense, $\left| \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right\rangle = \left| \frac{r_{\text{ol}}}{\sqrt{4N_0}} e^{i\theta} \right\rangle$ avec $r_{\text{ol}} \gg \sqrt{N_0}$, et le signal est dans un état quelconque $|\psi\rangle$, d'amplitude très inférieure à r_{ol} . La statistique de l'opérateur $\delta\hat{I}$ étant définie par sa fonction caractéristique $\langle e^{i\Omega\delta\hat{I}} \rangle$ [108]. Cette dernière est déterminée de manière unique par les moments³ $\langle \delta\hat{I}^n \rangle$ qu'on peut calculer dans ces conditions par le produit scalaire

$$(4N_0)^n \left\langle \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right| \otimes \langle \psi | (\hat{a}_{\text{ol}}^\dagger \hat{a}_s + \hat{a}_s^\dagger \hat{a}_{\text{ol}})^n | \psi \rangle \otimes \left| \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right\rangle \quad (5.12)$$

Le calcul exact est peu évident en raison de la non commutation des opérateurs. Cependant, dans le cas de l'oscillateur brillant, $\hat{a}_{\text{ol}}^\dagger$ et $\hat{a}_{\text{ol}}$ commutent presque. En effet, pour calculer le produit haut, il faut ordonner les termes de façon à pouvoir utiliser le fait que l'état cohérent est un état propre de l'opérateur annihilation (2.40), d'où découlent les équations

$$\hat{a}_{\text{ol}} \left| \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right\rangle = \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \left| \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right\rangle \quad \text{et} \quad \left\langle \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right| \hat{a}_{\text{ol}}^\dagger = \frac{\mathcal{E}_{\text{ol}}^*}{\sqrt{4N_0}} \left\langle \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right| \quad (5.13)$$

Le réordonnement des termes utilise la relation

$$\hat{a}_{\text{ol}} \hat{a}_{\text{ol}}^\dagger = \hat{a}_{\text{ol}}^\dagger \hat{a}_{\text{ol}} + 1. \quad (5.14)$$

Une fois tous les termes remis dans l'ordre normal, la propagation de la relation précédente nous permet donc de remplacer $\hat{a}_{\text{ol}} \hat{a}_{\text{ol}}^\dagger$ par $\frac{\mathcal{E}_{\text{ol}}^* \mathcal{E}_{\text{ol}}}{4N_0} + 1$. Si le faisceau est très brillant, le 1 est négligeable devant $\frac{\mathcal{E}_{\text{ol}}^* \mathcal{E}_{\text{ol}}}{4N_0}$ et on a

$$(4N_0)^n \left\langle \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right| \otimes \langle \psi | (\hat{a}_{\text{ol}}^\dagger \hat{a}_s + \hat{a}_s^\dagger \hat{a}_{\text{ol}})^n | \psi \rangle \otimes \left| \frac{\mathcal{E}_{\text{ol}}}{\sqrt{4N_0}} \right\rangle \simeq \langle \psi | (4N_0)^{\frac{n}{2}} (\mathcal{E}_{\text{ol}}^* \hat{a}_s + \mathcal{E}_{\text{ol}} \hat{a}_s^\dagger)^n | \psi \rangle \quad (5.15)$$

On a donc

$$\delta\hat{I} \simeq \sqrt{4N_0} \left(\mathcal{E}_{\text{ol}}^* \hat{a}_s + \mathcal{E}_{\text{ol}} \hat{a}_s^\dagger \right) = r_{\text{ol}} \sqrt{4N_0} \left(e^{-i\theta} \hat{a}_s + e^{i\theta} \hat{a}_s^\dagger \right) = 2r_{\text{ol}} \hat{Q}_\theta \quad (5.16)$$

L'observable $\delta\hat{I}$ est donc bien proportionnelle à l'observable $\hat{Q}_\theta$. Elles ont donc bien la même statistique, à un facteur d'échelle $2r_{\text{ol}} = 2\sqrt{I_{\text{ol}}}$ près, exactement comme on l'a vu par des arguments classiques section 5.2.1.

5.3 Équilibrage d'une détection homodyne

La section 5.3.1 est consacrée à l'étude théorique des effets d'un déséquilibre sur une détection homodyne, qui sont plus importants sur une détection homodyne impulsienne que sur les détections homodynes habituelles, qui fonctionnent en régime continu. Nous illustrerons, dans la section 5.3.2, le rôle crucial joué par ce paramètre en exposant certaines difficultés expérimentales rencontrées lors de l'équilibrage de notre système.

³Nous supposons ici que ces moments sont définis, ce qui est en général le cas.

5.3.1 Effet d'un déséquilibre

Jusqu'ici, nous avons supposé que la lame semi-réfléchissante était parfaitement équilibrée. En pratique, elle sera légèrement déséquilibrée et son coefficient de transmission vaudra $T = \frac{1}{2} + \varepsilon$ avec $|\varepsilon| \ll \frac{1}{2}$. L'approche classique que nous avons utilisée section 5.2.1 sera suffisante pour connaître la valeur maximale tolérable de $|\varepsilon|$.

L'équation (5.2) devient alors

$$\mathcal{E}_{\pm} = \mathcal{E}_{\pm}^{\text{eq}} \mp \varepsilon \mathcal{E}_{\mp}^{\text{eq}} \quad (5.17)$$

où $\mathcal{E}_{\pm}^{\text{eq}}$ représente la valeur de $\mathcal{E}_{\pm}$ lorsque la lame semi-réfléchissante est parfaitement équilibrée. On a alors

$$I_{\pm} = I_{\pm}^{\text{eq}} \pm 2\varepsilon(I_{\text{ol}} - I_s) + \varepsilon^2 I_{\mp}^{\text{eq}}. \quad (5.18)$$

Le changement de signe du terme proportionnel à ε montre qu'il va contaminer δI , qui devient, lorsqu'on néglige la petite contribution εI_s ,

$$\delta I \simeq (1 - \varepsilon^2)\delta I^{\text{eq}} + 4\varepsilon I_{\text{ol}} \simeq 2\sqrt{I_{\text{ol}}}(q_{\theta} + 2\varepsilon\sqrt{I_{\text{ol}}}). \quad (5.19)$$

Dans les détectations homodynes usuelles, résolues en fréquence, la composante continue de $\sqrt{I_{\text{ol}}}$ est éliminée par un filtrage spectral, de sorte que seules les fluctuations de I_{ol} contribuent au terme $2\varepsilon\sqrt{I_{\text{ol}}}$. Par contre, dans une détection impulsionnelle, la composante brillante $\sqrt{I_{\text{ol}}}$ a une structure temporelle trop proche de celle du signal pour pouvoir être filtrée, optiquement ou électroniquement. Dans ce cas, le terme $2\varepsilon\sqrt{I_{\text{ol}}}$ est dominé par la composante brillante de $\sqrt{I_{\text{ol}}}$, ce qui rend particulièrement critique l'équilibrage des deux voies de la détection homodyne.

Pour mesurer des signaux d'amplitude typique q_{θ} , il faut donc

$$|\varepsilon| \ll \frac{q_{\theta}}{2\sqrt{I_{\text{ol}}}}. \quad (5.20)$$

Pour exploiter une détection homodyne dans le domaine de l'information quantique, il faut avoir une résolution meilleure que $\sqrt{N_0}$, c'est-à-dire qui puisse mesurer les fluctuations du vide. En d'autres termes, on doit avoir

$$\boxed{|\varepsilon| \ll \frac{1}{4\sqrt{n_{\text{ol}}}}}, \quad (5.21)$$

où $n_{\text{ol}} = \frac{I_{\text{ol}}}{4N_0}$ désigne le nombre de photons de l'oscillateur local. Comme nous utilisons pour l'oscillateur local des impulsions d'environ 10^8 photons, cela exige un équilibrage à mieux que $2 \cdot 10^{-5}$ près.

En pratique, pour vérifier si la détection est suffisamment bien équilibrée, on mesure la variance de δI pour différentes valeurs de I_{ol} en mettant le vide pour tout signal. On a alors

$$\langle \delta I^2 \rangle = N_B + 4N_0 I_{\text{ol}} + 16 \langle \varepsilon^2 I_{\text{ol}}^2 \rangle, \quad (5.22)$$

où N_B représente le bruit électronique de la détection. Il faut choisir I_{ol} pour que le terme dominant dans cette variance soit celui qui nous intéresse, c'est-à-dire celui qui dépend linéairement de I_{ol} .

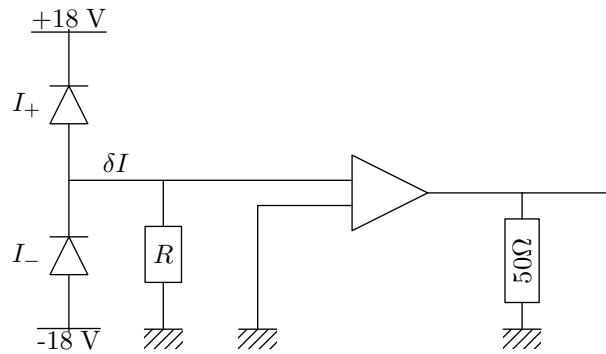


FIG. 5.2 – Schématisation du système utilisé pour la soustraction des photocourants

5.3.2 De la difficulté expérimentale d'équilibrer une détection homodyne impulsionnelle

Malheureusement, au cours du développement expérimental, il nous est fréquemment arrivé de constater que la détection n'était pas suffisamment équilibrée, et que le terme linéaire de l'équation (5.22) était largement négligeable devant le terme quadratique dû au déséquilibre et le terme constant dû au bruit électronique, quand la courbe était assez régulière pour que l'on puisse en déduire quoi que ce soit. Nous avons découvert au cours de nos deux premières années de thèse de multiples causes de déséquilibre de la détection homodyne ; nous en mentionnerons quelques-unes ici.

La première question qui se pose est le choix du dispositif utilisé pour la soustraction des photocourants mesuré par les deux photodiodes. Plusieurs raisons nous ont conduit à effectuer cette soustraction avant que le signal ne soit amplifié, par une simple loi des nœuds, comme schématisé FIG. 5.2. En effet, l'impossibilité de filtrer électroniquement la composante I_{0l} qui domine des photocourants proportionnels à $I_{\pm}$ élimine en pratique la possibilité de les amplifier séparément avant d'en extraire la petite composante $2\sqrt{I_{0l}} \hat{Q}_{\theta}$, environ dix mille fois plus petite. Cela exigerait deux amplificateurs identiques entre eux à 10^{-5} près sur l'ensemble du spectre, sans parler des problèmes éventuels de saturation.

De plus, tout déséquilibre ε de l'ordre de 10^{-4} risque donc de se traduire par une saturation de l'amplificateur, ce qui induit au mieux un biais dans la statistique observée de Q_{θ} . Nous avons donc choisi, dans un lot de vingt photodiodes *Centronix* BPX 65, les deux photodiodes dont les réponses se ressemblaient le plus, et dont l'efficacité quantique est maximale. Ces photodiodes sont polarisées en inverse par des piles, afin de réduire leur temps de réponse et d'assurer la linéarité de leur réponse.

Il faut également soigner l'équilibrage optique des deux faisceaux. Il faut notamment que le coefficient de réflexion soit stable à 10^{-5} près pendant l'intervalle de temps où l'on fait des mesures. Une première (mauvaise) idée avait été d'utiliser une lame demi-onde suivi d'un cube polariseur comme lame partiellement réfléchissante réglable, ce qui avait pour inconvénient de transformer toute variation de polarisation du faisceau en variation du coefficient de réflexion, et nous étions probablement loin d'avoir des faisceaux avec une polarisation constante à 10^{-5} près, notamment en modulant électriquement une diode laser.

Un effet analogue est un mouvement transversal du faisceau, qui induit des mouvements corrélés des taches lumineuses sur les photodiode. L'efficacité d'une photodiode n'étant pas identique sur toute sa surface, cet effet déséquilibre la détection homodyne, et le bruit de

position n'est pas soustrait mais ajouté car il se traduit par des fluctuations de ε .

Pour éviter ces bruits, il importe de soigneusement préparer le faisceau de l'oscillateur local. Nous avons finalement utilisé une diode laser SDL 5412 en cavité étendue, dont le faisceau continu est « haché » en impulsions de 120 ns par un modulateur acousto-optique. Ces impulsions sont ensuite filtrées spatialement par une fibre monomode à maintien de polarisation et un cube polariseur afin d'éviter les effets mentionnés dans les paragraphes précédents. L'équilibrage des intensités dans les voies $+$ et $-$ est assuré par deux cubes polariseurs précédés de lames demi-onde, qui permettent de régler l'intensité optique sans déplacer le faisceau.

5.4 Aspects temporels et fréquentiels

Notre système de détection homodyne est particulier car il est résolu en temps et non en fréquence. Nous étudierons ici les aspects temporels spécifiques d'un tel système. Nous commencerons par rappeler (section 5.4.1) le rôle du temps dans les détections homodynes usuelles, résolues en fréquence, avant de présenter les avantages d'une détection homodyne résolue en temps (section 5.4.2). Les aspects expérimentaux, notamment liés à l'électronique, seront abordés section 5.4.3. Nous étudierons ensuite cette dualité temps–fréquence plus quantitativement, sous certaines conditions de linéarité, exposées section 5.4.4, en définissant la réponse percutive (section 5.4.5) et ses liens avec la fonction de transfert (section 5.4.6), plus couramment utilisée en régime continu.

5.4.1 Détections homodyne résolues en fréquence

Les détections homodynes habituellement utilisées en optique quantique sont résolues en fréquence : l'oscillateur local et le signal sont des faisceaux continus monomodes. On analyse ensuite une bande étroite du spectre de la différence des photocourants δI , qui correspond à un mode étroit du signal. Cela permet en général de s'affranchir de nombreux bruits expérimentaux à basse fréquence et de choisir la bande du spectre la plus adaptée à la mesure. Cette approche peut également s'appliquer à des faisceaux impulsifs périodiques, où l'on isole une harmonique de la période de répétition, ce qui correspond à un mode réparti sur plusieurs impulsions.

De plus, un analyseur de spectre mesure directement des variances, ce qui permet de reconstituer la matrice de covariance. Un état gaussien étant complètement caractérisé par cette matrice, cette approche est très utile en optique quantique, où elle permet de mettre en évidence la production d'états comprimés ou intriqués [116, 117].

Dans l'étude théorique de la section 5.2, nous avons implicitement supposé des modes infiniment étroits en fréquence, ce qui correspond à une seule mesure de quadrature, effectuée en un temps infini. En pratique, lorsque les modes ont une largeur spectrale $\delta\nu$, une mesure prendra un temps au moins égal à $\frac{1}{2\delta\nu}$, et la mesure de données statistiquement significatives nécessite par définition un grand nombre de mesures.

Un système de communication est par définition résolu en temps, le « message » à transmettre n'étant pas connu à l'avance. Les systèmes de communication quantique n'échappent pas à cette règle, et l'étroitesse d'une bande du spectre se traduit directement sur son débit d'information [50], ce qui nous conduit à considérer des modes larges du faisceau.

5.4.2 Détection homodyne résolue en temps

La limite des modes très larges $\delta\nu \rightarrow \infty$ correspond à des modes résolus en temps. Si ce mode très large est limité par Fourier, cela correspond au mode impulsif extrême, dont l'enveloppe est une impulsion de Dirac $\delta(t)$. Ces impulsions infiniment brèves sont bien entendu aussi fictives que les modes infiniment étroits que nous avons utilisés pour exposer le principe de fonctionnement de la détection homodyne, mais elles sont aussi utiles pour exposer les principes de fonctionnement des détections impulsives résolues en temps que ces derniers le sont pour celles résolues en fréquence.

Si l'oscillateur local est constitué par une impulsion brève, modélisée par une fonction de Dirac, il échantillonnera la valeur d'une quadrature du champ signal à l'instant où l'impulsion arrive sur la lame semi-réfléchissante. Si on envoie un train d'impulsion, et que l'électronique est assez rapide pour discriminer les impulsions, on a une mesure par impulsion, ce qui permet d'envisager un grand nombre de mesures dans des temps relativement brefs.

En pratique, ces mesures sont échantillonnées par un ordinateur, ce qui permet d'en faire un traitement statistique *a posteriori* qui n'est pas limité à une mesure de variances. Cela permet de faire de la tomographie quantique où l'on doit avoir accès à l'ensemble de la distribution de probabilité des mesures de quadratures, et non seulement leur variances. Bien sûr, les systèmes continus permettent aussi de faire de la tomographie quantique, mais un système impulsif peut permettre d'accumuler une quantité significative de données en un temps plus bref. Ainsi, notre système expérimental accumule les échantillons au rythme de 800 kHz et peut accumuler 10 000 échantillons en 12,5 ms, c'est à dire accumuler des données statistiquement significatives en un temps assez bref pour négliger la plupart des dérives lentes du système expérimental, notamment les dérives de phases.

De plus, le seul état du champ expérimentalement réalisable aujourd'hui dont la fonction de Wigner présente des valeurs négatives est un état de Fock à un photon. Or ces états sont en général réalisés par production de faisceaux jumeaux, la détection d'un photon dans l'un des faisceaux indiquant la présence d'un photon dans le faisceau frère à l'instant du « clic » du premier détecteur. Le photon est donc dans un mode spatiotemporel bien précis, ce qui exige une adaptation de la détection homodyne [7, 113, 115] qui doit être également résolue en temps.

Comme nous l'avons mentionné plus haut, la résolution temporelle de la détection homodyne est fondamentale dans les protocoles de communication quantique, chaque mesure contenant en quelque sorte une information à transmettre ; une bonne résolution temporelle peut donc permettre de construire des systèmes avec des grands débits. L'approche impulsif permet d'étudier les protocoles de cryptographie quantique sur une base plus claire : chaque mode du champ étant une impulsion lumineuse, les actions possibles de l'espion sont plus clairement définies. Cet avantage n'a rien de fondamental et des mesures résolues en temps sur des faisceaux continus à large bande pourraient se révéler avantageuses à terme, mais elles exigent une étude soignée des correspondances temps-fréquences. Il faudra notamment étudier la structure spectrale des modulateurs éventuellement utilisés par Alice, toute bande latérale non-mesurée par la détection homodyne de Bob pouvant être capturée par Ève.

Par contre, une détection homodyne impulsif résolue en temps est plus délicate à réaliser expérimentalement. En effet, si on note Ω la fréquence d'échantillonnage et $T = \frac{n}{\Omega}$ l'intervalle de temps entre le premier et le dernier échantillon, les mesures recouvrent un in-

tervalle de fréquence très large, compris entre $\frac{1}{T}$ et Ω . Ce système sera donc très sensible aux bruits expérimentaux à basse fréquence, qui ne peuvent pas être filtrés sans que la résolution temporelle du système n'en pâtisse.

5.4.3 Électronique utilisée

Comme la mesure est impulsionnelle, la détection se fait nécessairement avec une bande passante plus large que le taux de répétition des impulsions, sans qu'il soit possible de filtrer les bruits indésirables. De plus, des courants à haute fréquence circulent dans d'autres parties de l'électronique pour générer des impulsions brèves, et ne manquent pas de générer des parasites qui ne demandent qu'à être captées par le circuit de détection. Toute longueur excessive de câblage, qui peut servir d'antenne doit donc être bannie, et l'électronique doit être blindée. Il faut également tenir compte des adaptations d'impédance, afin d'éviter des réflexions d'impulsions qui perturberaient le signal.

Le signal mesuré étant extrêmement faible avant amplification et très large spectralement, il était sensible à toute perturbation extérieure, et nous avons rapidement constaté que tout câble ajouté dans le circuit pour essayer d'en suivre le fonctionnement servait d'antenne pour introduire des parasites égarés cherchant un circuit à perturber. Elles changeaient également les caractéristiques électriques du système, en introduisant un léger changement dans la fonction de transfert (parfois juste un déphasage), ce qui en changeait la réponse percutive...

Nous avons utilisé des amplificateurs à bas bruit OEI AH0013CA, de 6 MHz de bande passante et de bruit équivalent ramené à l'entrée de $2 \text{ nV}/\sqrt{\text{Hz}}$. Ces amplificateurs étaient utilisés dans le groupe pour des mesures de bruit quantique en régime continu⁴ (mesures QND [82, 83], étude du bruit quantique des diodes laser [87, 88]).

Cet amplificateur étant un amplificateur de tension, les photocourants générés par les photodiodes sont convertis en tension par une résistance R de $4,7 \text{ k}\Omega$. Cette résistance induit un bruit thermique de $\sqrt{4kTR} = 9 \text{ nV}/\sqrt{\text{Hz}}$, proportionnel à $\sqrt{R}$, alors qu'elle multiplie les photocourants par R . Il semble donc intéressant d'avoir une résistance la plus grande possible, pour diminuer l'importance relative du bruit thermique. Cependant, lorsque la résistance devient trop grande, elle diminue la bande passante du premier amplificateur, ce qui diminue l'amplitude de la réponse percutive du signal. Nous avons donc choisi empiriquement la valeur de la résistance ($4,7 \text{ k}\Omega$) qui donnait le meilleur rapport signal à bruit.

Ce premier étage d'amplification est suivi d'un second amplificateur qui nous permet d'avoir des signaux assez importants (de l'ordre de 10 mV) pour être échantillonnés par une carte d'acquisition PCI-6111E, tout en ajoutant un bruit pratiquement négligeable. En pratique, la source essentielle du bruit électronique dans ce système est le bruit thermique de la résistance R .

Un soin tout particulier doit également être apporté pour éviter la saturation des amplificateurs, la composante brillante de l'oscillateur local ne pouvant pas être filtré par un simple condensateur comme dans le régime continu, puisqu'elle n'est justement pas continue.

⁴ Pour illustrer les difficultés posées par le régime impulsionnel, il suffit de mentionner que lorsque nous avons utilisé ces amplificateurs pour monter notre système impulsionnel, nous avons dû corriger quelques défauts de câblages qui étaient passés inaperçus pendant des années de régime continu !

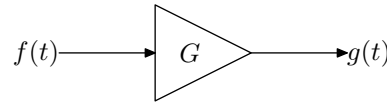


FIG. 5.3 – Schématisation d'un système linéaire

Il nous est ainsi apparu que les chemins optiques entre les photodiodes et la lame semi-réfléchissante devaient être identiques à quelques millimètres près. En effet, une différence de chemin optique induit une petite différence temporelle δt entre les intensités I_+ et I_- , ce qui conduit à mesurer

$$I_+(t) - I_-(t + \delta t) \simeq \delta I(t) + \frac{1}{2}(I_{0l}(t) - I_{0l}(t + \delta t)) \simeq \delta I(t) + \delta t \frac{dI_{0l}}{dt} \quad (5.23)$$

Ce terme supplémentaire est, suivant la durée de l'impulsion, proportionnel à la dérivée de l'impulsion ou de la réponse percutieuse des photodiodes. Si δt n'est pas assez petit, il semble induire des saturations dans la chaîne amplificatrice, qui faussent les mesures.

5.4.4 Notations

Comme nous l'avons dit plus haut, les approches temporelles et fréquentielles ne sont pas indépendantes l'une de l'autre. Nous reviendrons ici sur ces relations dans l'électronique du dispositif expérimental⁵. En effet, la différence des photocourants δI est trop faible (environ 10 nA pendant 100 ns) pour être directement mesurée par un dispositif expérimental et doit passer à travers une chaîne amplificatrice qui a une bande passante finie. Nous étudierons ici la réponse d'un amplificateur linéaire à un signal quelconque, avec les deux points de vue, fréquentiel et temporel.

Un amplificateur est un système qui à une entrée $f(t)$ associe une sortie $g(t)$. Nous négligerons ici les problèmes de bruit, et la sortie $g(t)$ sera donc entièrement déterminée par $f(t)$. Nous nous limiterons aux amplificateurs indépendants du temps et linéaires. L'indépendance par rapport au temps signifie que si on avance ou retarde l'entrée d'un intervalle de temps t_0 , la sortie sera décalée dans le temps du même délai, c'est à dire que la sortie $g(t - t_0)$ sera associée à l'entrée $f(t - t_0)$ quelque soit t_0 .

Par définition, un amplificateur linéaire associe la sortie $g_1(t) + \lambda g_2(t)$ à l'entrée $f_1(t) + \lambda f_2(t)$ pour tout $\lambda \in \mathbb{R}$ s'il associe les sorties $g_1(t)$ et $g_2(t)$ aux entrées $f_1(t)$ et $f_2(t)$. Cette propriété n'est jamais exactement vérifiée en pratique, les amplificateurs saturant pour des signaux trop importants. Cependant, ils sont en général conçus pour rester linéaires dans une certaine plage de fonctionnement. Dans ce régime, un amplificateur peut être caractérisé, soit par sa *réponse percutieuse* $G(t)$, ou *fonction de Green*, soit par sa *fonction de transfert* $G[\omega]$. La première privilégie un point de vue temporel et la dernière un point de vue fréquentiel : elles sont transformées de Fourier l'une de l'autre, comme nous l'établirons dans la suite.

⁵Les aspects quantiques de cette dualité temps-fréquence seront étudiés section 6.2.

5.4.5 Réponse percuttionnelle

5.4.5.1 Définition et propriétés

La *réponse percuttionnelle* $G(t)$, ou *fonction de Green*, est, comme son nom l'indique, la réponse de l'amplificateur à une impulsion très brève. Cette impulsion peut être modélisée par une fonction de Dirac $f(t - t_0) = \mathcal{I} \delta(t - t_0)$, où $\mathcal{I}$ est l'intégrale de cette impulsion et t_0 est l'instant où cette impulsion arrive dans l'amplificateur. La réponse $g(t)$ de l'amplificateur sera donc proportionnelle à $\mathcal{I}$ et dépendra de l'instant d'arrivée t_0 de l'impulsion. La réponse percuttionnelle $G(t)$ est alors définie par

$$G(t) \equiv \frac{g(t - t_0)}{\mathcal{I}}, \quad (5.24)$$

et ne dépendra plus de ces paramètres.

Les relations entre l'entrée et la sortie de l'amplificateur peuvent être représentées dans ce cas par les relations suivantes, pour des impulsions brèves.

$$f(t) = \mathcal{I} \delta(t - t_0) \quad \Longrightarrow \quad g(t) \equiv \mathcal{I} G(t - t_0). \quad (5.25)$$

La linéarité de l'amplificateur permet alors d'exprimer la réponse de l'amplificateur à une entrée quelconque $f(t)$ à l'aide de cette réponse percuttionnelle, en décomposant $f(t)$ en une somme de fonctions de Dirac :

$$f(t) = \int dt' f(t') \delta(t - t') \quad \Longrightarrow \quad g(t) = \int dt' f(t') G(t - t'). \quad (5.26)$$

La réponse à une fonction quelconque $f(t)$ est donc la convolution de celle-ci avec la réponse percuttionnelle $G(t)$.

Si on envoie dans l'amplificateur une impulsion de durée δt , la réponse de l'amplificateur sera donc la réponse percuttionnelle $G(t)$, où les détails d'une durée inférieure à δt sont moyennés. Ainsi, quant δt est petit devant le temps d'évolution de $G(t)$, la réponse de l'amplificateur est pratiquement identique à sa réponse percuttionnelle, à un facteur de proportionnalité près. Au contraire, quand δt est grand devant la durée totale de la réponse percuttionnelle, la structure de $G(t)$ est complètement effacée et l'amplificateur fournit un signal quasiment proportionnel à l'impulsion $f(t)$. Les impulsions de durée intermédiaire produisent des signaux qui ressemblent à la fonction $G(t)$, plus ou moins lissée selon la valeur de δt .

5.4.5.2 Mesures expérimentales

La FIG. 5.4 représente ainsi la réponse percuttionnelle de la chaîne d'amplificateurs que nous avons utilisée lors de nos premières expériences. Nous avons obtenu cette courbe en éclairant l'une des photodiodes par un train d'impulsions de 7 ns à 800 kHz environ, généré par modulation électrique d'une diode laser SDL 5412 à 780 nm montée sur un réseau. Pour pouvoir calibrer cette courbe, nous avons ajouté un commutateur entre les photodiodes et l'amplificateur, qui nous permettait d'observer les impulsions de courant traversant la photodiode et d'en déduire que la charge totale délivrée par ces impulsions valait environ 1,2 picocoulombs.

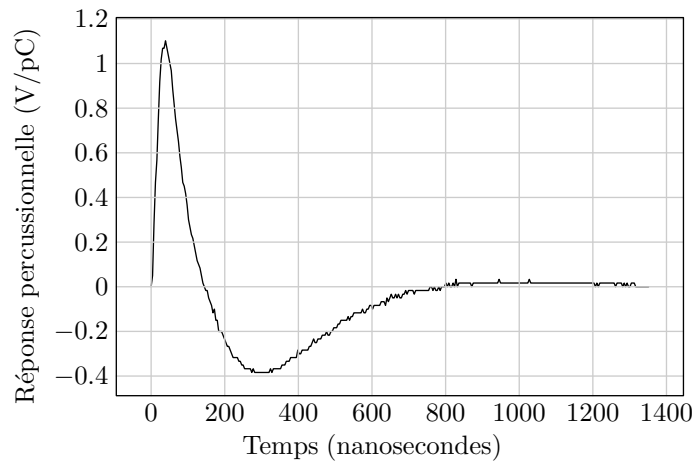


FIG. 5.4 – Réponse percussionnelle expérimentale

Nous avons également vérifié expérimentalement que cette réponse était proportionnelle à l'énergie des impulsions, jusqu'à ce que des phénomènes de saturation des amplificateurs se manifestent.

Cette réponse percussionnelle est pratiquement nulle après $1\ \mu\text{s}$, ce qui permet d'envoyer des impulsions au rythme de 1 MHz sans que les réponses ne se recouvrent. Nous avons choisi le rythme légèrement inférieur de 800 kHz, c'est-à-dire une impulsion toutes les $1,25\ \mu\text{s}$, car il correspond à celui du laser femtoseconde que nous utiliserons pour des expériences futures.

Il est alors possible de séparer la contribution des différentes impulsions d'entrée dans le signal électronique de sortie, le train d'impulsion optique ayant généré un train de réponses percussionnelles électroniques.

Pour avoir la valeur de la contribution de chaque impulsion, il suffit de mesurer la valeur de la réponse percussionnelle correspondante, toujours au même instant. On choisit en général un extremum de la courbe, ce qui permet d'avoir un signal plus fort, mais surtout de diminuer les exigences quant à la précision de l'instant d'échantillonnage. Ainsi, l'instant auquel il faut échantillonner la réponse percussionnelle présentée FIG. 5.4 correspond à un retard de 300 ns.

Une carte d'acquisition *National Instruments* PCI-6111E nous permet d'acquérir jusqu'à 2,5 millions d'échantillons par seconde. Nous nous contenterons ici d'en acquérir 800 000 par seconde, en en prenant un par impulsion. La synchronisation avec le train d'impulsion doit être directe et ne pas passer par l'horloge interne de la carte, contrairement aux procédures standard⁶. En effet, celle-ci est cadencée à 20 MHz, ce qui correspond à une imprécision de 50 ns dans l'instant d'échantillonnage lorsqu'on se repose sur elle pour la synchronisation. Un bruit important dans les mesures est alors constaté. La forme de la réponse percussionnelle expliquerait la présence d'un bruit de l'ordre de 5 %, mais la cause essentielle du bruit est la présence de courants transitoires rapides induits par les signaux de synchronisation dans la carte d'acquisition.

⁶Pour pouvoir effectuer cette synchronisation, nous avons dû demander à la société *Wavemetrics* produisant notre logiciel d'acquisition (*Igor*) d'ajouter une option supplémentaire aux procédures d'acquisition de données.

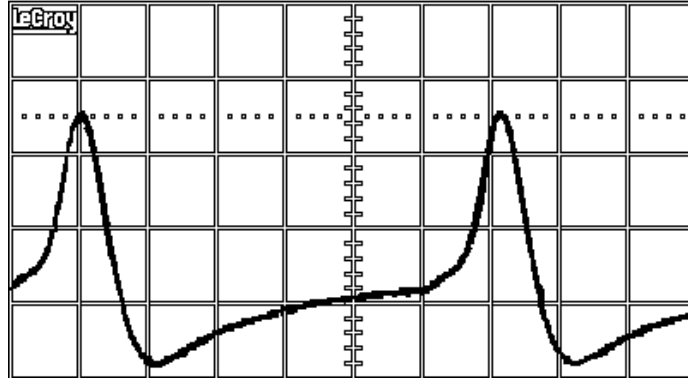


FIG. 5.5 – Réponse du système de détection homodyne finalement utilisé à des impulsions de 120 ns (200 ns par carreau)

Comme nous le verrons ci-dessous, le commutateur que nous avons ajouté pour pouvoir calibrer la réponse percutuelle changeait les caractéristiques électriques de la chaîne amplificatrice. Nous avons donc dû l'enlever. Nous avons également modifié le câblage de l'amplificateur final, de sorte que la courbe FIG. 5.4 n'est pas exactement la réponse percutuelle du système finalement utilisé, qui restera cependant similaire.

De plus, pour avoir des impulsions assez brèves pour mesurer une réponse percutuelle, nous modulons électriquement une diode laser, ce qui rendait extrêmement difficile l'équilibrage de la détection homodyne (voir section 5.3.2). Nous avons finalement utilisé un modulateur acousto-optique externe pour produire des impulsions satisfaisantes. Ces impulsions durent 120 ns, ce qui est trop long pour que l'on puisse considérer que la réponse de l'amplificateur comme une réponse percutuelle. En effet, conformément à l'équation (5.26), les structures plus brèves que 120 ns seront lissées ce qui est visible si on compare la réponse mesurée dans ce cas, représentée FIG. 5.5, à la réponse percutuelle FIG. 5.4.

La linéarité de l'amplificateur nous garantit toujours que, si les impulsions ont toujours la même structure temporelle, la réponse restera proportionnelle à cette réponse percutuelle lissée. Cela nous permettra donc d'utiliser le dispositif d'échantillonnage décrit ci-dessus pour mesurer l'intensité d'une impulsion, même si elle n'est pas brève devant le temps d'évolution de $G(t)$.

5.4.6 Lien avec la fonction de transfert

5.4.6.1 Définition

Les amplificateurs sont souvent caractérisés par leur *fonction de transfert*, qui décrit leur réponse à une entrée sinusoïdale. Nous pouvons appliquer la relation (5.26) à l'exponentielle complexe $\mathcal{Z}e^{i\omega t}$:

$$f(t) = \mathcal{Z} e^{i\omega t} \quad \implies \quad g(t) = \mathcal{Z} \int dt' e^{i\omega t'} G(t - t') = \sqrt{2\pi} G[\omega] \mathcal{Z} e^{i\omega t}, \quad (5.27)$$

où $G[\omega]$ est la transformée de Fourier de la fonction de Green $G(t)$. Comme $G(t) \in \mathbb{R}$ par définition, $G[-\omega] = G[\omega]^*$.

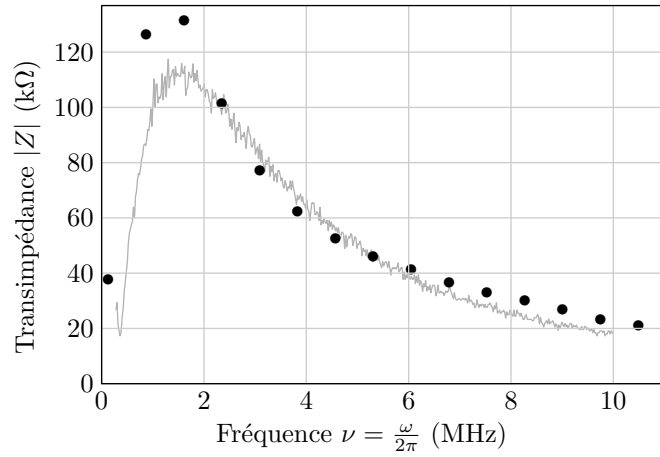


FIG. 5.6 – Transimpédance déduite du bruit (trait plein) et transformée de Fourier de la réponse percussionnelle (•)

La réponse de l'amplificateur à une entrée quelconque $f(t)$ peut alors s'exprimer simplement en fonction de sa transformée de Fourier $f[\omega]$:

$$f(t) = \frac{1}{\sqrt{2\pi}} \int d\omega f[\omega] e^{i\omega t} \quad \Rightarrow \quad g(t) = \int d\omega G[\omega] f[\omega] e^{i\omega t}, \quad (5.28)$$

c'est-à-dire, directement exprimé en coefficients de Fourier,

$$f[\omega] \quad \Rightarrow \quad g[\omega] = \sqrt{2\pi} G[\omega] f[\omega]. \quad (5.29)$$

5.4.6.2 Mesure à l'analyseur de spectre

Un analyseur de spectre permet de mesurer facilement les amplitudes $|f[\omega]|$ et $|g[\omega]|$. Lorsque l'entrée de l'amplificateur est un bruit blanc, comme le bruit thermique d'une résistance ou le bruit de photons d'un faisceau continu, l'amplitude $|f[\omega]|$ est indépendant de ω et la mesure de $|g[\omega]|$ est proportionnelle à $|G[\omega]|$. Un analyseur de spectre permet donc de mesurer rapidement la valeur absolue de la fonction de transfert, c'est à dire le gain.

Dans le cas de la chaîne amplificatrice que nous utilisons, qui associe une tension de sortie à une intensité d'entrée, le facteur $Z[\omega] \equiv \sqrt{2\pi} G[\omega]$ est une transimpédance complexe, et se mesure en ohms. Nous pouvons alors déduire le module de cette transimpédance en examinant le spectre produit par le bruit thermique de la résistance R d'entrée de l'amplificateur ou par le bruit de photons d'un faisceau continu éclairant une photodiode. Le résultat de ces mesures est représenté FIG. 5.6 et comparé à la transformée de Fourier⁷ de la réponse percussionnelle représentée FIG. 5.4. Ces deux estimations de la transimpédance sont similaires, la différence entre les deux courbes s'expliquant largement par les calibrations approximatives de la réponse percussionnelle et de l'analyseur de spectre.

⁷En fait, nous avons effectué une transformée de Fourier *discrète*, au moyen de l'algorithme usuel de FFT (Fast Fourier Transform)[128] directement appliqué aux mesures expérimentales de la transimpédance complexe, même si l'usage de cet algorithme n'est *pas recommandé* pour l'évaluation de transformée de Fourier continue [129].

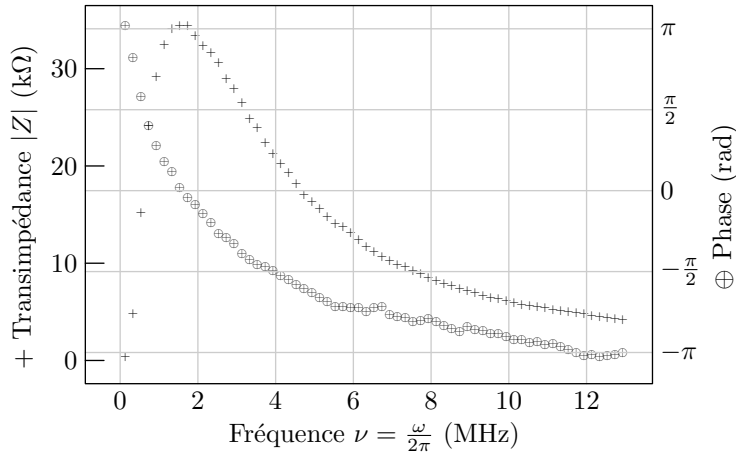


FIG. 5.7 – Transimpédance complexe $Z = |Z| e^{i\psi}$ de la chaîne amplificatrice

Si cette mesure apporte beaucoup d'information, notamment en ce qui concerne la bande passante, cette mesure n'apporte aucune information sur la phase, ce qui ne permet pas de reconstituer la réponse de la chaîne amplificatrice à une entrée quelconque, et notamment à une impulsion.

5.4.6.3 Modulation sinusoïdale

Si $G[\omega]$ est presque constant sur le domaine où $f[\omega]$ est non nul, c'est-à-dire si le mode du signal est suffisamment étroit, ou si la fonction de transfert de l'amplificateur est suffisamment plate (en gain et en phase), $G[\omega]$ peut être sorti de l'intégrale de l'équation (5.28) et être considéré comme un gain global.

Nous pouvons par exemple calculer la réponse à une entrée sinusoïdale réelle

$$f(t) = |Z| \cos(\omega t + \varphi) \quad \Rightarrow \quad g(t) = \sqrt{2\pi} |G[\omega]| |Z| \cos(\omega t + \varphi + \psi(\omega)), \quad (5.30)$$

où on a noté

$$Z = |Z| e^{i\varphi} \quad G[\omega] = |G[\omega]| e^{i\psi(\omega)}. \quad (5.31)$$

Une entrée sinusoïdale de fréquence ω induit donc bien une réponse à la même fréquence, puisque le système est linéaire et indépendant du temps. $\sqrt{2\pi} |G[\omega]|$ représente alors le gain avec lequel elle est amplifiée et $\psi(\omega)$ représente le déphasage de la sortie par rapport à l'entrée. Ces relations nous donnent la possibilité de mesurer expérimentalement $G[\omega]$ en mesurant le gain et le déphasage subit par une entrée sinusoïdale de fréquence ω .

5.4.6.4 Mesure de la transimpédance complexe

Nous avons mesuré le module et la phase de la transimpédance complexe de la chaîne amplificatrice en injectant des courants sinusoïdaux de différentes fréquences *via* le même commutateur qui nous avait permis de calibrer les impulsions avec lesquelles nous avons mesuré la réponse percussionnelle. L'amplitude et la phase de cette transimpédance complexe sont représentées FIG. 5.7.

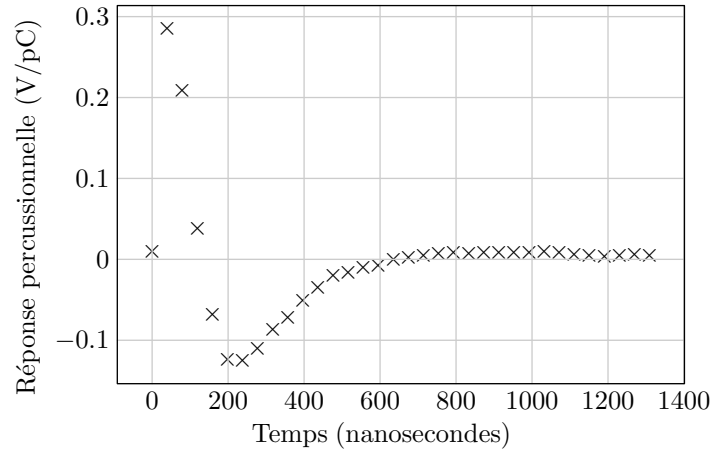


FIG. 5.8 – Transformée de Fourier inverse de la transimpédance complexe

Cette transimpédance est nettement différente de celles représentées FIG. 5.6 car nous avons changé le deuxième amplificateur de la chaîne entre les deux mesures.

Si on calcule la transformée de Fourier inverse de la transimpédance complexe, on retrouve bien une courbe qui ressemble à la réponse percussive (voir FIG. 5.8). On ne peut malheureusement pas directement comparer les deux courbes, en raison du changement de chaîne amplificatrice, mais les ordres de grandeur sont cohérents.

5.5 Résultats expérimentaux

Pour vérifier que ces problèmes sont réglés (ou, plus souvent, pour constater qu'il reste encore beaucoup de travail), on effectue une mesure homodyne du vide quantique. La première vérification à faire est que la statistique est bien gaussienne.

Il nous est arrivé d'avoir des distributions de probabilité présentant plusieurs bosses, qui correspondaient probablement à des sauts de modes dans la diode laser qui changeaient la polarisation ou la position des faisceaux et déséquilibraient la détection homodyne. Des effets similaires ont également été provoqués par des interférences électroniques dans l'électronique de détection. Elle peut également être biaisée par des phénomènes de saturations, ce qui se manifeste par un moment d'ordre 3 (la *skewness* [130] en anglais) nettement différent de 0.

Une statistique gaussienne n'est bien entendu pas suffisante pour garantir que la détection homodyne mesure bien les fluctuations du vide. Il faut alors vérifier l'évolution de cette variance en fonction de l'intensité de l'oscillateur local I_{0l} . L'équation (5.22) nous permet alors d'affirmer que la détection homodyne est équilibrée et limitée au bruit de photons si le terme linéaire de (5.22) domine le bruit électronique et le terme quadratique dû au déséquilibre.

Les mesures, représentées FIG. 5.9, s'alignent sur la droite théorique. Les écarts à la droite s'expliquent quantitativement par des erreurs statistiques et erreurs systématiques. Nous avons obtenu une détection homodyne équilibrée pour des oscillateurs locaux allant jusqu'à près de 500 millions de photons par impulsion. La détection est probablement assez bien équilibrée pour aller au delà, mais la diode laser que nous utilisons change brutalement de

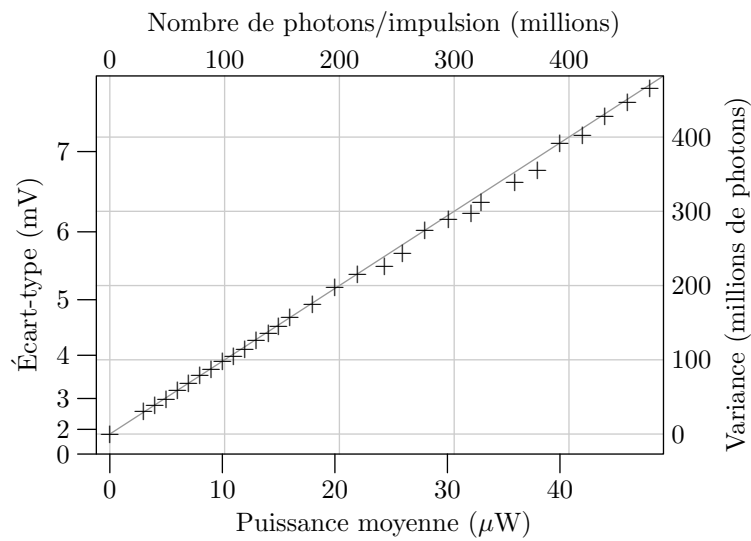


FIG. 5.9 – Mesures expérimentales (+) et estimation théorique (droite) du bruit de photons

régime pour des faisceaux plus brillants ce qui augmente brutalement la variance.

Le bruit électronique du système a en effet un écart-type de 1,8 mV, ce qui correspond à 5 600 électrons par impulsion, ou environ 6 000 photons par impulsion.

Nous avons par la suite légèrement changé le câblage de l'amplificateur une dernière fois. Dans toute la suite, on utilisera ce nouvel amplificateur, qui a un bruit électronique de 2,1 mV. Pour une valeur typique de 130 millions de photons par impulsion, l'écart-type du bruit de photons vaut 4,2 mV et est deux fois plus important (en amplitude) que celui du bruit électronique.

Chapitre 6

Mesure homodyne d'un état cohérent

Si la mesure homodyne du vide nous a permis de vérifier que notre système de mesure était bien équilibré et limité au bruit de photons, d'autres défauts se manifestent dès que l'on veut mesurer « quelque chose » (c'est-à-dire autre chose que le vide). En effet, le vide est indépendant de la phase et, surtout, présent partout, ce qui facilite sa mesure, mais gêne la mesure d'autres états.

En effet, la plupart des défauts d'une détection homodyne peuvent se ramener à une inefficacité, qui peut être modélisée par un mélange avec le vide (section 6.1). Nous étudierons plus en détail section 6.2, comment une différence entre les structures temporelles du signal et de l'oscillateur local se traduisent ainsi par une inefficacité¹. Nous étudierons ensuite (section 6.3) les effets d'un petit bruit de phase sur la mesure d'états brillants et montrerons les résultats de la mesure d'états cohérents section 6.4².

6.1 Inefficacité de la détection homodyne

La plupart des imperfections d'une détection homodyne (autres qu'un déséquilibre) peuvent se traduire par une efficacité quantique η .

6.1.1 Définition de l'efficacité

Si le faisceau signal subit des pertes avant d'arriver sur la lame semi-réfléchissante du détecteur homodyne, ces pertes peuvent être modélisées par une lame partiellement réfléchissante de transmission η , placée en amont d'une détection homodyne parfaite, suivant le schéma FIG. 6.1. Comme nous l'avons vu section 3.5.5, ces pertes induisent un bruit gaussien ajouté de variance $(1 - \eta)N_0$ sur les deux quadratures et une atténuation d'un facteur $\sqrt{\eta}$, notamment sur celle qui sera mesurée par la détection homodyne. L'atténuation représente un simple changement d'échelle et peut être compensée en multipliant simplement le résultat des mesures par $\frac{1}{\sqrt{\eta}}$. Le bruit gaussien est alors ramené à l'entrée et a une variance de $\frac{1-\eta}{\eta} N_0$.

¹Le contenu de la section 6.2 a été publié [7], accompagné d'études plus détaillées sur la structure temporelle de photons uniques générés par fluorescence paramétrique.

²Les résultats expérimentaux de ce chapitre ont été obtenus en collaboration avec Jérôme Wenger, qui effectue sa thèse dans notre groupe [131].

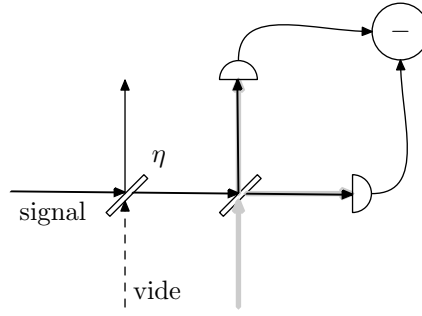


FIG. 6.1 – Modélisation de l'inefficacité d'une détection homodyne

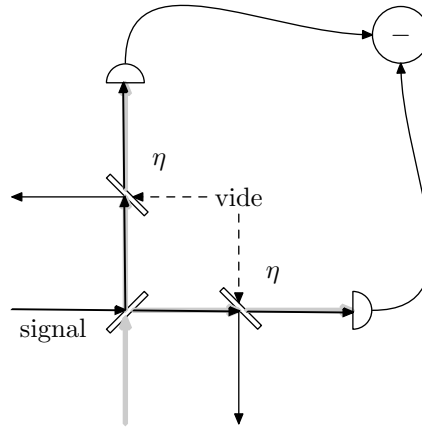
Réciproquement, toute imperfection produisant un bruit gaussien sur les mesures des quadratures peut être modélisée par des pertes équivalentes qui s'ajoutent aux précédentes. Ainsi, le bruit thermique dans l'électronique de détection pourrait être incorporé à l'efficacité globale de la détection homodyne : s'il induit un bruit gaussien de variance χN_0 sur la mesure de la quadrature, son effet est indiscernable de celui d'une lame semi-réfléchissante de transmission $\eta_{\text{élec}} = \frac{1}{\chi+1}$ placée en amont d'une détection parfaite, à un facteur d'échelle près. Cependant, ce bruit a une amplitude constante (en l'occurrence 2,1 mV), alors que le signal et les autres bruits sont proportionnels à l'amplitude de l'oscillateur local $\sqrt{I_{\text{ol}}}$ et aux autres efficacités de la détection $\sqrt{\eta'_{\text{hom}}}$. La valeur de χ serait donc inversement proportionnelle à $\sqrt{\eta'_{\text{hom}} I_{\text{ol}}}$, et l'efficacité $\eta_{\text{élec}}$ qu'on en déduirait sera donc dépendante de l'intensité de l'oscillateur local et ces autres efficacités. C'est pourquoi nous préférons en général caractériser le bruit électronique par sa variance ou son écart-type que par l'efficacité qui y est associée.

Une autre imperfection des détections homodynes est une mauvaise adaptation éventuelle des modes de l'oscillateur local et du signal. Si les deux modes ne sont pas parfaitement superposés, l'amplificateur local amplifiera un mélange du mode signal et du vide, exactement comme si ce mélange était fait par une lame partiellement réfléchissante. Nous démontrerons cela section 6.2 pour les modes temporels.

L'efficacité associée à l'adaptation spatiale des modes peut être mesurée en faisant interférer l'oscillateur local avec un signal de même intensité, en régime continu. Une variation de la phase relative nous permet de mesurer des franges, alternativement brillantes et sombres, dont le contraste est la racine carrée de l'efficacité associée. En pratique, nous avons régulièrement obtenu des contrastes de 96,5 %, et parfois de 99 %, ce qui correspond à des efficacités respectives de 93,1 % et 98 %.

6.1.2 Inefficacités après la lame semi-réfléchissante

L'efficacité limitée η d'une photodiode est en général modélisée par une lame partiellement réfléchissante fictive placée devant une photodiode parfaite. Les pertes optiques sur le trajet du faisceau entre la lame semi-réfléchissante et la photodiode peuvent être modélisées de la même manière. Pour un dispositif de détection homodyne, cela se traduit par la modélisation représentée FIG. 6.2, où l'on a supposé que les pertes η dans les voies + et – étaient égales. Si on note $q_{\pm 0}$ et $p_{\pm 0}$ les quadratures du vide injecté dans chacune des photodiodes

FIG. 6.2 – Modélisation d’une détection homodyne avec des photodiodes d’efficacité η

par ces lames partiellement réfléchissantes fictives, l’expression de $\mathcal{E}_{\pm}$ devient

$$\mathcal{E}_{\pm} = \left[\sqrt{\eta} \frac{q_{\theta} + ip_{\theta} \pm r_{ol}}{\sqrt{2}} + \sqrt{1-\eta} (q_{\pm 0} + ip_{\pm 0}) \right] e^{i\omega t + i\theta}. \quad (6.1)$$

Les intensités mesurées par les deux photodiodes s’écrivent alors, au premier ordre en $\sqrt{I_{ol}}$

$$I_{\pm} = \frac{I_{ol}}{2} \pm \sqrt{\eta I_{ol}} \left(\sqrt{\eta} q_{\theta} + \sqrt{\frac{1-\eta}{2}} q_{\pm 0} \right) \quad (6.2)$$

et leur différence vaut

$$\delta I = 2\sqrt{\eta I_{ol}} \left(\sqrt{\eta} q_{\theta} + \sqrt{1-\eta} \frac{q_{+0} + q_{-0}}{\sqrt{2}} \right). \quad (6.3)$$

Le terme $\frac{q_{+0} + q_{-0}}{\sqrt{2}}$ correspond à la quadrature q_0 d’un mode fictif que l’on aurait obtenu en mélangeant les deux modes vides sur une lame semi réfléchissante. On peut donc écrire cette différence

$$\delta I = 2\sqrt{\eta I_{ol}} \left(\sqrt{\eta} q_{\theta} + \sqrt{1-\eta} q_0 \right) \quad (6.4)$$

où q_0 est la quadrature d’un mode vide, exactement comme si on avait modélisé l’imperfection des photodiodes par une lame partiellement réfléchissante de transmission η placée en amont d’une détection homodyne parfaite, comme FIG. 6.1, et une autre sur le trajet de l’oscillateur local, qui se traduit par un oscillateur local effectif d’intensité ηI_{ol} .

6.1.3 Efficacité de notre dispositif expérimental

Les photodiodes *Centronix* BPX65 que nous utilisons ont une efficacité quantique de 92 %. Nous avons choisi dans un lot de photodiodes de ce modèle celles qui avaient la meilleure efficacité. Nous avons également retiré leur vitre de protection, qui diminuait l’efficacité de ces photodiodes de 4 %. Les pertes optiques après la lame semi-réfléchissante valent 92 %, ce qui correspond à une efficacité après la lame semi réfléchissante de 85 %. Combinée à l’adaptation des modes de 96,5 %, cela correspond à une efficacité globale de la détection

homodyne (bruit électronique exclu) $\eta'_{\text{hom}} = 79 \%$ (83 % pour l'adaptation de modes de 99 %).

Cette inefficacité génère donc un bruit

$$N'_{\text{hom}} = \chi'_{\text{hom}} N_0 = 0,27 N_0 \quad (0,20 N_0 \text{ pour un contraste de } 99 \%). \quad (6.5)$$

Lorsque l'oscillateur local est constitué par 130 millions de photons, le bruit de photons de $N_{\text{ph}} = 4,23 \text{ mV}$ tient à la fois compte des fluctuations du signal N_0 et de celles du vide N'_{hom} qui se mélange à celui-ci. On a donc

$$\sqrt{N_0} = \sqrt{\eta'_{\text{hom}} N_{\text{ph}}} = 3,76 \text{ mV} \quad (3,85 \text{ mV pour un contraste de } 99 \%) \quad (6.6)$$

$$\sqrt{N'_{\text{hom}}} = \sqrt{1 - \eta'_{\text{hom}} N_{\text{ph}}} = 1,95 \text{ mV} \quad (1,75 \text{ mV pour un contraste de } 99 \%). \quad (6.7)$$

Le bruit ajouté par ces inefficacités est donc du même ordre de grandeur que le bruit électronique $\sqrt{N_{\text{élec}}} = 2,15 \text{ mV}$. Ce dernier peut alors s'écrire

$$N_{\text{élec}} = \chi_{\text{élec}} N_0 = 0,33 N_0 \quad (0,31 N_0 \text{ pour un contraste de } 99 \%). \quad (6.8)$$

Le bruit total ajouté par la détection homodyne est la somme (en variance) des deux contributions et on a

$$\sqrt{N_{\text{hom}}} = \sqrt{N'_{\text{hom}} + N_{\text{élec}}} = 2,90 \text{ mV} \quad (2,77 \text{ pour un contraste de } 99 \%) \quad (6.9)$$

$$\chi_{\text{hom}} = \chi'_{\text{hom}} + \chi_{\text{élec}} = 0,60 \quad (0,51 \text{ pour un contraste de } 99 \%) \quad (6.10)$$

$$\eta_{\text{hom}} = \frac{1}{1 + \chi_{\text{hom}}} = 62 \% \quad (66 \% \text{ pour un contraste de } 99 \%). \quad (6.11)$$

6.2 Adaptation des modes temporels

6.2.1 Introduction

Nous avons montré [7]³ que l'adaptation entre le mode temporel du signal mesuré et l'oscillateur local peut aussi se traduire par une efficacité quantique équivalente. Nous comparerons dans cette section le détecteur homodyne monomode imparfait, représenté FIG. 6.1, avec un détecteur impulsionnel imparfait, l'analogie formelle entre les deux étant symbolisée par une double flèche.

Détecteur homodyne monomode imparfait

$\longleftrightarrow$

Détecteur homodyne impulsionnel parfait (6.12)

³Les calculs présentés ici diffèrent parfois d'un facteur $\sqrt{2\pi}$ des mêmes calculs effectués dans [7] en raison de conventions différentes dans la position des facteurs $\sqrt{2\pi}$ dans la définition de la transformée de Fourier.

6.2.2 État du champ et opérateurs de mesure

Tout état quantique peut se décomposer sur la base des états de Fock, dans laquelle nous nous placerons dans la suite. Pour comparer un état de Fock monomode $|n\rangle$ au paquet d'ondes correspondant $|\psi_n\rangle$ à n photons, nous devons définir l'opérateur annihilation du paquet d'ondes $\hat{A}_s$, analogue à l'opérateur annihilation du signal de la détection homodyne monomode $\hat{a}_s$. Ces opérateurs doivent vérifier

$$|n\rangle = \frac{\hat{a}_s^{+n}}{\sqrt{n!}} |0\rangle \iff |\psi_n\rangle = \frac{\hat{A}_s^{+n}}{\sqrt{n!}} |0\rangle, \quad (6.13)$$

où $|0\rangle$ représente le vide dans les deux cas.

Comme l'opérateur $\hat{E}(t)$ définissant le champ signal du détecteur homodyne impulsionnel est la transformée de Fourier de l'opérateur $\hat{a}_\omega$ définissant l'amplitude, on a

$$\hat{E}(t) = \frac{1}{\sqrt{2\pi}} \int d\omega \hat{a}_\omega e^{i\omega t}, \quad (6.14)$$

et il est naturel de définir

$$\hat{a}_s^+ \iff \hat{A}_s^+ = \frac{\int dt \mathcal{E}_s(t) \hat{E}^+(t)}{\sqrt{\langle \mathcal{E}_s | \mathcal{E}_s \rangle}} = \frac{\int dt \mathcal{E}_s[\omega] \hat{a}_\omega^+}{\sqrt{\langle \mathcal{E}_s | \mathcal{E}_s \rangle}}, \quad (6.15)$$

où la fonction $\mathcal{E}_s(t)$ définit l'enveloppe du paquet d'onde, et $\langle \cdot | \cdot \rangle$ désigne le produit scalaire usuel dans l'espace vectoriel des fonctions de carré sommable. La normalisation par $\sqrt{\langle \mathcal{E}_s | \mathcal{E}_s \rangle}$ permet de garder la relation de commutation usuelle $[\hat{A}_s, \hat{A}_s^+] = 1$.

Dans les deux cas, le champ oscillateur local sera quasi-classique et nettement plus grand que celui du signal mesuré, et sera supposé classique, c'est-à-dire défini par un nombre complexe $\mathcal{E}_{ol}$, commutant avec lui-même (voir section 5.2.3). Dans le cas du détecteur homodyne imparfait, le champ mesuré est un mélange du signal $\hat{a}_s$ et d'un mode vide $\hat{a}_0$ défini par

$$\hat{a} = t \hat{a}_s + r \hat{a}_0. \quad (6.16)$$

avec

$$|t|^2 + |r|^2 = 1 \quad \text{et} \quad |t|^2 \equiv \eta \quad (6.17)$$

D'un autre côté, le champ de l'oscillateur local $\mathcal{E}_{ol}(t)$ varie en général trop rapidement pour être suivi par le système de détection qui ne mesure que des intégrales sur la durée des impulsions. Si les impulsions sont plus lentes, on se contente en général de mesurer un paramètre proportionnel à cette intégrale (voir section 5.4). Les termes à comparer sont donc

$$\mathcal{E}_{ol}^* \hat{a} \iff \int dt \mathcal{E}_{ol}^*(t) \hat{E}(t) = \int d\omega \mathcal{E}_{ol}^*[\omega] \hat{a}_\omega. \quad (6.18)$$

Nous noterons $\delta \hat{f}$ le signal de la détection homodyne monomode et $\delta \hat{I}$ le signal correspondant de la détection impulsionnelle. On a donc

$$\delta \hat{f} = \mathcal{E}_{ol}^* \hat{a} + \mathcal{E}_{ol} \hat{a}^+ \iff \delta \hat{I} = \int d\omega \mathcal{E}_{ol}^*[\omega] \hat{a}_\omega + \int d\omega \mathcal{E}_{ol}[\omega] \hat{a}_\omega^+ \quad (6.19)$$

6.2.3 Correspondance entre les détecteurs homodynes monomode imparfait et impulsionnel parfait

Pour établir une analogie, il faut que les observables $\delta\hat{f}$ et $\delta\hat{I}$ aient la même distribution de probabilité. Elles doivent donc avoir les mêmes moments⁴ $\langle \delta\hat{f}^p \rangle$ et $\langle \delta\hat{I}^p \rangle$ pour tout entier p . Cette condition est suffisante, car elle garantit l'égalité des fonctions caractéristiques $\langle e^{i\Omega\delta\hat{f}} \rangle$ et $\langle e^{i\Omega\delta\hat{I}} \rangle$, qui sont, à une normalisation près, les transformées de Fourier des distributions de probabilité [108].

Pour calculer ce p -ième moment pour tout état de Fock à n photons, nous pouvons définir des « opérateurs de moment » par

$$\langle n | \hat{f}^p | n \rangle = \langle 0 | \delta\tilde{f}_{p,n} | 0 \rangle \iff \langle n | \hat{I}^p | n \rangle = \langle 0 | \delta\tilde{I}_{p,n} | 0 \rangle \quad (6.20)$$

$$\begin{aligned} \delta\tilde{f}_{p,n} &= \frac{1}{n!} \hat{a}_s^n \left[\mathcal{E}_{ol}^* \hat{a} + \mathcal{E}_{ol} \hat{a}^\dagger \right]^p \hat{a}_s^{+n} \iff \\ \delta\tilde{I}_{p,n} &= \frac{1}{n!} \hat{A}_s^n \left[\int d\omega \mathcal{E}_{ol}^*[\omega] \hat{a}_\omega + \int d\omega \mathcal{E}_{ol}[\omega] \hat{a}_\omega^\dagger \right]^p \hat{A}_s^{+n} \end{aligned} \quad (6.21)$$

Les expressions entre crochets doivent être développées pour calculer la valeur moyenne de ces opérateurs dans le vide. Ce développement sera identique pour les deux détecteurs homodynes, aux analogies (6.15) et (6.18) près. Pour simplifier le calcul des valeurs moyennes, les opérateurs création et annihilation peuvent être réécrits dans l'ordre normal. Les deux résultats seront alors aisément reliés par comparaison des coefficients multiplicatifs.

Si l'on veut récrire les expressions ci-dessus dans l'ordre normal sans écrire explicitement tous les termes, il est en fait suffisant de réaliser que ceux-ci peuvent être obtenus à partir des commutateurs suivants :

$$\left[\mathcal{E}_{ol}^* \hat{a}, \mathcal{E}_{ol} \hat{a}^\dagger \right] = \mathcal{E}_{ol}^* \mathcal{E}_{ol} \iff \left[\int d\omega \mathcal{E}_{ol}^*[\omega] \hat{a}_\omega, \int d\omega \mathcal{E}_{ol}[\omega] \hat{a}_\omega^\dagger \right] = \langle \mathcal{E}_{ol} | \mathcal{E}_{ol} \rangle \quad (6.22a)$$

$$\left[\mathcal{E}_{ol}^* \hat{a}, \hat{a}_s^\dagger \right] = \mathcal{E}_{ol}^* t \iff \left[\int d\omega \mathcal{E}_{ol}^*[\omega] \hat{a}_\omega, \hat{A}_s^\dagger \right] = \frac{\langle \mathcal{E}_{ol} | \mathcal{E}_s \rangle}{\sqrt{\langle \mathcal{E}_{ol} | \mathcal{E}_{ol} \rangle}}. \quad (6.22b)$$

En raison de l'analogie entre les développements précédents, des commutateurs analogues apparaîtront à la même place avec les mêmes coefficients dans les deux développements.

On peut alors relier les paramètres des deux types de détecteurs homodynes en utilisant les équations (6.22) :

$$\mathcal{E}_{ol}^* t t^* \mathcal{E}_{ol} \iff \frac{\langle \mathcal{E}_{ol} | \mathcal{E}_s \rangle \langle \mathcal{E}_s | \mathcal{E}_{ol} \rangle}{\langle \mathcal{E}_s | \mathcal{E}_s \rangle} \quad (6.23)$$

$$|t|^2 \equiv \eta \iff \frac{\langle \mathcal{E}_{ol} | \mathcal{E}_s \rangle \langle \mathcal{E}_s | \mathcal{E}_{ol} \rangle}{\langle \mathcal{E}_s | \mathcal{E}_s \rangle \langle \mathcal{E}_{ol} | \mathcal{E}_{ol} \rangle} = \frac{|\langle \mathcal{E}_{ol} | \mathcal{E}_s \rangle|^2}{\langle \mathcal{E}_s | \mathcal{E}_s \rangle \langle \mathcal{E}_{ol} | \mathcal{E}_{ol} \rangle} \quad (6.24)$$

$$t \iff \frac{\langle \mathcal{E}_{ol} | \mathcal{E}_s \rangle}{\sqrt{\langle \mathcal{E}_s | \mathcal{E}_s \rangle \langle \mathcal{E}_{ol} | \mathcal{E}_{ol} \rangle}} \quad (6.25)$$

⁴Nous n'étudierons ici que le cas où ces moments sont finis.

Ces expressions nous montrent comment trouver les paramètres d'un détecteur homodyne imparfait afin d'avoir le même développement des opérateurs de moments qu'un détecteur homodyne impulsif donné, c'est-à-dire

$$|\mathcal{E}_{\text{ol}}|^2 = \langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{ol}} \rangle \quad \text{et} \quad \eta = \frac{|\langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{s}} \rangle|^2}{\langle \mathcal{E}_{\text{s}} | \mathcal{E}_{\text{s}} \rangle \langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{ol}} \rangle} \quad (6.26a)$$

$$\text{ou } t = \frac{\langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{s}} \rangle}{\sqrt{\langle \mathcal{E}_{\text{s}} | \mathcal{E}_{\text{s}} \rangle \langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{ol}} \rangle}} \quad (6.26b)$$

Une fois ces substitutions effectuées, les deux développements des opérateurs dans l'ordre normal sont strictement analogues. Leur valeur moyenne dans le vide n'aura que des termes qui peuvent s'écrire sous la forme

$$[\mathcal{E}_{\text{ol}}^* \hat{a}]^k [\hat{a}_{\text{s}}]^l |0\rangle = \delta_{k,0} \delta_{l,0} |0\rangle \iff \left[\int d\omega \mathcal{E}_{\text{ol}}^*[\omega] \hat{a}_{\omega} \right]^k [\hat{A}_{\text{s}}]^l |0\rangle = \delta_{k,0} \delta_{l,0} |0\rangle \quad (6.27)$$

ou sous la forme hermitienne conjuguée. Comme ces termes analogues sont égaux et précédés par des coefficients complexes identiques, les deux valeurs moyennes sont égales :

$$\langle 0 | \delta \tilde{f}_{p,n} | 0 \rangle = \langle 0 | \delta \tilde{I}_{p,n} | 0 \rangle \quad \text{c'est-à-dire} \quad \langle n | \delta \hat{f}^p | n \rangle = \langle n | \delta \hat{I}^p | n \rangle \quad (6.28)$$

pour toutes les valeurs de $p \in \mathbb{N}$ et de $n \in \mathbb{N}$.

Comme tous leurs moments sont égaux, les deux distributions de probabilité sont identiques : une détection homodyne impulsif mesurant un paquet d'onde à n photons est équivalente à un détecteur homodyne monomode imparfait mesurant un état de Fock à n photons avec une efficacité donnée par

$$\boxed{\eta_{\text{eff}} = \frac{|\langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{s}} \rangle|^2}{\langle \mathcal{E}_{\text{s}} | \mathcal{E}_{\text{s}} \rangle \langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{ol}} \rangle}} \quad (6.29)$$

ce qui correspond à une adaptation de modes temporels entre l'enveloppe de l'oscillateur local et la forme du paquet d'ondes.

6.2.4 Généralisations

Si l'état mesuré n'est pas un état de Fock, mais une superposition linéaire de $|\psi_n\rangle$, l'équivalence entre les deux types de détecteurs homodynes tient toujours puisque η_{eff} est indépendant de n . Elle tient aussi bien entendu pour des mélanges statistiques de tels états. Ainsi des impulsions cohérentes pourront être décrites par un tel formalisme.

Cependant, la valeur de η_{eff} dépend de la forme $\mathcal{E}_{\text{s}}(t)$ du paquet d'onde et l'approche précédent ne s'applique pas à une superposition cohérente de paquets d'ondes avec des formes différentes, car plusieurs modes temporels sont alors impliqués. De telles superpositions ne seront pas considérées ici.

Un détecteur homodyne impulsif a bien entendu d'autres imperfections que les problèmes temporels dont il a été question ici. Ces imperfections peuvent en général être modélisées par des pertes η_{pertes} provenant d'une lame partiellement réfléchissante fictive placée comme sur la FIG. 6.1. Le champ mesuré est alors

$$\hat{A} = \tau \hat{A}_{\text{s}} + \rho \hat{A}_{\text{vide}} \quad (6.30)$$

où $|\tau|^2 = \eta_{\text{pertes}}$ et $|\tau|^2 + |\rho|^2 = 1$. La relation de commutation pertinent est alors

$$\left[\int d\omega \mathcal{E}_{\text{ol}}^*[\omega] \hat{a}_\omega, \hat{A}^\dagger \right] = \tau^* \left[\int d\omega \mathcal{E}_{\text{ol}}^*[\omega] \hat{a}_\omega, \hat{A}_s^\dagger \right], \quad (6.31)$$

et l'efficacité totale devient alors

$$\eta_{\text{totale}} = |\tau|^2 \frac{|\langle \mathcal{E}_{\text{ol}} | \mathcal{E}_s \rangle|^2}{\langle \mathcal{E}_s | \mathcal{E}_s \rangle \langle \mathcal{E}_{\text{ol}} | \mathcal{E}_{\text{ol}} \rangle} = \eta_{\text{pertes}} \eta_{\text{eff}}. \quad (6.32)$$

L'efficacité effective totale d'un détecteur homodyne impulsif imparfait est donc le produit de l'efficacité due aux pertes et de celle due à la forme des impulsions. En pratique, si cette inefficacité sera aisément rendue négligeable pour des impulsions de 120 ns, cet effet deviendra important lorsqu'on utilisera des impulsions de 100 fs, qui auront une longueur de 30 μm , ou des impulsions de structures temporelles différentes de l'oscillateur local, comme les photons uniques générés par fluorescence paramétrique [7, 115].

6.3 Bruit de phase

En mesurant des états cohérents brillants, nous avons que notre détection homodyne présentait un bruit de phase sur des fréquences inférieures à 500 Hz. Il provient très certainement de vibrations acoustiques (c'est évident lorsqu'on parle à proximité de l'expérience) qui induisent une différence de chemin optique entre le signal et l'oscillateur local de l'ordre de $\frac{\lambda}{250} \simeq 3 \text{ nm}$ ce qui correspond à un déphasage $\frac{2\pi}{250} = 25 \text{ mrad}$.

Ce déphasage, bien que petit, devient gênant pour les états brillants. En effet, il correspond à une rotation dans l'espace des phases, ce qui signifie que son effet est amplifié par une sorte de bras de levier, cette rotation correspondant à un déplacement d'amplitude $r_s \delta\theta$ dans l'espace des phases pour un état d'amplitude r_s . Pour que ces déplacements soient négligeables devant le bruit quantique standard $\sqrt{N_0}$, il faut que

$$\frac{r_s}{\sqrt{N_0}} \ll \frac{1}{\delta\theta} \simeq 40. \quad (6.33)$$

Comme $r_s = \sqrt{2N_0 n}$, cela correspond à des impulsions ayant nettement moins que 3 200 photons, ce qui correspond à nettement moins de $80 \cdot 10^{-15} \text{ J}$ par impulsion, ou 640 pW avec un taux de répétition de 800 kHz.

Si cette limite semble faible, elle laisse plus d'un ordre de grandeur en amplitude au delà du bruit de photon pour travailler, ce qui sera largement suffisant pour les protocoles de cryptographie quantique que nous étudierons dans la partie IV.

Ce bruit de phase dépend donc de l'amplitude de l'état mesuré et ne peut pas être mesuré par une efficacité quantique effective. Son importance dépend aussi du facteur de compression éventuel des états mesurés, ses états comprimés en phase étant bien entendu plus sensibles à ce bruit que des états cohérents. Pour que les effets de ce bruit soient indépendants du facteur de compression, à amplitude fixée, il suffit qu'il soit nettement inférieur au bruit de la détection homodyne N_{hom} , ce qui multiplie le nombre maximum de photons par $\chi_{\text{tot}} = 0,60$ ou $0,51$, dans la configuration décrite section 6.1.3.

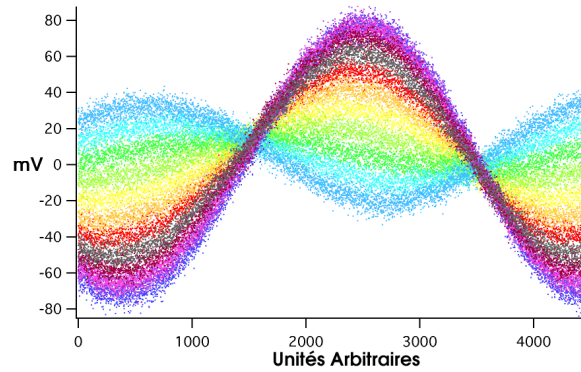


FIG. 6.3 – Mesure expérimentale d'un état cohérent

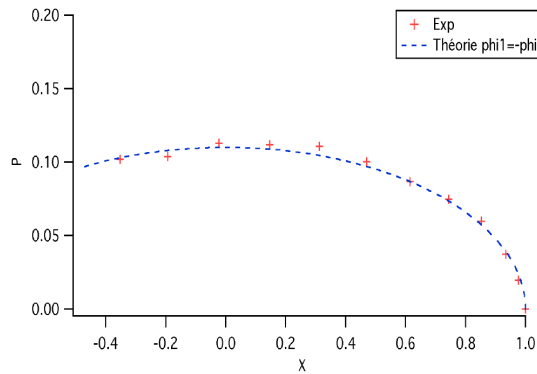


FIG. 6.4 – Représentation polaire du couplage amplitude-phase du modulateur (unités arbitraires)

6.4 Mesure expérimentale d'un état cohérent

Nous avons représenté FIG. 6.3 le résultat de la mesure homodyne d'une série d'états cohérents très peu intenses (de 1,2 à 97 photons par impulsion) générés avec un modulateur électro-optique, en fonction de la phase relative entre l'oscillateur local et le signal. Les fluctuations dues au bruit de photon sont bien visibles, ainsi que la dépendance en phase, qui montre qu'on mesure bien le champ et non l'intensité du signal.

Le déphasage relatif des différentes sinusoïdes est dû à un couplage entre la phase et l'amplitude dans notre modulateur électro-optique. C'est en effet un interféromètre de Mach-Zender déséquilibré dans lequel on induit des déphasages opposés dans les deux bras. Comme nous avons mesuré le champ transmis par ce modulateur nous avons pu reconstituer (FIG. 6.4) la dépendance amplitude-phase de ce modulateur, avec des signaux extrêmement faibles, jusqu'à 1 photon par impulsion.

Ces mesures illustrent la capacité des détectations homodynes à mesurer des signaux extrêmement faibles (de l'ordre de quelques photons) avec des photodiodes ordinaires. Nous verrons aussi plus loin qu'elle constitue la base expérimentale des protocoles de cryptographie quantique qui seront décrits dans la partie IV.

Troisième partie

Communication quantique

Chapitre 7

Théorie de l'information

Avant d'aborder la communication quantique avec des variables continues proprement dite dans les chapitres suivants, nous allons ici étudier la transmission d'information avec des variables continues. Nous commencerons ce chapitre par un rappel des bases [50] de la théorie de l'information avec des variables discrètes, section 7.1, et des variables continues, section 7.2. Nous appliquerons ensuite, dans la section 7.3, cette théorie à la communication d'information avec des états gaussiens. Nous montrerons, dans la section 7.4, la contribution des fluctuations de différentes amplitudes à l'information mutuelle dans le cas du canal additif gaussien. Nous terminerons ce chapitre en introduisant, dans la section 7.5, les coefficients de transfert d'information T , très utilisés dans le domaine de l'information quantique avec des variables continues.

7.1 Brefs rappels sur les variables discrètes

Shannon a montré en 1948 [50] que l'on pouvait définir la quantité d'information contenue dans une séquence aléatoire par son entropie. Pour une séquence ergodique constituée d'un grand nombre N de symboles A successifs tirées d'un alphabet discret $\{a_k\}$, cette entropie vaut $N H(A)$ bits, avec

$$H(A) = - \sum_k \mathcal{P}_{A=a_k} \log_2 \mathcal{P}_{A=a_k} \quad \text{bits/symbole,} \quad (7.1)$$

où $\log_2$ désigne le logarithme en base 2. $\mathcal{P}_{A=a_k}$, qui désigne la probabilité pour chaque symbole d'être égal à a_k , est bien défini lorsque la séquence est ergodique. La quantité $H(A)$, toujours positive ou nulle, désigne le nombre minimum de bits nécessaires en moyenne pour coder chaque symbole de la séquence avec une probabilité d'erreur arbitrairement faible.

Si on a deux séquences A et B , on peut bien entendu définir leur entropie conjointe par

$$H(A, B) = \sum_{k,l} \mathcal{P}_{a_k, b_l} \log_2 \mathcal{P}_{a_k, b_l}. \quad (7.2)$$

Il n'est pas difficile de se convaincre que

$$H(A, B) \leq H(A) + H(B). \quad (7.3)$$

Cette inégalité est saturée si et seulement si les variables aléatoires A et B sont indépendantes. Si elle ne sont pas indépendantes, l'inégalité signifie qu'une partie de l'information

apportée par la connaissance de B était déjà contenue dans A , et qu'il ne faut ajouter à l'entropie totale que l'information de B que l'on ignore lorsqu'on connaît A . Cette « ignorance » est quantifiée par l'entropie conditionnelle

$$H(B|A) = - \sum_{k,l} \mathcal{P}_{B=b_l} \mathcal{P}_{A=a_k|B=b_l} \log_2 \mathcal{P}_{B=b_l|A=a_k} \quad (7.4a)$$

$$= - \sum_{k,l} \mathcal{P}_{A=a_k, B=b_l} \log_2 \mathcal{P}_{B=b_l|A=a_k}. \quad (7.4b)$$

L'inégalité (7.3) devient alors l'égalité

$$H(A, B) = H(A) + H(B|A) = H(B) + H(A|B) \quad (7.5a)$$

avec

$$H(B|A) \leq H(B). \quad (7.5b)$$

$H(B|A)$ représente la quantité d'information qu'il nous manque encore pour spécifier B lorsqu'on connaît déjà A , et l'inégalité ci-dessus signifie simplement que la connaissance de A ne peut pas diminuer l'information dont on dispose sur B .

Dans le cas où les séquences A et B représentent le signal qu'Alice envoie dans un canal de communication bruité et la mesure que Bob fait de ce signal, la quantité intéressante n'est pas ce que Bob ignore du signal A d'Alice, $H(A|B)$, mais ce qu'il apprend de ce signal en observant B , c'est à dire

$$I_{AB} = H(A) - H(A|B). \quad (7.6)$$

Il est facile de se convaincre que

$$I_{AB} = H(A) - H(A|B) = H(B) - H(B|A) = H(A) + H(B) - H(A, B). \quad (7.7)$$

Comme $I_{AB} = I_{BA}$, Alice en sait autant sur la mesure B de Bob que ce dernier n'en sait sur le signal A d'Alice ; c'est pourquoi cette quantité est appelée information mutuelle. Shannon a montré [50] que Alice peut alors transmettre à Bob en moyenne jusqu'à I_{AB} bits par symbole avec un taux d'erreur arbitrairement faible.

Pour évaluer la quantité d'information maximale C que peut transmettre un canal bruité, défini par les probabilités conditionnelles $\mathcal{P}_{B=b_l|A=a_k}$, il faut choisir la distribution de probabilité pour le signal d'Alice A qui optimise I_{AB} . On a alors

$$C = \max_A \{I_{AB}\}. \quad (7.8)$$

7.2 Application aux variables continues

7.2.1 Définition de l'entropie différentielle

Si A et B ne sont plus des séquences tirées d'un alphabet discret, mais des séquences de nombres réels, continus, il faudra remplacer les sommes de la section précédente par des intégrales. Pour le montrer, commençons par découper $\mathbb{R}$ en tranches d'épaisseur δa et δb assez petites pour qu'on ait

$$\mathcal{P}_{A \in [a, a+\delta a], B \in [b, b+\delta b]} \simeq \mathcal{P}_{A=a, B=b} \delta a \delta b, \quad (7.9)$$

où $\mathcal{P}_{A,B}$ est la densité de probabilité de $\left(\frac{A}{B}\right)$ sur $\mathbb{R}^2$. Expérimentalement, δa peut représenter la résolution du modulateur et δb celle de l'appareil de mesure de Bob.

On se ramène alors à un cas discret et l'équation (7.1) devient

$$H_{\delta a}(A) = - \sum_{k \in \mathbb{Z}} \mathcal{P}_{A=k\delta a} \delta a \log_2(\mathcal{P}_{A=k\delta a} \delta a) \quad (7.10a)$$

$$= \sum_{k \in \mathbb{Z}} \delta a \mathcal{P}_{A=k\delta a} \log_2 \mathcal{P}_{A=k\delta a} - \log_2 \delta a \quad (7.10b)$$

$$\simeq \underbrace{\int da \mathcal{P}_{A=a} \log_2 \mathcal{P}_{A=a}}_{S(A)} - \log_2 \delta a \quad (7.10c)$$

L'entropie $H(A)$ tend vers l'infini lorsque δa tend vers 0, à cause du terme $-\log_2 \delta a$, parce qu'il faut un nombre infini de bits pour spécifier un nombre réel avec une précision arbitraire. Par contre, le terme $S(A)$, l'entropie différentielle, restera fini.

L'entropie différentielle d'une distribution uniforme sur un segment de longueur ΔX , par exemple, vaut $\log_2 \Delta X$.

Celle d'une distribution de probabilité gaussienne de variance $\langle A^2 \rangle$, peut également se calculer

$$S_{\sigma}^{\text{Gauss}}(A) = \int da \mathcal{P}_{A=a} \log_2 \sqrt{2\pi \langle A^2 \rangle} + \int da \mathcal{P}_{A=a} \frac{x^2}{2 \langle A^2 \rangle} \log_2 e \quad (7.11a)$$

$$= \frac{1}{2} \log_2(2\pi \langle A^2 \rangle) + \frac{\langle A^2 \rangle}{2 \langle A^2 \rangle} \log_2 e \quad (7.11b)$$

$$= \frac{1}{2} \log_2(2\pi e \langle A^2 \rangle) \quad (7.11c)$$

On peut définir de même des entropies différentielles $S(B)$ et $S(A, B)$, en remplaçant simplement les sommes du cas discret par les intégrales adéquates. On a alors

$$H_{\delta b}(B) \simeq S(B) - \log_2 \delta b \quad (7.12a)$$

$$H_{\delta a, \delta b}(A, B) \simeq S(A, B) - \log_2 \delta a - \log_2 \delta b \quad (7.12b)$$

On peut ainsi calculer l'entropie d'un vecteur aléatoire de $\mathbb{R}^d$ distribué uniformément sur une région de volume $\mathcal{V}$. Il n'est pas difficile de se convaincre que $S(\vec{A}) = \log_2 \mathcal{V}$.

Pour un vecteur aléatoire gaussien $\vec{A}$ de dimension d et de matrice de covariance K , on a

$$S(\vec{A}) = \frac{1}{2} \log_2 \left((2\pi e)^d |K| \right). \quad (7.13)$$

7.2.2 Propriétés de l'entropie différentielle

$S(A)$ présente beaucoup d'analogies avec l'entropie $H(A)$ du cas discret. Avant d'étudier ces analogies, commençons par souligner les différences. $S(A)$ peut être aussi bien négatif que positif et ne représente plus le nombre de bits (infini) qu'il faudrait pour représenter parfaitement A . En particulier $S(A) = 0$ ne correspond pas au cas où A est parfaitement déterminé, contrairement à $H(A) = 0$, mais à celui où A est réparti dans un volume unité.

De plus, si on multiplie A par une constante quelconque $\mathcal{Z}$, en changeant d'unités par exemple, on change la valeur de l'entropie différentielle :

$$S(\mathcal{Z}A) = S(A) - \log_2 \mathcal{Z}. \quad (7.14)$$

L'entropie différentielle est donc toujours définie à une constante près, et ce sont les différences d'entropies qui ont un sens, plutôt que l'entropie différentielle elle-même.

En effet, les équations nous permettant de nous débarrasser de ces constantes $\mathcal{Z}$ arbitraires sont également celles où on a une équivalence entre l'entropie différentielle et l'entropie discrète, car les termes $\log_2 \delta a$ disparaissent avec ces constantes dès que la résolution δa est assez bonne. L'équation (7.3), par exemple, devient

$$H_{\delta a, \delta b}(A, B) \leq H_{\delta a}(A) + H_{\delta b}(B) \quad (7.15a)$$

$$S(A, B) - \log_2 \delta a - \log_2 \delta b \leq S(A) - \log_2 \delta a + S(B) - \log_2 \delta b \quad (7.15b)$$

$$S(A, B) \leq S(A) + S(B), \quad (7.15c)$$

comme dans le cas discret.

On peut définir l'entropie différentielle conditionnelle $S(A|B)$ en remplaçant la somme discrète par l'intégrale

$$S(A|B) = \iint da db \mathcal{P}_{A=a, B=b} \log_2 \mathcal{P}_{A=a|B=b} \quad (7.16)$$

On a alors

$$H_{\delta a, \delta b}(A|B) \simeq S(A|B) - \log_2 \delta a \quad (7.17)$$

Il faut faire attention d'éviter le cas particulier $A = B$. En effet, dans ce cas

$$H_{\delta a, \delta b}(A|B) = 0 \quad \text{et} \quad S(A|B) \simeq \log_2 \delta a \xrightarrow{\delta a \rightarrow 0} -\infty \quad (7.18)$$

Ce cas pathologique correspond à une densité de probabilité conditionnelle en forme de pic de Dirac

$$\mathcal{P}_{B=b|A=a} = \delta(b - a). \quad (7.19)$$

Il est facile de se convaincre, par un cheminement analogue à celui des équations (7.15), que les équations (7.5) deviennent

$$S(A, B) = S(A) + S(B|A) = S(B) + S(A|B) \quad (7.20a)$$

$$S(B|A) \leq S(B). \quad (7.20b)$$

7.2.3 Information mutuelle et canaux additifs

De même, l'information mutuelle est bien définie et a le même sens que dans le cas discret. En effet,

$$I_{AB} = H_{\delta a}(A) - H_{\delta a}(A|B) = S(A) - S(A|B) \quad (7.21)$$

$$I_{AB} = S(A) - S(A|B) = S(B) - S(B|A) = S(A) + S(B) - S(A, B). \quad (7.22)$$

Ainsi, une modulation continue dans un canal continu bruité ne peut transmettre qu'un nombre fini de bits, I_{AB} .

Les canaux continus les plus fréquents sont des canaux additifs, où un bruit N statistiquement indépendant du signal A y est ajouté :

$$B = A + N. \quad (7.23)$$

On a donc

$$S(A, B) = S(A, N) = S(A) + S(N), \quad (7.24)$$

en utilisant l'indépendance de A et de N . L'injection de ce résultat dans l'équation (7.22) nous donne

$$\boxed{I_{AB} = S(B) - S(N)}. \quad (7.25)$$

Pour trouver la capacité d'un tel canal, il faut maximiser I_{AB} à N fixé, c'est à dire maximiser $S(B)$. Si on ne se fixe pas d'autre contrainte, la capacité du canal est infinie, car on peut envoyer un signal infiniment large, d'entropie $S(A) \rightarrow \infty$ et $S(B) \rightarrow \infty$.

Si on limite la variance de A , c'est à dire l'énergie moyenne du signal, on fixe la variance de B , car

$$\langle B^2 \rangle = \langle A^2 \rangle + \langle N^2 \rangle. \quad (7.26)$$

Il faut donc maximiser

$$S(B) = \int db \mathcal{P}_{B=b} \log_2 \mathcal{P}_{B=b} \quad (7.27a)$$

sous les contraintes

$$1 = \int db \mathcal{P}_{B=b} \quad \text{et} \quad \langle B^2 \rangle_{\max} \geq \int db \mathcal{P}_{B=b} b^2. \quad (7.27b)$$

Il faut donc maximiser l'intégrale

$$\int db \left[-\mathcal{P}_{B=b} \log_2 \mathcal{P}_{B=b} - \lambda \mathcal{P}_{B=b} - \mu \mathcal{P}_{B=b} b^2 \right], \quad (7.28)$$

où λ et μ sont des multiplicateurs de Lagrange. Cette intégrale est maximale lorsque la dérivée fonctionnelle de son intégrande par rapport à $\mathcal{P}_B$ est nulle [98], c'est-à-dire si

$$-\log_2 \mathcal{P}_{B=b} - 1 - \lambda - \mu b^2 = 0 \quad (7.29)$$

$$\mathcal{P}_{B=b} = 2^{-1-\lambda-\mu b^2} \quad (7.30)$$

$\mathcal{P}_B$ est donc une gaussienne. Si on ajuste les multiplicateurs de Lagrange pour satisfaire les contraintes, on a

$$\mathcal{P}_{B=b} = \frac{1}{\sqrt{2\pi \langle B^2 \rangle}} e^{-\frac{b^2}{2\langle B^2 \rangle}} \quad (7.31)$$

$$S(B) = \frac{1}{2} \log_2 \left(2\pi e \langle B^2 \rangle \right). \quad (7.32)$$

Si la structure du bruit N le permet, le signal que doit envoyer Alice pour maximiser I_{AB} à $\langle A^2 \rangle$ fixée est tel que B ait la structure d'un bruit blanc gaussien de variance $\langle B^2 \rangle = \langle A^2 \rangle + \langle N^2 \rangle$, et on a

$$\boxed{I_{AB}^{\text{opt}} = \frac{1}{2} \log_2 \frac{\langle B^2 \rangle}{\langle N^2 \rangle} = \frac{1}{2} \log_2 \left(1 + \frac{\langle A^2 \rangle}{\langle N^2 \rangle} \right)} \quad (7.33)$$

Le rapport $\frac{\langle A^2 \rangle}{\langle N^2 \rangle}$ est appelé *rapport signal à bruit* (*signal to noise ratio* ou *SNR* en anglais). Si le signal reçu par Bob n'est pas gaussien, $I_{AB} < I_{AB}^{\text{opt}}$.

Si N est un bruit blanc gaussien, la valeur de I_{AB}^{opt} donnée par l'équation (7.33) peut être atteinte avec une distribution de probabilité $\mathcal{P}_A$ gaussienne.

Si notre canal additif a un gain G , c'est à dire si

$$B = \sqrt{G} A + N, \quad (7.34)$$

celui-ci ne change pas I_{AB} . En d'autres termes, $\frac{1}{\sqrt{G}}B$ apporte autant d'information sur A que B . Pour se retrouver dans le domaine d'application de l'équation (7.33) il faut donc considérer

$$\frac{1}{\sqrt{G}}B = A + N_{\text{entrée}}, \quad (7.35)$$

où $N_{\text{entrée}} \equiv \frac{1}{\sqrt{G}}N$ est le *bruit équivalent ramené à l'entrée*, de variance $\langle N_{\text{entrée}}^2 \rangle = \frac{1}{G} \langle N^2 \rangle$. L'équation (7.33) devient alors

$$I_{AB}^{\text{opt}} = \frac{1}{2} \log_2 \left(1 + \frac{\langle A^2 \rangle}{\langle N^2 \rangle_{\text{entrée}}} \right). \quad (7.36)$$

Dans le cas des états gaussiens, il est facile de se convaincre que cette formulation est équivalent à celle qu'on obtiendrait en combinant les équations (7.22) et (7.13). Si

$$K_{A,B} = \begin{bmatrix} \langle A^2 \rangle & \langle A B \rangle \\ \langle A B \rangle & \langle B^2 \rangle \end{bmatrix}, \quad (7.37)$$

est la matrice de covariance de $\begin{pmatrix} A \\ B \end{pmatrix}$, on a en effet

$$I_{AB} = \frac{1}{2} \log_2 \frac{\langle A^2 \rangle \langle B^2 \rangle}{|K_{AB}|} = \frac{1}{2} \log_2 \frac{\langle A^2 \rangle \langle B^2 \rangle}{\langle A^2 \rangle \langle B^2 \rangle - \langle AB \rangle^2} = \frac{1}{2} \log_2 \frac{1}{1 - \rho_{AB}^2}, \quad (7.38)$$

où ρ_{AB} est le coefficient de corrélation des variables A et B .

7.3 Communication avec des états gaussiens

La théorie de l'information brièvement exposée dans les sections précédentes permet d'étudier la quantité d'information qui peut être transmise avec des impulsions lumineuses.

Nous nous limiterons ici à la communication avec des états gaussiens, ce qui nous fera notamment négliger l'utilisation d'états de Fock. Dans ce cas, qui est le plus pertinent expérimentalement à défaut d'être le plus efficace en théorie, tous les bruits sont gaussiens. Les faisceaux gaussiens constituent donc des canaux additifs gaussiens. Les équations (7.36) et (7.38) seront donc systématiquement utilisées pour calculer des taux de transfert d'information.

Comme nous l'avons démontré ci-dessus, ces taux optimum de transfert d'information s'obtiennent avec une modulation qui présente une statistique gaussienne. Je supposerai donc en général que les modulations sont gaussiennes, et, pour simplifier, que leurs valeurs moyennes (classiques) sont nulles.

Tous les taux d'informations que nous calculerons ici seront exprimés en bits par symbole. Un symbole correspondra à un mode spatiotemporel du champ électromagnétique, limité par Fourier. Un tel mode est défini par l'opérateur annihilation A_s défini par l'équation (6.15). Deux symboles différents correspondront à des modes orthogonaux du champ, comme des impulsions qui ne se recouvrent pas. Pour un faisceau continu de largeur spectrale $\delta\nu$, il est ainsi inutile de prendre des mesures avec un pas d'échantillonnage plus petit que $\frac{1}{2\delta\nu}$, car les échantillons successifs ne sont plus indépendants et correspondent à des modes qui ne sont pas orthogonaux. Par contre un échantillonnage plus espacé que $\frac{1}{2\delta\nu}$ garantit l'indépendance des mesures successive [50].

7.3.1 États cohérents

Supposons qu'Alice envoie un état cohérent modulé à Bob par une ligne sans pertes, et que celui-ci en mesure la quadrature Q avec une détection homodyne. Si Q_A représente la modulation d'Alice, le formalisme introduit section 4.2 nous donne

$$Q = Q_A + A_Q, \quad (7.39)$$

où A_Q représente le bruit quantique en position associé à l'état cohérent.

Ce bruit a une valeur moyenne nulle et sa variance vaut $\langle A_Q^2 \rangle = N_0$, ce qui limite la quantité d'information mutuelle entre Alice et Bob, pour une variance de modulation $\langle Q_A^2 \rangle = V_A N_0$ donnée, à

$$I_{AB}^{\text{coh}} = \frac{1}{2} \log_2 (1 + V_A) = \frac{1}{2} \log_2 V, \quad (7.40)$$

où $V = \frac{\langle Q^2 \rangle}{N_0}$ représente la variance du signal reçu par Bob, normée au bruit quantique standard N_0 .

7.3.2 États comprimés

La possibilité de comprimer le bruit quantique en dessous du bruit de photons nous fournit une possibilité de transmettre plus d'information pour une même modulation. Dans ce cas, on a $\langle A_Q^2 \rangle = s N_0$, où s est un paramètre strictement positif.

Alice envoie alors à Bob du vide comprimé déplacé de Q_A . La quantité d'information qu'elle peut transmettre est alors limitée par

$$I_{AB}^{\text{comp}} = \frac{1}{2} \log_2 \left(\frac{s + V_A}{s} \right) = \frac{1}{2} \log_2 \frac{V}{s}, \quad (7.41a)$$

$$\text{avec } V = V_A + s. \quad (7.41b)$$

On voit que la limite des états infiniment comprimés $s \rightarrow 0$, qui correspond à une ligne sans bruit, permettrait une transmission infinie d'information malgré une variance de modulation limitée. Cependant, cette limite n'est pas physique. En effet, il ne faut pas oublier

que l'autre quadrature, P , transporte de l'énergie, même si elle reste inutilisée en terme de transfert d'information ($P_A = 0$). On a en effet

$$\langle P^2 \rangle = \langle A_P^2 \rangle = \frac{1}{s} N_0. \quad (7.42)$$

Lorsque $s \rightarrow 0$, cette variance tend vers l'infini, ce qui implique un coût énergétique infini.

Pour calculer la quantité d'information échangée, il faut donc limiter le facteur de compression s , soit d'une manière *ad hoc*, soit en limitant l'énergie totale [132]. Si la dernière solution est plus satisfaisante d'un point de vue théorique et fondamental, elle est peu réaliste en raison des limitations technologiques actuelles. Les raisons qui limitent les facteurs de compression s atteints expérimentalement aujourd'hui n'ont en effet absolument rien à voir avec le coût de l'énergie.

De plus, une approche purement énergétique conduit à la conclusion que la méthode la plus efficace de transmission d'information est l'envoi et la mesure d'états de Fock. Cette approche semble encore plus délicate expérimentalement que des forts facteurs de compression, et elle ne fait plus intervenir de variable continue.

Si Alice et Bob partagent deux faisceaux jumeaux intriqués décrits par l'équation (3.41), nous avons vu section 4.2.4 qu'on se trouvait dans un cas équivalent aux faisceaux comprimés déplacés avec une variance $V = \frac{1}{s}$. On a alors

$$I_{AB} = -\frac{1}{2} \log_2 s^2 = -\log_2 s. \quad (7.43)$$

Cependant, si cette quantité d'information mutuelle quantifie les corrélations entre A et B , elle ne peut pas servir directement à envoyer un message car Alice ne peut exercer aucun contrôle sur A . Par contre, comme ces faisceaux sont corrélés simultanément en Q et en P , cette information commune pourra être utilisée dans des protocoles de communication quantique, comme le codage dense [132, 124] et la téléportation quantique [90, 91].

7.3.3 Transmission d'information en présence de pertes

En présence de pertes, l'atténuation du signal due au gain de la ligne peut être compensée par un simple changement d'échelle et n'influe pas sur le taux d'information. Par contre, le bruit inévitablement ajouté sur la ligne dégrade le taux d'information. Le bruit total qu'il faut utiliser dans l'équation (7.36) est la somme des variances du bruit quantique A_Q et du bruit B_Q ajouté sur la ligne :

$$\langle A_Q^2 \rangle + \langle B_Q^2 \rangle = (s + \chi) N_0. \quad (7.44)$$

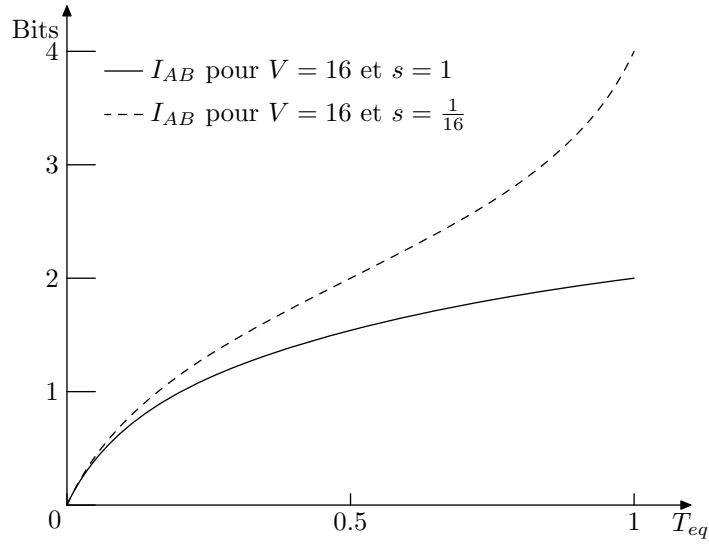
On a donc, dans le cas général des états comprimés,

$$I_{AB} = \frac{1}{2} \log_2 \frac{V_A + s + \chi}{s + \chi} = \frac{1}{2} \log_2 \frac{V + \chi}{s + \chi}, \quad (7.45)$$

expression dans laquelle il suffit de poser $s = 1$ pour se ramener aux états cohérents.

Si la ligne a une transmission $T < 1$, et que le bruit ajouté est seulement dû à cette atténuation, l'inégalité (4.22) est saturée (avec $\bar{T} \equiv G$). On a donc

$$\chi = \frac{1}{T} - 1. \quad (7.46)$$

FIG. 7.1 – Information mutuelle I_{AB} en fonction de la transmission équivalente T_{eq}

Dans le cas plus général, on peut définir une transmission équivalente, qui correspond à la transmission d'une ligne qui a le même bruit, par

$$T_{eq} = \frac{1}{\chi + 1}. \quad (7.47)$$

Si on envoie dans la ligne un état cohérent, cette transmission équivalente est égale au coefficient de transmission d'information défini section 7.5. L'équation (7.45) peut alors s'écrire

$$I_{AB} = \frac{1}{2} \log_2 \frac{1 + T(V - 1)}{1 - T(1 - s)}. \quad (7.48)$$

Le bruit ajouté (ou les pertes équivalentes) masque assez vite la différence entre les états comprimés et les états cohérents comme on peut le voir sur la FIG. 7.1 ; les états comprimés sont fragiles et ne résistent pas aux pertes. Ainsi, le taux d'information mutuelle infini mentionné plus haut pour des états infiniment comprimés devient fini pour des pertes infinitésimales et s'approche, lui aussi, assez vite de celui des états cohérents.

7.4 Contributions à l'information mutuelle

On peut essayer de comprendre comment l'information se répartit entre les petites fluctuations, codées sur les bits les moins significatifs, et les grandes fluctuations codées sur les bits les plus significatifs. Une approche rigoureuse reposerait sur le calcul de l'entropie d'une vraie description binaire, mais les calculs sont compliqués, et les résultats seront sous une forme discrète. Comme notre but ici est de chercher à comprendre intuitivement ce qui ce passe, nous choisirons une autre approche, continue.

Les fluctuations du signal dont l'amplitude correspond au bit σ sont celles qui ont une variance $2^{2\sigma}$. De même, la variance du bruit sera notée $2^{2\beta} \equiv \langle N_{\text{entrée}}^2 \rangle$. Nous avons ici omis un facteur multiplicatif correspondant au choix des unités, qui se traduit par le choix des

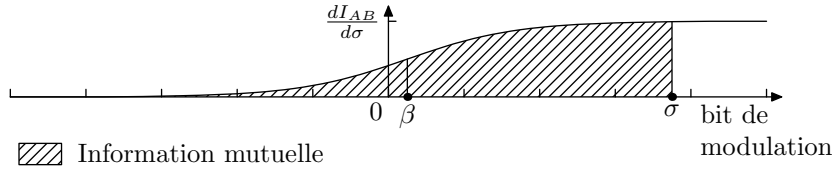


FIG. 7.2 – Contributions des diverses fluctuations à l'information mutuelle

origines pour σ et β . L'équation (7.36) peut alors se récrire sous la forme

$$I_{AB} = \frac{1}{2} \log_2(2^{2\sigma} + 2^{2\beta}) - \frac{1}{2} \log_2(2^{-\infty} + 2^{2\beta}) \quad (7.49)$$

$$= \int_{-\infty}^{\sigma} d\sigma' \frac{1}{1 + 2^{2(\beta-\sigma')}} \quad (7.50)$$

La valeur en σ de la fonction

$$\frac{dI_{AB}}{d\sigma} = \frac{1}{1 + 2^{2(\beta-\sigma)}}, \quad (7.51)$$

représentée FIG. 7.2, peut s'interpréter comme la contribution des fluctuations d'amplitude 2^σ dans la modulation à l'information mutuelle. σ correspond au bit modifié par ces fluctuations, les bits les moins significatifs correspondant à $\sigma \rightarrow \infty$ et le bit le plus significatif correspond à $\sigma = \frac{1}{2} \log_2 \langle A^2 \rangle$.

Cette fonction nous permet de quantifier notre intuition sur la contribution des différentes amplitudes de fluctuations. On voit aisément que les fluctuations nettement plus grandes que le bruit transportent l'essentiel de l'information. Dans ce cas, $2^{2(\beta-\sigma)} \ll 1$ et

$$\frac{dI_{AB}}{d\sigma} \simeq 1 - 2^{2(\beta-\sigma)} \simeq 1 \quad (7.52)$$

La contribution des grandes fluctuations est donc d'environ un bit d'information mutuelle par bit d'amplitude. On est déjà à 94% ($\frac{15}{16}$) de cette valeur à deux bits au delà du bruit (c'est à dire lorsque $\sigma = \beta + 2$.)

Par contre les fluctuations beaucoup plus petites que le bruit, qui correspondent à $2^{2(\sigma-\beta)} \ll 1$ contribuent à peine à l'information mutuelle : dans ce cas, on a

$$\frac{dI_{AB}}{d\sigma} \simeq 2^{2(\sigma-\beta)} \ll 1. \quad (7.53)$$

Ainsi, les fluctuations inférieures de deux bits au bruit ne contribuent déjà plus qu'à hauteur de 0.06 ($\frac{1}{16}$) bits d'information mutuelle par bit de fluctuation.

Les fluctuations d'amplitude égale au bruit sont à moitié masquées et contribuent à hauteur d'un demi bit d'information mutuelle par bit de fluctuation. On a donc une zone d'environ 4 bits autour du bruit β dans laquelle les fluctuations sont partiellement masquées par le bruit et ne transmettent qu'une partie de l'information qu'elle portent. En dessous, les petites fluctuations sont complètement masquées et ne transmettent que très peu d'information, et au delà, les grandes fluctuations transmettent quasiment toute leur information.

7.5 Coefficient de transfert d'information

Lorsqu'on utilise une chaîne de canaux additifs gaussiens, chacun d'entre eux peut être caractérisé par son coefficient de transfert d'information qui est le rapport des rapports signal à bruit entre l'entrée et la sortie du canal.

$$T \equiv \frac{\left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_{\text{sortie}}}{\left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_{\text{entrée}}} = \frac{G \langle \text{Bruit}^2 \rangle_{\text{entrée}}}{\langle \text{Bruit}^2 \rangle_{\text{sortie}}} \leq 1, \quad (7.54)$$

où G est le gain (en puissance) de cette chaîne de canaux

Supposons qu'Alice encode un message avec un rapport signal à bruit donné et qu'elle l'envoie à Bob par un canal de coefficient de transfert d'information T_A . Celui-ci le fait passer à Charles par un canal de coefficient T_C , qui le fait passer de proche en proche jusqu'à Zaccharie, qui le reçoit de Yolande par un canal de coefficient T_Y . Le coefficient de transfert d'information du canal ainsi constitué entre Alice et Zaccharie est donc le produit $T_A T_B \cdots T_Y$. On en déduit que l'information mutuelle entre Alice et Zaccharie vaut

$$I_{AZ} = \frac{1}{2} \log_2 \left(1 + T_A T_B \cdots T_Y \left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_A \right) \quad (7.55)$$

Appelons I_A l'information initialement accessible chez Alice.

$$I_A = \frac{1}{2} \log_2 \left(1 + \left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_A \right) \quad (7.56)$$

Si le rapport signal à bruit est encore grand devant 1 à l'arrivée chez Zaccharie, on a

$$I_{AZ} \simeq \frac{1}{2} \log_2 \left(T_A T_B \cdots T_Y \left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_A \right) \quad (7.57)$$

$$\simeq I_A + \frac{1}{2} \log_2 T_A + \frac{1}{2} \log_2 T_B + \cdots + \frac{1}{2} \log_2 T_Y \quad (7.58)$$

Si ce rapport est au contraire petit devant 1, on a

$$I_{AZ} \simeq \frac{1}{2 \ln 2} T_A T_B \cdots T_Y \left[\frac{\langle \text{Signal}^2 \rangle}{\langle \text{Bruit}^2 \rangle} \right]_A, \quad (7.59)$$

et si le rapport signal à bruit était déjà petit au départ, on a

$$I_{AZ} \simeq T_A T_B \cdots T_Y I_A. \quad (7.60)$$

Ces exemples justifient la dénomination de coefficient de transfert d'information donnée à ce coefficient. On y voit bien son rôle dans la dégradation de la quantité d'information extractible.

Remarquons pour conclure que la quantité $\frac{1}{T}$ est connue et souvent utilisée en électronique sous le nom de *figure de bruit* (*Noise figure*), notée NF . On a en effet

$$NF = \frac{1}{T} = \frac{\langle \text{Bruit}^2 \rangle_{\text{entrée}} + \langle \text{Bruit}^2 \rangle_{\text{ajouté}}}{\langle \text{Bruit}^2 \rangle_{\text{entrée}}} = 1 + \frac{\langle \text{Bruit}^2 \rangle_{\text{ajouté}}}{\langle \text{Bruit}^2 \rangle_{\text{entrée}}}, \quad (7.61)$$

donc NF mesure le bruit ramené à l'entrée du signal, normé au bruit déjà présent à l'entrée. En optique quantique, ce bruit est usuellement le bruit quantique standard N_0 , mais pour caractériser un amplificateur électronique, c'est souvent le bruit thermique de Johnson-Nyquist correspondant à une impédance d'entrée de $50 \, \Omega$.

Chapitre 8

Clonage quantique

8.1 Introduction

L'une des conséquences de la mécanique quantique et du principe d'incertitude de Heisenberg est l'impossibilité de copier parfaitement un objet quantique. Cette impossibilité, souvent désignée sous le nom de *théorème de non-clonage* est essentielle dans la plupart des protocoles de communication quantique. Elle est à la base même des protocoles de cryptographie quantique [65, 66, 133, 134], et la téléportation quantique est justement surprenante parce qu'elle permet de contourner ce théorème [76, 10].

Si l'impossibilité de copier un objet quantique est conjecturée depuis la connaissance du principe d'incertitude de Heisenberg, elle n'a été démontrée qu'en 1982 [135, 136] pour des variables discrètes. Ce théorème a récemment été étendu aux variables continues d'états gaussiens [137, 138, 134]. Le but de ce chapitre est de retrouver certains résultats de ces références avec le formalisme linéarisé introduit section 4.2.

Nous introduisons le duplicateur quantique section 8.2, dont nous étudierons les limites dans les sections 8.3 et 8.4. La section 8.5 présentera une implémentation optique d'un tel duplicateur 8.6. Les limites du clonage symétrique $1 \rightarrow M$ sont démontrées section 8.7.¹

8.2 Définition et description d'un duplicateur quantique

Un duplicateur quantique est un dispositif qui prend un état quantique inconnu en entrée et en fournit deux copies à la sortie. Pour fixer les idées, on supposera qu'Alice fabrique l'état quantique et que Bob et Charles récupèrent les deux copies. Un tel duplicateur est la plus simple des machines cloneuses non triviales et est également appelé *cloneuse* $1 \rightarrow 2$. On parlera de duplicateur optimal lorsque les copies de Bob et de Charles sont aussi bonnes que possible.

Un duplicateur peut être décrit comme un canal à une entrée et deux sorties, comme schématisé FIG. 8.1. Nous noterons $\begin{bmatrix} Q \\ P \end{bmatrix}$ la position et l'impulsion de l'état fourni en entrée du duplicateur par Alice, et $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ et $\begin{bmatrix} Q_C \\ P_C \end{bmatrix}$ celles des états fournis par le duplicateur à Bob et

¹Les démonstrations des sections 8.3 et 8.7 ont été publiées dans la section II de [10].

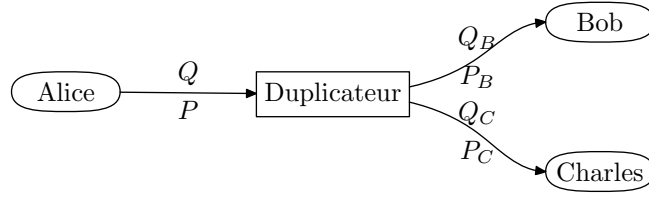


FIG. 8.1 – Schema d'un duplicateur quantique

Charles. On a alors

$$Q_B = g_Q(Q + B_Q) \quad Q_C = h_Q(Q + C_Q) \quad (8.1a)$$

$$P_B = g_P(P + B_P) \quad P_C = h_P(P + C_P), \quad (8.1b)$$

où on a introduit les bruits équivalents $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ et $\begin{bmatrix} C_Q \\ C_P \end{bmatrix}$ qui, par définition, ne sont pas corrélés au signal $\begin{bmatrix} Q \\ P \end{bmatrix}$ et commutent avec lui. Ces équations sont une généralisation de celles utilisées pour caractériser les mesures quantiques non destructives (QND) [79, 80, 81, 84, 85]. Un système de mesures QND correspond au cas où $g_Q = 1$ et $h_Q \rightarrow \infty$, une mesure QND parfaite correspondant à la limite $B_Q \rightarrow 0$ et $C_Q \rightarrow 0$.

Comme les modes de Bob et de Charles sont distincts, les observables d'un mode commutent avec celles de l'autre mode. En particulier,

$$[Q_B, P_C] = [Q_C, P_B] = 0. \quad (8.2a)$$

De plus, les quadratures des modes d'Alice, Bob et Charles vérifient les relations de commutation usuelles :

$$[Q, P] = 2iN_0 \quad (8.2b)$$

$$[Q_B, P_B] = [Q_C, P_C] = 2iN_0. \quad (8.2c)$$

8.3 Clonage optimal

Les variances des bruits équivalents $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ et $\begin{bmatrix} C_Q \\ C_P \end{bmatrix}$ permettent de mesurer la qualité d'un duplicateur : plus ces variances sont importantes et moins le duplicateur sera bon. Il faut donc calculer les contraintes que la mécanique quantique impose à ces variances.

Les relations (8.2a), calculées avec les expressions (8.1), et le commutateur (8.2b) nous permettent de calculer les commutateurs des bruits équivalents dans le cas assez général des gains non nuls. On a alors

$$[B_Q, C_P] = [B_P, C_Q] = -2iN_0. \quad (8.3)$$

Ces relations de commutation croisées entre les modes de Bob et Charles permet d'écrire des inégalités de Heisenberg croisées sur les variances des bruits équivalents à l'entrée.

$$\langle B_Q^2 \rangle \langle C_P^2 \rangle \geq N_0^2 \quad \langle B_P^2 \rangle \langle C_Q^2 \rangle \geq N_0^2. \quad (8.4)$$

Ces relations de Heisenberg croisées empêchent Bob et Charles de mesurer conjointement l'état d'Alice avec une précision meilleure que celle autorisée par le principe d'incertitude de Heisenberg (2.4). En effet Bob peut toujours mesurer parfaitement Q_B et Charles mesurer P_C , ce qui leur permet de mesurer simultanément Q et P , mais avec les bruits ajoutés B_Q et C_P . On retrouve là exactement l'équation (4.28), à partir des mêmes relations de commutation. Cette possibilité de Bob et de Charles de faire une mesure conjointe à partir de leurs clones respectifs est d'ailleurs l'argument central de la démonstration [138] du théorème de non-clonage quantique.

Un duplicateur sera considéré comme optimal lorsque ces deux inégalités sont saturées. On peut alors caractériser la répartition du bruit entre Bob et Charles par le coefficient $\mu > 0$ et la répartition entre les quadratures Q et P par le coefficient $\beta > 0$. On peut écrire

$$\langle B_Q^2 \rangle = \mu\beta N_0 \quad \langle B_P^2 \rangle = \frac{\mu}{\beta} N_0 \quad (8.5a)$$

$$\langle C_Q^2 \rangle = \frac{\beta}{\mu} N_0 \quad \langle C_P^2 \rangle = \frac{1}{\mu\beta} N_0. \quad (8.5b)$$

Dans le cas parfaitement symétrique, on a $\mu = \beta = 1$.

8.4 Fidélité

Une mesure fréquemment utilisée pour caractériser un canal est sa fidélité moyenne $\overline{\mathcal{F}}$ [109], comme nous l'avons vu section 4.4.2. De même, nous pouvons étudier la fidélité des deux canaux du duplicateur, $\overline{\mathcal{F}}_B$ désignant la fidélité du canal d'Alice à Bob et $\overline{\mathcal{F}}_C$ celle du canal d'Alice à Charles.

Pour optimiser cette valeur, on doit encore régler le gain, comme dans la section 4.4.3, en utilisant μ^2 ou $\frac{1}{\mu^2}$ comme valeur du paramètre $\chi_{\min}$. Si le duplicateur est symétrique, $\mu = 1$ et $\chi_{\min} = 1$. On a donc $\overline{\mathcal{F}}_{B,\text{opt}} = \overline{\mathcal{F}}_{C,\text{opt}} = \overline{\mathcal{F}}_{\text{opt}}$.

Si les états à cloner sont des états cohérents avec une répartition gaussienne de variance $\sigma^2 N_0$, la fidélité moyenne de cette cloneuse optimale est alors donnée par l'équation (4.46) où l'on a effectué la substitution $\chi_{\min} = 1$ et vaut

$$\overline{\mathcal{F}}_{\text{opt}} = \frac{4 + 2\sigma^2}{2 + 3\sigma^2} \xrightarrow{\sigma^2 \rightarrow \infty} \frac{2}{3}. \quad (8.6)$$

Conformément à l'équation (4.43), cette relation n'est valable que pour

$$\sigma^2 \geq 2(1 + \sqrt{2}) \simeq 4,83, \quad (8.7)$$

le gain optimal correspondant à des variances σ^2 inférieures étant limité à $\frac{1}{1+\chi_{\min}} = \frac{1}{2}$, sans quoi le duplicateur ne serait plus optimal, comme on le verra section 8.6. Pour les petits σ^2 , il faut alors utiliser l'équation (4.45), qui devient

$$\overline{\mathcal{F}}_{\text{opt}} = \frac{4}{4 + \sigma^2(3 - 2\sqrt{2})} \simeq \frac{4}{4 + 0,17\sigma^2} \quad (8.8)$$

On s'intéresse en général au cas où aucune information classique n'est disponible sur l'état envoyé par Alice, c'est à dire à la limite $\sigma^2 \rightarrow \infty$. Dans ce cas, le gain optimal vaut 1, et toute cloneuse symétrique est moins bonne que la cloneuse symétrique optimale :

$$\overline{\mathcal{F}}_B = \overline{\mathcal{F}}_C \leq \frac{2}{3}. \quad (8.9)$$

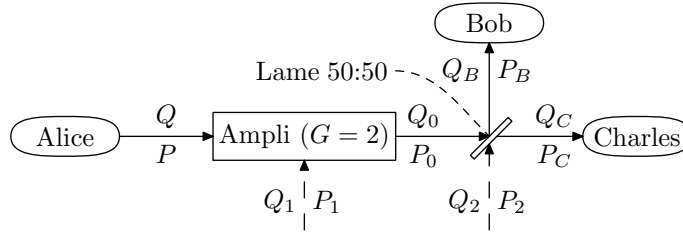


FIG. 8.2 – Implémentation d'un duplicateur quantique

8.5 Comment fabriquer un duplicateur ?

Un duplicateur quantique peut atteindre ces limites. Un schéma théorique a été proposé dans [137], une implémentation optique approximative a été proposée dans [139] et une implémentation exacte a été proposée simultanément dans [140, 141].

Un duplicateur quantique symétrique peut être construit avec un amplificateur indépendant de la phase de gain 2 suivi d'une lame semi réfléchissante, comme nous l'avons représenté FIG. 8.2. À la sortie de l'amplificateur, on a

$$Q_0 = \sqrt{2} \left(Q + \frac{1}{\sqrt{2}} Q_1 \right) \quad P_0 = \sqrt{2} \left(P - \frac{1}{\sqrt{2}} P_1 \right), \quad (8.10)$$

où $\begin{bmatrix} Q_1 \\ P_1 \end{bmatrix}$ représente le mode vide à l'origine du bruit de l'amplificateur. La lame semi-réfléchissante mélange alors le mode $\begin{bmatrix} Q_0 \\ P_0 \end{bmatrix}$ à un autre mode vide $\begin{bmatrix} Q_2 \\ P_2 \end{bmatrix}$ pour redistribuer les deux sorties $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ et $\begin{bmatrix} Q_C \\ P_C \end{bmatrix}$ à Bob et Charles.

$$Q_B = \frac{1}{\sqrt{2}} (Q_0 + Q_2) = Q + \frac{1}{\sqrt{2}} (Q_1 + Q_2) \quad (8.11a)$$

$$P_B = \frac{1}{\sqrt{2}} (P_0 + P_2) = P - \frac{1}{\sqrt{2}} (P_1 - P_2) \quad (8.11b)$$

$$Q_C = \frac{1}{\sqrt{2}} (Q_0 - Q_2) = Q + \frac{1}{\sqrt{2}} (Q_1 - Q_2) \quad (8.11c)$$

$$P_C = \frac{1}{\sqrt{2}} (P_0 - P_2) = P - \frac{1}{\sqrt{2}} (P_1 + P_2). \quad (8.11d)$$

On déduit alors aisément l'expression des bruits ajoutés par comparaison avec les équations (8.1) :

$$B_Q = \frac{1}{\sqrt{2}} (Q_1 + Q_2) \quad B_P = -\frac{1}{\sqrt{2}} (P_1 - P_2) \quad (8.12a)$$

$$C_Q = \frac{1}{\sqrt{2}} (Q_1 - Q_2) \quad C_P = -\frac{1}{\sqrt{2}} (P_1 + P_2). \quad (8.12b)$$

Il est alors aisé de vérifier que le cas où les modes $\begin{bmatrix} Q_1 \\ P_1 \end{bmatrix}$ et $\begin{bmatrix} Q_2 \\ P_2 \end{bmatrix}$ sont vides, les variances de tous les bruits ajoutés valent 1 et on a bien construit un duplicateur symétrique optimal pour des états cohérents distribués avec une forte variance.

Les expressions apparaissant dans (8.12) sont similaires aux sommes et différences impliquées dans le paradoxe EPR. Cela suggère que l'on peut construire un duplicateur asymétrique en injectant des faisceaux corrélés quantiquement (3.41) dans les modes $\begin{bmatrix} Q_1 \\ P_1 \end{bmatrix}$ et $\begin{bmatrix} Q_2 \\ P_2 \end{bmatrix}$, suivant un schéma similaire au dispositif de téléportation tout-optique proposé par Ralph

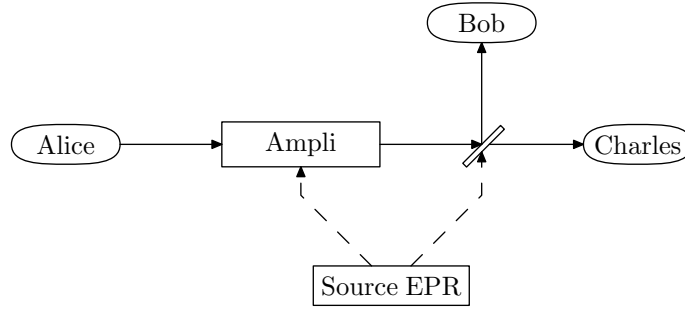


FIG. 8.3 – Schéma d'un duplicateur asymétrique

[142]. Il est alors aisé de voir que le degré d'intrication entre les deux faisceaux permet de jouer sur la distribution du bruit entre Bob et Charles.

Une autre cloneuse quantique asymétrique, sans doute expérimentalement plus réaliste, a été proposée par Fiurášek [141].

8.6 Les duplicateurs à gain différent de 1

Il peut être utile d'étudier les duplicateurs dont les gains dévient significativement de 1. La fidélité de tels duplicateurs sera inévitablement faible car les états produits seront différents des états fournis par Alice. Ils porteront néanmoins une information sur les états fournis en entrée, qui pourra être caractérisée par la variance des bruits équivalents.

De telles cloneuses peuvent correspondre à des mesures partielles, similaires à des mesures QND [81], où certains gains sont grands devant 1. Nous pouvons également étudier la production d'*anticlones*, qui correspondent au cas où le produit des gains $g_Q g_P$ est négatif. Un anticlon correspond au complexe conjugué de l'état d'entrée.

Nous montrerons que la possibilité de réaliser un duplicateur optimal, c'est à dire saturant les inéquations (8.3), n'est possible que dans certaines régions de l'espace des paramètres, décrit par le plan $\begin{pmatrix} G \\ H \end{pmatrix}$, où l'on a introduit $G \equiv g_Q g_P$ et $H \equiv h_Q h_P$. Nous montrerons ensuite que, même dans le domaine où une cloneuse optimale est réalisable, seules certaines valeurs du paramètre μ décrivant la répartition du bruit entre Bob et Charles sont accessibles, sauf dans le cas $G = H = 1$ où toutes les valeurs sont accessibles.

8.6.1 Conditions nécessaires à un duplicateur optimal

En effet, il ne faut pas oublier l'existence des relations de commutation (8.2c), qui deviennent

$$[B_Q, B_P] = \left(\frac{1}{G} - 1\right) 2iN_0 \quad [C_Q, C_P] = \left(\frac{1}{H} - 1\right) 2iN_0, \quad (8.13)$$

On en déduit les conditions supplémentaires sur les variances des bruits équivalents

$$\langle B_Q^2 \rangle \langle B_P^2 \rangle \geq \left(\frac{1}{G} - 1\right)^2 N_0^2 \quad (8.14a)$$

$$\langle C_Q^2 \rangle \langle C_P^2 \rangle \geq \left(\frac{1}{H} - 1\right)^2 N_0^2. \quad (8.14b)$$

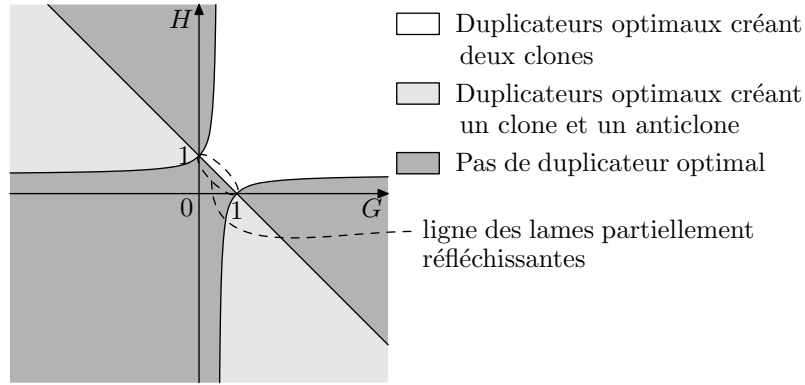


FIG. 8.4 – Domaines d'existence des duplicateurs optimaux

Le produit des équations (8.3) d'une part et celui des équations (8.14) d'autre part limitent les valeurs accessibles par le produit des bruits équivalents :

$$\langle B_Q^2 \rangle \langle B_P^2 \rangle \langle C_Q^2 \rangle \langle C_P^2 \rangle \geq N_0^4 \quad (8.15a)$$

$$\langle B_Q^2 \rangle \langle B_P^2 \rangle \langle C_Q^2 \rangle \langle C_P^2 \rangle \geq \left[\frac{(1-G)(1-H)}{GH} \right]^2 N_0^4. \quad (8.15b)$$

Seule une de ces inégalités est contraignante à la fois, et on ne peut avoir une cloneuse optimale au sens du bruit équivalent que si la première peut être saturée, c'est à dire si

$$\left| \frac{(1-G)(1-H)}{GH} \right| \leq 1 \quad (8.16)$$

$$|1 - G - H + GH| \leq |GH|. \quad (8.17)$$

8.6.2 Domaine d'existence des duplicateurs optimaux

L'inégalité (8.17) a quatre expressions possibles suivant le signe de GH et de $(1-G)(1-H)$. Nous étudierons ces différents cas dans cette section, ce qui nous permettra de tracer la FIG. 8.4, qui définit les zones de l'espace des paramètres où il est possible de construire des cloneuses optimales.

Il est aisé de constater que les lignes de l'espace des paramètres d'équations $G = 1$ et $H = 1$ vérifient toujours cette inégalité, et qu'il est donc toujours possible de fabriquer un duplicateur optimal dont l'une des sorties à un gain de 1.

$GH > 0$, correspond soit au cas où les deux copies sont des anticlones, soit au cas où ce sont des clones. $(1-G)(1-H) > 0$ correspond au cas où G et H sont du même côté de 1 : soit ils sont tous les deux supérieurs à 1 soit ils sont tous les deux inférieurs à 1. Dans ce cas, (8.17) devient

$$G + H \geq 1. \quad (8.18)$$

Cette relation n'est jamais vérifiée lorsque G et H sont négatifs et elle est toujours vérifiée lorsqu'ils sont tous les deux supérieurs à 1. On en déduit qu'une cloneuse qui produit deux anticlones n'est jamais une cloneuse optimale, et qu'il est toujours possible de réaliser une cloneuse optimale amplificatrice

Le cas le plus intéressant est celui où G et H sont compris entre 0 et 1. La ligne $G + H = 1$, où les deux inéquations sont simultanément saturées correspond, entre autres, aux lames partiellement réfléchissantes. Une lame partiellement réfléchissante est donc le duplicateur quantique optimal minimum, et aucun duplicateur qui crée deux clones plus atténués que s'ils sortaient d'une lame partiellement réfléchissante n'est optimal.

Si $GH > 0$ et $(1 - G)(1 - H) < 0$, (8.17) devient

$$\left(G - \frac{1}{2}\right) \left(H - \frac{1}{2}\right) \geq -\frac{1}{4}. \quad (8.19)$$

En particulier si H et G sont supérieurs à $\frac{1}{2}$, il est toujours possible d'avoir une cloneuse optimale.

Il reste à examiner le cas où $GH < 0$, qui correspond à un duplicateur qui produit un clone et un anticlon. Nous nous limiterons au cas $G < 0$ et $H > 0$, l'autre cas s'en déduisant par symétrie. Si $H > 1$, (8.17) devient

$$\left(G - \frac{1}{2}\right) \left(H - \frac{1}{2}\right) \leq -\frac{1}{4}. \quad (8.20)$$

Si $H > 1$, elle devient

$$G + H \leq 1. \quad (8.21)$$

Ces équations définissent trois domaines dans le plan $\left(\frac{G}{H}\right)$ dans lesquels on peut fabriquer des duplicateurs optimaux (cf FIG. 8.4) : un domaine qui comprend tous les duplicateurs produisant deux clones, et deux domaines symétriques contenant les duplicateurs produisant un couple clone-anticlon.

8.6.3 Distribution du bruit entre Bob et Charles

Là où des cloneuses optimales sont réalisables, on peut caractériser les bruits équivalents par les paramètres μ et β introduits dans les équations (8.5). Il n'est en général pas possible à G et H fixés de choisir arbitrairement le coefficient μ indiquant la répartition du bruit entre Bob et Charles. C'est assez intuitif si on est sur la ligne des lames partiellement réfléchissantes ($G + H = R + T = 1$), où μ est fixé et vaut $\mu = \frac{H}{G} = \frac{T}{R}$. Plus généralement, les équations (8.14) s'écrivent alors

$$\mu \geq \left| \frac{1}{G} - 1 \right| \qquad \frac{1}{\mu} \geq \left| \frac{1}{H} - 1 \right|. \quad (8.22)$$

On voit bien que ces équations seront contraignantes pour μ , sauf dans le cas où $G = H = 1$.

Il est aisé de se convaincre que, dans le cas symétrique ($\mu = 1$), ces inéquations sont équivalentes à

$$G \geq \frac{1}{2} \qquad H \geq \frac{1}{2}. \quad (8.23)$$

Il n'est donc possible de faire des cloneuses optimales répartissant symétriquement le bruit que pour des gains supérieurs à $\frac{1}{2}$.

Si $\mu > 1$ (l'autre cas se déduit par symétrie), Bob reçoit plus de bruit que Charles et l'on a

$$G \in \left] -\infty, -\frac{1}{\mu-1} \right] \cup \left[\frac{1}{\mu+1}, +\infty \right[\qquad H \in \left[\frac{\mu}{\mu+1}, \frac{\mu}{\mu-1} \right] \quad (8.24)$$

On voit en particulier que lorsque la copie de Bob est moins bonne que celle de Charles ($\mu \geq 1$), Charles ne peut avoir d'anticlon car $H > 0$. Un anticlon est donc toujours moins bon que le clone associé par une cloneuse optimale.

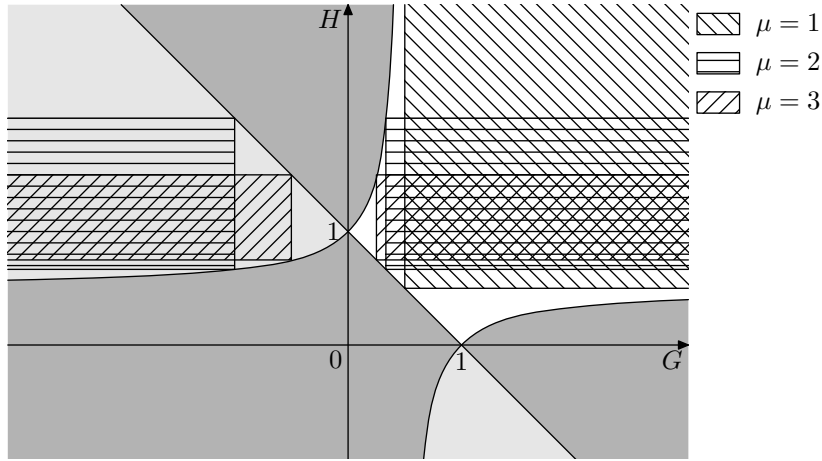


FIG. 8.5 – Exemples de valeurs possibles du coefficient μ et zones de l'espace des paramètres où ils sont réalisables.

8.7 Clonage symétrique $1 \rightarrow M$

Les références [138, 134] généralisent le problème du clonage aux cloneuses $\left(\begin{smallmatrix} N \\ N^* \end{smallmatrix}\right) \rightarrow \left(\begin{smallmatrix} M \\ M^* \end{smallmatrix}\right)$, prenant N copies de l'état initial et N^* copies de son complexe conjugué en entrée, et produisant M clones et M^* anticlones à la sortie. [140, 141] proposent des implémentations de ces machines cloneuses avec un amplificateur et un réseau de lames partiellement réfléchissantes.

Nous retrouverons ici une démonstration de la limite de fidélité d'une cloneuse symétrique $1 \rightarrow M$ ($M > 2$) de gain 1, toujours en utilisant les notations linéarisées.

Les quadratures des M sorties d'une telle cloneuse s'écriront

$$\forall 1 \leq k \leq M \quad Q_k = Q + B_{Q_k} \quad (8.25a)$$

$$P_k = P + B_{P_k}. \quad (8.25b)$$

Si elle est symétrique, nous pouvons définir les trois quantités $\mathcal{C}_Q$, χ_Q et χ_P par

$$\forall k \neq l \quad \langle B_{Q_k} B_{Q_l} \rangle = \mathcal{C}_Q N_0 \quad (8.26a)$$

$$\forall k \quad \langle B_{Q_k}^2 \rangle = \chi_Q N_0 \quad \langle B_{P_k}^2 \rangle = \chi_P N_0. \quad (8.26b)$$

Comme précédemment, nous avons les relations de commutation

$$\forall k \neq l \quad [B_{Q_k}, B_{P_l}] = -2iN_0 \quad (8.27a)$$

$$\forall k \quad [B_{Q_k}, B_{P_k}] = 0. \quad (8.27b)$$

Pour tout $\lambda \in \mathbb{R}$, nous pouvons définir

$$\Lambda = B_{Q_1} + \lambda \sum_{k=2}^M B_{Q_k}. \quad (8.28)$$

Le commutateur de cet opérateur avec le bruit équivalent B_{P_1} vaut alors

$$[\Lambda, B_{P_1}] = -\lambda (M - 1) 2iN_0, \quad (8.29)$$

d'où on déduit la relation

$$\langle \Lambda^2 \rangle \langle B_{P_1}^2 \rangle \geq \lambda^2 (M-1)^2 N_0^2. \quad (8.30)$$

La variance de Λ peut se calculer directement à partir de sa définition (8.28)

$$\langle \Lambda^2 \rangle = \langle B_{Q_1}^2 \rangle + \lambda^2 \sum_{k=2}^M \langle B_{Q-k}^2 \rangle + 2\lambda \sum_{k>1} \langle B_{Q_1} B_{Q_k} \rangle + 2\lambda^2 \sum_{k>1} \sum_{l>k} \langle B_{Q_k} B_{Q_l} \rangle \quad (8.31)$$

$$= \left\{ [1 + \lambda^2(M-1)]\chi_Q + 2\lambda \left[1 + \lambda \frac{M-2}{2}\right] (M-1)\mathcal{C}_Q \right\} N_0 \quad (8.32)$$

Cette expression est compliquée par les coefficients de corrélations $\mathcal{C}_Q$, qui ne nous intéressent pas. Nous pouvons nous en débarrasser en choisissant judicieusement le paramètre λ : si l'on pose

$$\lambda = -\frac{2}{M-2}, \quad (8.33)$$

l'équation (8.32) se simplifie et devient

$$\langle \Lambda^2 \rangle = \frac{M^2}{(M-2)^2} \chi_Q N_0. \quad (8.34)$$

Cette expression nous permet d'écrire la relation (8.30) sous la forme familière d'une relation « de Heisenberg »

$$\chi \equiv \sqrt{\chi_Q \chi_P} \geq 2 \frac{M-1}{M}. \quad (8.35)$$

Cette relation reste valable pour le cas trivial ($M=1$) et le duplicateur quantique ($M=2$).

Comme le gain de cette machine cloneuse vaut 1 pour toutes ses sorties, on peut calculer la fidélité de celles-ci en utilisant l'équation (4.39). Si $\chi_Q = \chi_P = \chi$ et qu'elle est utilisée pour cloner des états cohérents, on a

$$\mathcal{F} \leq \frac{2}{2+\chi} = \frac{M}{2M-1}. \quad (8.36)$$

Il est aisé de voir que la limite est la même pour des états comprimés si on adapte la répartition du bruit entre les quadratures au facteur de compression.

Cette limite est bien sûr vraie pour les grands facteurs de modulation seulement, car nous avons limité notre étude au cas de gain 1.

Ces limites de clonages nous serviront pour discuter la signification physique de différents critères de téléportation quantique avec de variables continues, au chapitre suivant, mais aussi pour évaluer les attaques possibles des protocoles de cryptographie quantiques développés dans la partie IV.

Chapitre 9

Critères de téléportation quantique

Nous commencerons ce chapitre, consacré à la téléportation quantique, par une définition de la téléportation (section 9.1) et du rôle que joue ce protocole dans le domaine de l'information quantique (section 9.2). Nous décrirons ensuite, dans les sections 9.3 et 9.4, deux critères employés pour évaluer des expériences de téléportation quantique¹.

9.1 Définition de la téléportation

En 1993, Bennett et ses collaborateurs ont découvert un protocole de communication quantique qu'ils ont appelé, en s'inspirant de la science-fiction, *téléportation quantique* [76]. D'une manière plus générale, la téléportation est en général définie [76, 143, 144] comme un processus au cours duquel un objet disparaît d'un endroit pour apparaître ailleurs, d'une manière plus ou moins instantanée. En d'autre terme, il s'agit de transporter un objet de A en B sans parcourir la distance entre les deux points.

Pour accomplir cette tâche fantastique, les auteurs de science-fiction avaient imaginé de nombreux processus plus ou moins fantaisistes [143] que nous ne détaillerons pas ici. Nous présenterons dans la section 9.1.1 les limites de la méthode intuitive par mesure et de reconstruction, souvent appelée *téléportation classique*, avant d'exposer rapidement, dans la section 9.1.2, comment la téléportation quantique contourne ces limitations. Nous détaillerons ce processus pour les variables continues dans la section 9.1.3.

9.1.1 Mesure et reconstruction

Une manière intuitive pour « téléporter » un objet consiste à procéder par mesure et reconstruction : toutes les caractéristiques de l'objet sont mesurées et transmises au point B en utilisant un système de télécommunications, où elles servent à reconstituer cet objet [143, 144]. Ce type de processus est parfois appelé *téléportation classique*. Cependant, il s'agit plus d'une copie à distance que d'un véritable transport. Il est ainsi plus naturel de désigner un fax, qui fonctionne sur ces principes, par le nom de *télécopieur* que de l'appeler « téléporteur de documents » !

La physique impose des limites à ce qui est réalisable par ce type de système. La transmission des informations nécessaires à la reconstitution de l'objet téléporté est bien entendu

¹Le contenu de ce chapitre a fait l'objet des publications [8, 9, 10]

limitée à la vitesse de la lumière c , mais la limitation qui nous intéressera le plus est celle imposée par la physique quantique. En effet le principe d'incertitude de Heisenberg limite la qualité des mesures qui peuvent être faites sur l'état de l'objet et la qualité de la reconstruction, même si cette limite n'est pas gênante pour téléporter un objet macroscopique à température ambiante [143, 145, 144]. Nous avons étudié quantitativement ces limites dans le cas où l'objet à téléporter est un mode du champ électromagnétique dans un état cohérent dans les sections 4.3.4 et 4.4.4.

Au delà de la science-fiction, la téléportation classique, est intéressante comme point de comparaison avec la téléportation quantique. Ainsi, Braunstein et ses collaborateurs [146] considèrent que l'on peut parler de téléportation quantique à partir du moment où l'on fait mieux que n'importe quel système classique.

9.1.2 Téléportation quantique

Plus généralement, un système de téléportation étant assimilé à un système de copie à distance, le théorème de non-clonage quantique [135, 136] semblait exclure la possibilité d'une téléportation parfaite. Cette difficulté a été contournée dans le cas des qubits par Bennett et ses collaborateurs [76], qui ont imaginé un protocole où l'original est nécessairement détruit *avant* que la copie ne soit recréée [46, 143, 145].

Supposons qu'Alice veuille téléporter un qubit dans un état quantique inconnu à Bob. Le protocole de téléportation nécessite qu'ils disposent chacun de la moitié d'une paire de particules dans un état maximalement intriqué. L'état de chacune des particules n'est pas défini, mais elles sont parfaitement corrélées. En d'autres termes, le résultat de toute mesure faite sur la particule d'Alice ou de Bob est aléatoire, mais si Alice et Bob font la même mesure sur leurs particules respectives, ils auront exactement le même résultat. L'astuce consiste alors pour Alice à faire une mesure conjointe, ou *mesure de Bell* sur sa particule de la paire intriquée et le qubit à téléporter. Cette mesure la renseignera sur la différence entre le qubit et la demi-paire, mais ne lui dira rien sur l'état du qubit lui-même. Elle détruit nécessairement le qubit au cours du processus de mesure. Comme les deux moitiés de la paire intriquée sont identiques, la différence entre la qubit et la demi-paire d'Alice est égale à la différence entre le qubit et la demi-paire de Bob. Si Alice transmet cette information à Bob, celui-ci pourra alors modifier l'état de sa demi-paire par la modulation adéquate, qui se trouvera alors dans l'état exact du qubit de Bob, sans que ni Alice, ni Bob n'aient rien appris de cet état au cours du processus.

Ce processus ressemble donc énormément à la téléportation de la science-fiction, que nous avons décrit plus haut, d'où son nom : le qubit disparaît en effet de chez Alice avant de réapparaître chez Bob. De plus le théorème de non-clonage empêche que plusieurs copies existent simultanément, ce qui renforce l'impression que c'est vraiment le qubit d'Alice qui a été transporté plutôt que copié.

9.1.3 Téléportation quantique de variables continues

Le protocole de téléportation proposé en 1993 par Bennett et ses collaborateurs pour des qubits [76] a été étendu par Vaidman en 1994 aux variables continues [89]. Braunstein et Kimble ont ensuite proposé une implémentation expérimentale de ce protocole [90], qu'ils ont réalisé avec Furusawa et leurs collaborateurs [91].

Un système de téléportation quantique pour des variables continues est similaire au système de téléportation classique décrit section 4.3.4, et peut en particulier être décrit par des équations analogue aux équations (4.29). On a donc

$$Q_C = g_Q (Q_B + C_Q) = g_Q (Q + B_Q + C_Q) \quad (9.1a)$$

$$P_C = g_P (P_B + C_P) = g_P (P + B_P + C_P), \quad (9.1b)$$

où $\begin{bmatrix} Q \\ P \end{bmatrix}$ représentent la position et l'impulsion de l'objet à téléporter, $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ représentent les résultat des mesures et $\begin{bmatrix} Q_C \\ P_C \end{bmatrix}$ la position et l'impulsion de l'objet téléporté à l'arrivée chez Bob. $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ correspond au bruit ajouté lors de la mesure et $\begin{bmatrix} C_Q \\ C_P \end{bmatrix}$ à celui ajouté lors de la reconstruction. Comme dans le cas de la téléportation classique, on a les relations de commutation

$$\begin{aligned} [B_Q, B_P] &= -2iN_0 & [C_Q, C_P] &= \frac{1}{G} 2iN_0, \end{aligned} \quad (9.2)$$

où $G \equiv g_Q g_P$. Ces relations de commutation imposent une variance minimale aux bruits ajoutés à chaque étape. On a cependant

$$[B_Q + C_Q, B_P + C_P] = \frac{1-G}{G} 2iN_0, \quad (9.3)$$

ce qui signifie en particulier que si $G = 1$, la variance de la somme des bruits ajoutés peut être arbitrairement faible. Il faut pour cela que les bruits $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ et $\begin{bmatrix} C_Q \\ C_P \end{bmatrix}$ soient corrélés quantiquement.

Le bruit de mesure $\begin{bmatrix} B_Q \\ B_P \end{bmatrix}$ vient de la mesure simultanée de l'impulsion et de la position. Une telle mesure simultanée peut-être modélisée par le système suivant : le faisceau à mesurer $\begin{bmatrix} Q \\ P \end{bmatrix}$ est séparé en deux par une lame semi-réfléchissante ce qui permet de mesurer la position et l'impulsion avec deux détecteurs homodyne. Cependant, le champ électromagnétique $\begin{bmatrix} Q_- \\ P_- \end{bmatrix}$ présent à l'autre entrée de la lame semi-réfléchissante se mélange au champ mesuré et introduit du bruit. On a donc

$$B_Q = -Q_- \quad \text{et} \quad B_P = P_-. \quad (9.4)$$

Dans le cas de la téléportation classique, $\begin{bmatrix} Q_- \\ P_- \end{bmatrix}$ est le vide quantique, mais rien n'interdit d'utiliser un autre faisceau, comme nous le verrons plus bas.

Le bruit de reconstruction $\begin{bmatrix} C_Q \\ C_P \end{bmatrix}$ correspond au bruit quantique du mode de Bob. On peut modéliser la reconstitution de Bob comme le déplacement d'un état quantique, décrit par $\begin{bmatrix} Q_+ \\ P_+ \end{bmatrix}$ en fonction du résultat de la mesure d'Alice. On a alors

$$g_Q C_Q = Q_+ \quad \text{et} \quad g_P C_P = P_+. \quad (9.5)$$

Dans le cas de la téléportation classique, $\begin{bmatrix} Q_+ \\ P_+ \end{bmatrix}$ est usuellement le vide quantique, mais ça ne sera pas le cas pour la téléportation quantique.

Nous nous limiterons dans la suite au cas où les gains valent 1. Dans ce cas le bruit total ajouté par la téléportation pour chaque quadrature peut s'écrire

$$B_Q + C_Q = Q_+ - Q_- \quad \text{et} \quad B_P + C_P = P_+ + P_-. \quad (9.6)$$

La téléportation quantique avec des variables continues consiste à injecter dans les modes $+$ et $-$ les faisceaux EPR définis par l'équation (3.41). Le bruit total ajouté par le processus de téléportation peut alors être arbitrairement faible, si ces états sont générés avec de forts facteurs de compression. Plus le facteur de compression est fort, plus le bruit ajouté au cours de la mesure est important, et moins Alice en apprend sur l'état à téléporter. Par contre le bruit ajouté par Bob au cours de la reconstruction est fortement anticorrélé au bruit de mesure et le compense presque exactement.

Cela permet donc à Bob de reconstituer presque exactement l'état initial à partir de la mesure d'Alice, sans que ni Alice, ni Bob, n'aient appris d'information significative sur l'état à téléporter. Comme Alice doit détruire cet état pour effectuer ses mesures, l'état en question n'existe jamais en double exemplaire et le théorème de non-clonage quantique est respecté.

9.2 La téléportation quantique en pratique

9.2.1 Intérêts de la téléportation quantique

Même si elle n'est pas envisagée comme le moyen de transport de la science-fiction, la téléportation quantique présente de nombreux intérêts.

Il s'agit en effet d'un moyen « d'exploiter » l'intrication d'un système quantique, quelque soit sa nature. Ainsi, après les photons uniques [147, 148] et les quadratures du champ électromagnétique [91], des protocoles de téléportation quantique sont proposés pour quasiment tous les systèmes étudiés dans le domaine de l'information quantique, des atomes uniques en cavité, aux systèmes supraconducteurs, en passant par les nuages d'atomes froids [149], pour n'en citer que quelques uns. La réalisation expérimentale d'un protocole de téléportation quantique est considérée comme une étape dans le contrôle de l'intrication, qui permet de comparer entre eux des systèmes expérimentaux très différents. Il est intéressant de noter que les inégalités de Bell ont en pratique un statut similaire, étant donné que peu de monde s'attend aujourd'hui à ce que ce test de la mécanique quantique échoue.

Sur un plan plus théorique, les protocoles de téléportation quantique sont un outil théorique important pour étudier les liens entre les informations classique et quantique. C'est également un outil de base dans d'autres protocoles de communication quantique, notamment dans les registres des ordinateurs quantiques.

Le plus prometteur est sans doute le transfert de variables quantiques de natures différentes. Cela peut être des photons de longueur d'onde différente utilisés pour d'autres protocoles de communication quantique : l'état quantique d'un photon visible, dont la longueur d'onde correspond au maximum d'efficacité quantique des détecteurs peut ainsi être téléporté sur un photon de longueur d'onde télécom, adapté à la transmission dans une fibre optique. Cela peut aussi correspondre au transfert d'un état quantique de photons à atomes et réciproquement [149], ce qui permet notamment d'envisager la fabrication de mémoires et de répéteurs quantiques.

9.2.2 Quand y a-t-il eu téléportation ?

Si tout le monde est d'accord pour dire ce que serait une expérience parfaite de téléportation quantique, les avis divergent [150, 151, 152, 153] sur la définition du seuil au delà duquel une expérience imparfaite peut être qualifiée de système de téléportation quantique.

En effet, la téléportation quantique a peu d'utilité par elle-même, contrairement à la cryptographie quantique, par exemple. Il est donc possible de définir un critère différent en fonction de l'utilisation envisagée de la téléportation.

Dans le cas de la téléportation de qubits, les expériences réalisées à ce jour permettent de téléporter presque parfaitement un état quantique, mais avec une probabilité de succès de 25 % seulement [147, 148]. Le succès ou l'échec de la tentative de téléportation ne pouvant être déterminé qu'après la destruction de l'état d'Alice, ce type de téléportation est en général désigné sous le nom de téléportation *a posteriori* [150, 151]. Le protocole de téléportation quantique de variables continues proposé par Braunstein et Kimble [90], qu'ils ont réalisé avec Furusawa et leurs collaborateurs [91], ne présente pas cet inconvénient : la téléportation est toujours réussie. En revanche, la qualité de l'état téléporté est en général dégradée par rapport à l'état initial.

La téléportation quantique crée un canal quantique entre Alice et Bob, qui peut être décrit avec le formalisme que nous avons introduit dans la section 4.3. Elle est souvent évaluée, comme un canal quantique, par la fidélité moyenne $\overline{\mathcal{F}}$ que nous avons définie section 4.4.2. Cette fidélité moyenne dépend de l'ensemble des états à téléporter dans le cas des variables continues. Braunstein et ses collaborateurs [146] soulignent qu'il est assez simple de trouver deux états orthogonaux qui exigent une fidélité arbitrairement proche de 1 pour être téléportés de façon satisfaisante. Les états généralement considérés dans la téléportation avec des variables continues sont des états cohérents, distribués suivant une gaussienne.

Si Alice connaît précisément l'état quantique à téléporter, elle peut le décrire à Bob qui peut le recréer parfaitement. Ce scénario ne correspond pas à la téléportation quantique, où Alice doit mesurer un état qu'elle ne connaît pas. Pour évaluer un système de téléportation quantique, Braunstein, Fuchs et Kimble [146] ont décrit un scénario faisant intervenir un troisième personnage en plus d'Alice et de Bob : Victor, le vérificateur. C'est lui qui choisit l'état à téléporter et le donne à Alice. Il le récupère également de Bob afin de vérifier la qualité de la téléportation. En répétant l'opération un certain nombre de fois, Victor peut mesurer la fidélité moyenne de la téléportation $\overline{\mathcal{F}}$, et peut ainsi vérifier qu'Alice et Bob n'ont pas « triché », en utilisant un système de téléportation classique au lieu d'un système de téléportation quantique.

9.3 Le seuil classique

9.3.1 Définition du seuil

Braunstein, Fuchs et Kimble ont proposé de considérer qu'il y avait téléportation quantique dès lors que Victor, le vérificateur, pouvait déduire qu'Alice et Bob n'utilisaient pas un protocole de téléportation classique, par simple mesure et reconstruction [146, 154]. Comme nous l'avons vu dans la section 4.3.4, le bruit ajouté par un protocole de téléportation classique est d'au moins deux fois le bruit quantique standard ($\chi \geq 2$), dans le cas d'un gain unité. Pour reprendre une terminologie utilisée par Furusawa et ses collaborateurs [91], Alice et Bob doivent payer deux qutaxes (*qutaxes*, en anglais) : une fois lors de la mesure et une fois lors de la reconstruction. Comme nous l'avons vu dans la section 4.4.4, cela correspond à une fidélité moyenne maximale $\overline{\mathcal{F}}_{\text{class}} \leq \frac{1}{2}$, pour des états cohérents fortement modulés.

Furusawa et ses collaborateurs ont téléporté des états cohérents en 1998 [91] avec une

fidélité de 58 %, au delà de ce seuil. Des améliorations postérieures de cette expérience leur ont permis d'atteindre une fidélité de 61 % cette année [155]. Bowen et ses collaborateurs ont également téléporté des états cohérents avec une fidélité moyenne de 64 % en 2002 [156].

9.3.2 Lien entre le seuil classique et l'intrication

Ce seuil correspond au niveau minimal pour lequel les faisceaux d'Alice et de Bob sont intriqués [116, 117]. Alice et Bob peuvent alors en théorie utiliser un protocole de purification [157] sur leurs faisceaux EPR imparfaitement intriqués pour en extraire des faisceaux parfaitement intriqués, et réaliser une téléportation arbitrairement proche de la perfection. Ces protocoles sont loin d'être évident à implémenter, et ils restent même à imaginer si l'on veut rester dans le domaine des variables continues.

Si Victor dispose de faisceaux jumeaux fortement intriqués, avec un facteur de compression arbitrairement proche de 0, il peut faire téléporter l'un d'eux par Alice et Bob et conserver l'autre. L'application des critères donnés dans les références [116, 117] nous permet de voir aisément que le faisceau téléporté chez Bob est intriqué au faisceau conservé par Victor si et seulement si le bruit ajouté par la téléportation est inférieur à deux fois le bruit quantique standard ($\chi < 2$), c'est à dire si la fidélité de téléportation d'états cohérents est meilleure que la limite classique $\overline{\mathcal{F}} > \frac{1}{2}$.

9.3.3 Rôle des pertes et du facteur de compression

Comme ce seuil classique correspond à une démonstration de l'utilisation d'intrication, il est en particulier atteignable avec une intrication infinitésimale. Il suffit en effet en théorie de générer les faisceaux EPR à l'aide de faisceaux à peine intriqués et de les transmettre à Alice et Bob à travers un canal qui peut correspondre à de fortes pertes pour pouvoir franchir ce seuil. Bien sûr, c'est plus difficile expérimentalement que cela, car l'intrication permet de compenser un certain nombre d'imperfections expérimentales.

Cette limite peut sembler « trop facile » à atteindre. En tout cas, un certain nombre de propriétés de la téléportation quantique, notamment celles qui lui ont valu son nom, ne sont pas remplies lorsqu'on franchit le seuil $\overline{\mathcal{F}} = \frac{1}{2}$ et le seront lorsqu'on atteint un autre seuil $\overline{\mathcal{F}} = \frac{2}{3}$, comme nous le verrons dans la section suivante.

9.4 Le seuil $\overline{\mathcal{F}} = \frac{2}{3}$

Le seuil $\overline{\mathcal{F}} = \frac{2}{3}$ a un certain nombre de propriétés et a été trouvé de plusieurs manières différentes. Il correspond à un bruit équivalent à l'entrée d'une fois le bruit quantique standard $\chi = 1$, et est bien entendu plus difficile à atteindre que le précédent. A ce jour, aucune expérience de téléportation quantique n'a atteint cette limite, même si l'expérience de Bowen et ses collaborateurs publiée cette année est très proche de ce seuil, avec une fidélité moyenne de 64 %.

Ralph et Lam ont trouvé cette limite [158], qu'ils ont exprimé différemment, simultanément et indépendamment de Braunstein, Fuchs et Kimble qui trouvaient la limite classique discutée ci-dessus. Leur formulation de ce critère est similaire aux critères utilisés dans le contexte des mesures quantiques non destructives [84].

9.4.1 Paradoxe EPR

Pour atteindre un tel seuil, Alice et Bob doivent pouvoir fabriquer des états intriqués avec un facteur de compression d'au moins 3 dB et transmettre l'un des faisceaux EPR avec moins de 3 dB de pertes. Ces conditions sont nécessaires, mais nullement suffisantes ; la condition générale s'exprime

$$T(1 - s) > \frac{1}{2}, \quad (9.7)$$

où s correspond au facteur de compression et T représente la transmission de la ligne.

Cette limite correspond également à la limite où le paradoxe EPR se manifeste pour les faisceaux jumeaux d'Alice et de Bob, dans une formulation proche de la formulation originale d'Einstein [37, 152]. En effet, elle correspond à la limite où les positions et impulsions des bruits ajoutés chez Alice et Bob sont nécessairement mieux corrélés que le bruit quantique standard [8].

De plus, elle correspond à la limite pour que Victor puisse faire du transfert d'intrication, dans des conditions réalistes. Supposons en effet que Victor ne dispose pas d'états mieux intriqués qu'Alice et Bob et qu'il essaie de faire téléporter cette intrication jusqu'à Bob. Pour que l'état reçu par Bob soit toujours intriqué avec l'état que Victor a gardé, il faut que les différents facteurs d'intrication utilisés soient d'au moins 3 dB dans le cas sans pertes.

9.4.2 Fonction de Wigner

Si on téléporte un état dont la fonction de Wigner est localement négative, la fonction de Wigner de l'état reçu par Bob ne peut être négative que si la fidélité moyenne de téléportation d'états cohérents est supérieure à $\overline{\mathcal{F}} > \frac{2}{3}$. Il est aisé de s'en convaincre par une comparaison avec la section 3.5.6 si l'on se souvient que cette condition est équivalente à $\chi < 1$.

9.4.3 Non-Clonage

Ce critère a initialement été trouvé par des considérations sur les coefficients T de transfert d'information introduits dans la section 7.5 [158, 8]. Ces considérations nous ont amené à faire le lien avec le clonage quantique, en examinant la possibilité de fuites d'information au cours du protocole de téléportation [9, 10].

En effet le lien entre téléportation quantique et le clonage quantique est profond. Bennett et ses collaborateurs [76] mentionnent ainsi le théorème de non-clonage dans le paragraphe où ils expliquent pourquoi ils ont appelé *téléportation* leur protocole : c'est le théorème de non clonage qui oblige Alice à détruire son état quantique avant que Bob ne puisse le reconstruire, ce qui fait étrangement ressembler ce processus à la téléportation de la science fiction. Ce théorème de non-clonage est d'ailleurs souvent cité lorsque l'on parle de téléportation [144, 155]. Braunstein et ses collaborateurs [153, 154] vont plus loin en soulignant que « dans une théorie qui autoriserait de cloner des états, il n'y aurait aucune utilité à parler de téléportation—des états inconnus pourraient être clonés et transmis, avec une fidélité arbitrairement proche de 1 »²

²« In a theory which allows states to be cloned, there would be no need to discuss teleportation at all—unknown states could be cloned and transmitted, with a fidelity arbitrarily close to 1. »

	Nombre de copies autorisées	Opérations locales	Opérations distantes
1	1	Transfert Quantique	Téléportation Quantique
2	2	Copie Quantique	Télécopie Quantique
3	3		
1	1		
2	∞	Copie Classique	Télécopie Classique
0			

TAB. 9.1 – Régimes de téléportation en fonction de la fidélité $\overline{\mathcal{F}}$

Si Alice garde une copie de l'état qu'elle téléporte chez Bob, le protocole mérite plus le nom de *télécopie quantique* que de téléportation. Ainsi, si $\overline{\mathcal{F}} \leq \frac{2}{3}$, Alice peut garder une copie de meilleure qualité que celle que Bob reçoit, car cette fidélité est inférieure à celle du clonage optimal donnée par l'équation (8.9).

Ainsi, la fidélité de 58 % ($\chi = 1,45$) obtenue par Furusawa et ses collaborateurs [91] pourrait autoriser Alice à garder une copie de meilleure qualité avec une fidélité de 74 %. L'expérience de Zhang et ses collaborateurs [155] (respectivement Bowen et ses collaborateurs [156]) a une fidélité de 61 % (respectivement 64 %), ce qui autoriserait Alice à garder une copie de fidélité 71 % (respectivement 69 %). On voit bien que si aucune expérience n'a franchi à ce jour la limite $\overline{\mathcal{F}} > \frac{2}{3}$, elle semble proche.

Au delà de cette limite, la copie conservée par Alice sera de moins bonne qualité que celle reçue par Bob. À la limite $\overline{\mathcal{F}} \rightarrow 1$, le théorème de non clonage empêche Alice (ou qui que ce soit) d'apprendre quoi que ce soit sur l'état transféré à Bob.

L'équation (8.36) nous permet de définir une série de seuils en fonction du nombre $M - 1$ de copies qu'Alice peut garder. Ainsi, une fidélité $\overline{\mathcal{F}} \leq \frac{3}{5}$ permettrait à Alice de conserver deux copies de l'état téléporté, au moins aussi bonnes que celle reçue par Bob. La limite $M \rightarrow \infty$, où Alice conserve une infinité de copies correspond à la limite classique $\overline{\mathcal{F}} \leq \frac{1}{2}$.

En conclusion, si la limite $\overline{\mathcal{F}} > \frac{1}{2}$ ne peut pas être atteinte par des protocoles de téléportation classique et qu'elle est fortement liée à l'intrication quantique, la limite, plus stricte, $\overline{\mathcal{F}} > \frac{2}{3}$ a de nombreuses propriétés intéressantes. Elle est notamment fortement liée au clonage quantique. Comme il semble difficile d'appeler téléportation une opération où l'état initial n'est pas nécessairement détruit, nous avons proposé l'appellation *télécopie quantique* (*quantum fax* en anglais) pour le régime quantique $\frac{1}{2} < \overline{\mathcal{F}} \leq \frac{2}{3}$, où Alice ne détruit pas nécessairement l'état quantique à téléporter (voir TAB. 9.1).

Quatrième partie

Cryptographie Quantique

Chapitre 10

Généralités sur la cryptographie quantique

Nous avons étudié au cours de cette thèse des protocoles de cryptographie quantique avec des variables continues, que nous présenterons dans les trois chapitres suivants. Mais, auparavant, nous rappellerons dans la section 10.1 les principes de Kerckhoffs, qui sont fondamentaux en cryptologie depuis près de 120 ans. Nous aborderons plus spécifiquement la cryptographie quantique dans la section 10.2, où nous définirons rapidement la distribution quantique de clefs et où nous dresserons un bref état des lieux des protocoles de cryptographie quantique avec des variables continues dont nous avons connaissance au moment où nous commençons cette étude (juillet 2001).

10.1 Les principes de Kerckhoffs

Avant de définir un nouveau système cryptographique, il convient d'examiner les critères habituellement utilisés en cryptologie pour évaluer un système de codage. Il ne suffit pas en effet de produire un texte crypté qui ait l'air d'un charabia sans relation évidente avec le texte clair pour pouvoir légitimement prétendre avoir réalisé un protocole cryptographique sûr. Le déchiffrement des codes secrets utilise des procédés en général plus astucieux que véritablement complexes, qui permettent de déchiffrer la plupart des codes « naïfs » sans avoir à essayer toutes les combinaisons possibles [159, 160, 161, 162, 163, 58].

Kerckhoffs a établi en 1883 [160] un certain nombre de conditions que doit remplir un système de cryptographie opérationnel. Ces critères sont toujours considérés en cryptologie comme des principes qui doivent présider à la conception de tout nouveau système cryptologique ; ils sont énumérés TAB. 10.1, dans leur formulation originale. Les trois premiers sont les plus souvent cités seuls aujourd'hui, alors que les trois derniers sont généralement considérés comme obsolètes¹. Ces six principes nous fourniront donc les règles à suivre dans l'étude d'un système de cryptographie quantique, ainsi que des critères nous permettant de comparer un tel système avec les protocoles algorithmiques couramment employés aujourd'hui.

¹C'est assez curieux quand on pense que Kerckhoffs disait lui-même : « *Tout le monde est d'accord pour admettre la raison d'être des trois derniers desiderata ; on ne l'est plus, lorsqu'il s'agit des trois premiers.* »

1. *Le système doit être matériellement, sinon mathématiquement, indéchiffrable ;*
2. *Il faut qu'il n'exige pas le secret, et qu'il puisse sans inconvénient tomber aux mains de l'ennemi ;*
3. *La clef doit pouvoir en être communiquée et retenue sans l'aide de notes écrites, et être changée ou modifiée au gré de ses correspondants ;*
4. *Il faut qu'il soit applicable à la correspondance télégraphique ;*
5. *Il faut qu'il soit portatif ; et que son maniement ou son fonctionnement n'exige pas le concours de plusieurs personnes ;*
6. *Enfin, il est nécessaire, vu les circonstances qui en commandent l'application, que le système soit d'un usage facile, ne demandant ni tension d'esprit, ni la connaissance d'une longue série de règles à observer.*

TAB. 10.1 – Les principes de Kerckhoffs [160, page 12]

10.1.1 Impossibilité du déchiffrement

Le premier critère exige de comparer le système à l'état de l'art du déchiffrement. Les systèmes informatiques de cryptologie couramment utilisés aujourd'hui sont de fait matériellement indéchiffrables mais ne sont pas à l'abris d'éventuels progrès brutaux de l'informatique.

Par exemple, il y a eu beaucoup de discussions cette année autour une proposition de Bernstein [166] de construire des machines massivement parallèles spécialisées qui permettraient de factoriser des nombres entiers trois fois plus longs que des ordinateurs habituels. Si la pertinence pratique, et même la justesse, de cette proposition est discutée [167], un tel système serait, selon certains, à la portée d'agences gouvernementales et pourrait permettre de déchiffrer des messages codés avec des clefs dont la longueur était jusque là considérée comme amplement suffisante. Bien sûr, se fonder sur cet article pour affirmer que l'algorithme de cryptographie le plus utilisé aujourd'hui, RSA, n'est plus sûr, serait grandement exagéré.

Cependant, c'est sur ce premier critère de Kerckhoffs que les systèmes de cryptographie quantique pourraient se révéler avantageux par rapport aux systèmes informatiques, car ils peuvent être « mathématiquement » indéchiffrables tout en respectant (presque) tous les autres critères.

10.1.2 Système connu de l'ennemi

Le second principe, peu intuitif, est l'un des plus importants.

En effet, il est difficile de garder un secret, et il serait illusoire d'imaginer qu'un système de cryptographie réellement utilisé ne sera pas analysé sous toutes ses coutures, sans parler des risques de fuites quand un grand nombre de personnes l'a entre les mains. Quantifier exactement ce risque est délicat, et supposer que le système peut être parfaitement connu de l'adversaire n'est en général pas si loin de la vérité qu'on pourrait le croire.

Le fait de supposer que le système soit connu de l'ennemi nous amène à concentrer tout le secret dans la clef. Cela présente plusieurs avantages. Le premier d'entre eux est que s'il est difficile de garder secrètes toutes les caractéristiques d'un système cryptographique, garder

la clef secrète est beaucoup plus simple, car le secret est beaucoup plus petit (voir section 10.1.3).

De plus, cette séparation claire entre ce qui est secret et ce qui est connu permet de faciliter l'analyse de la sécurité d'un protocole cryptographique. S'il n'est pas absolument nécessaire de communiquer les détails de l'algorithme à ses adversaires, notamment dans le cas des systèmes militaires pour des raisons plus stratégiques que cryptologiques, c'est devenu la norme dans les systèmes civils de cryptographie. En effet, les concepteurs d'un système cryptographique ne sont pas toujours les mieux placés pour en trouver les faiblesses, vu qu'ils l'ont *a priori* conçu en ayant déjà évité toutes les faiblesses qu'ils connaissaient. Les algorithmes civils qui sont en général considérés comme les plus sûrs sont ceux qui, publics, ont déjoué les efforts de nombreux cryptographes pendant de nombreuses années.

Ce principe, s'il est respecté, permet également de vendre un système cryptographique à quelqu'un qui ne nous fait pas *a priori* confiance. Ce client peut alors inspecter lui même le système, ou le faire inspecter par un expert de confiance, pour être sûr de ne pas être floué.

L'algorithme de cryptage A5/1 utilisé dans les téléphones GSM est un bon contre-exemple de ce principe : il était gardé secret, mais son principe de fonctionnement a assez vite été analysé et publié sur Internet [58]. Cet algorithme présentait certaines faiblesses qui auraient sans doute été évitées s'il avait été rendu public et soumis à la critique.

Nous devons supposer dans l'analyse de la sécurité de tout système de cryptographie quantique que l'espion, Ève, connaît parfaitement toutes les caractéristiques physiques des dispositifs des deux correspondants, Alice et Bob, à l'exception, bien sûr, de la partie bien définie qui constitue la clef.

10.1.3 La clef

Le troisième principe de Kerckhoffs est étroitement lié au précédent. En effet, si le deuxième principe traitait de la partie publique du système, le troisième traite de sa partie secrète.

Le fait que cette « *clef puisse être changée et modifiée au gré de ses correspondants* » permet de restreindre le secret nécessaire aux personnes concernées par la communication, et de limiter le risque de fuite.

La première partie de ce principe est devenue obsolète avec l'avènement des ordinateurs, car toute clef assez courte pour être « *retenue et communiquée sans l'aide de notes écrites* » peut être rapidement retrouvée par essais systématiques avec l'aide d'un ordinateur. Elle pose cependant le problème crucial de la distribution de la clef. En effet, plus la clef est longue, plus il sera difficile en pratique pour Alice et Bob de se mettre d'accord sur cette clef.

Par exemple, il existe un système de cryptographie qui répond à tous les critères de Kerckhoffs sauf celui-ci : il s'agit du *code de Vernam*, ou *one-time-pad*, qui utilise une clef aléatoire aussi longue que le message. Il s'agit alors de mettre en regard chaque lettre du message avec une lettre de la clef, et de les additionner. Ainsi, le message BONJOURX, codé avec la clef AVDOYGVQ donnera le texte crypté AXWRTWN. Ce système est absolument inviolable, car le message crypté AXWRTWN pourrait tout aussi bien correspondre au clair AUREVOIR, si la clef avait été ZJJIFCWM, ou à n'importe quel message clair de la même longueur. Cette sécurité, combinée à sa facilité d'emploi avait conduit ce code à être massivement utilisé pendant la seconde guerre mondiale et la guerre froide. La difficulté consiste alors à transmettre en sécurité des clefs aussi longues que les messages aux espions. Les époux Rosenberg ont ainsi été découverts au début des années 1950 car ils avaient utilisé plusieurs fois la même clef secrète dans un code de Vernam. La preuve de sécurité tombe en effet dès qu'on utilise la clef

plusieurs fois, et Shannon a montré [168] qu'il n'existait aucun algorithme mathématique absolument sûr fondé sur une clef dont l'entropie est inférieure à celle du message.

Les algorithmes à clef secrètes utilisés aujourd'hui utilisent des clefs plus courtes que le message. Si c'est insuffisant pour permettre un procédé inconditionnellement sûr, cela peut rendre les messages cryptés suffisamment difficiles à décrypter pour être hors de portée des systèmes informatiques actuels. Si la distribution des clefs est simplifiée par rapport au code de Vernam, elle doit toujours se faire par un canal sûr.

La cryptographie à clef publique a fourni un moyen de construire ce canal sûr, sans qu'il soit nécessaire de se fier à des coursiers. Il s'agit de protocoles mathématiques fondés sur le fait que la clef de décodage, secrète, est différente de la clef de codage, publique, qui permet en pratique à Alice et Bob de se mettre d'accord sur une clef à utiliser dans un protocole de cryptographie à clef secrète, tout en communiquant par un canal accessible à l'espion. Cependant, la sécurité du protocole le plus fréquemment utilisé, RSA, repose sur une conjecture mathématique non-démontrée (mais que les mathématiciens croient vraie), et n'est pas à l'abri de progrès mathématiques et techniques brutaux [166, 167, 53].

La cryptographie quantique se place dans le même créneau que ces systèmes à clef publique, la distribution de clef. En effet, la mécanique quantique nous permet de construire un canal sûr, l'espionnage étant empêché par le principe d'incertitude de Heisenberg, ce qui permettra de générer une clef utilisable pour un code de Vernam [67, 68].

C'est là l'avantage essentiel de la cryptographie quantique, qui permet de résoudre le problème de la distribution des clefs avec une sécurité absolue, qui ne repose que sur les lois de la mécanique quantique.

10.1.4 La « correspondance télégraphique »

Le quatrième principe se traduit aujourd'hui par la nécessité de produire des messages binaires, adapté aux systèmes de communication modernes. Si ce point est trivial pour les algorithmes informatiques et les protocoles de cryptographie quantique fondés sur des variables discrètes, il est beaucoup plus délicat pour les variables continues.

10.1.5 Système portable

Ce cinquième critère est évidemment important pour l'utilisation réelle d'un système de cryptographie quantique. Si la cryptographie quantique ne pourra jamais être compétitive sur ce point avec les systèmes informatiques, qui ne sont que des logiciels et s'installent aisément sur n'importe quel ordinateur, des progrès importants ont été faits ces dernières années dans le domaine de la cryptographie à photons uniques [68], et des prototypes de systèmes commerciaux.

Bien entendu, le système de démonstration que nous avons construit au cours de cette thèse [14] est loin d'obéir à ce critère, mais il ne s'agit que d'une expérience de faisabilité.

10.1.6 Ergonomie

Le dernier critère serait aujourd'hui désigné sous l'ergonomie du système. Si elle se réduit pour les systèmes algorithmiques à l'interface du logiciel, l'ergonomie d'un système de cryptographie quantique concerne aussi tous les réglages physiques indispensables à une

expérience de cryptographie. Ils doivent donc être automatisés autant que possible pour un véritable système opérationnel.

Ce critère d'ergonomie n'est bien entendu pas plus respecté que le précédent par le système que nous avons construit.

10.2 La cryptographie quantique

10.2.1 La distribution quantique de clefs

Un système de cryptographie quantique est un système qui exploite les lois de la mécanique quantique pour accomplir une tâche cryptographique. En général, mais pas toujours, cette tâche est la distribution d'une clef cryptographique, ce qui conduit souvent à désigner la *distribution quantique de clef* par le terme plus général de *cryptographie quantique*. Nous utilisons cet abus de langage dans cette thèse, la distribution quantique de clef étant le seul type de protocole de cryptographie quantique considéré ici.

Avant de définir un protocole de cryptographie, quantique ou non, les principes de Kerckhoffs nous disent qu'il faut clairement séparer ce qui est public, et supposé connu de l'espion, Ève, de ce qui est secret. Cette séparation est à peu près la même pour tous les protocoles de distribution quantique de clef, qu'ils soient fondés sur des variables discrètes [68] ou continues.

Dans ces protocoles, les deux protagonistes, Alice et Bob, disposent de deux canaux pour communiquer : un canal quantique, bien sûr, mais aussi un canal classique. Le canal quantique permet de transmettre des objets quantiques, c'est à dire des objets suffisamment petits pour que les effets quantiques soient mesurables. La nature précise de ces objets dépendra du protocole considéré, même s'il s'agira en pratique toujours d'impulsions lumineuses, transmises par des fibres optiques ou par un télescope. L'espion est censé avoir accès à ce canal, mais la nature quantique du canal limitera ses actions.

Le canal classique doit juste permettre à Alice et Bob de communiquer et peut être une ligne téléphonique ordinaire ou une fréquence radio. Ce canal est supposé écouté par Ève, qui écoute toutes les conversations entre Alice et Bob, mais ne peut pas le modifier. En d'autres termes, le canal doit être authentifié, ce qui est possible par des algorithmes de cryptographie classique, du moment qu'Alice et Bob partagent déjà une (relativement petite) clef secrète.

Dans un premier temps Alice et Bob s'échangent des objets quantiques par le canal quantique, ou mesurent des objets quantiques intriqués, ce qui est équivalent. Ces objets sont préparés de sorte que toute tentative pour Ève d'acquérir de l'information sur eux se traduirait, en vertu des lois de la mécanique quantique, par une perturbation du signal qu'Alice et Bob peuvent mesurer en comparant certaines de leurs mesures par le canal classique.

En pratique, un certain nombre de bruits expérimentaux seront indiscernables de l'action d'un espion et l'on devra toujours considérer qu'un espion est présent. Alice et Bob peuvent déduire des perturbations qu'ils ont mesuré la quantité d'information qu'Ève aurait pu extraire de ses mesures du canal quantique. Ils utilisent alors le canal classique et des algorithmes de réconciliation et d'amplification de confidentialité [70, 71, 72, 73, 74] pour pouvoir extraire une clef secrète commune, et rendre inutiles les mesures éventuelles de l'espion.

10.2.2 Cryptographie quantique avec des variables continues

Depuis le premier protocole de distribution quantique de clef élaboré en 1984 par Bennett et Brassard [66], les progrès ont été nombreux, tant sur les plans expérimentaux que théoriques [68]. Sur le plan théorique, il est notamment difficile ne serait-ce que de compter les différents protocoles à base d'états discrets de la littérature. Bennett a montré en 1992 [169] que la condition minimale pour faire un protocole de cryptographie quantique était de disposer de deux états non-orthogonaux, ce qui montrait que, au moins en principe, la cryptographie était possible avec des variables continues.

Un certain nombre de protocoles de cryptographie quantique utilisant les variables continues décrites dans la partie I de cette thèse ont été proposées depuis 1999 [170, 171, 172, 173, 174, 175, 176, 177, 178, 179, 180, 181, 133, 134, 11, 12, 13]. Ils nécessitent tous l'utilisation d'une détection homodyne analogue à celle que nous avons présenté dans la partie II. Nous présenterons ici rapidement ces protocoles, qui représentent la plupart des protocoles de cryptographie quantique avec des variables continues qui étaient publiés en juillet 2001, au moment où nous avons commencé à étudier les protocoles de cryptographie quantique présentés dans les chapitres suivants.

L'une des difficultés conceptuelles posée par l'utilisation des variables continues, liée au quatrième principe de Kerckhoffs, est le passage du continu au discret. De nombreux systèmes étaient ainsi fondés sur une modulation binaire des variables continues [170, 171, 172, 174, 175], ce qui n'était pas optimal et semblait exiger l'usage d'états comprimés. Ralph avait considéré l'usage d'états cohérents dans [171, 172], mais ils semblaient peu sûrs jusqu'à ce qu'il montre, dans [173]², qu'ils pouvaient être utilisés conjointement avec des algorithmes de réconciliation et d'amplification de confidentialité.

Ces protocoles étaient tous conçus comme des analogues de BB84, et cette analogie a été poussée le plus loin par Gottesmann et Preskill [176] qui ont trouvé le moyen d'encoder un qubit dans des variables continues [182], par une technique fondée sur la théorie des codes correcteurs d'erreurs quantiques. Ils en ont déduit un protocole fondé sur la modulation continue d'états comprimés, qui est, à notre connaissance, le seul protocole de cryptographie quantique avec des variables continues dont il existe une démonstration inconditionnelle de sécurité. Malheureusement, ce protocole a une portée limitée à des pertes inférieures à 1,6 dB et exige un facteur de compression d'au moins 2,5 dB.

D'autres protocoles extraient directement un bit des fluctuations gaussiennes de faisceaux gaussiens quantiquement corrélés [178, 179, 180, 177]. Il s'agit alors de déduire une variable binaire à partir de l'observation d'une variable aléatoire continue. Cette réduction directe d'une variable gaussienne à une variable binaire fait disparaître une partie de l'information disponible qui pourrait être utile pour évaluer l'espionnage d'Ève. Cependant, Silberhorn et ses collaborateurs ont récemment réussi à étendre ces protocoles à de fortes pertes et à des états cohérents [181].

Cerf et ses collaborateurs, en revanche, ont proposé un protocole [133, 134] fondé sur la théorie de l'information avec des variables continues, que nous avons présenté au chapitre précédent, où le bit n'est qu'une unité de mesure de la quantité d'information. Cette approche présente l'avantage de pouvoir séparer l'étude des performances théoriques optimales de ces protocoles de la production concrète d'une clef [183, 14, 15].

Ces protocoles exigeaient l'emploi d'états comprimés, et nous avons montré qu'ils pouvaient être généralisés aux états cohérents ([11, 12] et chapitre 11 de cette thèse), puis aux

²Cette publication est juste postérieure à notre article [12], mais indépendante.

fortes pertes ([13] et chapitre 12). Nous avons ensuite, en collaboration avec Gilles Van Assche et Nicolas Cerf, démontré expérimentalement la faisabilité de ces protocoles ([14, 15] et chapitre 13).

Chapitre 11

Protocoles directs

11.1 Introduction

Nous commencerons ce chapitre par un bref rappel du chapitre 7 sur les protocoles de communication avec des variables quantiques continues (section 11.2). Nous examinerons ensuite (section 11.3) les limites imposées à l'espionnage de ces protocoles par le théorème de non-clonage, et (sections 11.4 et 11.5) comment Alice et Bob peuvent se protéger de cet espionnage en symétrisant leur protocole. Nous calculerons, section 11.6, la quantité d'information secrète qu'Alice et Bob et examinerons le rôle que jouent l'amplitude des modulations d'une part (section 11.7), et le facteur de compression d'autre part (section 11.8). Nous récapitulerons d'une manière plus formelle section 11.9 les protocoles définis dans les parties précédentes ¹.

11.2 Transfert d'information avec des variables gaussiennes

Ces protocoles sont fondés sur des protocoles de communication similaires à ceux présentés au chapitre 7. Si on reprend directement ce protocole, Alice envoie à Bob des états comprimés de facteur de compression s ($s = 1$ pour des états cohérents) centrés en $\begin{bmatrix} Q_A \\ 0 \end{bmatrix}$. Bob mesure alors la quadrature Q_B de ce faisceau, qui a traversé un canal bruité. On a alors

$$Q_B = g_Q(Q + B_Q) = g_Q(Q_A + A_Q + B_Q), \quad (11.1)$$

¹Ce chapitre reprend le contenu des articles [12, 11].

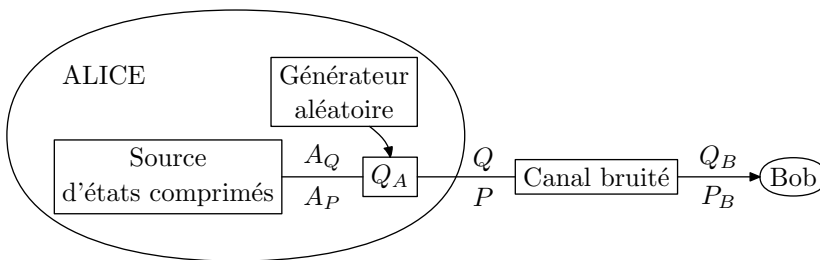


FIG. 11.1 – Système de communication quantique entre Alice et Bob

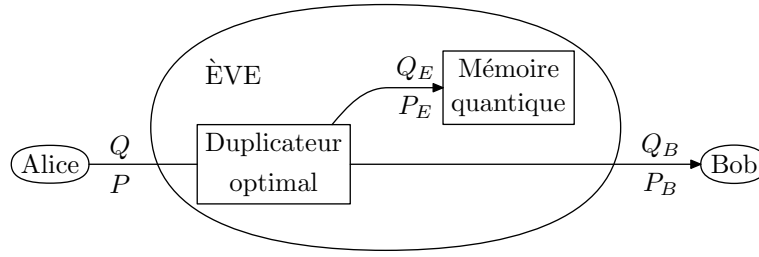


FIG. 11.2 – Stratégie possible d’espionnage

où A_Q et B_Q sont des bruits additifs gaussiens de variances respectives sN_0 et $\chi_{B,Q}N_0$. A_Q représente alors le bruit quantique du faisceau envoyé par Alice et B_Q le bruit équivalent à l’entrée du canal bruité.

Si Alice choisit de moduler Q_A avec une distribution aléatoire gaussienne de variance $V_A N_0$, le taux d’information mutuelle entre Alice et Bob est donné par la formule de Shannon (7.36)

$$I_{AB} = \frac{1}{2} \log_2 \frac{V_A + s + \chi_{B,Q}}{s + \chi_{B,Q}} = \frac{1}{2} \log_2 \frac{V + \chi_{B,Q}}{s + \chi_{B,Q}}, \quad (11.2)$$

où $V \equiv \langle Q^2 \rangle = V_A + s$ est la variance totale du faisceau en sortie de chez Alice, normée au bruit de photons N_0 .

11.3 L’espionnage d’Ève

Pour espionner la ligne, Ève doit faire une copie de ce qu’Alice a envoyé. Cette problématique est similaire à celle des cloneuses quantiques, que nous avons étudié au chapitre 8. Si nous nous limitons aux attaques individuelles, nous pouvons exploiter les résultats de ce chapitre, et en particulier les équations (8.3), en considérant que Charles et Ève sont une seule et même personne. Si on appelle $\left[\frac{Q}{P} \right]$ les quadratures du champ à la sortie de chez Alice, on a des équations quasiment identiques à (8.1) :

$$Q_B = g_Q(Q + B_Q) \quad Q_E = h_Q(Q + E_Q) \quad (11.3a)$$

$$P_B = g_P(P + B_P) \quad P_E = h_P(P + E_P). \quad (11.3b)$$

Les variances des bruits ajoutés chez Ève, normées à N_0 , seront notées $\chi_{E,Q}$ et $\chi_{E,P}$.

Alice et Bob ignorent ce que fait Ève, qui par définition, n’est pas coopérative. Ils doivent donc supposer son action la plus gênante possible au vu de ce qu’ils savent.

Le bruit ajouté chez Ève obéit aux relations de Heisenberg croisées (8.3) qui s’écrivent ici

$$\chi_{E,Q} \chi_{B,P} \geq 1 \quad \chi_{E,P} \chi_{B,Q} \geq 1. \quad (11.4)$$

Comme Alice et Bob peuvent évaluer $\chi_{Q,B}$ et $\chi_{P,B}$, ils peuvent déduire de ces équations des limites pour $\chi_{E,Q}$ et $\chi_{E,P}$. Il doivent supposer que ces limites sont saturées, car Ève peut les atteindre, si elle n’a pas de contrainte technologique. Elle peut en effet utiliser l’attaque schématisée FIG. 11.2, qui consiste à remplacer le canal bruité par un duplicateur optimal

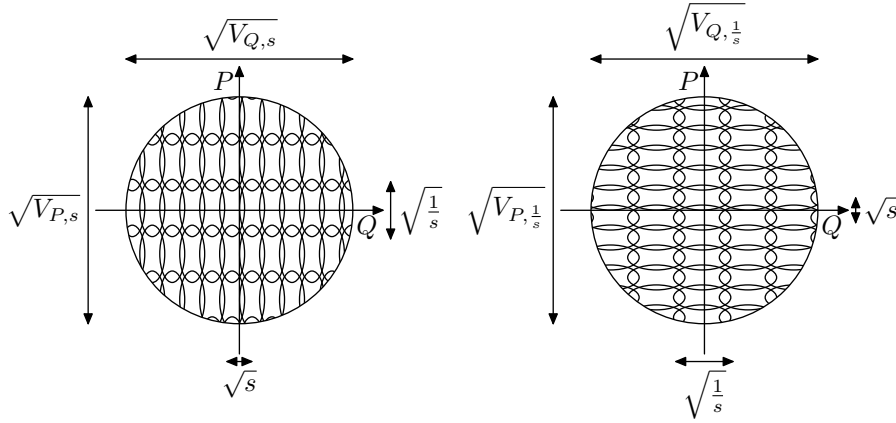


FIG. 11.3 – Schématisation des différentes variances de la modulation symétrisée d'Alice

induisant le même bruit chez Bob. Ève peut alors garder son clone dans une mémoire quantique pour effectuer ses mesures après avoir pris connaissance des communications classiques entre Alice et Bob.

Ces derniers doivent donc supposer

$$\chi_{E,Q} = \frac{1}{\chi_{B,P}} \quad \chi_{E,P} = \frac{1}{\chi_{B,Q}}. \quad (11.5)$$

Si les bruits ajoutés chez Ève sont gaussiens, restriction dans laquelle on se placera dans la suite, on peut utiliser la formule de Shannon (7.36) pour calculer l'information qu'Ève peut apprendre sur la modulation d'Alice.

$$I_{AE} = \frac{1}{2} \log_2 \frac{V + \chi_{E,Q}}{s + \chi_{E,Q}} = \frac{1}{2} \log_2 \frac{V + \frac{1}{\chi_{B,P}}}{s + \frac{1}{\chi_{B,Q}}} \quad (11.6)$$

11.4 Émission d'Alice

Si Alice et Bob utilisent le protocole défini plus haut, rien n'empêchera Ève de faire une mesure QND arbitrairement précise de Q avec une petite valeur de $\chi_{E,Q}$. En effet, elle peut dans ce cas mettre un fort bruit $\chi_{B,P}$ sur la quadrature P_B que Bob ne vérifie pas. Pour limiter l'intérêt pour Ève de faire une mesure QND, Alice et Bob peuvent symétriser leur protocole, en échangeant aléatoirement les rôles de Q et de P .

L'état envoyé par Alice sera donc comprimé la moitié du temps suivant Q et l'autre moitié du temps comprimé suivant P . L'amplitude des déplacements doit être telle que la matrice densité de l'état soit indépendante de la direction de la compression. Cette condition garantit l'impossibilité pour Ève de déduire la direction de compression des états, ce qui l'incite à choisir une attaque symétrique.

Pour des états gaussiens, déplacés avec une statistique gaussienne centrée en 0, cette condition revient à ce que la matrice de covariance de l'état envoyé par Alice soit indépen-

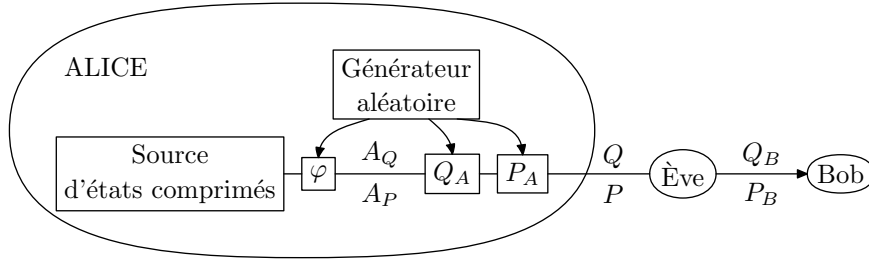


FIG. 11.4 – Alice

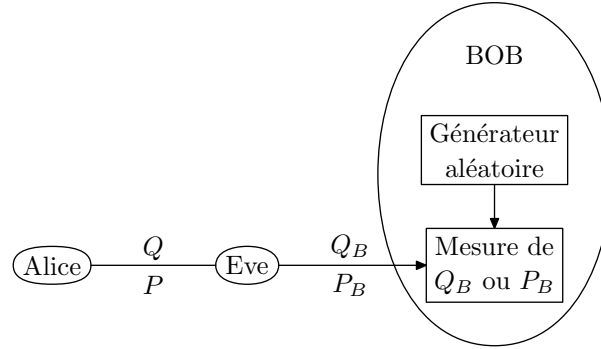


FIG. 11.5 – Bob

dante de la direction de la compression, ce qui se traduit par les égalités suivantes :

$$V_Q \equiv V_{Q,s} = V_{Q,\frac{1}{s}} \quad V_P \equiv V_{P,s} = V_{P,\frac{1}{s}} \quad (11.7a)$$

$$V_{A,Q,s} + s = V_{A,Q,\frac{1}{s}} + \frac{1}{s} \quad V_{A,P,s} + \frac{1}{s} = V_{A,P,\frac{1}{s}} + s. \quad (11.7b)$$

Pour alléger les notations, nous nous placerons dans le cas symétrique, où $V_Q = V_P \equiv V \geq \frac{1}{s}$ et où $\chi_{B,Q} = \chi_{B,P} \equiv \chi$.

La dernière supposition correspond aux canaux bruités les plus courants. Cependant, pour une implémentation réelle, il faudra mesurer séparément $\chi_{B,Q}$ et $\chi_{B,P}$, qui ne seront sans doute jamais complètement identiques.

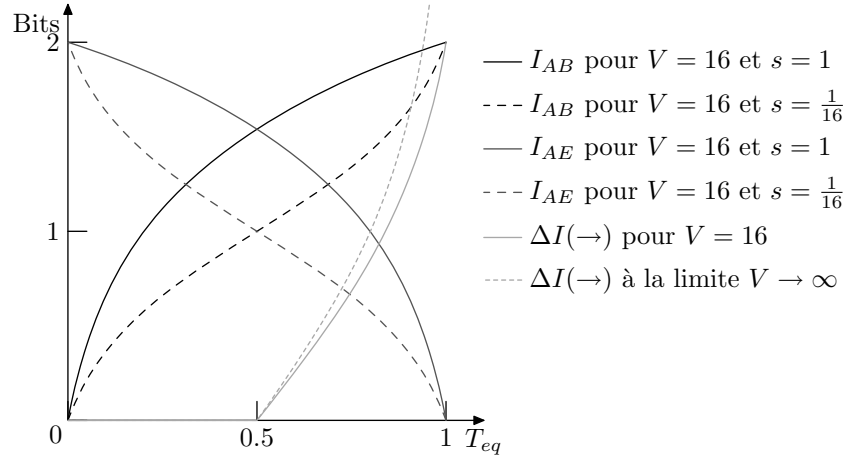
11.5 Mesures de Bob

Bob choisit aléatoirement de mesurer Q_B ou P_B . S'il a de la chance, il mesure selon la direction comprimée, la moins bruitée et

$$I_{AB,s} = \frac{1}{2} \log_2 \frac{V + \chi}{s + \chi}. \quad (11.8a)$$

S'il n'en a pas, il mesurera selon l'autre direction, la plus bruitée, et

$$I_{AB,\frac{1}{s}} = \frac{1}{2} \log_2 \frac{V + \chi}{\frac{1}{s} + \chi}. \quad (11.8b)$$

FIG. 11.6 – Évolution des informations mutuelles en fonction des pertes équivalentes T_{eq}

À V fixé, le facteur de compression maximum est $s = \frac{1}{V}$. Dans ce cas, Alice envoie un état comprimé déplacé orthogonalement à sa direction de compression seulement, et $I_{AB} = 0$ lorsque Bob mesure la mauvaise quadrature. C'est ce cas qui était considéré par Cerf et ses collaborateurs [133, 134]. Ce cas est également équivalent au cas où Alice prépare deux faisceaux EPR et mesure la position ou l'impulsion de l'un d'entre eux, comme nous l'avons montré section 4.2.4.

L'information acquise par Bob, qui se trompera une fois sur deux sera alors égale à la moyenne

$$I_{AB} = \frac{1}{2}I_{AB,s} + \frac{1}{2}I_{AB,\frac{1}{s}} = \frac{1}{2}\log_2 \frac{V + \chi}{\sqrt{1 + (s + \frac{1}{s})\chi + \chi^2}}. \quad (11.9)$$

Il est intéressant de noter que l'utilisation de faisceaux comprimés *diminue* l'information mutuelle entre Alice et Bob. Cette dégradation est due aux pertes d'information lorsque Bob mesure la mauvaise quadrature.

L'information qu'Alice avait éventuellement mise sur la quadrature que Bob n'a pas mesurée sera perdue, et ne servira pas à l'élaboration de la clef. Bob doit donc annoncer à Alice quelle quadrature il a mesuré.

Contrairement à Bob, Ève peut toujours mesurer la « bonne » quadrature. En effet, elle peut toujours enregistrer son clone dans une mémoire quantique, et ne le mesurer qu'après que Bob a révélé son choix de quadrature. Si Ève ne dispose pas de mémoire quantique, elle peut parfois avoir, sous les réserves évoquées dans la section 8.6, un clone amplifié avec un gain arbitrairement grand, ce qui lui permet de mesurer simultanément les deux quadratures.

Ève peut donc toujours mesurer la même quadrature que Bob, et on a alors

$$I_{AE} = \frac{1}{2}\log_2 \frac{V + \frac{1}{\chi}}{\sqrt{1 + (s + \frac{1}{s})\frac{1}{\chi} + \frac{1}{\chi^2}}} = \frac{1}{2}\log_2 \frac{V\chi + 1}{\sqrt{1 + (s + \frac{1}{s})\chi + \chi^2}}. \quad (11.10)$$

Il est intéressant de noter que la dépendance de cette information mutuelle au facteur de compression est exactement la même que celle de I_{AB} : elle est diminuée de la même quantité que celle-ci lorsqu'on comprime les faisceaux.

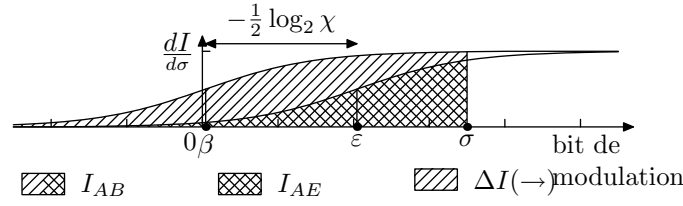


FIG. 11.7 – Contributions des différentes amplitudes de modulation aux différentes informations mutuelles

Comme on le verra ci-dessous, Alice et Bob peuvent générer une clef secrète tant que Bob en sait plus sur Alice qu'Ève, c'est-à-dire tant qu' $I_{AB} > I_{AE}$, ce qui correspond à $\chi < 1$, c'est-à-dire à des pertes équivalentes inférieures à $\frac{1}{2}$ ou 3 dB. Cette limitation de portée, assez intuitive, se retrouve dans beaucoup de protocoles de cryptographie quantique à variables continues. Elle sera discuté plus en détail section 12.1.1 .

11.6 Information secrète

Notons $S_{AB||E}$ la quantité d'information secrète qu'Alice et Bob peuvent extraire de leurs mesures $\left[\begin{smallmatrix} Q_A \\ P_A \end{smallmatrix} \right]$ et $\left[\begin{smallmatrix} Q_B \\ P_B \end{smallmatrix} \right]$ par communication publique. Il est possible [184, 71] de borner inférieurement $S_{AB||E}$ par la différence des informations mutuelles $\underline{\Delta I}$:

$$S_{AB||E} \geq \underline{\Delta I} \equiv I_{AB} - I_{AE}. \quad (11.11)$$

Cette borne quantifie l'intuition selon laquelle Bob peut toujours exploiter l'avantage qu'il a sur Ève en terme d'information. Il est possible d'approcher cette borne en pratique par des techniques informatiques dont nous parlerons plus en détail dans le chapitre 13.

Le calcul de $\underline{\Delta I}$ nous donne

$$\underline{\Delta I} = \frac{1}{2} \log_2 \frac{V + \chi}{\chi V + 1} \xrightarrow{V \rightarrow \infty} -\frac{1}{2} \log_2 \chi. \quad (11.12)$$

Cette valeur est indépendante du facteur de compression, et tend vers une limite finie pour les grandes modulations.

11.7 Rôle des petites et des grandes modulations

L'équation (11.12) tend vers une limite finie lorsque V augmente. Cela se comprend aisément en reprenant l'approche que nous avons développé dans la section 7.4. sur la FIG. 11.7, tracée pour des états cohérents, on peut constater que les bits de poids faibles sont essentiellement inconnus pour Bob et Ève, et n'apportent donc aucune contribution à $\underline{\Delta I}$. *A contrario*, les bits de poids forts, présents lors des grandes modulations sont parfaitement connus de Bob, mais presque aussi bien connus d'Ève et ne contribuent pas non plus à $\underline{\Delta I}$. $\underline{\Delta I}$ est donc essentiellement constitué des bits de poids moyens, compris entre environ 2 bits sous le niveau β du bruit chez Bob et 2 bits au dessus du niveau ε du bruit chez Ève, avec

$$\varepsilon \equiv \frac{1}{2} \log_2 \left(1 + \frac{1}{\chi} \right) \quad \beta \equiv \frac{1}{2} \log_2 (1 + \chi). \quad (11.13)$$

L'aire comprise entre les deux courbes correspond à la limite asymptotique de $\underline{\Delta I}$ pour les grandes modulations. Comme ces deux courbes, de hauteur 1, se déduisent l'une de l'autre par une translation horizontale de $\varepsilon - \beta$, l'aire entre les deux courbes se calcule aisément : elle vaut

$$1 \times (\varepsilon - \beta) = -\frac{1}{2} \log_2 \chi \quad (11.14)$$

Cette interprétation graphique nous permet aussi d'évaluer assez aisément la variance V au delà de laquelle $\underline{\Delta I}$ sera essentiellement constant. Si on pose

$$\sigma \equiv \frac{1}{2} \log_2 V, \quad (11.15)$$

cette limite correspond à $\sigma \simeq \varepsilon + 2$, c'est à dire $V \simeq 16(1 + \frac{1}{\chi})$.

On peut également calculer la variance au bout de laquelle $\underline{\Delta I}$ vaudra la moitié de sa valeur asymptotique. Cela correspond à $\sigma = \frac{1}{2}(\varepsilon + \beta)$, c'est à dire à $V = \sqrt{2 + \chi + \frac{1}{\chi}}$.

11.8 Rôle du facteur de compression

L'équation (11.12) semble rendre l'utilisation d'états comprimés inutiles. Cela n'est pas tout à fait exact et l'utilisation d'états comprimés peut être utile de plusieurs manières.

D'un point de vue pratique, on peut considérer que la différence entre $\underline{\Delta I}$ et I_{AB} représente la quantité d'information dont il faut se débarrasser par des protocoles informatiques. Or, cette différence est inférieure dans le cas de l'utilisation d'états comprimés, donc l'extraction de clef semble plus facile. Cet argument heuristique se trouve conforté par le fait que la plupart des protocoles à variables continues publiés au moment où nous avons commencé cette étude n'utilisaient pas les techniques d'amplification de confidentialité et exigeaient l'utilisation d'états comprimés. Cela ressort particulièrement dans les articles [171, 172], où Ralph envisage un protocole avec des états cohérents, qu'il rejette en raison de son trop fort taux d'erreurs. Dans une version ultérieure de cet article [173], rédigée indépendamment de nos articles [12, 11], il montre que ce protocole fonctionne s'il est complété par des protocoles de réconciliation et d'amplification de confidentialité.

Un autre rôle, plus trivial, qui peut être attribué au facteur de compression, c'est celui de la génération aléatoire de la clef, si Alice utilise des faisceaux EPR. Alice n'a plus à choisir un nombre aléatoire, avec toute la difficulté que représente la génération d'un « vrai » nombre aléatoire de qualité cryptographique [185], car celui-ci est directement généré par le bruit quantique.

D'un point de vue plus fondamental, si Bob disposait d'une mémoire quantique, il pourrait stocker l'impulsion envoyée par Alice. Celle-ci pourrait dire ensuite à Bob quelle observable mesurer. L'impulsion n'étant plus dans le domaine d'Ève, il est alors trop tard pour qu'elle puisse implémenter une mesure QND. L'utilisation d'états maximales comprimés ($s = \frac{1}{V}$) permettrait alors de doubler I_{AB} et I_{AE} , donc $\underline{\Delta I}$, comme représenté FIG. 11.8.

Même en l'absence de mémoire quantique, ce protocole peut être approché dans le cas de faibles pertes. Alice peut en effet utiliser une clef binaire secrète préalablement partagée avec Bob pour choisir la direction de compression de ses états. Bob utilisera alors cette clef et ne se trompera jamais de quadrature et peut ainsi aller jusqu'à doubler la quantité d'information

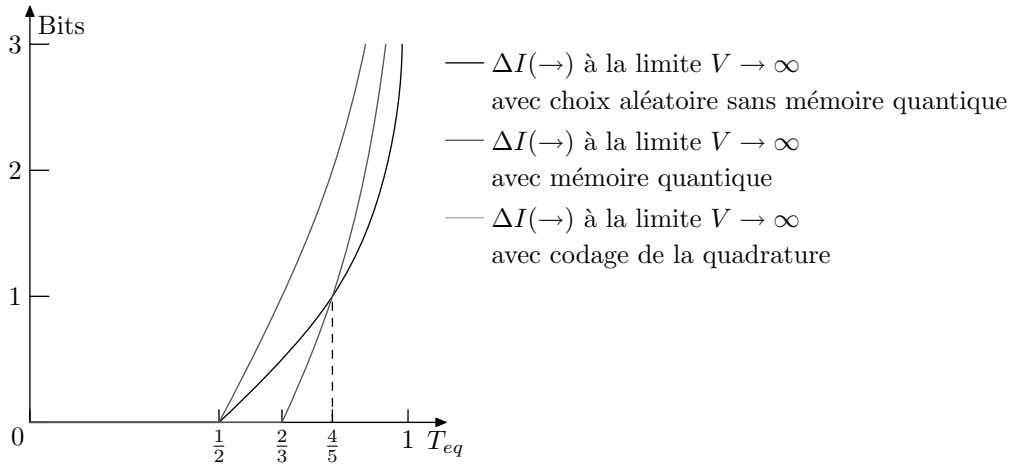


FIG. 11.8 – Taux d’information secrète que l’on peut atteindre avec différents protocoles directs (états maximalelement comprimés)

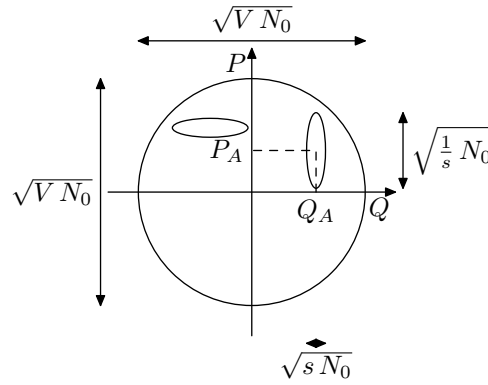


FIG. 11.9 – Schématisation de la modulation d’Alice dans le cas général

dont il dispose à la limite des états comprimés. Pour calculer le taux net de génération de clef, il ne faut pas oublier de soustraire le bit utilisé pour le choix de la base, et on a $2\Delta\vec{I} - 1$.

Ces protocoles ne fonctionnent, bien entendu, que si l’on consomme moins de bits secrets pour coder les quadratures que l’on en crée, c’est à dire si $2\Delta\vec{I} > 1$. À la limite des fortes modulations, et pour $\chi < \frac{1}{2}$, c’est-à-dire pour des pertes équivalentes inférieures à $\frac{1}{3}$ ou 1,8 dB et dépasse le protocole présenté plus haut pour $\chi < \frac{1}{4}$ c’est-à-dire pour des pertes équivalentes inférieures à $\frac{1}{5}$ ou 1,0 dB.

11.9 Protocoles

11.9.1 Cas général

Ce qui précède nous permet d’écrire explicitement une famille de protocoles de cryptographie quantique avec des variables continues :

1. Alice et Bob se mettent publiquement d’accord sur les différents paramètres du protocole à

utiliser, notamment le facteur de compression ($s \leq 1$) des états envoyés par Alice, ainsi que la variance ($V > \frac{1}{s}$) de modulation.

2. Alice et Bob répètent un grand nombre ($n + k$) de fois les opérations suivantes :
 - (a) Alice choisit aléatoirement d'envoyer des états comprimés suivant Q ou P . Nous appellerons X la quadrature choisie et Y l'autre quadrature.
 - (b) Alice choisit deux nombres réels aléatoires X_A et Y_A . X_A et Y_A sont choisis selon des lois gaussiennes de variances respectives $V - s$ et $V - \frac{1}{s}$.
 - (c) Alice envoie à Bob l'état comprimé de facteur de compressions (suivant la direction X), centré en $X = X_A$ et $Y = Y_A$.
 - (d) Bob reçoit cet état et choisit aléatoirement de mesurer Q ou P .
 - (e) Bob informe Alice de son choix de mesure par le canal classique public. Alice peut alors « jeter » la valeur de sa modulation selon l'autre quadrature, qui ne servira plus dans la suite.
 - (f) Alice informe Bob de la direction de compression X ou Y par le même canal.

À l'issue de ce processus, Alice et Bob partagent un grand nombre ($n + k$) de variables aléatoires gaussiennes corrélées, dont ils vont devoir extraire leur clef privée.

3. Alice et Bob choisissent aléatoirement k éléments de clef. Alice et Bob utilisent alors le canal classique pour comparer la modulation d'Alice et le résultat de la mesure de Bob. Ces comparaisons constituent une sorte de sondage, qui permet à Alice et Bob de déterminer les paramètres du canal ($g_Q, g_P, \chi_{B,Q}$ et $\chi_{B,P}$). Si k doit être grand devant 1 pour que la loi des grands nombres puisse garantir la validité du sondage, il n'est pas proportionnel au nombre total $n + k$ d'états envoyés et devient négligeable quand ce dernier est grand. Les paramètres $\chi_{B,Q}$ et $\chi_{B,P}$ permettent à Alice et Bob de déduire les paramètres $\chi_{E,Q}$ et $\chi_{E,P}$ de la meilleure attaque possible d'Ève.
4. Connaissant les paramètres du canal, ils effectuent alors un protocole de réconciliation sur les n autres éléments de clef pour en extraire une clef binaire commune.
5. Cette clef est partiellement connue d'Ève. Comme Alice et Bob connaissent $\chi_{E,Q}$ et $\chi_{E,P}$, ils peuvent quantifier la quantité d'information qui a fui vers Ève pendant la transmission et la réconciliation. Ils peuvent alors utiliser des algorithmes d'amplification de confidentialité pour extraire une clef binaire commune, plus courte que la précédente, mais inconnue d'Ève.

Dans ce cas la quantité d'information mutuelle entre Alice et Bob vaut

$$I_{AB} = \frac{1}{2} \log_2 \frac{V + \chi}{\sqrt{1 + (s + \frac{1}{s})\chi + \chi^2}}, \quad (11.16)$$

où l'on a effectué la moyenne entre les moments où Bob a mesuré la bonne quadrature et a acquis beaucoup d'information et ceux où il a mesuré la mauvaise quadrature et en a acquis moins. Celle dont dispose l'espion Ève vaut

$$I_{AE} = \frac{1}{2} \log_2 \frac{V\chi + 1}{\sqrt{1 + (s + \frac{1}{s})\chi + \chi^2}}, \quad (11.17)$$

et la quantité d'information secrète qu'Alice et Bob peuvent extraire par réconciliation directe vaut

$$\boxed{\Delta I = \frac{1}{2} \log_2 \frac{V + \chi}{V\chi + 1} \xrightarrow{V \rightarrow \infty} -\frac{1}{2} \log_2 \chi = \frac{1}{2} \log_2 \frac{T_{eq}}{1 - T_{eq}}.} \quad (11.18)$$

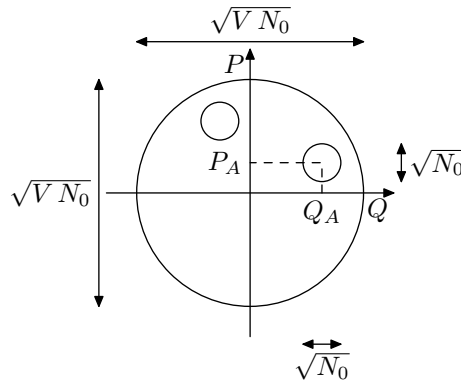


FIG. 11.10 – Schématisation de la modulation d'Alice pour un protocole à états cohérents

Cette expression est indépendante de la valeur du facteur de compression, et tend vers une limite finie pour les grandes modulations ($V \gg 1$). Elle est positive pour $\chi < 1$, c'est à dire pour des pertes équivalentes inférieures à 3 dB ($T_{eq} > \frac{1}{2}$).

11.9.2 États cohérents

Ce protocole est légèrement simplifié dans le cas où Alice et Bob utilisent des états cohérents ($s = 1$), en raison de la symétrie de ces états :

1. Alice et Bob se mettent publiquement d'accord sur les différents paramètres du protocole à utiliser, notamment la variance ($V > 1$) de modulation.
2. Alice et Bob répètent un grand nombre ($n + k$) de fois les opérations suivantes :
 - (a) Alice choisit deux nombres réels aléatoires Q_A et P_A selon des lois gaussiennes de variance $V - 1$.
 - (b) Alice envoie à Bob l'état cohérent centré en $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$
 - (c) Bob reçoit cet état et choisit aléatoirement de mesurer Q ou P .
 - (d) Bob informe Alice de son choix de mesure par le canal classique public. Alice peut alors « jeter » la valeur de sa modulation selon l'autre quadrature.

À l'issue de ce processus, Alice et Bob partagent un grand nombre ($n + k$) de variables aléatoires gaussiennes corrélées, dont ils vont devoir extraire leur clef privée.

3. Alice et Bob choisissent aléatoirement k éléments de clef. Alice et Bob utilisent alors le canal classique pour comparer la modulation d'Alice et le résultat de la mesure de Bob afin de déterminer $g_Q, g_P, \chi_{B,Q}$ et $\chi_{B,P}$ et d'en déduire $\chi_{E,Q}$ et $\chi_{E,P}$.
4. Ils effectuent alors un protocole de réconciliation sur les n autres éléments de clef pour en extraire une clef binaire commune.
5. Ils utilisent alors des algorithmes d'amplification de confidentialité pour extraire une clef binaire secrète commune.

La symétrie des états cohérents simplifie légèrement l'expression des informations mutuelles. En effet, contrairement au cas plus général des états comprimés, Bob n'a plus le choix entre une bonne et une mauvaise quadrature, et il acquiert la même quantité d'information quelque soit la quadrature choisie. On obtient alors, en posant $s = 1$ dans les équations

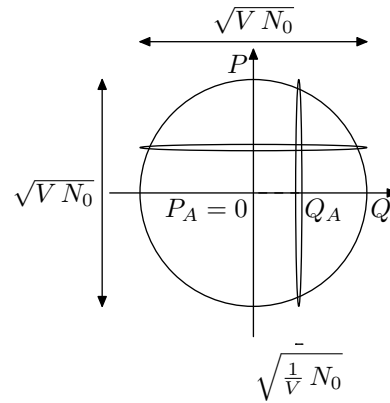


FIG. 11.11 – Schématisation des états maximalement comprimés envoyés à Bob dans le cas d'un protocole utilisant des faisceaux intriqués

précédentes,

$$I_{AB,\text{coh}} = \frac{1}{2} \log_2 \frac{V + \chi}{1 + \chi}, \quad (11.19)$$

et, pour Ève

$$I_{AE,\text{coh}} = \frac{1}{2} \log_2 \frac{V\chi + 1}{1 + \chi}. \quad (11.20)$$

Par contre, le taux d'information secrète ne change pas et vaut toujours

$$\boxed{\Delta I \xrightarrow{\rightarrow} = \frac{1}{2} \log_2 \frac{V + \chi}{V\chi + 1} \xrightarrow{V \rightarrow \infty} -\frac{1}{2} \log_2 \chi = \frac{1}{2} \log_2 \frac{T_{\text{eq}}}{1 - T_{\text{eq}}}.} \quad (11.21)$$

11.9.3 États intriqués

En pratique, les implémentations de ces protocoles avec des états comprimés seront sans doute réalisées avec des faisceaux intriqués. La différence essentielle avec le protocole décrit section 11.9.1 est le fait que la « modulation » aléatoire n'est pas choisie mais mesurée par Alice. On peut considérer après cette mesure que le faisceau envoyé à Bob est dans un état maximalement comprimé ($s = \frac{1}{V}$). On peut alors écrire le protocole comme suit :

1. Alice et Bob se mettent publiquement d'accord sur les différents paramètres du protocole à utiliser, notamment la variance ($V > 1$) de modulation.
2. Alice et Bob répètent un grand nombre ($2n + 2k$) de fois les opérations suivantes :
 - (a) Alice prépare deux faisceaux jumeaux de variance V .
 - (b) Alice envoie à Bob l'un des faisceaux et garde l'autre.
 - (c) Alice choisit aléatoirement la direction X (Q ou $-P$) de sa mesure
 - (d) Alice mesure X sur son faisceau. Le résultat de sa mesure sera un nombre réel distribué selon une loi gaussienne de variance V .
 - (e) Bob reçoit cet état et choisit aléatoirement de mesurer Q ou P sur son faisceau.
 - (f) Bob informe Alice de son choix de mesure par le canal classique public.

(g) Alice informe Bob de son choix de mesure par le même canal. Ils jettent alors leurs mesures s'ils n'ont pas choisi la même quadrature, et la conservent s'ils ont choisi la bonne.

À l'issue de ce processus, Alice et Bob partagent un grand nombre ($\sim n + k$) de variables aléatoires gaussiennes corrélées, dont ils vont devoir extraire leur clef privée.

3. Alice et Bob choisissent aléatoirement k éléments de clef. Alice et Bob utilisent alors le canal classique pour comparer la modulation d'Alice et le résultat de la mesure de Bob afin de déterminer $g_Q, g_P, \chi_{B,Q}$ et $\chi_{B,P}$ et d'en déduire $\chi_{E,Q}$ et $\chi_{E,P}$.
4. Ils effectuent alors un protocole de réconciliation sur les n autres éléments de clef pour en extraire une binaire commune.
5. Ils utilisent alors des algorithmes d'amplification de confidentialité pour extraire une clef binaire secrète commune.

Dans ce cas Alice et Bob ne disposent d'information mutuelle qu'une fois sur deux. On obtient alors, en posant $s = \frac{1}{V}$,

$$I_{AB,EPR} = \frac{1}{4} \log_2 \frac{V(V + \chi)}{1 + V\chi}. \quad (11.22)$$

et, pour Ève,

$$I_{AE,EPR} = \frac{1}{4} \log_2 \frac{V(1 + V\chi)}{V + \chi}. \quad (11.23)$$

La différence de ces deux taux d'information mutuelle est toujours inchangée, ce qui permet toujours de générer une clef secrète avec un débit donné par

$$\boxed{\Delta I = \frac{1}{2} \log_2 \frac{V + \chi}{V\chi + 1} \xrightarrow{V \rightarrow \infty} -\frac{1}{2} \log_2 \chi = \frac{1}{2} \log_2 \frac{T_{eq}}{1 - T_{eq}}.} \quad (11.24)$$

Ce débit est nul pour une transmission équivalente T_{eq} inférieure à $\frac{1}{2}$. Nous allons étudier, dans le chapitre suivant, une variante de ce protocole qui n'a pas cette limitation, mais fonctionne pour des grandes pertes.

Chapitre 12

Protocoles inverses

Les protocoles présentés au chapitre précédent présentent l'inconvénient d'être limités à des pertes inférieures à 3 dB, comme tous les protocoles à variables continues publiés avant avril 2002. Nous présenterons section 12.1.1 les raisons de l'omniprésence de cette limite et le principe de la *réconciliation inverse*, qui permet de la franchir par une variante des protocoles précédents. Nous présenterons ensuite (section 12.2) les attaques d'Ève dans le cas de ces protocoles inverses. Nous verrons, section 12.3, les conditions dans lesquels ces protocoles inverses fonctionnent, et nous calculerons (section 12.4) la quantité d'information secrète qu'ils permettent de générer.¹

12.1 Comment franchir la limite des 3 dB

12.1.1 La limite des 3 dB

Comme nous l'avons déjà mentionné plus haut, la totalité des protocoles de cryptographie quantique avec des variables continues publiés avant avril 2002 étaient limités à des pertes équivalentes de 3 dB. La question se posait de savoir si cette limite était fondamentale.

Cette limite peut en effet être justifiée par un argument assez simple : si les pertes sont supérieures à 50 %, Ève peut remplacer la ligne de transmission par une lame partiellement réfléchissante. Comme elle récupère plus de la moitié du faisceau, elle en a une plus grande partie que Bob ; elle peut donc en extraire plus d'information que lui. L'étude des cloneuses quantiques de variables continues permet d'étendre ce raisonnement aux pertes équivalentes.

Malgré ce raisonnement simple, la question qui se posait n'était pas tant la nature fondamentale de cette limitation, que les moyens à mettre en oeuvre pour franchir cette limite. En effet, deux faits montraient que ce raisonnement était insuffisant, concernant d'une part l'étude du degré d'intrication des variables continues, et la résistance aux pertes de BB84 d'autre part.

Des résultats théoriques montraient en effet que l'intrication persistait malgré des pertes arbitraires [116, 117]. Duan et ses collaborateurs ont même proposé un protocole de purification d'intrication [157] qui extrait des paires de qubits dans des états maximalement intriqués à partir de faisceaux intriqués gaussiens ayant subi des pertes. Ces qubits peuvent

¹Les résultats présentés dans ce chapitre ont été publiés dans [13]

alors être utilisés pour des protocoles de cryptographie quantique discrets. Malheureusement, ces protocoles de purifications sont très complexes et paraissent peu réalistes expérimentalement.

Si le résultat précédent montrait que « quelque chose » de quantique passait avec les variables continues pour des pertes arbitraires, ce quelque chose étant inaccessible à Ève et exploitable pour la cryptographie, il ne fournissait pas en lui-même un moyen réaliste de construire un protocole robuste aux pertes de cryptographie quantique avec des variables continues. Cependant, par sa généralité même, le raisonnement mentionné plus haut montrait ses limites. En effet, l'argument est indépendant de la nature des états envoyés sur le faisceau : si Ève dispose de plus de la moitié du faisceau, elle dispose de la majeure partie de l'information, que celle-ci soit codée sous la forme de variables continues ou discrètes. Le même argument, s'il était valable, devrait également limiter le fonctionnement du premier protocole de cryptographie quantique, BB84, à 3 dB de pertes.

Pourtant, BB84, avec des détecteurs parfaits, si la seule source d'erreur est constituée par les pertes, a une portée infinie. Si les pertes diminuent son débit, elle n'entravent en rien sa sécurité. En effet, si les photons perdus sont récupérés par Ève, qui peut alors apprendre la polarisation qu'Alice leur avait donné, Bob ne les reçoit jamais et ils ne servent pas à l'élaboration de la clef. Si Ève apprend effectivement plus d'information que Bob, celle-ci lui est donc totalement inutile.

L'analyse précédente n'était pas pertinente en ce qu'elle négligeait le retour d'information de Bob, qui annonce à Alice les moments où il n'a pas reçu de photons, ce qui permet à Alice et Bob de ne sélectionner après coup que les moments où le photon est passé et n'a pas été intercepté par Ève. Le rôle de cette post sélection peut être interprété de deux manières différentes, qui conduisent à deux adaptations différentes aux variables continues.

On peut, comme Silberhorn, Ralph, Lütkenhaus et Leuchs [181], attribuer un rôle particulier à l'acte de sélectionner après coup les événements rares où Ève n'a pas pu apprendre beaucoup d'information. Dans le cas de BB84 avec de fortes pertes, cela correspond aux quelques photons qui traversent la ligne de bout en bout et qu'Ève n'a pas interceptés ; dans le cas du protocole à variables continues proposé dans [181], cela correspond aux fluctuations exceptionnelles de variables continues, ce qui permet de calculer numériquement un taux d'information mutuelle secrète entre Alice et Bob qui reste positif au delà de 3 dB de pertes.

12.1.2 Réconciliation inverse

Indépendamment et simultanément, nous avons étudié une autre approche, qui insiste plus sur le flux d'information venant de Bob, à rebours de la transmission physique des objets quantiques, comme une rétroaction. En effet, la formule (11.11) repose sur des considérations purement mathématiques, indépendante de la direction de transmission physique entre Alice et Bob. Elle ne dépend que des corrélations entre les mesures d'Alice, de Bob et d'Ève, et, si on est tout à fait libre d'invertir les rôles de Bob et d'Alice, le taux d'information secrète entre Alice et Bob en présence d'Ève est indépendant de cette interversion. Son expression, donnée dans [71], est donc bornée par

$$S_{AB||E} \geq \max\{I_{AB} - I_{AE}, I_{AB} - I_{BE}\} \equiv \max\{\underline{\Delta I}; \overline{\Delta I}\}. \quad (12.1)$$

Le premier terme, $\underline{\Delta I}$, correspond aux protocoles où Alice envoie à Bob suffisamment d'information pour qu'il corrige ses erreurs. Comme la clef alors construite est exclusive-

ment fonction de la modulation d'Alice, et que tous les flux d'information vont dans le même sens, nous appellerons *réconciliation directe* ces protocoles. Le second terme, ΔI correspond au cas où c'est Bob qui envoie les indications à Alice, pour que celle-ci puisse deviner ses mesures, et reproduire en quelque sorte les bruits survenus au cours de la transmission. La clef construite est alors fonction des résultats des mesures de Bob, et le flux d'information au cours de la réconciliation se propage en sens inverse de la transmission physique des objets quantiques. Nous parlerons donc de protocoles de *réconciliation inverse*.

Pour attaquer ces protocoles inverses, Ève devra essayer de deviner le résultat de la mesure de Bob, en aval de son dispositif, plus que les caractéristiques du faisceau préparé par Alice. Même lorsqu'Ève intercepte plus de la moitié du faisceau, il est difficile de comparer intuitivement l'information dont elle dispose sur Bob à celle dont Alice dispose.

12.1.3 Interprétation de BB84 en terme de protocoles inverses

BB84 peut être interprété en termes de réconciliation inverse, ce qui nous permettra de mettre en évidence les raisons pour lesquelles le raisonnement tenu plus haut sur les pertes ne s'applique pas à BB84. En effet, une fois le choix de base rendu publique, un canal de transmission η peut être décrit par le tableau suivant :

Alice envoie 0	$\Rightarrow$	Bob reçoit	$\begin{cases} 0 & \text{avec la probabilité } \eta \\ \otimes & \text{avec la probabilité } 1 - \eta \end{cases}$
Alice envoie 1	$\Rightarrow$	Bob reçoit	$\begin{cases} 1 & \text{avec la probabilité } \eta \\ \otimes & \text{avec la probabilité } 1 - \eta \end{cases}$

où $\otimes$ désigne l'absence de photons. Tous les photons non reçus par Bob doivent être considérés comme capturés, donc connus, par Ève.

Par exemple, on peut avoir la situation suivante pour un canal à de transmission 0,25 (6 dB de pertes) :

Alice envoie	1	1	0	0	1	1	1	1	0	1	1	0	1	0	1	0
Bob reçoit	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	1	$\otimes$	$\otimes$	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	$\otimes$	1	$\otimes$
Ève sait d'Alice	1	1	0		1		1	1	0	1	1		1	0		0
Ève sait de Bob	$\otimes$	$\otimes$	$\otimes$		$\otimes$		$\otimes$	$\otimes$	$\otimes$	$\otimes$	$\otimes$		$\otimes$	$\otimes$		$\otimes$

Lorsque Bob annonce à Alice les photons qu'il a reçus, il ne fait que transmettre une information déjà connue d'Ève, et ne lui apprend rien de plus. Après cette étape de réconciliation inverse, Alice connaît parfaitement la clef de Bob, et on a

Clef d'Alice	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	1	$\otimes$	$\otimes$	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	$\otimes$	1	$\otimes$
Clef de Bob	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	1	$\otimes$	$\otimes$	$\otimes$	$\otimes$	$\otimes$	0	$\otimes$	$\otimes$	1	$\otimes$
Ève connaît	$\otimes$	$\otimes$	$\otimes$		$\otimes$		$\otimes$	$\otimes$	$\otimes$	$\otimes$	$\otimes$		$\otimes$	$\otimes$		$\otimes$

Ensuite, Alice et Bob éliminent de leur clé commune tous les $\otimes$, déjà connus d'Ève. Après cette étape, qui peut être considérée comme une amplification de confidentialité triviale, Alice et Bob partagent une clé secrète, en l'occurrence 0 1 0 1 .

On peut comparer ce protocole à celui que donnerait un protocole de réconciliation direct appliqué à BB84. Dans ce cas, Alice doit donner à Bob suffisamment d'information pour qu'il puisse remplacer ses $\otimes$ par le 0 ou le 1 qui figurait à sa place dans la clef initiale d'Alice.

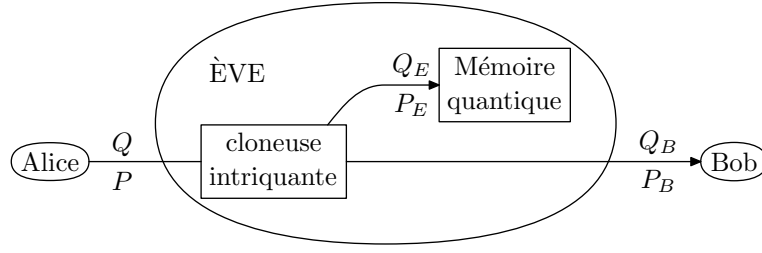


FIG. 12.1 – Espionnage par une cloneuse intriquante

Cette information peut par exemple prendre la forme de parités de sous blocs de la clef. Si ce protocole marche pour de faibles pertes, il n'est pas difficile de se convaincre que, si Alice ignore la position des $\otimes$, elle est obligée de transmettre suffisamment d'information par le canal classique pour permettre une reconstitution intégrale de la clef par Ève, si les pertes sont supérieures à 3 dB.

12.2 Les cloneuses intriquantes

Pour appliquer ces idées aux variables continues, nous supposons qu'Alice et Bob utilisent les mêmes protocoles physiques qu'au chapitre précédent, Alice envoyant des états comprimés centrés en $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$, dont la direction de compression choisie aléatoirement et Bob mesurant aléatoirement la quadrature Q_B ou P_B . Ils ne se comporteront différemment qu'au moment d'effectuer le protocole de réconciliation, où la clef sera fondée sur les mesures de Bob et non les modulations d'Alice.

Avant de calculer le taux d'information secrète ΔI qu'Alice et Bob peuvent obtenir par réconciliation inverse, il convient d'étudier les stratégies d'attaque d'Ève, qui sont différentes par rapport au protocole de réconciliation directe. Elles exigent l'utilisation de *cloneuses intriquantes*, que nous définirons section 12.2.1. Nous caractériserons ces cloneuses intriquantes, dont les sorties obéissent à une inégalité de Heisenberg (section 12.2.2), par les variances conditionnelles d'Alice (section 12.2.3) et d'Ève (section 12.2.4). Nous montrerons section 12.2.5 qu'il est possible de réaliser une telle cloneuse intriquante optimale.

12.2.1 Définition

En effet, si Alice et Bob utilisent un protocole de réconciliation inverse, Ève devra utiliser un dispositif qui lui permettra de deviner le résultat de la mesure de Bob. Nous appellerons *cloneuse intriquante* un tel dispositif. C'est en effet une cloneuse dont les sorties sont corrélées entre elles, de préférence quantiquement.

Comme pour les cloneuses utilisées pour attaquer les protocoles directs, nous noterons $\begin{bmatrix} Q \\ P \end{bmatrix}$ les quadratures du champ à l'entrée de la cloneuse intriquante, $\begin{bmatrix} Q_E \\ P_E \end{bmatrix}$ celles du champ qu'Ève stocke dans sa mémoire quantique et $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ celles du champ qu'elle renvoie à Bob.

Comme précédemment, le canal d'Alice à Bob peut être décrit comme un canal bruité dont ils peuvent mesurer les paramètres par sondage :

$$Q_B = g_Q(Q + B_Q) \qquad P_B = g_P(P + B_P). \qquad (12.2)$$

Nous utiliserons les mêmes notations que dans le chapitre précédent :

$$\langle Q^2 \rangle = \langle P^2 \rangle \equiv V N_0 \quad V \geq 1 \quad (12.3a)$$

$$\langle B_Q^2 \rangle \equiv \chi_Q N_0 \quad \langle B_P^2 \rangle \equiv \chi_P N_0 \quad (12.3b)$$

$$\langle Q B_Q \rangle = \langle P B_P \rangle = 0 \quad (12.3c)$$

12.2.2 Inégalité de Heisenberg sur les variances conditionnelles

Pour un protocole de réconciliation inverse, Alice a besoin d'estimer Q_B . On peut noter αQ_A son estimateur, avec $\alpha \in \mathbb{R}$. L'estimateur d'Ève pour P_B sera noté de même εP_E , avec $\varepsilon \in \mathbb{R}$. Leurs erreurs respectives seront donc

$$Q_{B|A,\alpha} \equiv Q_B - \alpha Q_A \quad (12.4a)$$

$$P_{B|E,\varepsilon} \equiv P_B - \varepsilon P_E. \quad (12.4b)$$

Ce sont les variances de ces quantités qui quantifient la précision des mesures, Comme toujours, un calcul de commutateur nous renseignera sur ces variances :

$$[Q_{B|A,\alpha}, P_{B|E,\varepsilon}] = \underbrace{[Q_B, P_B]}_{2iN_0} - \alpha \underbrace{[Q_A, P_B]}_0 - \varepsilon \underbrace{[Q_B, P_E]}_0 - \alpha \varepsilon \underbrace{[Q_A, P_E]}_0 \quad (12.5a)$$

$$= 2iN_0. \quad (12.5b)$$

En effet, $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$, $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ et $\begin{bmatrix} Q_E \\ P_E \end{bmatrix}$ sont des modes distincts et commutent entre eux. Cette relation de commutation se traduit aisément en terme de variances par l'inégalité de Heisenberg

$$\langle Q_{B|A,\alpha}^2 \rangle \langle P_{B|E,\varepsilon}^2 \rangle \geq N_0^2. \quad (12.6)$$

Alice et Ève sont bien entendu libres d'optimiser les paramètres réels α et ε pour avoir la meilleure estimation possible des résultats de Bob. L'information dont ils disposent est donc caractérisée par les variances conditionnelles $V_{Q_B|Q_A}$ et $V_{P_B|P_E}$ définies par les relations

$$V_{Q_B|Q_A} \equiv \min_{\alpha \in \mathbb{R}} \left\{ \langle Q_{B|A,\alpha}^2 \rangle \right\} \quad V_{P_B|P_E} \equiv \min_{\varepsilon \in \mathbb{R}} \left\{ \langle P_{B|E,\varepsilon}^2 \rangle \right\}. \quad (12.7)$$

L'inégalité (12.6) s'applique aux variances conditionnelles et devient

$$V_{Q_B|Q_A} V_{P_B|P_E} \geq N_0^2 \quad V_{P_B|P_A} V_{Q_B|Q_E} \geq N_0^2. \quad (12.8)$$

La deuxième inégalité s'obtient en effectuant le même raisonnement, mais en intervertissant les rôles d'Alice et d'Ève. Ces inégalités de Heisenberg croisées nous disent encore une fois qu'Alice et Ève ne peuvent pas connaître les quadratures du faisceau de Bob mieux que ne l'autorise le principe d'incertitude de Heisenberg, même si elles se mettent d'accord pour mesurer deux quadratures différentes.

Ces inégalités limitent la connaissance d'Ève et peuvent s'écrire sous la forme

$$V_{P_B|P_E} \geq \frac{N_0^2}{V_{Q_B|Q_A}} \quad V_{Q_B|Q_E} \geq \frac{N_0^2}{V_{P_B|P_A}}. \quad (12.9)$$

12.2.3 Variance conditionnelle d'Alice

Si Alice crée le champ $\begin{bmatrix} Q \\ P \end{bmatrix}$, on peut écrire

$$Q \equiv Q_A + A_Q \quad P \equiv P_A + A_P \quad (12.10)$$

où $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$ est proportionnel à la meilleure estimation qu'Alice a de $\begin{bmatrix} Q \\ P \end{bmatrix}$ et $\begin{bmatrix} A_Q \\ A_P \end{bmatrix}$ est le bruit quantique, non-corrélé à $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$. On a

$$\langle A_Q^2 \rangle = sN_0 \quad \text{ou} \quad \langle A_P^2 \rangle = sN_0, \quad (12.11)$$

où s représente le facteur de compression utilisé par Alice pour générer ce champ. Comme Alice change aléatoirement la quadrature comprimée de compression de son faisceau, Ève et Bob ne peuvent pas savoir dans quelle direction elle a comprimé son faisceau. Si on ne garde que les cas où Bob a mesuré la bonne quadrature, tout se passe comme si le faisceau était comprimé sur les deux quadratures. On traiterait de la même manière les cas où Bob a mesuré la mauvaise quadrature, considérant que le faisceau est anticomprimé sur les deux quadratures, en remplaçant s par $\frac{1}{s}$. On a nécessairement

$$s \geq \frac{1}{V}. \quad (12.12)$$

La matrice de covariance entre Alice et Bob, restreinte aux cas où ils ont choisi la même quadrature, est donc définie par

$$\langle Q_A^2 \rangle = (V - s)N_0 \quad \langle P_A^2 \rangle = (V - s)N_0 \quad (12.13a)$$

$$\langle Q_B^2 \rangle = G_Q(V + \chi_Q)N_0 \quad \langle P_B^2 \rangle = G_P(V + \chi_P)N_0 \quad (12.13b)$$

$$\langle Q_A Q_B \rangle = g_Q \langle Q_A^2 \rangle \quad \langle P_A P_B \rangle = g_P \langle P_A^2 \rangle, \quad (12.13c)$$

où on a utilisé les notations usuelles $G_Q \equiv g_Q^2$ et $G_P \equiv g_P^2$. Nous pouvons alors calculer les variances conditionnelles d'Alice sur les mesures de Bob :

$$V_{Q_B|Q_A} = \langle Q_B^2 \rangle - \frac{\langle Q_A Q_B \rangle^2}{\langle Q_A^2 \rangle} \quad (12.14a)$$

$$= G_Q V N_0 + G_Q \chi_Q N_0 - G_Q V N_0 + G_Q s N_0 \quad (12.14b)$$

$$= G_Q(\chi_Q + s)N_0 \quad (12.14c)$$

$$V_{P_B|P_A} = G_P(\chi_P + s)N_0 \quad (12.14d)$$

Dans le cas du canal symétrique, ces équations deviennent

$$V_{B|A} = G(\chi + s)N_0 \quad (12.14e)$$

Ces équations, combinées à la contrainte (12.12) nous permettent de borner les variances conditionnelles d'Alice par

$$V_{Q_B|Q_A} \geq V_{Q_B|Q_A, \min} = G_Q(\chi_Q + \frac{1}{V})N_0 \quad (12.15a)$$

$$V_{P_B|P_A} \geq V_{P_B|P_A, \min} = G_P(\chi_P + \frac{1}{V})N_0, \quad (12.15b)$$

et, dans le cas du canal symétrique,

$$V_{B|A} \geq V_{B|A, \min} = G(\chi + \frac{1}{V})N_0. \quad (12.15c)$$

12.2.4 Variance conditionnelle d'Ève

Pour calculer la variance conditionnelle d'Ève, on serait tenté d'injecter la valeur obtenue en (12.14) dans les relations de Heisenberg croisées (12.9). L'inégalité obtenue serait bien vérifiée, mais ce mode de calcul ne donne en général pas la limite optimale.

En effet, une cloneuse intriquante est un dispositif physique, et la matrice densité de ses sorties $\begin{bmatrix} Q_B \\ P_B \end{bmatrix}$ et $\begin{bmatrix} Q_E \\ P_E \end{bmatrix}$ ne dépendent que de la matrice densité des états fournis en entrée $\begin{bmatrix} Q \\ P \end{bmatrix}$, et pas de la manière dont ils ont été préparés. En particulier, les corrélations sortie-sortie, décrites en l'occurrence par $V_{Q_B|Q_A}$ et $V_{P_B|P_A}$, ne dépendent que de la variance V de Q et de P et pas du facteur de compression utilisé pour générer ces états. Au fond, il importe peu qu'Alice utilise réellement des états comprimés ; ce qui importe, c'est qu'elle *aurait pu* le faire sans qu'Ève et Bob n'aient aucun moyen de s'en rendre compte.

Ainsi, les inégalités de Heisenberg croisées (12.9) doivent être vérifiées pour toutes les valeurs physiquement accessibles de $V_{Q_B|Q_A}$ et de $V_{P_B|P_A}$. Il faut donc prendre les valeurs les plus contraignantes, données par 12.15. Nous pouvons donc écrire

$$V_{P_B|P_E} \geq \frac{N_0^2}{V_{Q_B|Q_A, \min}} \quad V_{Q_B|Q_E} \geq \frac{N_0^2}{V_{P_B|P_A, \min}}, \quad (12.16)$$

ce qui nous donne les expressions suivantes pour les variances conditionnelles

$$V_{P_B|P_E} \geq V_{P_B|P_E, \min} = \frac{N_0}{G_P(\chi_P + \frac{1}{V})} \quad (12.17a)$$

$$V_{Q_B|Q_E} \geq V_{Q_B|Q_E, \min} = \frac{N_0}{G_Q(\chi_Q + \frac{1}{V})}. \quad (12.17b)$$

Dans un système pratique de cryptographie, Alice et Bob donneront le même rôle à Q et P . On peut donc raisonnablement se limiter à l'étude du cas symétrique où $G_Q = G_P = G$ et $\chi_Q = \chi_P = \chi$. Les deux équations précédentes se réduisent alors à une seule :

$$V_{B|E} \geq V_{B|E, \min} = \frac{N_0}{G(\chi + \frac{1}{V})}. \quad (12.17c)$$

12.2.5 Implémentation

La limite (12.17) peut être saturée, comme nous le montrerons ici en construisant une cloneuse intriquante optimale. Notre construction ne concerne que le cas $G < 1$, symétrique en $\begin{bmatrix} Q \\ P \end{bmatrix}$, mais les mêmes principes permettent de construire des cloneuses intriquantes optimales plus générales.

Pour construire une telle cloneuse intriquante, schématisée FIG. 12.2, Ève utilise une lame partiellement réfléchissante de transmission G pour récupérer une fraction du faisceau transmis d'Alice à Bob. Elle injecte avec cette lame semi-réfléchissante le champ $\begin{bmatrix} Q_{E1} \\ P_{E1} \end{bmatrix}$, qui a la variance adéquate pour induire un bruit de variance totale $G\chi N_0$ chez Bob. On a donc

$$\langle Q_{E1}^2 \rangle = \langle Q_{E2}^2 \rangle = \frac{G\chi}{1-G} N_0. \quad (12.18)$$

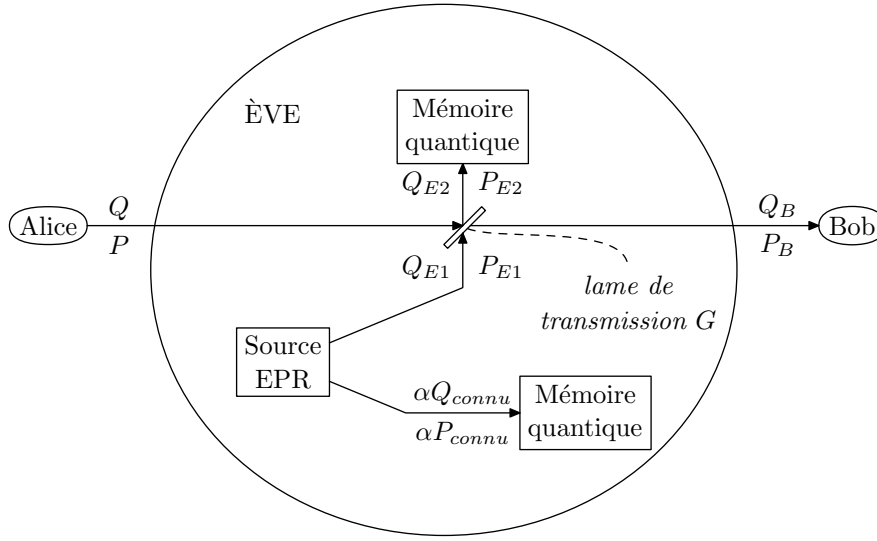


FIG. 12.2 – Implémentation d'une cloneuse intriquante

Pour connaître le mieux possible la modulation de Bob, Ève cherche à connaître le bruit $\begin{bmatrix} Q_{E1} \\ P_{E1} \end{bmatrix}$ qu'elle injecte. Le mieux qu'elle puisse faire est de préparer deux faisceaux jumeaux EPR, injecter l'un d'eux dans $\begin{bmatrix} Q_{E1} \\ P_{E1} \end{bmatrix}$ et stocker l'autre dans une mémoire quantique, jusqu'à ce qu'elle sache quelle quadrature de Bob elle veut connaître. Le contenu de sa mémoire quantique est alors quantiquement corrélé au faisceau de Bob, ce qui donne le caractère « intriquant » de l'attaque.

Pour fixer les idées, supposons que Bob ait mesuré Q_B . Une fois qu'il en a averti Alice par le canal classique, écouté par Ève, celle-ci peut mesurer la quadrature correspondante du faisceau conservé dans sa mémoire quantique, le « frère jumeau » de celui qu'elle a injecté chez Bob. Elle en déduit une estimation Q_{connu} de Q_{E1} , et on peut écrire

$$Q_{E1} = Q_{\text{connu}} + Q_{\text{inconnu}}, \quad (12.19)$$

où Q_{inconnu} désigne le bruit qu'elle ne peut pas connaître. On a alors

$$\langle Q_{\text{inconnu}}^2 \rangle = \frac{N_0^2}{\langle Q_{E1}^2 \rangle} = \frac{1-G}{G\chi} N_0 \quad (12.20a)$$

$$\langle Q_{\text{connu}}^2 \rangle = \langle Q_{E1}^2 \rangle - \langle Q_{\text{inconnu}}^2 \rangle. \quad (12.20b)$$

Ève stocke également l'autre sortie $\begin{bmatrix} Q_{E2} \\ P_{E2} \end{bmatrix}$ de la lame semi-réfléchissante dans une mémoire quantique. Elle peut ensuite mesurer Q_{E2} pour obtenir de l'information sur le champ d'entrée :

$$Q_{E2} = \sqrt{G} Q_{E1} - \sqrt{1-G} Q. \quad (12.21a)$$

Elle peut annuler une partie du bruit induit par Q_{E1} en soustrayant la partie proportionnelle à Q_{connu} . Ainsi, elle connaît

$$Q'_{E2} = \sqrt{G} Q_{\text{inconnu}} - \sqrt{1-G} Q. \quad (12.21b)$$

On peut remarquer au passage que le bruit équivalent d'Ève sur la quadrature d'Alice vaut

$$\chi_B = \frac{1-G}{G} \langle Q_{\text{inconnu}}^2 \rangle = \frac{N_0}{\chi}, \quad (12.22)$$

c'est à dire qu'Ève en apprend autant sur Alice que si elle utilisait un duplicateur optimal. On voit donc bien ici qu'une cloneuse intrigante est un cas particulier de cloneuse, dont les sorties sont en plus corrélées entre elles.

D'autre part, la quadrature de Bob a pour expression

$$Q_B = \sqrt{G} Q + \sqrt{1-G} Q_{E1}, \quad (12.23a)$$

dont Ève connaît déjà la partie proportionnelle à Q_{connu} , injectée avec Q_{E1} . Il ne lui reste donc plus qu'à deviner

$$Q'_B = \sqrt{G} Q + \sqrt{1-G} Q_{\text{inconnu}} \quad (12.23b)$$

à partir de sa mesure de Q_{E2} . On a donc

$$V_{Q_B|Q_{\text{connu}}, Q_{E2}} = V_{Q'_B|Q'_{E2}} = \langle Q_B'^2 \rangle - \frac{\langle Q'_B Q'_{E2} \rangle^2}{\langle Q_{E2}'^2 \rangle}. \quad (12.24)$$

Les équations (12.21b) et (12.23b) nous permettent de calculer les éléments de la matrice de covariance de Q'_{E2} et de Q'_B :

$$\langle Q_B'^2 \rangle = \left(GV + \frac{(1-G)^2}{G\chi} \right) N_0 \quad (12.25a)$$

$$\langle Q_{E2}'^2 \rangle = (1-G) \left(V + \frac{1}{\chi} \right) N_0 \quad (12.25b)$$

$$\langle Q'_B Q'_{E2} \rangle = \sqrt{G(1-G)} \left(\frac{1-G}{G\chi} - V \right) N_0, \quad (12.25c)$$

d'où on déduit

$$\langle Q_B'^2 \rangle = \left(GV + \frac{1}{G\chi} - \frac{2}{\chi} + \frac{G}{\chi} \right) N_0 \quad (12.26a)$$

$$\frac{\langle Q'_B Q'_{E2} \rangle^2}{\langle Q_{E2}'^2 \rangle} = \left(\frac{1}{G\chi(\chi V + 1)} - \frac{2}{\chi} + GV + \frac{G}{\chi} \right) N_0 \quad (12.26b)$$

$$V_{Q'_B|Q'_{E2}} = \frac{1}{G(\chi + \frac{1}{V})} N_0 \quad (12.26c)$$

On a donc

$$V_{Q_B|Q_{\text{connu}}, Q_{E2}} = \frac{1}{G(\chi + \frac{1}{V})} N_0 = V_{B|E \text{ min}}, \quad (12.27a)$$

et des calculs similaires nous donnent

$$V_{P_B|P_{\text{connu}}, P_{E2}} = \frac{1}{G(\chi + \frac{1}{V})} N_0 = V_{B|E \text{ min}}. \quad (12.27b)$$

Les inégalités (12.17) sont donc saturées, et on a bien construit une cloneuse intriquante optimale.

Il est intéressant d'examiner ce que devient cette cloneuse intriquante optimale lorsque le bruit est du aux pertes seules, c'est à dire lorsque

$$\chi = \frac{1-G}{G} \quad (V < 1). \quad (12.28)$$

Dans ce cas, le champ $\begin{bmatrix} Q_{E1} \\ P_{E1} \end{bmatrix}$ est vide. Ève ne peut donc rien connaître de ses fluctuations. En effet, dans ce cas, on a

$$\langle Q_{E1}^2 \rangle = \langle Q_{\text{inconnu}}^2 \rangle = N_0 \quad (12.29)$$

et l'équation (12.20b) devient

$$\langle Q_{\text{connu}}^2 \rangle = 0. \quad (12.30)$$

La cloneuse intriquante optimale dans ce cas est donc une simple lame partiellement réfléchissante. Dans ce cas, Ève ne peut pas injecter de bruit connu, et sa seule source d'information sur Bob est le mode $\begin{bmatrix} Q_{E2} \\ P_{E2} \end{bmatrix} = \begin{bmatrix} Q'_{E2} \\ P'_{E2} \end{bmatrix}$. Par contre, un excès de bruit sur la ligne (bruit supérieur aux fluctuations du vide associées aux pertes) est immédiatement exploitable par Ève.

12.3 Canaux bruités permettant la cryptographie inverse

Connaissant la stratégie optimale d'Ève, nous pouvons étudier les conditions sous lesquelles Alice et Bob peuvent utiliser des protocoles de réconciliation inverse (section 12.3.1). Nous étudierons ensuite les cas particulier constitués par les états maximalement comprimés (section 12.3.2) et les états cohérents (section 12.3.3), avant d'étudier le cas plus général des facteurs de compression intermédiaires (section 12.3.4).

12.3.1 Condition générale

Si Alice et Bob veulent utiliser un protocole de réconciliation inverse, ils doivent supposer qu'Ève utilise une cloneuse intriquante optimale, définie par les équations (12.17). Il leur est possible d'extraire une clé secrète si Alice en sait plus sur Bob qu'Ève, c'est à dire si

$$V_{Q_B|Q_A} < V_{Q_B|Q_E, \min} \quad \text{ou} \quad V_{P_B|P_A} < V_{P_B|P_E, \min}. \quad (12.31)$$

Lorsqu'on tient compte de (12.14) et (12.17), les conditions ci-dessus deviennent

$$G_Q(\chi_Q + s)G_P(\chi_P + \frac{1}{V}) < 1 \quad \text{ou} \quad G_P(\chi_P + s)G_Q(\chi_Q + \frac{1}{V}) < 1 \quad (12.32a)$$

Dans la suite, nous nous limiterons au cas symétrique en Q, P , où ces équations s'écrivent

$$G^2(\chi + s)(\chi + \frac{1}{V}) < 1 \quad (12.32b)$$

Contrairement aux protocoles inverses, pour lesquels la condition équivalente s'écrit simplement $\chi < 1$, cette condition fait intervenir les deux paramètres du canal (G et χ), mais aussi les paramètres de la modulation d'Alice (s et V). On voit en particulier qu'une grande modulation et/ou une forte compression accroît la robustesse du protocole au bruit, puisqu'elles permettent d'accroître la valeur de χ tolérée par l'inégalité (12.32b).

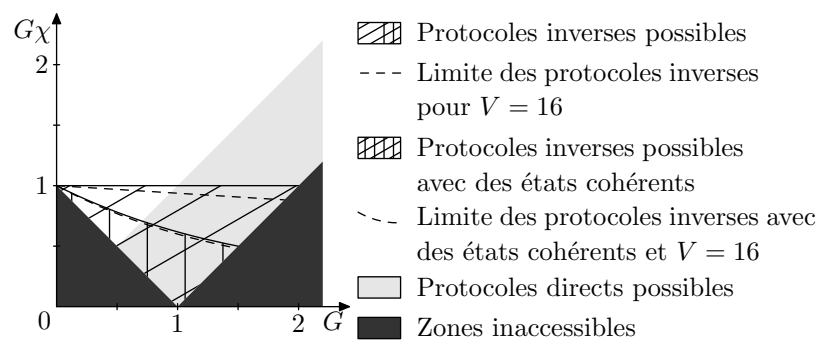


FIG. 12.3 – Canaux où il est possible de faire de la cryptographie inverse

12.3.2 États maximalement comprimés

Les paramètres G et $G\chi$ des canaux symétriques où un protocole de cryptographie inverse est possible sont représentés FIG. 12.3.

Pour des états maximalement comprimés, on a $s = \frac{1}{V}$, et l'inégalité (12.32b) devient

$$G(\chi + \frac{1}{V}) < 1 \quad \text{c'est à dire} \quad G\chi < 1 - \frac{G}{V} \quad (12.33)$$

À la limite des fortes modulations cette inégalité devient

$$G\chi < 1. \quad (12.34)$$

Cette inégalité définit l'ensemble des canaux à travers lesquels il est possible d'utiliser des protocoles inverses. Contrairement aux protocoles directs, le facteur limitant est le bruit total à la sortie de canal $G\chi$, et non le bruit équivalent ramené à l'entrée χ . Cela est dû à la possibilité d'utiliser des attaques intriquantes, qui contrôlent le bruit à la sortie. Si celui-ci est trop petit, Ève n'a en quelque sorte « pas la place » d'y dissimuler des corrélations quantiques.

Dans le cas où le bruit est dû aux seules pertes $G\chi = 1 - G < 1$, et il est possible de générer une clé secrète par réconciliation inverse pour des pertes arbitrairement grandes, alors que les protocoles de réconciliation directs ne fonctionnent pas au delà de 3 dB de pertes. Par contre, ceux ci fonctionnent dans certains régimes où les protocoles inverses ne fonctionnent pas, lorsque $G > 2$ par exemple, et que $\chi < 1$. Dans ce cas, on a en effet $G\chi > G - 1 > 1$.

12.3.3 États cohérents

Pour des états cohérents, $s = 1$ et l'inégalité (12.32b) devient

$$(G\chi)^2 + G(1 + \frac{1}{V})G\chi + \frac{G^2}{V} - 1 < 0. \quad (12.35)$$

Le discriminant de ce trinôme en $G\chi$ vaut

$$\Delta = G^2(1 + \frac{1}{V})^2 - 4\frac{G^2}{V} + 4 = G^2(1 - \frac{1}{V}) + 4, \quad (12.36)$$

ce qui est toujours positif. Cette condition devient donc

$$G\chi < G\chi_{\max, \text{coh}}^\infty = \frac{1}{2} \left[\sqrt{G^2(1 - \frac{1}{V})^2 + 4} - G(1 - \frac{1}{V}) \right]. \quad (12.37)$$

Cette inégalité devient, à la limite des fortes modulations,

$$G\chi < G\chi_{\max, \text{coh}}^{\infty} = \frac{1}{2} \left[\sqrt{G^2 + 4} - G \right]. \quad (12.38)$$

Une comparaison des inégalités (12.37) (12.33) nous montre que les états comprimés permettent de faire de la cryptographie inverse dans des canaux plus bruités. En effet, il n'est pas difficile de se convaincre que

$$G\chi_{\max, \text{coh}} < 1 - \frac{G}{V}. \quad (12.39)$$

Contrairement aux protocoles directs, les protocoles inverses font donc une différence entre les états cohérents et les états comprimés. Ces derniers sont notamment plus robustes au bruit ajouté.

On peut quantifier cette robustesse au bruit ajouté en scindant le bruit équivalent χ en deux parties :

$$\chi = \chi_0 + \xi, \quad (12.40)$$

où $\chi_0 = \frac{1-G}{G}$ est le bruit du aux pertes (on se contente ici d'étudier le cas $G \leq 1$) et ξ est le bruit ajouté proprement dit. L'équation (12.32b) donne alors la condition suivante pour ξ (si $G < 1$) :

$$G\xi^2 + \left[2 - G(2 - s - \frac{1}{V})\right]\xi - 2 + s + \frac{1}{V} + G(1 - s)(2 - \frac{1}{V}) < 0. \quad (12.41)$$

À la limite des faibles transmissions ($G \ll 1$), ce trinôme se simplifie et devient

$$\boxed{\xi < 1 - \frac{1}{2}(s + \frac{1}{V})}. \quad (12.42)$$

Il est alors aisé de voir qu'un protocole de cryptographie inverse avec une forte modulation ($V \rightarrow \infty$) et une forte compression ($s \rightarrow 0$) tolère deux fois plus de bruit ajouté ($\xi = 1$) qu'un protocole avec des états cohérents ($s = 1$ et $\xi = \frac{1}{2}$)

12.3.4 Facteurs de compression intermédiaires

L'inégalité (12.32b) peut s'écrire pour faire ressortir la quantité minimale de compression en fonction du gain G et du bruit χ :

$$s < s_{\max} \equiv \frac{1}{G^2(\chi + \frac{1}{V})} - \chi = 1 - \xi + \frac{1 - \xi - \frac{1}{V}}{1 - G(1 - \xi - \frac{1}{V})} \quad (12.43)$$

Si le bruit ajouté ξ est trop important, $s_{\max}$ peut être négatif. Dans ce cas il est impossible de faire de la cryptographie inverse, même avec des états maximalelement comprimés. Au contraire, il peut-être supérieur à 1, ce qui autorise à faire de la cryptographie inverse avec des états non minimaux.

Dans le cas particulier d'un canal avec des pertes seules, $\xi = 0$ et cette inégalité devient

$$s < s_{\max}^0 \equiv 1 + \frac{1 - \frac{1}{V}}{1 - G(1 - \frac{1}{V})} \quad (12.44)$$

Il est aisé de se convaincre que

$$s_{\max} < s_{\max}^0 < V \quad (12.45)$$

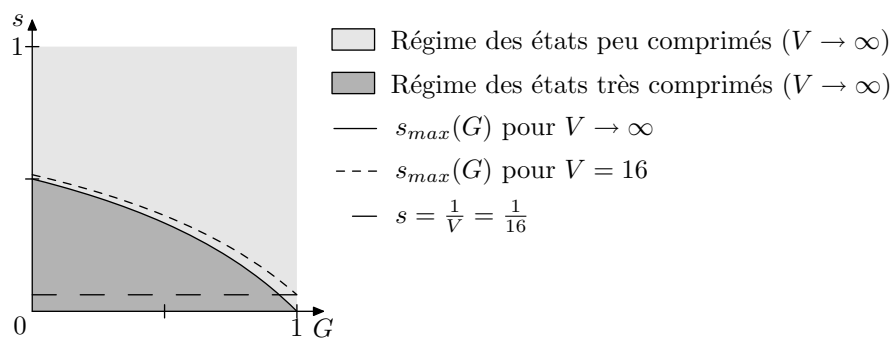


FIG. 12.4 – Les différents régimes utilisant des états comprimés dans des protocoles de réconciliation inverse (pertes seules)

On a $s_{\max}^0 > 1$ pour tout $G \in]0, 1]$ et pour tout $V > 1$. Il est donc toujours possible de faire de la cryptographie quantique avec des états cohérents (en l'absence de bruit ajouté), même pour des pertes arbitrairement élevées ($G \rightarrow 1$). C'est *a fortiori* possible pour des états comprimés.

Cependant, l'existence de cette limite $s_{\max}$, finie et toujours inférieure à V permet de distinguer plusieurs régimes dans le cas des états minimaux, en fonction du facteur de compression $s \leq 1$:

- Si $s_{\max} < 1$, pour les états insuffisamment comprimés ($s > s_{\max}$) Alice et Bob ne peuvent pas extraire d'information secrète par réconciliation inverse.
- Pour les états peu comprimés ($s_{\max} > 1$ et $\frac{1}{s_{\max}} < s < s_{\max}$), Alice et Bob peuvent extraire de l'information secrète qu'ils aient ou non choisi la même quadrature, comme pour la réconciliation directe.
- Pour les états très comprimés ($s < s_{\max}$ et $\frac{1}{s} > s_{\max}$) la variance de la quadrature « anticomprimée » est supérieure à $s_{\max}$. Alice connaît donc moins bien cette quadrature qu'Ève, et ne peut en extraire d'information secrète.

Nous avons représenté ces régimes dans le cas des pertes seules FIG. 12.4.

12.4 Quantités d'information

Nous avons maintenant tous les éléments pour calculer le taux d'information secrète que l'on peut extraire avec un protocole de réconciliation inverse. Après quelques calculs élémentaires (section 12.4.1), nous étudierons les cas particuliers des états maximalelement comprimés (section 12.4.2) et des états cohérents (section 12.4.3), avant de nous intéresser au cas plus général des facteurs de compression intermédiaires (section 12.4.4). Enfin, section 12.4.5, nous évaluerons les performances de ces protocoles dans le régime des fortes pertes et nous le comparerons à BB84.

12.4.1 Informations mutuelles et information secrète

Comme la modulation d'Alice est gaussienne, la formule de Shannon (7.36) nous permet de calculer aisément les quantités d'information mutuelles. Les équations (12.14) nous permettent de calculer l'information que Bob a sur Alice. Une comparaison avec (11.8a) nous

permet de vérifier qu'on a bien $I_{BA} = I_{AB}$.

$$I_{BA} = \frac{1}{2} \log_2 \frac{\langle Q_B^2 \rangle}{V_{B|A}} = \frac{1}{2} \log_2 \frac{GV + G\chi}{G(s + \chi)} = I_{AB} \quad (12.46)$$

On peut calculer de même la quantité d'information I_{BE} qu'Ève a sur la mesure de Bob :

$$I_{BE} = \frac{1}{2} \log_2 \frac{\langle Q_B^2 \rangle}{V_{B|E}} = \frac{1}{2} \log_2 \left[G^2(\chi + V)(\chi + \frac{1}{V}) \right] \quad (12.47)$$

On en déduit alors la quantité d'information secrète qu'on peut extraire par un protocole de réconciliation inverse :

$$\boxed{\Delta I = I_{BA} - I_{BE} = \frac{1}{2} \log_2 \frac{V_{B|E}}{V_{B|A}} = -\frac{1}{2} \log_2 [G^2(\chi + \frac{1}{V})(\chi + s)]} \quad (12.48)$$

Comme dans le cas de la réconciliation directe, ce taux d'information secrète tend vers une limite finie à la limite des grandes modulations ($V \rightarrow \infty$) :

$$\Delta I^\infty = -\frac{1}{2} \log_2 [G^2\chi(\chi + s)] \quad (12.49)$$

12.4.2 États maximalement comprimés

Lorsqu'Alice envoie des états maximalement comprimés ou, ce qui revient au même, (voir section 4.2.4) une demi-paire de faisceaux EPR, $s = \frac{1}{V}$. Comme le protocole est symétrisé, Alice et Bob ne choisissent la même quadrature qu'une fois sur deux. Si Alice a envoyé un état comprimé en Q et que Bob a mesuré P_B , par exemple, Alice n'a aucune information sur la mesure de Bob, contrairement à Ève. Cet élément de clef n'apportera donc aucune information secrète, et Alice et Bob jetteront au total la moitié des éléments de clef. On se trouve dans ce cas dans le régime des état très comprimés, pour reprendre la terminologie de la section 12.3.4, à moins bien sûr que le facteur de compression soit insuffisant ($s = \frac{1}{V} > s_{\max}$) compte tenu des caractéristiques du canal pour autoriser la génération d'une clef secrète.

On a donc

$$\Delta I_{\text{EPR}} = -\frac{1}{2} \frac{1}{2} \log_2 [G^2(\chi + \frac{1}{V})(\chi + \frac{1}{V})] = -\frac{1}{2} \log_2 [G(\chi + \frac{1}{V})] \quad (12.50)$$

et, à la limite des grandes modulations et des fortes compressions ($V \rightarrow \infty$),

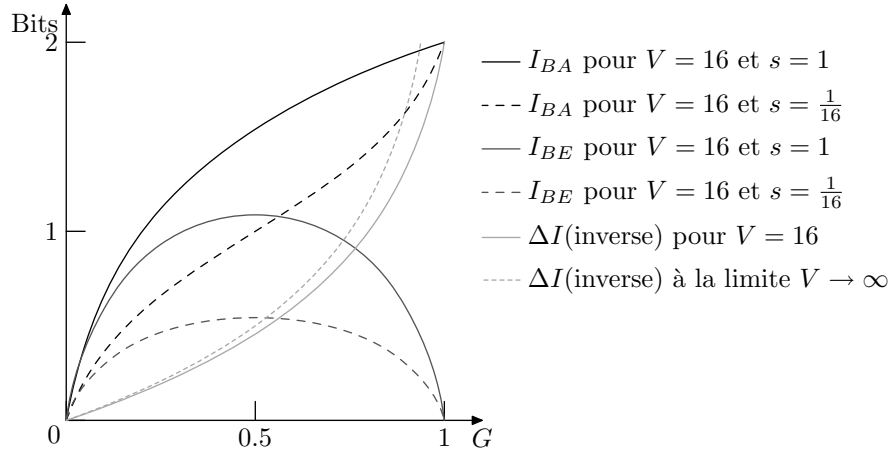
$$\Delta I_{\text{EPR}}^\infty = -\frac{1}{2} \log_2 G\chi. \quad (12.51)$$

Si on distingue dans le bruit équivalent χ la partie due aux pertes² $\chi_0 = \frac{1}{G} - 1$ et l'excès de bruit $\xi \equiv \chi - \chi_0 \geq 0$, l'équation (12.50)

$$\Delta I_{\text{EPR}} = -\frac{1}{2} \log_2 [1 - G(1 - \frac{1}{V} - \xi)] \quad (12.52)$$

Les limites de fortes modulations et des pertes seules s'en déduisent aisément, en faisant tendre respectivement $\frac{1}{V}$ et ξ vers 0.

²Dans le cas où $G > 1$, χ_0 est négatif et n'a plus de sens physique. $\xi > 2(G - 1)$ ne représente plus l'excès de bruit à proprement parler, mais les équations restent valides.

FIG. 12.5 – Informations mutuelles en fonction de G (pertes seules)

12.4.3 États cohérents

Lorsqu'Alice envoie des états cohérents modulés en Q et en P , $s = 1$ et la situation est symétrique. La quantité d'information récupérée par Bob est alors indépendante de la quadrature qu'il mesure. Pour reprendre les termes de la section 12.3.4, on est soit dans le régime des états insuffisamment comprimés, soit dans le régime des états faiblement comprimés, ce qui est somme toute logique, puisqu'on utilise des états qui ne sont pas comprimés.

On a alors

$$\underline{\Delta I}_{\text{coh}} = -\frac{1}{2} \log_2 \left[G^2 \left(\chi + \frac{1}{V} \right) (\chi + 1) \right] = \underline{\Delta I}_{\text{EPR}} - \frac{1}{2} \log_2 [G(\chi + 1)], \quad (12.53)$$

Réécrite en fonction de l'excès de bruit ξ , cette égalité devient

$$\underline{\Delta I}_{\text{coh}} = \underline{\Delta I}_{\text{EPR}} - \frac{1}{2} \log_2 [1 + G\xi] \leq \underline{\Delta I}_{\text{EPR}} \quad (12.54)$$

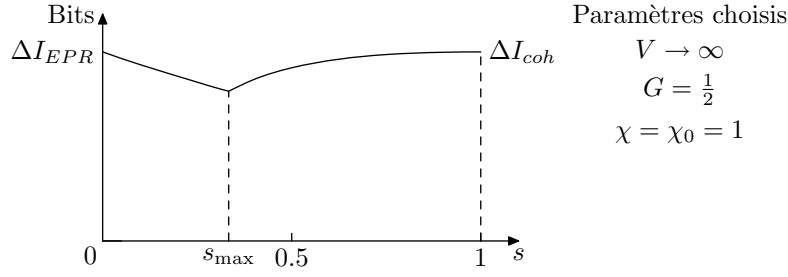
Les deux taux d'information secrètes sont donc égaux si et seulement si le bruit est dû aux seules pertes ($\xi = 0$). Cependant, contrairement au cas de la réconciliation directe, les protocoles à états fortement comprimés sont plus robustes à l'excès de bruit et leurs performances sont moins dégradées que ceux avec des états cohérents.

12.4.4 Facteurs de compression intermédiaires

Les états cohérents et les états maximalelement comprimés représentent les cas extrêmes, et il est bien entendu possible d'étudier les protocoles de cryptographie correspondant aux cas intermédiaires. Comme nous l'avons vu section 12.3.4, il convient de distinguer deux régimes, en fonction du paramètre de compression s et de s_{max} , défini par l'équation (12.43), lorsqu'une génération de clef est possible ($s < s_{\text{max}}$).

12.4.4.1 Régimes des états très comprimés

Il s'agit du régime où $s < s_{\text{max}}$ et où $\frac{1}{s} \geq s_{\text{max}}$, représenté à l'extrême par les états maximalelement comprimés ($s = \frac{1}{V} < \frac{1}{s_{\text{max}}}$) et du seul régime qu'on peut atteindre si $s_{\text{max}} \leq 1$,

FIG. 12.6 – ΔI en fonction du facteur de compression s

c'est-à-dire s'il est impossible de faire de la cryptographie inverse avec des états cohérents. Dans ce cas, Alice en sait plus qu'Ève sur une des deux quadrature de Bob, et une seulement. Comme le protocole est symétrisé, elle obtient le bon résultat une fois sur deux seulement et on a

$$\underline{\Delta I}_{\text{comp.}} = -\frac{1}{4} \log_2 [G^2(\chi + \frac{1}{V})(\chi + s)] = \frac{1}{2} \underline{\Delta I}_{\text{EPR}} - \frac{1}{4} \log_2 [1 - G(1 - s - \xi)]. \quad (12.55a)$$

Il est aisé de voir que le taux d'information secrète décroît lorsque les états sont moins comprimés.

12.4.4.2 Régimes des états peu comprimés

Cependant, cette décroissance s'arrête lorsque les états sont suffisamment peu comprimés pour qu'Alice en sache plus qu'Ève sur les deux quadratures de Bob, c'est à dire lorsque $s < s_{\max}$ et $\frac{1}{s} < s_{\max}$. On peut atteindre ce régime si $s_{\max} > 1$, c'est à dire s'il est possible de faire de la cryptographie avec des états cohérents, celle-ci représentant le cas extrême.

Dans ce cas, Alice et Bob peuvent également exploiter la quadrature « anticomprimée » et

$$\underline{\Delta I}_{\text{-comp.}} = -\frac{1}{2} \log_2 [G^2(\chi + \frac{1}{V}) \sqrt{(\chi^2 + (s + \frac{1}{s})\chi + 1)}] \quad (12.56a)$$

$$= \underline{\Delta I}_{\text{EPR}} - \frac{1}{4} \log_2 [1 - G(1 - s - \xi)] - \frac{1}{4} \log_2 [1 - G(1 - \frac{1}{s} - \xi)]. \quad (12.56b)$$

Comme $s + \frac{1}{s}$ est minimal pour $s = 1$, le taux d'information secrète décroît lorsqu'on comprime les états. Il est donc minimum à la transition entre ses deux régimes, lorsque $s = \frac{1}{s_{\max}}$.

12.4.5 Fortes pertes

À la limite des fortes pertes ($G \ll 1$), un développement limité nous donne la quantité d'information secrète qu'Alice et Bob peuvent extraire

$$\underline{\Delta I}_{\text{EPR}} \simeq \frac{1}{2 \ln 2} G(1 - \frac{1}{V} - \xi) \quad \underline{\Delta I}_{\text{coh}} \simeq \frac{1}{2 \ln 2} G(1 - \frac{1}{V} - 2\xi). \quad (12.57)$$

Si le bruit est dû aux seules pertes ($\xi = 0$), à la limite des fortes modulations ($V \rightarrow \infty$), nous avons

$$\underline{\Delta I}_{\text{EPR,pertes}} = \underline{\Delta I}_{\text{coh,pertes}} \simeq \frac{1}{2 \ln 2} G \simeq 0,721 G. \quad (12.58)$$

Ce taux peut-être comparé à au taux d'information secrète de BB84, qui vaut $S_{\text{BB84}} = \frac{1}{2}G\bar{n}$ dans le cas des seules pertes, où $\bar{n} = 1$ dans le cas de sources parfaites de photons uniques et $\bar{n} \ll 1$ pour les impulsions lumineuses atténuées. Ces deux valeurs sont donc du même ordre de grandeur dans le cas idéal (impulsions à un seul photon), le protocoles à variables continues prenant nettement l'avantage sur la cryptographie à impulsions atténuées. Même si une astuce permet de doubler le taux de BB84 à $S'_{\text{BB84}} = G\bar{n}$ [186], les protocoles inverses peuvent être considérés comme plus performants car ils n'utilisent qu'un mode du champ par élément de clef, alors que BB84 en utilise deux.

Si on examine par exemple le cas d'une fibre optique de 100 km avec 20 dB de pertes ($G = 0,01$) avec une modulation raisonnable $V = 16$, on obtient un taux d'information secrète de $6,8 \cdot 10^{-3}$ bits par symbole. Pour les mêmes paramètres, une variante de BB84, avec des sources de photons uniques parfaites et des détecteurs parfaits pourrait monter à 10^{-2} bits par symbole. L'utilisation d'impulsions lumineuses atténuées avec $\bar{n} = 0,1$ diminue cette valeur à 10^{-3} bits par symbole. De plus, le bruit des détecteurs de photons uniques actuels à 1550 nm réduisent en pratique ce taux à zéro. Pour les variables continues, il est parfaitement envisageable d'envoyer des impulsions à quelques mégahertz, ce qui aboutirait à une génération de clef secrète à un rythme de l'ordre de 10 kbits/s.

Malheureusement, si les protocoles à variables continues semblent plus efficaces que BB84, tant sur le plan théorique qu'expérimental, nous verrons dans le chapitre suivant que les imperfections des protocoles actuels de réconciliation en limitent la portée.

Chapitre 13

Expérience de cryptographie quantique

Nous avons implémenté les protocoles à états cohérents décrits dans les chapitres précédents en utilisant la détection homodyne décrite dans la partie II. Nous avons collaboré avec Gilles Van Assche et Nicolas Cerf, de l'Université Libre de Bruxelles, qui ont extrait de vraies clefs secrètes binaires à partir de nos signaux optiques.

Si cette expérience est une démonstration de faisabilité et pas un véritable système de cryptographie quantique, c'est essentiellement pour des raisons techniques aisément surmontables. Nous avons par contre abordé et surmonté tous les points délicats, que leur nature soit physique ou informatique, ce qui nous a permis de générer des clefs secrètes à des taux allant de 1,6 mégabits/s pour de faibles pertes, à 75 kilobit/s pour 3,1 dB de pertes, franchissant ainsi le seuil des 3 dB de pertes.

Nous étudierons dans la section 13.1 la structure des signaux envoyés par Alice, avant d'étudier la mise en œuvre expérimentale proprement dite dans la section 13.2. La partie optique du dispositif a été, pour une grande partie, réalisée par Jérôme Wenger et sera décrite plus en détail dans sa thèse[131]. La quantité d'information accessible à Ève dans le cadre d'un protocole expérimental sera évaluée section 13.3. La section 13.4 traitera des procédés informatiques qui ont été développés par Gilles Van Assche pour extraire une clef secrète à partir des données acquises par l'expérience. Cette partie sera traitée plus en détail et plus en profondeur dans la thèse de Gilles Van Assche [187]. Ce chapitre se conclura, section 13.5 par les résultats obtenus avec ce système¹.

13.1 Signaux envoyés par Alice

Alice fabrique l'oscillateur local à l'aide d'une diode laser monomode hachée en impulsions de 120 ns avec un taux de répétition de 800 kHz. Le faisceau est ensuite filtré spatialement et spectralement comme nous l'avons expliqué section 5.5. Alors que l'essentiel de l'oscillateur local est directement envoyé à Bob, une fraction en est prélevée et injectée dans un modulateur intégré qui fabrique l'état cohérent voulu. Comme expliqué section 6.3, un bruit de phase d'origine acoustique limitera l'intensité des états cohérents à quelques centaines de photons au maximum, ce qui ne nous gênera pas vraiment.

¹Les résultats présentés dans ce chapitre ont fait l'objet d'un article à paraître dans *Nature* [14, 15]

13.1.1 Génération aléatoire du signal

13.1.1.1 Introduction

Les protocoles à états cohérents décrits dans les deux chapitres précédents nécessitent une modulation aléatoire gaussienne de $\begin{bmatrix} Q_A \\ P_A \end{bmatrix}$ de variance $V_A N_0 = (V - 1)N_0$ dans l'espace des phases. Nous utiliserons le générateur pseudo-aléatoire standard du micro-ordinateur pilotant l'expérience même s'il n'est très certainement pas de qualité cryptographique [185]. Il est en effet probablement cyclique avec une période de 32 767 points [188], ce qui serait rédhibitoire si on cherchait à faire autre chose qu'une démonstration de principe, vu que les blocs de données de base sont constitués de 50 000 éléments de clefs. Pour un vrai système de cryptographie, il faudrait traiter la production de nombres aléatoires avec plus de soins, en recourant éventuellement à des systèmes physiques de génération aléatoire.

Il faut ensuite transformer les nombres fournis par le générateur pseudo-aléatoire, uniformément distribués sur l'intervalle $[0,1]$ en signaux à envoyer sur les modulateurs pour générer la distribution gaussienne d'états cohérents dont nous avons besoin pour nos protocoles de cryptographie quantique.

13.1.1.2 Distributions en amplitude et en phase

Comme les modulateurs ne permettent en général pas de changer séparément Q et P mais plutôt l'amplitude R et la phase Θ , il faut d'abord calculer les distributions en amplitude et en phase correspondant à la distribution de probabilité gaussienne bivariée de variance $V_A N_0$

$$\mathcal{P}_{Q=q, P=p} = \mathcal{P}_{Q=q} \mathcal{P}_{P=p} = \frac{1}{2\pi V_A N_0} e^{-\frac{q^2 + p^2}{2V_A N_0}} \quad (13.1)$$

Pour réexprimer cette probabilité en fonction de $\begin{pmatrix} R \\ \Theta \end{pmatrix}$, il ne faut pas égaliser les distributions, au changement de variable près, mais il faut tenir compte du jacobien du changement de variables, [108, 188] ce qui nous donne²

$$|\mathcal{P}_{q,p} dq dp| = |\mathcal{P}_{r,\theta} r d\theta dr|. \quad (13.2)$$

Comme l'élément de surface élémentaire $dq dp$ en coordonnées cartésiennes devient $r d\theta dr$, l'équation précédente devient

$$|\mathcal{P}_{r,\theta}| = |r \mathcal{P}_{Q=r \cos \theta, P=r \sin \theta}|. \quad (13.3)$$

La distribution gaussienne (13.1) devient donc

$$\mathcal{P}_{r,\theta} dr d\theta = \frac{r}{V_A N_0} e^{-\frac{r^2}{2V_A N_0}} dr \frac{d\theta}{2\pi} = \mathcal{P}_r dr \mathcal{P}_\theta d\theta \quad (13.4a)$$

$$\mathcal{P}_r dr = \frac{r}{V_A N_0} e^{-\frac{r^2}{2V_A N_0}} dr = -d \left\{ e^{-\frac{r^2}{2V_A N_0}} \right\} \quad (13.4b)$$

$$\mathcal{P}_\theta d\theta = \frac{d\theta}{2\pi} \quad (13.4c)$$

²On abrégera dans la suite, en l'absence d'ambiguïté, $\mathcal{P}_{Q=q, P=p}$ en $\mathcal{P}_{q,p}$, $\mathcal{P}_{Q=q}$ en $\mathcal{P}_q$, $\mathcal{P}_{P=p}$ en $\mathcal{P}_p$, $\mathcal{P}_{R=r, \Theta=\theta}$ en $\mathcal{P}_{r,\theta}$, $\mathcal{P}_{R=r}$ en $\mathcal{P}_r$, et $\mathcal{P}_{\Theta=\theta}$ en $\mathcal{P}_\theta$

Ces équations nous donnent les distributions de probabilité en amplitude et en phase correspondant à une distribution gaussienne bivariee. La distribution en amplitude (13.4b) n'est rien d'autre qu'une distribution de Maxwell à deux dimensions. Nous nous servirons plus loin de son expression différentielle. La distribution en phase (13.4c) est uniforme sur $[0, 2\pi[$ et ne pose pas de problème particulier. Nous n'y reviendrons plus dans la suite.

13.1.1.3 Amplitudes finies

La distribution (13.4b) donne des amplitudes distribuées sur $[0, +\infty[$, alors que les amplitudes que l'on fabriquera physiquement seront nécessairement bornés. Si on appelle $R_{\max}$ l'amplitude maximale accessible, nous devons restreindre la distributions de probabilité précédente à l'intervalle $[0, R_{\max}]$. Comme les distributions de probabilité gaussiennes décroissent rapidement, il suffit que $R_{\max}$ vaille quelques écart-types $\sqrt{V_A N_0}$ pour que les événements hors de l'intervalle soient très rares. Soit f le facteur

$$f = \frac{\sqrt{V_A N_0}}{R_{\max}}. \quad (13.5)$$

Une valeur couramment utilisée pour des distributions gaussiennes est $f = \frac{1}{3}$.

Pour être plus quantitatif, il faut calculer la probabilité $\mathcal{P}_{r>R_{\max}}$ qu'à la distribution idéale de tomber hors de la zone autorisée. Cette probabilité sera parfois exprimée en bits, une probabilité de $n_{r>R_{\max}}$ correspondant à

$$\mathcal{P}_{r>R_{\max}} = 2^{-n_{r>R_{\max}}} \quad (13.6)$$

L'expression différentielle (13.4b) nous permet de calculer simplement cette probabilité :

$$\mathcal{P}_{r>R_{\max}} = \int_{R_{\max}}^{\infty} \mathcal{P}_r dr = e^{-\frac{R_{\max}^2}{2V_A N_0}} = e^{-\frac{1}{2f^2}} \quad (13.7)$$

d'où on déduit la relation

$$\frac{1}{f} = \sqrt{-2 \ln \mathcal{P}_{r>R_{\max}}} = \sqrt{n_{r>R_{\max}} 2 \ln 2}. \quad (13.8)$$

L'approximation usuelle $\frac{1}{f} = 3$ correspond à une probabilité de

$$n_{r>R_{\max}} \simeq \frac{9}{1.39} \simeq 9 \frac{3}{4} = \frac{27}{4} \simeq 7 \text{ bits environ}, \quad (13.9)$$

de l'ordre du pour cent (un calcul exact donne 6,5 bits, soit $\mathcal{P}_{r>R_{\max}} = 1/90 = 1,1 \%$).

13.1.1.4 D'une distribution uniforme à la modulation adéquate

Pour passer de la variable aléatoire X , donnée par le micro-ordinateur et uniformément distribuée sur $[0, 1]$, à l'amplitude R , distribuée selon (13.4b), il faut [108, 188] utiliser une équation analogue à (13.2) :

$$|\mathcal{P}_{X=x} dx| = |\mathcal{P}_{R=r} dr|. \quad (13.10)$$

Comme $\mathcal{P}_{X=x} = 1$ si $x \in [0, 1]$, cette équation, combinée à (13.4b), devient

$$\left| d \left\{ e^{-\frac{r^2}{2V_A N_0}} \right\} \right| = dx \quad (13.11)$$

$$e^{-\frac{r^2}{2V_A N_0}} = x \quad (13.12)$$

$$\boxed{r = \sqrt{-2V_A N_0 \ln x}} \quad (13.13)$$

Les petites valeurs de x correspondant aux grandes valeurs r , et réciproquement.

Si on obtient l'amplitude $r(s)$ en envoyant le signal $s = r^{-1}(r)$ au modulateur, il n'est pas difficile de voir que la fonction à appliquer à x pour obtenir le signal s à envoyer au modulateur sera

$$\boxed{s = r^{-1} \left(\sqrt{-2V_A N_0 \ln x} \right)}. \quad (13.14)$$

Bien sûr, le modulateur ne pourra pas sortir de signaux supérieurs à $R_{\max}$ et la distribution sera tronquée. Cela nous amène à rejeter les tirages de X hors de l'intervalle $[X_{\min}, 1]$, avec

$$X_{\min} = e^{-\frac{R_{\max}^2}{2V_A N_0}} = e^{-\frac{1}{2f^2}} = \mathcal{P}_{r > R_{\max}} = 2^{-n_{r > R_{\max}}}. \quad (13.15)$$

On peut éviter de rejeter des tirages aléatoires en contractant X avec la transformation linéaire adéquate. On définit donc $x' \in [X_{\min}, 1]$ par

$$x' = X_{\min} + x(1 - X_{\min}) \quad (13.16a)$$

et s est donné par

$$\boxed{s = r^{-1} \left(\sqrt{-2V_A N_0 \ln x'} \right)}. \quad (13.16b)$$

13.1.2 Effets du codage discret des nombres réels

13.1.2.1 Généralité

Tous nos raisonnements sont fondés sur des variables continues alors que l'ordinateur qui commande les modulateurs ne peut traiter que des entiers. Nos calculs restent cependant valables si le « quadrillage » est suffisamment dense pour être effacé par le bruit de photons, ce qui rend la situation indiscernable en pratique pour l'espion Ève. Si on note $\delta \sqrt{N_0}$ la distance entre deux points consécutifs du quadrillage, cette condition devient

$$\delta \equiv 2^{-n_\delta} \ll 1, \quad (13.17)$$

où l'on a défini n_δ , expression « en bits » de δ . *A priori*, n_δ doit être au moins égal à $n_{\delta, \min} = 3$, ce qui correspond à $\delta \leq \delta_{\min} = 0,125$.

13.1.2.2 Modulation de la phase

On peut définir δ_Θ (respectivement δ_R), qui correspond à la distance entre deux points adjacents de même amplitude (respectivement de même phase) et $n_{\delta, \Theta}$ (respectivement $n_{\delta, R}$). On a alors

$$\delta = \max\{\delta_\Theta, \delta_R\} \quad n_\delta = \min\{n_{\delta, \Theta}, n_{\delta, R}\} \quad (13.18)$$

et on peut traiter séparément la numérisation de la phase et de l'amplitude.

Si la phase est définie par n_θ bits, la distance séparant deux points adjacents de même amplitude r vaut

$$2^{-n_\theta} 2\pi r. \quad (13.19)$$

Cette distance est maximale pour $r = R_{\max}$ et on a donc

$$\delta_\Theta = 2^{-n_{\delta,\Theta}} = \frac{2^{-n_\theta} 2\pi R_{\max}}{\sqrt{N_0}}. \quad (13.20)$$

d'où on a

$$n_\theta = n_{\delta,\Theta} + \log_2 2\pi + \frac{1}{2} \log_2 \frac{R_{\max}^2}{N_0} \quad (13.21a)$$

$$= n_{\delta,\Theta} + \log_2 2\pi + \frac{1}{2} \log_2 V_A - \log_2 f \quad (13.21b)$$

$$= n_{\delta,\Theta} + \underbrace{\log_2(8\pi \ln 2)}_{2,89} + \underbrace{\frac{1}{2} \log_2 V_A}_{\simeq I_{AB}^0} + \frac{1}{2} \log_2 n_{r>R_{\max}}, \quad (13.21c)$$

où I_{AB}^0 désigne le taux d'information mutuelle entre Alice et Bob en l'absence de pertes et de bruit. L'égalité entre I_{AB}^0 et $\frac{1}{2} \log_2 V_A$ est approximativement valable dès que $I_{AB}^0 \geq 1,5$ bits. Pour les variances plus faibles, I_{AB}^0 reste une borne supérieure.

Le terme $\frac{1}{2} \log_2 n_{r>R_{\max}}$ évolue extrêmement lentement en fonction de la probabilité $\mathcal{P}_{r>R_{\max}}$ de dépasser l'amplitude maximale. Pour $n_{r>R_{\max}} = 6$ bits ($\mathcal{P}_{r>R_{\max}} = 1/64$), il vaut 1,3, et il semble peu raisonnable d'envisager de tronquer la distribution au delà de 16 bits ($\mathcal{P}_{r>R_{\max}} = 1/65\,536$), où ce terme ne vaut que 2.

En prenant les estimations minimales pour les différents paramètres de l'équation (13.21), on obtient pour $I_{AB}^0 \simeq 3$ bits ($V_A = 64$), qui est supérieure aux modulations expérimentales que nous utiliserons.

$$\boxed{n_\theta \geq n_{\theta,\min} \equiv 3 + 2,9 + 3 + 1,3 = 10,2 \text{ bits.}} \quad (13.22)$$

11 bits de modulation en phase semblent donc nécessaires. Or, les sorties analogiques de la carte PCI-6111E transmet des signaux codés sur 16 bits : cela nous laisse 5,8 bits de marge, ce qui est amplement suffisant.

13.1.2.3 Modulation de l'amplitude

Pour l'amplitude, si le signal s de commande du modulateur est défini sur n_s bits, la distance radiale maximale vaut environ

$$\delta_R \sqrt{N_0} \simeq 2^{-n_s} S_{\max} \max_s \left\{ \left| \frac{dr}{ds} \right| \right\}, \quad (13.23)$$

où $S_{\max}$ est le signal à envoyer au modulateur pour que l'amplitude de sortie soit $R_{\max}$ (en supposant que $r(s=0) \simeq 0$).

Si le modulateur d'amplitude est un interféromètre de Mach-Zender bien équilibré, on a

$$r(s) = R_{\max} \sin \left(\frac{\pi}{2} \frac{s}{S_{\max}} \right) \quad (13.24)$$

Précision	ϵ_{ordi}	n_{ordi}
float	5,96 10^{-8}	24
double	1,11 10^{-16}	53

TAB. 13.1 – Précision des codages en virgule flottante obéissant à la norme IEEE 754–1985 [189]

et le maximum de la dérivée vaut

$$\max_s \left\{ \left| \frac{dr}{ds} \right| \right\} = \frac{\pi}{2} \frac{R_{\max}}{S_{\max}} \max_s \left\{ \left| \cos \left(\frac{\pi}{2} \frac{s}{S_{\max}} \right) \right| \right\} = \frac{\pi}{2} \frac{R_{\max}}{S_{\max}} \quad (13.25)$$

Cette valeur est supérieure à ce que donne un modulateur mal équilibré ou un modulateur linéaire. Nous nous en servons donc comme borne supérieure et (13.23) devient

$$\delta_R \sqrt{N_0} \simeq 2^{-n_s} \frac{\pi}{2} R_{\max} = \delta_\theta \sqrt{N_0} 2^{n_\theta - n_s - 2} \quad (13.26)$$

$$-n_{\delta,R} = -n_{\delta,\Theta} + n_\theta - n_s - 2. \quad (13.27)$$

Si on pose la condition $n_{\delta,R} = n_{\delta,\Theta} = n_{\delta,\min}$, on a alors

$$n_{s,\min} = n_{\theta,\min} - 2 = 8, 2 \text{ bits}, \quad (13.28)$$

sous les mêmes conditions que précédemment. Comme en général $n_s = n_\theta$, cette égalité implique que c'est le caractère discret de la phase qui risque d'être gênant, et pas celui de l'amplitude, qui a 2 bits de marge en plus.

13.1.2.4 Générateur aléatoire discret

Nous avons précédemment supposé que nous disposions d'une variable réelle uniformément répartie sur $[0, 1]$, ce qui est bien entendu faux, puisque x' , qui en tient lieu, est généré par un ordinateur. Si le générateur pseudo-aléatoire est bon, ce que l'on supposera dans la suite, on n'est limité que par le codage en virgule flottante de la machine³, et nos nombres « réels » entre sur $[0, 1]$ sont des multiples entiers de

$$\epsilon_{\text{ordi}} = 2^{-n_{\text{ordi}}}. \quad (13.29)$$

Les valeurs usuelles de ses paramètres sont récapitulées dans [190] : $\epsilon_{\text{ordi}} = \text{epsneg}$ (on est entre 0 et 1), d'où $n_{\text{ordi}} = \text{negexp}$ dans le cas usuel des codages en base 2 ($\text{ibeta} = 2$). Les valeurs pour les machines respectent la norme IEEE 754–1985 [189] sont rappelées TAB. 13.1.

Comme la distribution en phase est uniforme, et que cette précision (24 bits en simple précision) est nettement supérieure aux 16 bits de la sortie analogique de la carte PCI-6111E, le caractère discret du modulateur aléatoire n'est pas le facteur limitant en phase.

³Pour être vraiment précis, le générateur aléatoire sera supposé donner des nombres régulièrement espacés, dont l'espacement ne sera limité que par la plus mauvaise précision sur l'intervalle $[0, 1]$ du codage en virgule flottante. Les générateurs pseudo-aléatoires usuels n'exploitent en effet pas le caractère flottant du codage en virgule flottante, qui permettrait d'avoir une précision *relative* constante et donc une précision *absolue* accrue lorsqu'on approche de 0.

En revanche, la relation entre n_{ordi} et n_s n'est pas si directe, car les relations (13.16), qui permettent de passer du nombre pseudo-aléatoire x au signal s , sont loin d'être linéaires.

On peut être tenté de dériver les relations (13.16) pour trouver la valeur minimale de $n_{\text{ordi,min}}$ qui permet de couvrir tout le quadrillage défini par $n_{s,\text{min}}$, mais cette démarche nous donnerait des résultats excessivement pessimistes. En effet, près de R_{max} , les valeurs de r accessibles par des valeurs de s uniformément espacées sont inutilement resserrées (au moins sous l'hypothèse (13.24)) car

$$\left. \frac{dr}{ds} \right|_{s_{\text{max}}} = 0. \quad (13.30)$$

Nous devons donc regarder directement les effets sur R de la discrétisation de x , qui définit une sorte de quadrillage en amplitude. La distance radiale entre deux points successifs de ce quadrillage est donnée par

$$\delta_R^{x',\text{aléa}} \sqrt{N_0} = 2^{-n_{\delta,R}^{x',\text{aléa}}} \sqrt{N_0} = \left| \frac{dr}{dx'} \frac{dx'}{dx} \right| \epsilon_{\text{ordi}}, \quad (13.31)$$

avec $\frac{dx'}{dx} = 1 - \mathcal{P}_{r>R_{\text{max}}} \simeq 1$ dans tous les cas où la distribution de probabilité reste raisonnablement gaussienne. Si on dérive (13.13), on obtient

$$2^{-n_{\delta,R}^{x',\text{aléa}}} = \frac{1}{x} \sqrt{\frac{V_A}{-2 \ln x}} 2^{-n_{\text{ordi}}}, \quad (13.32)$$

Cette équation diverge quand $x \rightarrow 0$ ($r \rightarrow \infty$) et quand $x \rightarrow 1$ ($r \rightarrow 0$) : ce sont les extrémités qui risquent de nous poser des problèmes, des valeurs de X régulièrement espacées se traduisant par des valeurs de R de plus en plus espacées lorsqu'on s'approche des petites et des grandes amplitudes.

Près de X_{min} , cette équation reste valable si $X_{\text{min}} \gg \epsilon_{\text{ordi}}$. Une comparaison entre l'équation (13.15) et les valeurs de la TAB. 13.1 nous montre que cette supposition est très raisonnable, même en simple précision et en prenant une très grande marge de sécurité pour englober toute la gaussienne ($n_{r>R_{\text{max}}} = 16$). L'équation (13.32), évaluée en $X_{\text{min}} = 2^{-n_{r>R_{\text{max}}}}$, nous donne

$$n_{\text{ordi}} = n_{\delta,R}^{0,\text{aléa}} + n_{r>R_{\text{max}}} + \underbrace{\frac{1}{2} \log_2 V_A}_{\simeq I_{AB}^0} - \frac{1}{2} \log_2 n_{r>R_{\text{max}}} - \underbrace{\frac{1}{2} \log_2 (2 \ln 2)}_{0,24}, \quad (13.33)$$

Dans les conditions minimales évoquées plus haut, cette équation se traduit par la condition

$$n_{\text{ordi}} \geq n_{\text{ordi,min}}^0 \equiv 3 + 6 + 3 - 1,3 - 0,2 = 10,5. \quad (13.34)$$

Par contre, l'équation (13.32) n'est pas pertinente près de $x = 1$ car le bord, où la dérivée est infinie, est atteint. Il faut donc calculer directement la distance entre deux points successifs, sans passer par la dérivée

$$\delta_R^{1,\text{aléa}} \sqrt{N_0} = |r(x=1) - r(x=1 - \epsilon_{\text{ordi}})| = \left| 0 - \sqrt{-2V_A N_0 \ln(1 - \epsilon_{\text{ordi}})} \right| \quad (13.35a)$$

$$2^{-n_{\delta,R}^{1,\text{aléa}}} \simeq \sqrt{V_A 2^{1-n_{\text{ordi}}}} \quad (13.35b)$$

$$n_{ordi} \simeq 2 n_{\delta,R}^{1,\text{aléa}} + \underbrace{\log_2 V_A}_{\simeq 2 I_{AB}^0} + 1 \quad (13.35c)$$

Sous les mêmes conditions minimales que précédemment, cette égalité se traduit par la condition

$$n_{ordi} \geq n_{ordi,\min}^1 \equiv 2 \times 3 + 2 \times 3 + 1 = 13 \text{ bits.} \quad (13.36)$$

L'utilisation de nombres aléatoires en simple précision est donc largement suffisante, même en ajoutant une marge pour tenir compte des erreurs d'arrondis lors des calculs. Elle est même largement excessive si notre source de nombres aléatoires est une ressource limitée.

Si on change de point de vue et qu'on calcule, à discrétisations $n_{ordi} = 24$ et $n_s = n_\theta = 16$ fixées, les paramètres n_δ qui décrivent les différentes tailles des « quadrillages » dans l'espace des phases, on constate que le paramètre limitant est le nombre de bits de la commande de la phase, mais que le codage en simple précision du générateur aléatoire pourrait poser des problèmes pour de grandes valeurs de I_{AB}^0 .

13.2 Mise en œuvre expérimentale

L'ensemble du dispositif est présenté sur la FIG. 13.1. Nous allons en détailler ci-dessous les éléments principaux.

13.2.1 Modulateur d'amplitude d'Alice

Le modulateur d'amplitude est un interféromètre de Mach-Zender intégré en niobate de lithium (LiNbO_3). Il ne se fabrique plus de modulateur intégré à 780 nm depuis une dizaine d'années, époque où ces longueurs d'ondes ont cessé d'être utilisées pour les télécommunications au profit de la fenêtre de transmission maximale des fibres optiques autour de 1 550 nm. Celui-ci nous a été prêté par Thierry Debuisschert de la société Thalès TRT. Les fibres optiques permettant de faire entrer et sortir la lumière du modulateur étaient cassées, mais Jérôme Wenger a pu les cliver afin de les réparer.

Notre modulation étant à 800 kHz seulement, nous avons dû utiliser la voie DC, prévue pour la tension continue de réglage, notre signal étant fortement déformé par la voie Rf, conçue pour des modulations de quelques gigahertz.

Ce modulateur est alimenté avec la séquence aléatoire définie ci-dessus. En pratique, la variance de la modulation reste fixe, et la variance du signal envoyé par Alice est réglée par un atténuateur placé en amont du modulateur.

Comme nous l'avons montré section 6.4, ce modulateur a malheureusement un taux d'extinction maximum en amplitude de 10 % seulement, c'est-à-dire 1 % en intensité. Cela nous oblige à faire une modulation dans l'espace des phases présentant un trou en son centre. Nous aurions pu corriger ce taux d'extinction insuffisant par un dispositif interférométrique adéquat ou étudier les taux d'information secrète dans le cas de telles modulations non-gaussiennes. Ces deux approches nous ont paru compliquées et inutiles. En effet, notre but était de faire une expérience de principe, permettant de générer physiquement des données ayant la même structure que celles créées dans un vrai système de cryptographie.

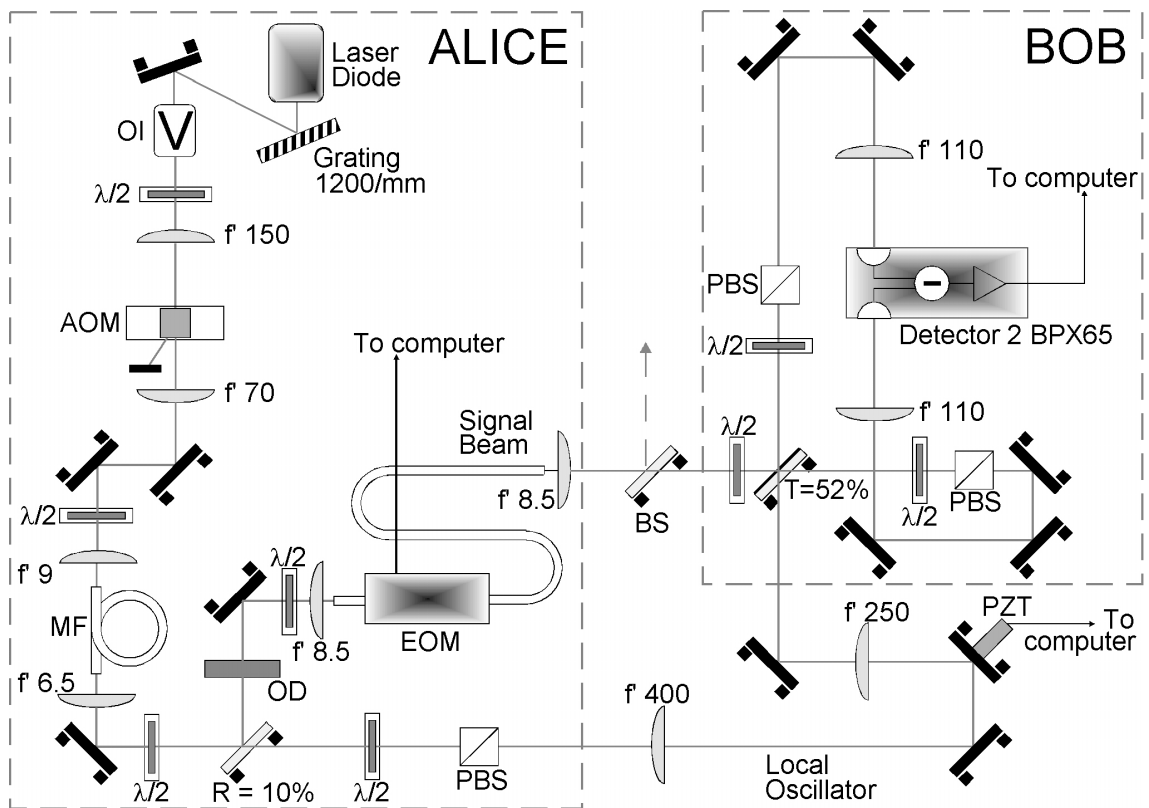


FIG. 13.1 – Dispositif expérimental

Or, un tel système serait construit à 1 550 nm, où de bons modulateurs sont disponibles commercialement et où ce problème ne se posera pas. Nous nous sommes donc contentés de corriger les données *a posteriori* par le traitement informatique *ad hoc*, ce qui nous permet de faire comme si notre modulateur avait un bon taux d'extinction.

13.2.2 Modulation de la phase

Pour pouvoir couvrir le plan $\left[\frac{Q}{P}\right]$, Alice a besoin de moduler la phase de ses états cohérents aléatoirement à un taux de 800 kHz. Malheureusement, si nous avons trouvé un modulateur d'amplitude intégré à 780 nm, nous ne sommes pas parvenus à trouver de modulateur de phase rapide à cette même longueur d'onde.

Nous avons donc créé ces déphasages par un miroir monté sur une cale piézoélectrique. La tension nécessaire pour faire décaler le miroir d'une demi-longueur d'onde vaut 90 V environ, et il est difficile de faire varier aléatoirement de telles tensions électriques à 800 kHz. Nous nous sommes donc contentés de faire varier la phase continuellement et uniformément entre 0 et $\frac{6}{5}2\pi$ pendant les 75 ms durant lesquelles Alice envoie ses 60 000 impulsions à Bob. Ils ne conservent que les 50 000 premières environ, de façon à couvrir exactement l'intervalle $[0, 2\pi[$ de phases. Ces impulsions sont ensuite mélangées *a posteriori* par un traitement informatique.

De plus, Alice et Bob partagent la même cale piézoélectrique, qui module la phase de l'oscillateur local. La modulation de phase d'Alice et le choix de quadrature de Bob sont donc tous les deux séparés artificiellement *a posteriori*.

Tout cela serait extrêmement gênant si l'on voulait construire un vrai système de cryptographie quantique, mais ne change pas la structure des données. En effet, ce qui compte alors, c'est la phase relative entre l'oscillateur local et le signal. Comme, en plus, des modulateurs de phase rapides sont commercialement disponibles aux longueurs d'ondes télécom, la manière dont on crée le déphasage relatif entre le signal et l'oscillateur local importe peu pour une expérience de principe comme celle-ci.

13.2.3 Bob

Pour mesurer l'état d'Alice, Bob utilise la détection homodyne impulsionnelle dont nous avons décrit la construction dans la partie II. Nous n'y reviendrons pas plus en détail ici, si ce n'est pour préciser que la procédure que nous utilisons pour fixer la phase relative entre Alice et Bob est très loin d'être optimale et dure 1,6 s, soit environ 20 fois plus longtemps que la transmission réelle des blocs de données. Lors du calcul du débit effectif, nous négligerons ce temps d'asservissement, facilement améliorable.

13.2.4 Ève

Les pertes qui se produiraient lors d'une véritable communication à longue distance sont introduites par une lame de verre placée sur le faisceau. En changeant l'angle de la lame par rapport au faisceau, on peut changer le coefficient de transmission. Un espion, Ève, pourrait utiliser cette lame de verre dans un dispositif analogue à celui décrit dans la section 12.2.5.

13.3 Hypothèse sur les attaques d'Ève

Pour évaluer les attaques possibles d'Ève, nous n'utiliserons pas directement les équations (11.10) et (12.47) des chapitres précédents. En effet, nous avons supposé en établissant ces équations que tout le bruit ajouté χ était dû à l'action d'Ève ; or, comme nous l'avons vu section 6.1.3, les imperfections de notre détection homodyne contribuent à ce bruit : on a

$$\chi = \chi_{\text{ligne}} + \frac{\chi_{\text{hom}}}{G_{\text{ligne}}} \quad (13.37)$$

avec $\chi_{\text{hom}} = 0,60$ (ou $0,51$ dans le meilleur des cas). G_{ligne} et χ_{ligne} représentent respectivement le gain et le bruit ajouté sur la ligne. Utiliser directement cette valeur de χ revient à supposer qu'Ève peut s'intriquer avec les sources de bruit de la détection homodyne, ce que nous avons considéré comme excessivement *paranoïaque*.

Nous avons préféré une approche *réaliste*, en considérant que le bruit ajouté par la détection homodyne était complètement inconnu d'Ève. L'équation 11.10 devient alors

$$I_{AE} = \frac{1}{2} \log_2 \frac{1 + V\chi_{\text{ligne}}}{1 + \chi_{\text{ligne}}}. \quad (13.38)$$

Pour calculer I_{BE} , il suffit de voir que la variance conditionnelle vaut

$$V_{B|E} = \frac{1}{G_{\text{ligne}}(\chi_{\text{ligne}} + \frac{1}{V})} + \chi_{\text{hom}}. \quad (13.39)$$

L'équation (12.47) devient alors

$$I_{BE} = \frac{1}{2} \log_2 \frac{G_{\text{ligne}}^2 (V + \chi_{\text{ligne}} + \frac{\chi_{\text{hom}}}{G_{\text{ligne}}}) (\chi_{\text{ligne}} + \frac{1}{V})}{1 + G_{\text{ligne}}^2 \chi_{\text{hom}} (\chi_{\text{ligne}} + \frac{1}{V})}. \quad (13.40)$$

Contrairement à ce que la désignation de *réaliste* pourrait laisser croire, il est difficile d'imaginer un espion ayant tous ces pouvoirs : Ève est en effet supposée disposer d'un ordinateur infiniment puissant, d'une fibre optique sans pertes, de mémoires quantiques parfaites, de faisceaux arbitrairement intriqués, etc. Il est possible de restreindre encore plus les pouvoirs de l'espion et, probablement, d'améliorer les performances de notre expérience. Nous n'avons cependant pas étudié ces attaques limitées.

Nous avons représenté FIG. 13.2 les différents taux d'information mutuelles dans l'approche paranoïaque et l'approche réaliste avec notre système de détection homodyne⁴, pour une variance $V = 40$. Les points expérimentaux représentent ce que l'on a déduit de la mesure directe de la variance sur 60 000 points. Les barres d'erreurs sont dues pour moitié environ aux incertitudes statistiques et aux erreurs systématiques (calibration et dérives).

Il semble *a priori* possible avec notre système expérimental, en disposant d'algorithmes de réconciliation parfaits, de faire de la cryptographie quantique avec des états cohérents jusqu'à 70 % de pertes (5 dB) environ, les incertitudes statistiques sur l'évaluation du bruit permettant peu d'aller au delà.

⁴Dans le cas habituel de l'adaptation de mode de 96,5 %

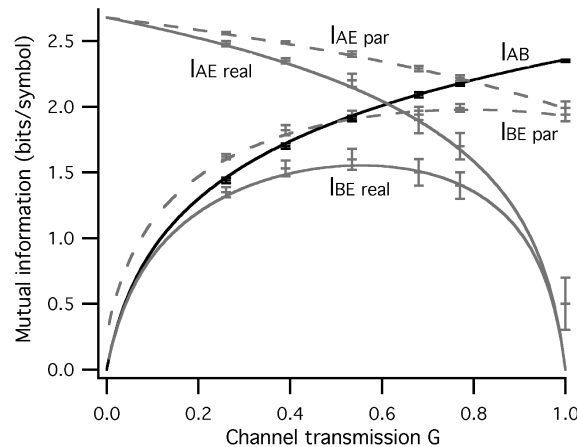


FIG. 13.2 – Taux d’informations mutuelles en fonction des pertes avec notre système de détection homodyne ($V = 40$)

13.4 Encore un peu d’informatique...

Alice et Bob doivent encore extraire une véritable clef secrète binaire à partir de leurs mesures corrélées. Cette étape exige un traitement informatique assez élaboré, qui a été fait par Gilles Van Assche et Nicolas Cerf de l’Université Libre de Bruxelles. Cette section a pour but de présenter rapidement cette étape informatique, qui sera traitée en détail dans [187].

13.4.1 Estimation

Pour estimer les paramètres du canal, et détecter l’espionnage éventuel d’Ève, Alice et Bob doivent sacrifier un certain nombre d’impulsions. Ils les choisissent aléatoirement, de manière à avoir un échantillon représentatif et imprévisible par Ève. Ils les comparent par le canal classique et ne doivent pas s’en servir pour générer leur clef secrète.

Nous avons triché lors de cette étape en évaluant ces paramètres sur la totalité des impulsions transmises.

13.4.2 Réconciliation par tranches

Une fois les paramètres du canal connu, il faut qu’Alice et Bob génèrent une clef binaire commune à partir de leur 36 800 variables gaussiennes corrélées, tout en révélant le moins possible d’information à l’espion. L’algorithme que nous utiliserons est l’algorithme de réconciliation par tranche, développé par Gilles Van Assche, Jean Cardinal et Nicolas Cerf, décrit en détail dans les références [183, 134]. Nous en présenterons ici rapidement son fonctionnement dans le cas de la réconciliation inverse.

Cet algorithme utilise comme élément de base un protocole de réconciliation binaire pour réconcilier successivement des tranches successives des variables gaussiennes. En pratique, la modulation d’Alice et la mesure de Bob sont codées sur quelques bits (ici 5). Les limites exactes des intervalles de cette binarisation, les tranches, sont précalculées par une optimisation numérique pour une dizaine de valeurs de χ (le gain G_{ligne} n’intervient ici que comme

un facteur d'échelle). La position des tranches est alors déterminée pour $\chi_{\text{ligne}}^{\text{est}} + \chi_{\text{hom}}$ par interpolation.

La première tranche sera ainsi constituée du bit le moins significatif de chaque impulsion, la seconde tranche du deuxième bit le moins significatif, et ainsi de suite jusqu'à la dernière tranche, qui est constituée par le bit le plus significatif. Comme Alice et Bob utilisent 36 800 impulsions par bloc, on aura donc 5 tranches de 36 800 bits chacune. Dans un protocole inverse, la clef binaire réconciliée sera donc constituée par les tranches de Bob, qu'Alice reconstituera.

Le principe général du protocole est le suivant, pour chaque tranche. Comme la tranche d'Alice n'est pas identique à celle de Bob, ils utilisent un protocole de réconciliation binaire pour corriger leurs erreurs. Même en supposant que ce protocole de réconciliation est parfait, l'information qui circule dans le canal classique peut être utilisée par Ève pour avoir une connaissance partielle de la clef. Il faut donc soigneusement tenir compte de ces fuites d'information afin de pouvoir l'éliminer ensuite, lors de l'étape d'amplification de confidentialité.

Une fois cette tranche réconciliée, Alice utilise alors les informations qu'elle a à sa disposition pour construire le meilleur estimateur de la tranche suivante de Bob. En d'autres termes, pour construire sa n -ième tranche, Alice utilise la connaissance qu'elle a des $n - 1$ tranches précédentes, ainsi que des valeurs $\{X_{A,k}\}_k$ de ses modulations. En pratique, les 2 ou 3 premières tranches, suivant la valeur de I_{AB} seront très peu corrélées entre elles, avec des taux d'erreurs proches de 50 %. L'algorithme de réconciliation utilisé sera alors l'algorithme de réconciliation trivial : Bob révélera directement la valeur de ses premières tranches à Alice (et Ève). Comme Alice possède une meilleure estimation *a priori* de X_B qu'Ève, cette information lui permet d'avoir une meilleure estimation de la tranche suivante qu'Ève.

D'une tranche à la suivante, tout se passe comme si on progressait de la gauche vers la droite le long des courbes de la FIG. 11.7⁵. Le taux d'erreur des tranches d'Alice chutera plus vite que ceux d'Ève, et donc la quantité d'information qu'elle apprend sur chaque tranche augmente plus vite que celle d'Ève.

Cet algorithme a une efficacité intrinsèque de l'ordre de 90 % de la limite de Shannon, et peut approcher arbitrairement de la limite de Shannon en utilisant des tranches multidimensionnelles [183]. Utilisé avec un véritable protocole de réconciliation binaire imparfait, comme *Cascade*, son efficacité est plus proche de 80 %.

13.4.3 *Cascade*

L'algorithme de réconciliation binaire que nous avons utilisé est *Cascade* [72, 70, 74]. Cet algorithme est fondé sur la comparaison de la parité de sous-blocs des clefs brutes (en l'occurrence les tranches) d'Alice et Bob, afin de trouver et de corriger les erreurs. L'implémentation de *Cascade* que nous utilisons a été optimisée par Kim-Chi Nguyen au cours de son travail de fin d'études. Une présentation plus détaillée de ce programme se trouve dans [74, 187].

Le choix des blocs dont Alice et Bob comparent la parité au cours de *Cascade* dépendent des différences de parités, c'est à dire des erreurs, révélées au fur et à mesure de l'avancement du protocole. En d'autres termes, si on note T_A et T_B les tranches respectives d'Alice

⁵Cette figure, ainsi que la section 7.4 sont d'ailleurs directement issues de réflexions sur le fonctionnement de cet algorithme de réconciliation par tranche.

et de Bob, et R la matrice rectangulaire qui spécifie les blocs dont la parité est comparée au cours de *Cascade*, R est construite ligne après ligne, la n -ième ligne de R dépendant des $n - 1$ premières lignes de R $T_A \oplus R T_B = R (T_A \oplus T_B)$. Cela permet à *Cascade* d'aller chercher les erreurs là où elles se trouvent en choisissant de ne redécouper pour analyse approfondie que les blocs où il y a des erreurs.

Cette interactivité de *Cascade* est à l'origine de son efficacité, mais pose un problème dans le cas des protocoles inverses. En effet, l'équation (12.1) n'est valable que si les protocoles de réconciliations sont unidirectionnels, c'est-à-dire si les communications entre Alice et Bob pour réconcilier leur clef ne révèlent d'information que sur l'un des deux. Comme les messages échangés par Alice et Bob sur le canal classique dépendent à la fois de X_A et X_B , ils risquent de permettre à Ève d'utiliser ses connaissances plus facilement acquises sur la modulation d'Alice pour exploiter au mieux les messages de correction d'erreur échangés entre Alice et Bob pour acquérir l'information qui lui manque.

Il nous faut donc évaluer l'information mutuelle $I_{K;E,M}$, où K représente la clef binaire obtenue à l'issue de *Cascade* ($K = T_B$ dans le cas de la réconciliation inverse), E les informations qu'Ève a pu obtenir lors du transfert physique des impulsions, et M les messages échangés entre Alice et Bob pendant le protocole de réconciliation. On a

$$I_{K;E,M} = I_{K;E} + I_{K;M|E}. \quad (13.41)$$

Le premier terme est borné par I_{BE} pour les protocoles de réconciliation inverse et par I_{AB} pour ceux de réconciliation directe. Le second terme désigne l'information supplémentaire apportée par les messages proprement dits. Il est bien entendu borné par $2l$, où l désigne le nombre de lignes de la matrice R , c'est-à-dire le nombre de parités qu'Alice et Bob comparent. Cette borne est beaucoup trop importante, et nous devons chercher à réduire ce nombre de bits révélés.

Une solution est alors de crypter les échanges d'Alice et de Bob par un code de Vernam, ou *one-time-pad*, en utilisant une partie de la clef générée précédemment [186]⁶. Alice et Bob peuvent ainsi utiliser une partie P de longueur l de la clef générée par le bloc précédent et communiquer les messages cryptés $R T_A \oplus P$ et $R T_B \oplus P$.

Appelons $\Delta \equiv T_A \oplus T_B$ la position des erreurs entre la tranche d'Alice et de Bob. Ève peut déduire Δ de l'écoute des communication classiques entre Alice et Bob. Pour ce faire, elle doit se débarrasser de la clef P en additionnant les deux messages :

$$(R T_A \oplus P) \oplus (R T_B \oplus P) = R T_A \oplus R T_B = R \Delta. \quad (13.42)$$

Elle peut ensuite déduire Δ de $R \Delta$. En effet, le bloc dont la parité est définie par la n -ième ligne de R est choisi en fonction des $n - 1$ parités précédentes, données par les $n - 1$ premiers coefficients. À la fin du protocole de réconciliation, Alice et Bob doivent connaître la position de toutes les erreurs. En d'autres termes, si on complète la matrice rectangulaire R par les autres parités linéairement indépendantes S , $S (T_A \oplus T_B) = 0$. Ève peut alors calculer

$$\Delta = \begin{bmatrix} R \\ S \end{bmatrix}^{-1} \begin{bmatrix} R \Delta \\ 0 \end{bmatrix}. \quad (13.43)$$

⁶Une méthode strictement équivalente en terme de débit net de clef serait de supprimer *a posteriori* par amplification de confidentialité le même nombre de bits. Nous l'exposerons ici en terme de cryptage pour deux raisons : c'est la solution que nous avons adoptée, et elle nous semble plus claire exposée ainsi.

Réciproquement, toute l'information connue par Ève est déterminée par Δ , car R est une fonction de Δ et que toute l'information disponible sur $R T_A$ ou $R T_B$ est contenue dans $R \Delta$, le reste étant protégé par le cryptage.

On a donc $I(K; M|E) = I(K; \Delta|E)$. Si *Cascade* est utilisé pour protéger BB84, $I(K; \Delta|E) = 0$ car Ève apprend déjà Δ en utilisant l'attaque optimale. Cela n'est malheureusement pas le cas pour les protocoles à variables continues et nous devons alors évaluer $I(K; \Delta, E)$ numériquement en effectuant l'intégrale

$$I(K; \Delta|E) = \int d\mathbf{e} \mathcal{P}_{E=\mathbf{e}} I(K; \Delta|E). \quad (13.44)$$

Cette intégrale a été évaluée dans le cas de la cloneuse intriquante décrite section 12.2.5. E représente alors la gaussienne bivariée constituée par le champ X_{connu} injecté par Ève et celui qu'elle mesure, X_{E2} . Pour chaque valeur de $E = \mathbf{e} = \begin{pmatrix} x_{\text{connu}} \\ x_{E2} \end{pmatrix}$, Ève peut déduire une distribution de probabilité gaussienne bivariée pour les valeurs de la modulation d'Alice X_A et de la mesure de Bob X_B . Or, le k -ième bit de $K = T_B$ est une fonction discrète qui ne dépend que de la valeur de la mesure de Bob de la k -ième impulsion $X_{B,k}$ et le k -ième bit de Δ en est une qui ne dépend que de $X_{A,k}$ et $X_{B,k}$, associés à la k -ième impulsion. On peut donc calculer leurs distributions de probabilité conditionnées à $E_k = \mathbf{e}$ et en déduire $I_{K;\Delta|E}$.

Malheureusement, Δ apporte à Ève de l'information concernant B , ce qui se traduit par une dégradation significative des performances de *Cascade* avec les pertes dans le cas de la réconciliation inverse. Nous avons alors envisagé d'utiliser un autre protocole de cryptage, où Alice révèle $R A$ sans aucun cryptage et Bob crypte son message $R B \oplus B$. Ce protocole est plus proche d'un protocole unidirectionnel, et devrait être meilleur. Malheureusement, l'évaluation de $I_{K;R A|E}$ est beaucoup plus complexe que la précédente : la valeur de chaque bit $R A$ est une fonction non triviale de toutes les valeurs de $\{X_{A,k}, X_{B,k}\}_{k=1}^{36\,800}$. Et l'intégrale qu'il faudrait calculer pour évaluer $I_{K;R A|E}$ n'est plus une intégrale à deux dimensions, mais une intégrale à $2 \times 36\,800$ dimensions !

Une autre possibilité pour rendre le protocole de réconciliation globalement plus unidirectionnel serait d'augmenter le nombre de dimensions des tranches en « mélangeant » des impulsions successives au niveau du protocole de réconciliation par tranche [183].

13.4.4 Amplification de confidentialité

Une fois la réconciliation achevée, Alice et Bob disposent d'une clef commune K d'1 bit par tranche et par impulsion. Sur cette clef, Ève peut connaître $I_{K;E} \leq I_{BE}$ bits qu'elle a déduit directement des résultats de sa cloneuse intriquante et $I_{K;\Delta|E}$ bits qu'elle a obtenus en écoutant le canal classique, soit $I_{K;E,\Delta}$ bits au total.

Alice et Bob doivent donc se débarrasser des bits connus d'Ève, au moyen d'une procédure d'amplification de confidentialité [73, 70]. Ils commencent par simplement jeter les deux ou trois premières tranches, qui ont été entièrement révélées au début du processus de réconciliation.

Pour « nettoyer » les autres tranches, partiellement connues d'Ève, ils doivent les mélanger en choisissant aléatoirement une fonction au sein d'une classe universelle de fonctions de hachage [191]. Ce mélange a en quelque sorte pour rôle de diluer les bits connus d'Ève en les répartissant uniformément dans la clef. Si Alice et Bob jettent ensuite $I(E; K, \Delta) + \mathcal{S}$

V	transmission		I_{BA} (bits)	I_{BE} (%) I_{BA}	I_{rec} (%) I_{BA}	Taux d'information (kbps)			
	(%)	(dB)				inverse		direct	
						idéal	réel	idéal	réel
41,7	100	0	2,39	0	88	1 920	1 690	1 910	1 660
38,6	79	-1,0	2,17	58	85	730	470	540	270
32,3	68	-1,7	1,93	67	79	510	185	190	—
27,0	49	-3,1	1,66	72	78	370	75	0	—
43,7	26	-5,9	1,48	93	71	85	—	0	—

TAB. 13.2 – Taux d'information secrète idéal et pratique

bits, où s est un paramètre de sécurité, Ève n'a plus qu'une probabilité 2^{-s} , essentiellement nulle, de connaître un bit de la clef secrète ainsi générée.

La classe universelle de fonctions de hachage généralement choisie est fondée sur l'utilisation de matrices aléatoires. Cette démarche nous a semblé peu pratique en raison de la grande taille des matrices dont elle nécessitait la manipulation.

Nous avons préféré utiliser la multiplication tronquée sur un corps fini (parfois appelé *champ de Galois*), en l'occurrence $GF(2^{110\,503})$, qui sera représenté par le corps des polynômes à coefficients binaires modulo le polynôme irréductible $p[X] = X^{110\,503} + X^{519} + 1$ [192]. Ces fonctions ont été choisies car la multiplication des polynômes peut être implémentée efficacement [193] au moyen de FFT de nombres entiers. Le polynôme p devait donc répondre à trois critères : il devait être irréductible, son degré doit être le plus proche possible d'une puissance de 2 par valeurs inférieures et devait être de l'ordre de 10 000 au moins pour pouvoir s'appliquer à un nombre intéressant de bits.

Le protocole d'amplification de confidentialité est alors le suivant : Alice et Bob regroupent les tranches à amplifier, issues de la réconciliation, en une chaîne de bits de longueur inférieure ou égale au degré du polynôme, 110 503. La taille conséquente de ce polynôme nous a autorisé à traiter en une seule étape 3 tranches de 36 834 bits ou 2 tranches de 55 251 bits. Ils choisissent ensuite, publiquement et aléatoirement, un polynôme sur dans $GF(2^{110\,503})$ et le multiplient modulo p au polynôme correspondant à la chaîne de bits réconciliés.

La clef secrète est le résultat de cette multiplication, amputée du nombre de bits déterminé plus haut. Avant d'utiliser ces bits, ils ne doivent pas oublier d'en conserver une partie pour crypter *Cascade* lors de la réconciliation du bloc suivant.

13.5 ...et Alice et Bob peuvent enfin discuter tranquillement

Après toutes ces opérations, Alice et Bob disposent (enfin !) d'une clef secrète, qu'ils peuvent utiliser pour leurs communications.

Nous avons représenté sur le TAB. 13.2 les taux d'information secrète⁷ obtenus pour différentes valeurs des pertes, en les comparant aux taux idéaux déduits de l'équation (13.40), qui auraient été obtenus avec des protocoles de réconciliations parfaits. Les valeurs équivalentes obtenues par réconciliation directe sont également indiquées pour comparaison.

Le protocole de réconciliation réellement implémenté a une efficacité limitée, qui limite la quantité nette d'information obtenue par Alice au cours de la réconciliation à 80 % environ

⁷Comme nous l'avons mentionné plus haut, le temps d'1,5 s utilisé entre deux blocs pour l'asservissement de la phase n'est pas compris dans ces estimations.

de la valeur de I_{BA} donnée par la limite de Shannon. Si cette efficacité, relativement élevée, permet des débits d'information secrète de plusieurs centaines de kilobits par seconde aux faibles pertes, elle limite les pertes que ces protocoles peuvent tolérer car Ève dispose d'un algorithme efficace à 100 %.

Si cette efficacité décroît lorsque les pertes augmentent, cela peut être compensé en réduisant la variance V de la modulation, ce qui augmente le rapport $\frac{I_{BA}}{I_{BE}}$. Cette perte d'efficacité de la réconciliation est nettement plus forte pour les protocoles de réconciliation inverse que pour les protocoles directs, pour les raisons mentionnées dans la section 13.4.3. Cependant, l'avantage quantique initial de ces protocoles, manifeste dans lorsqu'on compare les taux d'informations donnés par la formule de Shannon, n'est jamais rattrapé, et les protocoles inverses restent plus efficaces que les protocoles directs.

Nous avons ainsi pu obtenir un taux net de génération clef secrète de 75 kbits/s à 3,1 dB de pertes, montrant ainsi que cette limite, que beaucoup considéraient comme fondamentale il y a moins d'un an, pouvait être franchie par des protocoles réels. Si les données du point expérimental à 3,1 dB ont été acquises dans de très bonnes conditions, avec notamment une adaptation de modes de 99 % nous avons généré beaucoup de clefs autour de 3 dB de pertes, avec des taux d'information secrète typiquement au dessus de 55 kilobits/s.

Chapitre 14

Conclusion

Nous avons élaboré au cours de cette thèse un système de cryptographie quantique complet avec des variables continues, ce qui nous a permis d'aborder une variété surprenante de problèmes, tous reliés à la problématique de ce domaine en plein essor que sont les communications quantiques avec des variables continues.

Nous avons ainsi commencé cette thèse par la mise au point d'une détection homodyne impulsionnelle rapide, limitée au bruit de photons, élément indispensable à de nombreux protocoles de communication quantique.

Nous avons parallèlement étudié sur un plan théorique la signification physique des critères de téléportation quantique et leurs liens avec le clonage quantique. Cette étude du clonage quantique nous a naturellement conduit, en suivant Nicolas Cerf et ses collaborateurs, à la cryptographie quantique [133, 134]. Leur approche, fondée sur la théorie de l'information de Shannon, permet en effet de calculer les débits d'information secrète à partir de simples calculs de rapport signal à bruit. Si cela ne fournissait pas directement un protocole exploitable, cela a permis de diviser la partie théorique de la conception d'un protocole de cryptographie quantique avec des variables continues en deux parties bien distinctes que je désignerais pour simplifier sous les noms de partie *quantique* et de partie *informatique*.

La partie *quantique* du protocole consiste à utiliser les lois de la mécanique quantique pour calculer des rapports signal sur bruit et des variances conditionnelles, avec des méthodes utilisées depuis longtemps dans l'étude des mesures quantiques non-destructives, et d'utiliser ces grandeurs pour calculer des taux d'information mutuelles avec la formule de Shannon. Cette approche nous a permis d'étendre les protocoles de Cerf aux états cohérents, puis aux fortes pertes, deux domaines que l'on pensait difficilement accessibles aux protocoles de cryptographie quantique avec des variables continues. Cette « division du travail » nous a permis de proposer ces protocoles et d'évaluer leurs performances sans avoir la moindre idée de la manière concrète d'en extraire des clefs cryptographiques utiles !

La partie *informatique* du protocole concerne bien entendu cette extraction des clefs. Cette partie n'était en général pas clairement séparée de la précédente, ce qui conduisait à des protocoles peu efficaces. Pendant que nous généralisions la partie *quantique* de leurs protocoles, Gilles Van Assche et Nicolas Cerf ont développé des protocoles de réconciliation adaptés aux protocoles à variables continues utilisant une modulation gaussienne, ainsi qu'un protocole d'amplification de confidentialité efficace.

Avec la détection homodyne, qui constitue la troisième partie, la partie *optique*, indispensable à toute réalisation expérimentale, il devenait possible de construire la première démonstration expérimentale complète d'un système de cryptographie quantique avec des

variables continues [14, 15], d'autant plus que la partie optique était grandement simplifiée par la possibilité d'utiliser des états cohérents. Cette première démonstration expérimentale, qui franchit également la « limite » des 3 dB avec un taux conséquent (75 kbits/s) de génération de clef secrète, notamment grâce aux performances de notre détection homodyne, ouvre la voie pour des systèmes pratiques.

En particulier, la mise en œuvre d'un démonstrateur complet aux longueurs d'ondes télécom (1 550 nm) est actuellement à l'étude en collaboration avec la société *Thalès*.

Cinquième partie

Annexes

Annexe A

Formules mathématiques

A.1 Algèbre d'opérateurs

Les propriétés algébriques des opérateurs présentées ici sont démontrées plus en détail dans [194].

A.1.1 Propriétés élémentaires

L'équation qui suit se démontre aisément.

$$[\hat{A}, \hat{B}\hat{C}] = [\hat{A}, \hat{B}] \hat{C} + \hat{B} [\hat{A}, \hat{C}] \quad (\text{A.1})$$

La plupart des démonstration de cette section ne sont valable que pour les opérateurs qui commutent avec leurs commutateur, c'est à dire qui vérifient

$$[\hat{A}, [\hat{A}, \hat{B}]] = 0 \quad (\text{A.2a})$$

$$[\hat{B}, [\hat{A}, \hat{B}]] = 0. \quad (\text{A.2b})$$

Les opérateurs considérés dans cette thèse vérifient en général ces deux relations car leur commutateur est souvent scalaire.

A.1.2 Commutateurs de fonctions d'opérateurs

Nous nous placerons ici dans le cas où $\hat{B}$ commute avec le commutateur $[\hat{A}, \hat{B}]$. Nous démontrerons d'abord par récurrence que

$$[\hat{A}, \hat{B}^n] = n\hat{B}^{n-1} [\hat{A}, \hat{B}]. \quad (\text{A.3})$$

Cette propriété est évidemment vérifiée pour $n = 1$. Si on la suppose vérifiée pour n , la relation (A.1) nous permet de calculer le $n + 1$ ème terme :

$$[\hat{A}, \hat{B}^{n+1}] = [\hat{A}, \hat{B}\hat{B}^n] = [\hat{A}, \hat{B}] \hat{B}^n + \hat{B} [\hat{A}, \hat{B}^n] = [\hat{A}, \hat{B}] \hat{B}^n + n\hat{B}\hat{B}^{n-1} [\hat{A}, \hat{B}]. \quad (\text{A.4})$$

La condition (A.2b) nous permet alors de finir le calcul

$$[\hat{A}, \hat{B}^{n+1}] = (n + 1)\hat{B}^n [\hat{A}, \hat{B}]. \quad (\text{A.5})$$

On a donc démontré la relation (A.3) par récurrence.

On peut se servir du calcul précédent pour établir la valeur du commutateur $[\hat{A}, f(\hat{B})]$, où $f(B)$ est une fonction analytique quelconque. Si on appelle f_n les coefficients du développement en série entière de f , on a

$$[\hat{A}, f(\hat{B})] = \left[\hat{A}, \sum_n f_n \hat{B}^n \right] = \sum_n f_n [\hat{A}, \hat{B}^n] = \sum_n n f_n \hat{B}^{n-1} [\hat{A}, \hat{B}]. \quad (\text{A.6})$$

On reconnaît ici le développement en séries entière de $f'(\hat{B})$, où f' est la dérivée de f . On a donc

$$[\hat{A}, f(\hat{B})] = f'(\hat{B}) [\hat{A}, \hat{B}], \quad (\text{A.7})$$

sous la condition (A.2b).

A.1.3 Formule de Baker–Hausdorff

Dans le cas d’observables commutant avec leur commutateur, la formule de Baker–Hausdorff

$$e^{\hat{A}+\hat{B}} = e^{-\frac{1}{2}[\hat{A}, \hat{B}]} e^{\hat{A}} e^{\hat{B}} \quad (\text{A.8})$$

est vérifiée. Elle nous sera très utile.

A.2 Gaussiennes et paquets d’ondes

A.2.1 Intégrale d’une gaussienne

Une fonction gaussienne centrée s’écrit sous la forme

$$g(x) = \mathcal{Z} e^{-a x^2}. \quad (\text{A.9})$$

Si on appelle I l’intégrale de $g(x)$, on a

$$I \equiv \int dx \mathcal{Z} e^{-a x^2} \quad (\text{A.10a})$$

$$I^2 = \int dx dy \mathcal{Z}^2 e^{-a(x^2+y^2)} \quad (\text{A.10b})$$

$$= \mathcal{Z}^2 \int_0^\infty dr \int_0^{2\pi} r d\theta e^{-a r^2} \quad (\text{A.10c})$$

$$= \mathcal{Z}^2 \underbrace{\int_0^{2\pi} d\theta}_{2\pi} \underbrace{\int_0^\infty dr r e^{-a r^2}}_{\left[-\frac{1}{2a} e^{-a r^2}\right]_0^\infty} \quad (\text{A.10d})$$

Lorsque $\mathcal{R}(a) > 0$, la valeur à l’infini du crochet est nulle et l’intégrale se calcule aisément, sinon, l’intégrale diverge. Dans ce cas, on a

$$I \equiv \int dx \mathcal{Z} e^{-a x^2} = \mathcal{Z} \sqrt{\frac{\pi}{\mathcal{R}(a)}} \quad (\text{A.11})$$

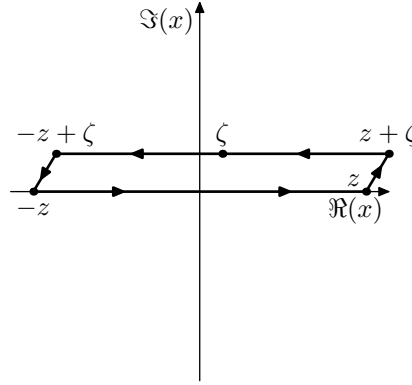


FIG. A.1 – Contour d'intégration dans le plan complexe

A.2.2 Définition et intégrale d'un paquet d'onde gaussien

Un paquet d'onde gaussien est défini par une fonction de la forme

$$h(X = x) = \mathcal{Z} e^{-\frac{x^2}{2\sigma^2} + \zeta x}, \quad (\text{A.12})$$

où $\mathcal{Z}$ et ζ sont des constantes complexes et σ est réel. Cette fonction représente une onde de nombre d'onde $\mathcal{I}(\zeta)$ dans une enveloppe gaussienne centrée en $X = \sigma^2 \mathcal{R}(\zeta)$ et d'écart-type σ .

Pour calculer l'intégrale de cette fonction, nous pouvons la récrire sous la forme

$$\int dx \mathcal{Z} e^{-\frac{x^2}{2\sigma^2} + \zeta x} = \mathcal{Z} e^{\frac{\sigma^2 \zeta^2}{2}} \int dx e^{-\frac{(x - \sigma^2 \zeta)^2}{2\sigma^2}}. \quad (\text{A.13})$$

Un simple changement de variable ne suffit pas à ramener l'intégrale ci-dessus à l'intégrale (A.11). En effet, le paramètre ζ peut être complexe. Cependant, la fonction $g(X)$ définie en (A.9) est holomorphe et ne présente aucun pôle. Toute intégration de $g(X)$ sur un circuit fermé du plan complexe est donc nulle. L'intégrale de $g(X)$ sur le quadrilatère défini par les quatre nombres complexes $(-z, z, z - \sigma^2 \zeta, -z - \sigma^2 \zeta)$, avec $z \in \mathbb{R}$ et $z \rightarrow \infty$ est donc nulle. Comme $g(X)$ tend vers zéro lorsque la partie réelle de X tend vers $\pm\infty$, l'intégrale sur les « petits » segments $[z, z - \sigma^2 \zeta]$ et $[-z - \sigma^2 \zeta, -z]$ tend donc vers zéro et l'intégrale des deux « grands » segments peut s'écrire

$$\int_{-z}^z dx e^{-\frac{x^2}{2\sigma^2}} + \int_{z - \sigma^2 \zeta}^{-z - \sigma^2 \zeta} dx e^{-\frac{x^2}{2\sigma^2}}, \xrightarrow{z \rightarrow \infty} 0 \quad (\text{A.14})$$

où les intégrales sont supposées prises en ligne droite dans le plan complexe. On a donc égalité entre les intégrales

$$\int dx e^{-\frac{(x - \sigma^2 \zeta)^2}{2\sigma^2}} = \int dx e^{-\frac{x^2}{2\sigma^2}} = \sqrt{2\pi \sigma^2} \quad (\text{A.15})$$

et on peut calculer l'intégrale (A.13) :

$$\int dx \mathcal{Z} e^{-\frac{x^2}{2\sigma^2} + \zeta x} = \mathcal{Z} e^{\frac{\sigma^2 \zeta^2}{2}} \sqrt{2\pi \sigma^2} \quad (\text{A.16})$$

A.2.3 Paquets d'ondes à variance complexe

Le même raisonnement que précédemment permet de calculer l'intégrale

$$\int dx \mathcal{Z} e^{-ax^2 + \zeta x} = \mathcal{Z} e^{-\frac{\zeta^2}{4a}} \sqrt{\frac{\pi}{\mathcal{R}(a)}}, \quad (\text{A.17})$$

si a est complexe de partie réelle positive.

A.3 Transformées de Fourier

A.3.1 Convention

Les transformées de Fourier seront notées avec des crochets autour de l'argument et les fonctions dans l'espace réel avec des parenthèses. La convention utilisée pour placer les facteurs 2π est la suivante :

$$f(x) = \frac{1}{\sqrt{2\pi}} \int d\kappa e^{i\kappa x} f[\kappa] \quad (\text{A.18a})$$

$$f[\kappa] = \frac{1}{\sqrt{2\pi}} \int dx e^{-i\kappa x} f(x). \quad (\text{A.18b})$$

Les variables de l'espace réel seront notées autant que possible par des lettres latine, alors que celles de l'espace de Fourier le seront par des lettres grecque. Ces conventions s'appliquent aussi aux fonctions caractéristiques qui ne sont rien d'autre que des transformées de Fourier de distributions de probabilité.

Les transformées de Fourier de fonction d'onde seront notées avec les conventions (B.30) et (B.32), légèrement différentes, tenant compte du facteur d'échelle N_0 . L'origine de ces conventions est donnée dans la section B.5, page 206.

A.3.2 Fonctions de Dirac et ondes planes

La transformée de Fourier d'une Fonction de Dirac se calcule aisément :

$$\frac{1}{\sqrt{2\pi}} \int dx e^{-i\kappa x} \delta(x - x_0) = \frac{e^{-i\kappa x_0}}{\sqrt{2\pi}}, \quad (\text{A.19})$$

d'où on déduit

$$\delta(x - x_0) = \frac{1}{\sqrt{2\pi}} \int d\kappa e^{i\kappa x} \frac{e^{-i\kappa x_0}}{\sqrt{2\pi}} = \frac{1}{2\pi} \int d\kappa e^{i\kappa(x-x_0)} \quad (\text{A.20})$$

A.3.3 Paquets d'onde gaussiens

La transformée de Fourier d'un paquet d'onde gaussien $h(X)$ défini en (A.12) s'écrit simplement sous la forme d'une intégrale gaussienne, qu'on peut calculer en utilisant l'égalité

(A.16)

$$f[K = \kappa] = \frac{\mathcal{Z}}{\sqrt{2\pi}} \int dx e^{-\frac{x^2}{2\sigma^2} + (\zeta - i\kappa)x} \quad (\text{A.21a})$$

$$= \frac{\mathcal{Z} e^{\frac{\sigma^2}{2}(\zeta - i\kappa)^2}}{\sqrt{2\pi}} \sqrt{2\pi\sigma^2} = \mathcal{Z} \sigma e^{\frac{\sigma^2\zeta^2}{2}} e^{-\frac{\sigma^2\kappa^2}{2} - i\sigma^2\zeta\kappa} \quad (\text{A.21b})$$

Cette transformée de Fourier est un paquet d'onde centré en $\mathcal{I}(\zeta)$, d'écart-type $\frac{1}{\sigma}$ et de nombre d'onde $-\sigma^2\mathcal{R}(\zeta)$.

A.3.4 Produit scalaire

Le produit scalaire usuel dans l'espace des fonctions de carré sommable sera noté avec la convention de Dirac :

$$\langle f | g \rangle \equiv \int dx f^*(x) g(x) \quad (\text{A.22a})$$

$$= \iint d\kappa d\lambda f^*[\kappa] g[\lambda] \frac{1}{2\pi} \int dx e^{i(\lambda - \kappa)x} \quad (\text{A.22b})$$

$$= \iint d\kappa d\lambda f^*[\kappa] g[\lambda] \delta(\lambda - \kappa) \quad (\text{A.22c})$$

$$= \int d\kappa f^*[\kappa] g[\kappa] \quad (\text{A.22d})$$

Annexe B

Propriétés quantiques élémentaires des variables continues

Le but de cette annexe est de retrouver les propriétés des observables P et Q , essentiellement à partir de la relation de commutation (2.3). Les démonstrations de cette section viennent essentiellement des références [195, 196].

B.1 Choix des observables

L'équation (2.8) montre clairement que, dans le cas de l'oscillateur harmonique, le choix arbitraire de l'origine des temps t_0 définit la quadrature observée. Plus généralement, on peut toujours redéfinir les observables P_θ et Q_θ par une rotation dans l'espace des phases :

$$\begin{bmatrix} Q_\theta \\ P_\theta \end{bmatrix} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} Q \\ P \end{bmatrix}. \quad (\text{B.1})$$

Le changement de variable réciproque s'obtient évidemment en changeant θ en $-\theta$ dans les équations ci-dessus. Il est facile de vérifier que

$$[\hat{Q}_\theta, \hat{P}_\theta] = [\hat{Q}, \hat{P}] \cos^2 \theta - [\hat{P}, \hat{Q}] \sin^2 \theta = [\hat{Q}, \hat{P}], \quad (\text{B.2})$$

et que les observables $\hat{Q}_\theta$ et $\hat{P}_\theta$ ont donc exactement les mêmes propriétés que $\hat{Q}$ et $\hat{P}$. Le choix de l'angle θ est donc uniquement une commodité expérimentale, purement arbitraire du point de vu théorique. Par exemple, le choix $\theta = \frac{\pi}{2}$ intervertit, à un signe près, les rôles de Q et de P .

B.2 États propres de $\hat{Q}$ et fonction d'onde

Les états propres des opérateurs $\hat{Q}$ et $\hat{P}$ associés aux valeurs propres q et p seront notés $|Q = q\rangle$ et $|P = p\rangle$, ou $|q\rangle$ et $|p\rangle$ en l'absence d'ambiguïté. On a donc, par définition,

$$\hat{Q} |q\rangle = q |q\rangle \quad (\text{B.3})$$

$$\hat{P} |p\rangle = p |p\rangle, \quad (\text{B.4})$$

En multipliant l'équation (B.3) à gauche par un autre vecteur propre $\langle q'|$ et en développant le produit $\langle q'|\hat{Q}|q\rangle$ de deux manières différentes, on a

$$\langle q'|\hat{Q}|q\rangle = \langle q'|\hat{Q}|q\rangle \quad (\text{B.5})$$

$$\langle q'|q'\rangle = \langle q'|q\rangle \quad (\text{B.6})$$

$$(q' - q) \langle q'|q\rangle = 0. \quad (\text{B.7})$$

Comme nous le verrons section B.7, la relation de commutation (2.3) assure que le spectre de $\hat{Q}$ est continu et recouvre $\mathbb{R}$ de $-\infty$ à $+\infty$. L'équation précédente a donc pour solution

$$\langle q'|q\rangle = \delta(q' - q) \quad (\text{B.8})$$

Les états propres de $\hat{Q}$ correspondant à des positions différentes sont donc orthogonaux entre eux. Ces états ont un statut un peu particulier : ils constituent une base d'états très pratique et très utilisée. Malheureusement, on a

$$\langle q|q\rangle = \delta(0) \neq 1 : \quad (\text{B.9})$$

à cause de la fonction de Dirac, ils ne sont pas normalisables et ne représentent pas de « vrais » états physiques.

Le commutateur assure aussi [195] que, si l'opérateur $\hat{Q}$ n'est pas dégénéré pour une valeur propre, alors il n'est dégénéré pour aucune valeur propre. Dans ce cas, l'ensemble des états propres de l'opérateur position constitue une base complète de l'ensemble des états :

$$\int dq |q\rangle \langle q| = \hat{\mathbb{I}}, \quad (\text{B.10})$$

où $\hat{\mathbb{I}}$ représente l'opérateur identité. Sinon, cet ensemble constitue une base complète d'un sous espace de Hilbert.

On peut ainsi définir la fonction d'onde $\psi(Q)$ associée à un état $|\psi\rangle$ par le produit scalaire

$$\psi(q) \equiv \langle q|\psi\rangle \quad (\text{B.11})$$

et en déduire qu'un état est parfaitement défini par sa fonction d'onde :

$$|\psi\rangle = \int dq |q\rangle \langle q|\psi\rangle = \int dq \psi(q) |q\rangle. \quad (\text{B.12})$$

En multipliant cette dernière équation à gauche par $\langle\psi|$, on obtient la condition de normalisation des fonctions d'ondes :

$$1 = \langle\psi|\psi\rangle = \int dq \psi^*(q) \psi(q) = \int dq |\psi(q)|^2. \quad (\text{B.13})$$

L'ensemble des états quantiques est donc l'espace de Hilbert des fonctions de carré sommable.

Bien sûr, ce que nous avons dit ici pour Q s'applique aussi à Q_θ et à toutes les variables continues.

B.3 Action de l'opérateur impulsion

Pour calculer l'action de l'opérateur impulsion sur un état quelconque $|\psi\rangle$ de fonction d'onde $\psi(Q)$, nous allons d'abord calculer son action sur un état propre $|q\rangle$ de la position.

Il nous faut donc calculer $\langle q' | \hat{P} | q \rangle$. Or on a

$$\langle q' | [\hat{Q}, \hat{P}] | q \rangle = \langle q' | (\hat{Q}\hat{P} - \hat{P}\hat{Q}) | q \rangle = (q' - q) \langle q' | \hat{P} | q \rangle \quad (\text{B.14})$$

$$\langle q' | [\hat{Q}, \hat{P}] | q \rangle = 2iN_0 \langle q' | q \rangle. \quad (\text{B.15})$$

En comparant ces deux équations et en tenant compte du produit scalaire (B.8), on obtient

$$\langle q' | \hat{P} | q \rangle = \frac{2iN_0 \delta(q - q')}{q' - q}. \quad (\text{B.16})$$

Cette quantité est donc nulle si $q' \neq q$ et est singulière dans le voisinage $q' \simeq q$.

Elle nous servira à calculer la fonction d'onde $\langle q | \hat{P} | \psi \rangle$ de l'état $\hat{P} | \psi \rangle$. En effet,

$$\langle q | \hat{P} | \psi \rangle = \int dq' \langle q | \hat{P} | q' \rangle \langle q' | \psi \rangle = \int dq' \frac{2iN_0 \delta(q' - q)}{q - q'} \psi(q'). \quad (\text{B.17})$$

Or l'intégrande est nul partout, sauf dans le voisinage de $q' = q$ où l'on a

$$\psi(q') = \psi(q) + \left[\frac{d}{dq} \psi(q) + \varepsilon(q' - q) \right] (q' - q), \quad (\text{B.18})$$

où $\varepsilon(q' - q)$ est une fonction qui tend vers zéro lorsque q' tend vers q . En injectant ce développement limité dans (B.17), on a

$$\begin{aligned} \langle q | \hat{P} | \psi \rangle = 2iN_0 \left[\psi(q) \int dq' \frac{\delta(q - q')}{q - q'} - \frac{d}{dq} \psi(q) \int dq' \delta(q - q') \right. \\ \left. + \int dq' \delta(q - q') \varepsilon(q' - q) \right]. \quad (\text{B.19}) \end{aligned}$$

Le premier terme est nul, car l'intégrande est impaire et le troisième l'est car on a $\varepsilon(q - q) = 0$. On a donc

$$\langle q | \hat{P} | \psi \rangle = -2iN_0 \frac{d}{dq} \psi(q) = -2iN_0 \frac{d}{dq} \langle q | \psi \rangle. \quad (\text{B.20})$$

On peut donc exprimer l'action de l'opérateur impulsion comme une dérivée de la fonction d'onde par rapport à la position.

$$\hat{P} = \frac{2N_0}{i} \frac{d}{dq}. \quad (\text{B.21})$$

B.4 Fonction d'onde de $|P\rangle = p$

Cette expression de l'opérateur $\hat{P}$ nous permettra d'exprimer la fonction d'onde $\langle q | p \rangle$ de ses vecteurs propres. On a en effet

$$p \langle q | p \rangle = \langle q | \hat{P} | p \rangle = \frac{2N_0}{i} \frac{d}{dq} \langle q | p \rangle. \quad (\text{B.22})$$

Cette équation différentielle se résout simplement et on a

$$\langle q | p \rangle = \mathcal{Z} e^{\frac{ipq}{2N_0}}. \quad (\text{B.23})$$

Reste à calculer le facteur de normalisation $\mathcal{Z}$. En combinant l'expression du produit scalaire (B.8) et la relation de complétude $\hat{\mathbb{I}} = \int dp |p\rangle \langle p|$, on a

$$\delta(q' - q) = \langle q' | q \rangle = \int dp \langle q' | p \rangle \langle p | q \rangle = |\mathcal{Z}|^2 \int dp e^{\frac{ip(q' - q)}{2N_0}} \quad (\text{B.24})$$

Or la définition intégrale (A.20) de la fonction de Dirac, après le changement de variable $\kappa \rightarrow \frac{p}{2N_0}$, devient

$$\delta(q' - q) = \frac{1}{4\pi N_0} \int dp e^{\frac{ip(q' - q)}{2N_0}}. \quad (\text{B.25})$$

On en déduit donc

$$\mathcal{Z} = \frac{e^{i\varphi}}{\sqrt{4\pi N_0}}. \quad (\text{B.26})$$

En laissant tomber le facteur de phase $e^{i\varphi}$ arbitraire, on a donc

$$\langle q | p \rangle = \frac{1}{\sqrt{4\pi N_0}} e^{\frac{ipq}{2N_0}}. \quad (\text{B.27})$$

Les états propre $|p\rangle$ de l'impulsion peuvent donc s'écrire

$$|p\rangle = \frac{1}{\sqrt{4\pi N_0}} \int dq e^{\frac{ipq}{2N_0}} |q\rangle. \quad (\text{B.28})$$

Ce sont des ondes planes de nombre d'onde $\frac{p}{2N_0}$ et le longueur d'onde $\frac{4\pi N_0}{p}$. Si Q et P représentent la position et l'impulsion d'une particule, ces ondes sont les ondes de de Broglie associées à cette particule. Comme leur amplitude est constante de $-\infty$ à $+\infty$, ces états ne sont pas plus normalisables que les états $|q\rangle$.

B.5 Relations de Fourier entre impulsion et position

La base des états propres de l'impulsion $|p\rangle$ est donc une base d'ondes planes. Si on développe un état $|\psi\rangle$ sur cette base, on aura donc la transformée de Fourier $\psi[P]$ de sa fonction d'onde. En d'autres termes sa fonction d'onde en impulsion $\psi[P]$ est la transformée de Fourier de sa fonction d'onde en position $\psi(Q)$.

En effet, on définit la fonction d'onde en impulsion par

$$\psi[P = p] \equiv \langle P = p | \psi \rangle. \quad (\text{B.29})$$

Si on intercale la relation de complétude (B.10), on a

$$\psi[p] = \int dq \langle p | q \rangle \langle q | \psi \rangle = \frac{1}{\sqrt{4\pi N_0}} \int dq e^{\frac{ipq}{2N_0}} \psi(q) \quad (\text{B.30})$$

Cette relation est identique à celle définissant les transformées de Fourier (A.18b), au changement de variable $\kappa = \frac{p}{2N_0}$ et au facteur $\frac{1}{\sqrt{2N_0}}$ près. La transformée de Fourier inverse (A.18a) devient alors

$$\frac{1}{\sqrt{2N_0}}\psi(q) = \frac{1}{\sqrt{2\pi}} \int \frac{dp}{2N_0} e^{-\frac{ipq}{2N_0}} \psi[p] \quad (\text{B.31})$$

$$\psi(q) = \frac{1}{\sqrt{4\pi N_0}} \int dp e^{-\frac{ipq}{2N_0}} \psi[p] \quad (\text{B.32})$$

Les équations (B.30) et (B.32) montrent bien que la fonction d'onde en impulsion $\psi[P]$ est la transformée de Fourier de la fonction d'onde en position $\psi(Q)$, à un facteur et un changement de variables près. Ce facteur sera sous-entendu dans la suite et l'on dira simplement que ces deux fonctions sont transformées de Fourier l'une de l'autre.

B.6 Relation d'incertitude

La relation d'incertitude de Heisenberg peut se comprendre comme conséquence directe des relations de Fourier (B.30)–(B.32) [36]. Elle peut également se démontrer directement à partir de la relation de commutation (2.3), et c'est cette démonstration, plus générale, que nous présenterons ici.

On peut considérer, pour tout état quantique $|\psi\rangle$ et pour tout $\lambda \in \mathbb{R}$, l'opérateur $\hat{\Lambda} \equiv \hat{P} + i\lambda\hat{Q}$ et son action sur $|\psi\rangle$. On a alors

$$\|\hat{\Lambda}|\psi\rangle\|^2 = \|(\hat{P} + i\lambda\hat{Q})|\psi\rangle\|^2 \quad (\text{B.33a})$$

$$= \langle \psi | (\hat{P} - i\lambda\hat{Q}) (\hat{P} + i\lambda\hat{Q}) | \psi \rangle \quad (\text{B.33b})$$

$$= \langle \psi | \hat{P}^2 | \psi \rangle + i\lambda \langle \psi | [\hat{P}, \hat{Q}] | \psi \rangle + \lambda^2 \langle \psi | \hat{Q}^2 | \psi \rangle \quad (\text{B.33c})$$

Il est aisé de se convaincre que le commutateur $[\hat{P}, \hat{Q}]$ ne dépend pas des valeurs moyennes $\langle \psi | \hat{P} | \psi \rangle$ et $\langle \psi | \hat{Q} | \psi \rangle$ de ces observables. On peut donc redéfinir $\hat{\Lambda}'$ à partir des observables $\hat{P}' \equiv \hat{P} - \langle \psi | \hat{P} | \psi \rangle$ et $\hat{Q}'$ centrées en zéro. L'équation (B.33c) peut alors se récrire

$$\|\hat{\Lambda}'|\psi\rangle\|^2 = \Delta P^2 + i\lambda \langle \psi | [\hat{P}, \hat{Q}] | \psi \rangle + \lambda^2 \Delta Q^2, \quad (\text{B.33d})$$

où ΔP et ΔQ sont les écarts types des observables $\hat{P}$ et $\hat{Q}$. Comme $\hat{P}$ et $\hat{Q}$ sont hermitiens, leur commutateur est antihermitien et sa valeur moyenne est donc imaginaire pure. Le côté droit de l'équation (B.33d) est bien réel. Le côté gauche étant une norme, il doit être positif ou nul quelque soit λ . Le discriminant Δ du trinôme en λ doit donc être négatif ou nul :

$$\Delta = |\langle \psi | [\hat{P}, \hat{Q}] | \psi \rangle|^2 - 4\Delta P^2 \Delta Q^2 \leq 0. \quad (\text{B.34})$$

Cette dernière équation se récrit simplement sous la forme de la relation d'incertitude

$$\Delta P \Delta Q \geq \frac{|\langle \psi | [\hat{P}, \hat{Q}] | \psi \rangle|}{2}. \quad (\text{B.35})$$

Lorsque le commutateur $[\hat{Q}, \hat{P}]$ est scalaire, suivant la relation (2.3), l'inégalité (B.35) ne dépend pas de l'état quantique $|\psi\rangle$. On a alors la relation d'incertitude de Heisenberg

$$\Delta Q \Delta P \geq N_0, \quad (\text{B.36})$$

qui n'est autre que (2.4).

B.7 Opérateur translation et spectre de $\hat{Q}$

Les relations de Fourier entre l'impulsion et la position nous permettent de définir pour tout $x \in \mathbb{R}$ l'opérateur

$$\hat{T}_x \equiv e^{-\frac{ix\hat{p}}{2N_0}}. \quad (\text{B.37})$$

qui effectue une translation de x . Cet opérateur nous permettra, par des translations arbitraires, de montrer que les valeurs propres de $\hat{Q}$ recouvrent complètement $\mathbb{R}$. Comme cette propriété était essentielle pour établir le produit scalaire (B.8), d'où découlent, entre autres, les relations de Fourier entre la position et l'impulsion, nous étudierons les propriétés de l'opérateur translation $\hat{T}_x$ en nous appuyant uniquement sur la relation de commutation (2.3), en suivant [195].

Il est aisé de vérifier la relation

$$\hat{T}_x \hat{T}_y = \hat{T}_{x+y}. \quad (\text{B.38})$$

$\hat{T}_x$ est donc unitaire, car $\hat{T}_x^\dagger = \hat{T}_{-x} = \hat{T}_x^{-1}$.

De plus, l'équation (A.7) nous permet de calculer le commutateur

$$[\hat{Q}, \hat{T}_x] = -\frac{ix}{2N_0} \hat{T}_x [\hat{Q}, \hat{p}] = x \hat{T}_x \quad (\text{B.39})$$

Cette relation nous permettra de montrer que l'état $\hat{T}_x |q\rangle$ est un état propre de l'opérateur position, de valeur propre $(q + x)$. En effet

$$\hat{Q} \hat{T}_x |q\rangle = \hat{T}_x \hat{Q} |q\rangle + [\hat{Q}, \hat{T}_x] |q\rangle = (q + x) \hat{T}_x |q\rangle. \quad (\text{B.40})$$

On peut ainsi définir tous les états propres de l'opérateur $\hat{Q}$ à partir de l'un d'entre eux :

$$\forall q \in \mathbb{R}, \quad |Q = q\rangle \equiv \hat{T}_q |Q = 0\rangle. \quad (\text{B.41})$$

On a donc bien démontré à partir de la seule relation de commutation (2.3) que le spectre de $\hat{Q}$ était continu et recouvrait $\mathbb{R}$.

Les relations

$$\hat{T}_x |q\rangle = |q + x\rangle \quad (\text{B.42})$$

$$\langle q | \hat{T}_x = \langle q | \hat{T}_{-x}^\dagger = \langle q - x | \quad (\text{B.43})$$

nous permettrons de calculer l'effet de l'opérateur $\hat{T}_x$ sur un état quelconque $|\psi\rangle$. La fonction d'onde de l'état $\hat{T}_x |\psi\rangle$ vaut en effet

$$\langle q | \hat{T}_x | \psi \rangle = \langle q - x | \psi \rangle = \psi(q - x). \quad (\text{B.44})$$

C'est la fonction d'onde de $|\psi\rangle$ translatée de $+x$, d'où le nom d'*opérateur translation* donné à $\hat{T}_x$.

B.8 Opérateur déplacement

Le même raisonnement s'applique pour $\hat{P}$ et permet de montrer que l'opérateur

$$\hat{T}'_p \equiv e^{i\frac{p\hat{Q}}{2N_0}} \quad (\text{B.45})$$

effectue une translation en impulsion de p , c'est à dire qu'il décale la fonction d'onde en P de tout état quantique de p . Il est donc très tentant de définir l'opérateur déplacement $\hat{D}_\alpha$ qui décalerait un état dans l'espace des phase de l'amplitude $\alpha \equiv \frac{1}{\sqrt{4N_0}}(q + ip)$ par l'opérateur

$$\hat{D}_{\frac{1}{\sqrt{4N_0}}(q+ip)} \equiv e^{\frac{i}{2N_0}(q\hat{P}-p\hat{Q})}. \quad (\text{B.46})$$

Le facteur $\frac{1}{\sqrt{4N_0}}$, arbitraire ici, se justifiera lorsque nous introduiront l'opérateur amplitude $\hat{a}$ dans la section suivante. La formule de Baker–Hausdorff (A.8) nous permet de calculer l'exponentielle :

$$\hat{D}_{\frac{1}{\sqrt{4N_0}}(q+ip)} = e^{-\frac{1}{2}\left[\frac{iq\hat{P}}{2N_0}, -\frac{ip\hat{Q}}{2N_0}\right]} e^{-i\frac{q\hat{P}}{2N_0}} e^{i\frac{p\hat{Q}}{2N_0}} \quad (\text{B.47})$$

$$= e^{i\frac{qp}{4N_0}} \hat{T}_q \hat{T}'_p = e^{-i\frac{qp}{4N_0}} \hat{T}'_p \hat{T}_q. \quad (\text{B.48})$$

C'est bien la combinaison d'une translation en Q et d'une translation en P , à un facteur de phase près, qui est dû à la non-commutation de ses translations.

Les formules (B.48) permettent de vérifier aisément l'unitarité de cet opérateur, puisque $\hat{D}_\alpha^\dagger = \hat{D}_{-\alpha} = \hat{D}_\alpha^{-1}$.

B.9 Opérateur amplitude

Malgré le principe d'incertitude, on peut définir l'opérateur $\hat{a}$ associé à l'amplitude complexe $\hat{Q} + i\hat{P}$ d'un état quantique. Cependant, cet opérateur n'est pas hermitien et n'est donc pas une observable. Il est défini par

$$\hat{a} = \frac{\hat{Q} + i\hat{P}}{\sqrt{4N_0}}. \quad (\text{B.49})$$

Le facteur de normalisation $\frac{1}{\sqrt{4N_0}}$ assure la relation de commutation

$$[\hat{a}, \hat{a}^\dagger] = 2\frac{-i[\hat{Q}, \hat{P}]}{4N_0} = 1. \quad (\text{B.50})$$

On peut bien sûr récrire les opérateurs position $\hat{Q}$ et impulsion $\hat{P}$ à partir de l'opérateur amplitude $\hat{a}$ et de son hermitien conjugué $\hat{a}^\dagger$:

$$\hat{Q} = \sqrt{N_0}(\hat{a} + \hat{a}^\dagger) \quad (\text{B.51})$$

$$\hat{P} = -i\sqrt{N_0}(\hat{a} - \hat{a}^\dagger). \quad (\text{B.52})$$

Il est aisé de se convaincre que les opérateurs $\hat{Q}_\theta$ et $\hat{P}_\theta$ définis en B.1 correspondent à un opérateur d'amplitude déphasé de θ :

$$\hat{a}_\theta = \hat{a} e^{i\theta}. \quad (\text{B.53})$$

Comme Le même raisonnement que celui de la section 2.1.3, nous conduit à l'expression de l'opérateur $\hat{R}(\theta)$ correspondant à une rotation dans l'espace des phases :

$$\hat{R}(\theta) = e^{-i\theta \frac{\hat{H}}{\hbar\omega}} = e^{-\frac{i}{4N_0} \left(\frac{\hat{Q}^2}{2} + \frac{\hat{P}^2}{2} \right)}, \quad (\text{B.54})$$

où $\hat{H}$ est l'hamiltonien (2.6) de l'oscillateur harmonique.

L'hamiltonien (2.6) de l'oscillateur harmonique s'écrit d'une manière simple en fonction des opérateurs $\hat{a}$ et $\hat{a}^\dagger$:

$$\hat{H} = \frac{\hbar\omega}{4N_0} N_0 \{ (\hat{a} + \hat{a}^\dagger)^2 - (\hat{a} - \hat{a}^\dagger)^2 \} \quad (\text{B.55})$$

$$= \frac{\hbar\omega}{2} (\hat{a} \hat{a}^\dagger + \hat{a}^\dagger \hat{a}) \quad (\text{B.56})$$

$$= \hbar\omega \left(\hat{a}^\dagger \hat{a} + \frac{1}{2} \right). \quad (\text{B.57})$$

L'opérateur $\hat{n} \equiv \hat{a}^\dagger \hat{a}$ qui intervient dans l'expression de cet hamiltonien est appelé opérateur *nombre de photons* si Q et P sont les quadratures d'un mode du champ électromagnétique. Il est appelé *nombre de phonons* si l'oscillateur harmonique est un oscillateur mécanique. Ces dénominations, ainsi que les appellations d'opérateur *annihilation* et *création* pour $\hat{a}$ et $\hat{a}^\dagger$ seront justifiées section 2.2.5.

L'opérateur déplacement défini par l'équation (B.46) s'écrit alors

$$\hat{D}_\alpha = e^{i \frac{N_0}{2N_0} \{ (\alpha + \alpha^*) i(\hat{a}^\dagger - \hat{a}) + i(\alpha - \alpha^*) (\hat{a} + \hat{a}^\dagger) \}} \quad (\text{B.58})$$

$$= e^{\alpha^* \hat{a} - \alpha \hat{a}^\dagger}. \quad (\text{B.59})$$

La formule de Baker–Hausdorf (A.8) permet d'écrire $\hat{D}_\alpha$ sous une forme qui nous sera plus utile :

$$\hat{D}_\alpha = e^{\frac{|\alpha|^2}{2}} e^{\alpha^* \hat{a}} e^{\alpha \hat{a}^\dagger} = e^{-\frac{|\alpha|^2}{2}} e^{\alpha \hat{a}^\dagger} e^{\alpha^* \hat{a}}. \quad (\text{B.60})$$

On peut en déduire l'effet de l'opérateur déplacement sur l'opérateur $\hat{a}$:

$$\hat{D}_\alpha^\dagger \hat{a} \hat{D}_\alpha = \hat{D}_{-\alpha} \hat{a} \hat{D}_\alpha \quad (\text{B.61})$$

$$= e^{-\frac{|\alpha|^2}{2}} e^{-\alpha \hat{a}^\dagger} e^{-\alpha^* \hat{a}} \hat{a} e^{\frac{|\alpha|^2}{2}} e^{\alpha^* \hat{a}} e^{\alpha \hat{a}^\dagger} \quad (\text{B.62})$$

$$= e^{-\alpha \hat{a}^\dagger} \hat{a} e^{\alpha \hat{a}^\dagger} \quad (\text{B.63})$$

$$= e^{-\alpha \hat{a}^\dagger} \left(\left[\hat{a}, e^{\alpha \hat{a}^\dagger} \right] + e^{\alpha \hat{a}^\dagger} \hat{a} \right). \quad (\text{B.64})$$

Le commutateur se calcule avec la formule (A.7) et on obtient

$$\hat{D}_\alpha^\dagger \hat{a} \hat{D}_\alpha = \hat{a} + \alpha. \quad (\text{B.65})$$

L'opérateur déplacement $\hat{D}_\alpha$ décale donc bien l'amplitude $\hat{a}$ de α .

Bibliographie

- [1] Adán Cabello. Bibliographic guide to the foundations of quantum mechanics and quantum information. *E-print quant - ph/0012089v5*, March 2002. Cette collection de 6 569 (!) références bibliographiques a été bien utile pour établir la présente bibliographie.
- [2] Counterpane. Index of cryptographic papers available online. <http://www.counterpane.com/biblio>.
- [3] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Mécanique Quantique*. Hermann, 1977.
- [4] Jean-Louis Basdevant and Jean Dalibard. *Mécanique Quantique*. École Polytechnique, 1995. Est également édité chez Ellipse.
- [5] Wolfgang P. Schleich. *Quantum Optics in Phase Space*. Wiley-VCH, 2001.
- [6] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Numerical Recipes in C. The Art of Scientific Computing*. Cambridge University Press, seconde édition, 1992. Disponible à <http://www.nr.com>.
- [7] Frédéric Grosshans and Philippe Grangier. Effective quantum efficiency in the pulsed homodyne detection of a n-photon state. *The European Physical Journal D*, 14(1) :119–125, April 2001.
- [8] Philippe Grangier and Frédéric Grosshans. Quantum teleportation criteria for continuous variables. *E-print quant - ph/0009079*, September 2000.
- [9] Philippe Grangier and Frédéric Grosshans. Evaluating quantum teleportation of coherent states. *E-print quant - ph/0010107*, October 2000.
- [10] Frédéric Grosshans and Philippe Grangier. Quantum cloning and teleportation criteria for continuous quantum variables. *Physical Review A (Rapid Communication)*, 64(1) :010301(R), June 2001. *E-print quant - ph/0012121*.
- [11] Frédéric Grosshans and Philippe Grangier. Continuous variable quantum cryptography using coherent states. *Physical Review Letters*, 88(5) :057902, January 2002. [12] en est une version plus longue.
- [12] Frédéric Grosshans and Philippe Grangier. Continuous variable quantum cryptography using coherent states. *E-print quant - ph/0109084*, September 2001. Version étendue de [11].
- [13] Frédéric Grosshans and Philippe Grangier. Reverse reconciliation protocols for quantum cryptography with continuous variables. In Jeffrey H. Shapiro and Osamu Hirota, editors, *Proceedings of the 6th International Conference on Quantum Communication, Measurement, and Computing (QCMC'02, Boston, 22–26 juillet 2002)*. Rinton Press, December 2002. *E-print quant - ph/0204127*.

- [14] Frédéric Grosshans, Gilles Van Assche, Jérôme Wenger, Rosa Brouri, Nicolas J. Cerf, and Philippe Grangier. High-rate quantum key distribution using gaussian-modulated coherent states. *Nature*, 421(6920) :238–241, January 2003. Complété par [15].
- [15] Frédéric Grosshans, Gilles Van Assche, Jérôme Wenger, Rosa Brouri, Nicolas J. Cerf, and Philippe Grangier. Supplementary information for “high-rate quantum key distribution using gaussian-modulated coherent states”. *Nature*, 2003. <http://www.nature.com/nature/journal/v421/n6920/supinfo/nature01289.html>. Complète [14].
- [16] Jérôme Wenger, Mohammad Hafezi, Frédéric Grosshans, Rosa Brouri, and Philippe Grangier. Maximal violation of Bell inequalities using continuous variables measurements. *Physical Review A*, 67(1) :012105, January 2003. *E-print quant - ph/0211067*.
- [17] Jean-Louis Basdevant and Roland Omnès. *Historique de la mécanique quantique*, chapter XXII, pages 431–450. Volume II of [4], 1995.
- [18] Sven Ortoli and Jean-Pierre Pharabod. *Le Cantique des Quantiques*. Number 4066 in Le Livre de Poche biblio, essais. Éditions La Découverte, 1984.
- [19] Max Tegmark and John Archibald Wheeler. 100 ans de mystères quantiques. *Pour la Science*, 282 :82–90, April 2001. Traduction française de [21].
- [20] David C. Cassidy. Werner Heisenberg et le principe d’incertitude. *Pour la Science*, 177 :86–92, July 1992.
- [21] Max Tegmark and John Archibald Wheeler. 100 years of the quantum. *Scientific American*, 284 :68–75, 2001. [19] en est une traduction française. [22] en est une version plus longue.
- [22] Max Tegmark and John Archibald Wheeler. 100 years of the quantum. *E-print quant - ph/0101077*, January 2001. Version longue (« director’s cut ») de [21].
- [23] George Gamow. *M. Tompkins au Pays des Merveilles, Histoire de c, G et h*. Dunod, 1953. Traduction de [25].
- [24] George Gamow. *M. Tompkins explore l’atome*. Dunod, 1954. Traduction de [26].
- [25] George Gamow. *Mr. Tompkins in Wonderland*. Cambridge University Press, 1939.
- [26] George Gamow. *Mr. Tompkins explores the atom*. Cambridge University Press, 1944.
- [27] Giulio Peruzzi. La constante de Planck. *Pour la Science*, 279 :14–16, January 2001.
- [28] Louis Leprince-Ringuet. *Cours de Physique, Tome II*. École Polytechnique, 1^e Division, 1940–1941.
- [29] Louis de Broglie. Waves and quanta. *Nature*, 112 :540, October 1923. Disponible à <http://www.nature.com/physics/lookingback/debroglie/>.
- [30] C. Davisson and L. H. Germer. The scattering of electrons by a single crystal of nickel. *Nature*, 119(2998) :558–560, April 1927. Disponible à <http://www.nature.com/physics/lookingback/davisson/>.
- [31] Josette Rey-Debove and Alain Rey, editors. *Le nouveau petit Robert, dictionnaire alphabétique et analogique de la langue française*. Dictionnaires le Robert, 1996.
- [32] Anatole Bailly. *Dictionnaire Grec–Français*. Librairie Hachette, 26^{ème} édition, 1963.

- [33] David C. Cassidy. L'expérience de pensée du microscope à rayons gamma. *Pour la Science*, 177 :91, July 1992. Encadré dans l'article [20].
- [34] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Ondes et particules. Introduction aux idées fondamentales de la mécanique quantique*, chapter I, pages 7–90. Volume I of [3], 1977.
- [35] Jean-Louis Basdevant and Jean Dalibard. *Mécanique ondulatoire I : Fonction d'onde, équation de Schrödinger*, chapter I, pages 23–67. Volume I of [4], 1995.
- [36] Jean-Michel Bony. *Théorie des distributions et analyse de Fourier*. École Polytechnique, Palaiseau (France), 1997.
- [37] Albert Einstein, B. Podolsky, and Nathan Rosen. Can quantum mechanical description of physical reality be considered complete? *Physical Review*, 47(10) :777–780, 1935.
- [38] John S. Bell. On the Einstein–Podolsky–Rosen paradox. *Physics*, 1 :195–200, 1964. Réédité dans [39, pages 14–21].
- [39] John S. Bell. *Speakable and unspeakable in Quantum Mechanics. Collected papers on quantum philosophy*. Cambridge University Press, 1993. Recueil d'articles publiés entre 1964 et 1986.
- [40] Michel Bitbol. *L'élision. Essai sur la philosophie de Schrödinger*, essai liminaire, pages 9–149. In [41], 1990.
- [41] Erwin Schrödinger. *L'esprit et la matière*. Éditions du Seuil, 1990. Traduction française de [42].
- [42] Erwin Schrödinger. *Mind and Matter*. Cambridge University Press, 1958. [41] en est une traduction française.
- [43] Hervé Zwirn. Du quantique au classique. *Pour la Science*, 182 :38–44, December 1992.
- [44] John Preskill. Quantum information and computation. Lecture notes for physics 229, California Institute of Technology, September 1998. Disponible à <http://www.theory.caltrch.edu/people/preskill/ph229>.
- [45] Jean-Paul Delahaye. Les ordinateurs quantiques. *Pour la Science*, 208 :100–104, February 1995.
- [46] Jean-Paul Delahaye. Les lois nouvelles de l'information quantique. *Pour la Science*, 250 :66–72, August 1998.
- [47] Nicolas J. Cerf and Nicolas Gisin. Les promesses de l'information quantique. *La Recherche*, 327 :46–53, January 2000.
- [48] Jérôme Ségala. *Théorie de l'information : sciences, techniques et société de la seconde guerre mondiale à l'aube du XXI^e siècle*. Thèse de doctorat, Université Lumière Lyon 2, 1998. Disponible à <http://theses.univ-lyon2.fr/Theses/jsegal/tm.html> et <http://www.mpiwg-berlin.mpg.de/staff/segal/thesis/>.
- [49] Jérôme Ségala. Le géomètre de l'information. *Pour la Science*, 295 :26–29, May 2002.
- [50] Claude Elwood Shannon. A mathematical theory of communication. *Bell System Technical Journal*, 27 :379–423 et 623–656, July et October 1948. Réédité à <http://cm.bell-labs.com/cm/ms/what/shannonday/paper.html>.
- [51] Richard Feynmann. Simulating physics with computers. *International Journal of Theoretical Physics*, 21(6–7) :467–488, 1982.

- [52] David Deutsch. Quantum theory, the church-turing principle and the universal quantum computer. *Proceedings of the Royal Society of London A*, 400 :97–117, 1985. Disponible à <http://www.qubit.org/oldsite/resource/deutsch85.pdf>.
- [53] Peter W. P. Shor. Polynomial time algorithms for prime factorization and discrete logarithm on a quantum computer. In *Proceedings of the 35th Annual Symposium on Foundations of Computer Science, Santa Fe*, pages 124–139. IEEE Computer Society Press, 1994. E-print quant - ph/9508027v2.
- [54] Lov K. Grover. A fast quantum mechanical algorithm for databas search. In *Proceedings of 28th Annual ACM Symposium on Theory of Computiong (STOC)*, pages 212–219, 1996. Disponible à <http://www.bell-labs.com/user/lkgrover>.
- [55] Lov K. Grover. Quantum mechanics helps in searching for a needle in a haystack. *Physical Review Letters*, 78(2) :325–328, July 1997. Disponible à <http://www.bell-labs.com/user/lkgrover>.
- [56] Lieven S. K. Vandersypen, Matthias Steffen, Gregory Breyta, Costantino S. Yannoni, Mark H. Sherwood, and Isaac L. Chuang. Experimental realization of Shor’s quantum factoring algorithm using nuclear magnetic resonance. *Nature*, 414 :883–887, December 2001.
- [57] François Morain. La factorisation d’entiers. In *Pour la Science* [58], pages 62–64.
- [58] L’art du secret : la cryptographie. *Pour la Science*, Dossier hors-série 36, July–October 2002.
- [59] Nicolas Schlosser. *Étude et réalisation de micro-pièges dipolaires optiques pour atomes neutres*. Thèse de doctorat, Université Paris XI, December 2001.
- [60] Georges-Olivier Reymond. *Etudes expérimentales d’atome dans un piège dipolaire microscopique*. Thèse de doctorat, Université Paris XI, November 2002.
- [61] Nicolas Schlosser, Georges-Olivier Reymond, Igor Protsenko, and Philippe Grangier. Sub-poissonian loading of single atoms in a microscopic dipole trap. *Nature*, 411 :1024–1027, 2001.
- [62] Igor Protsenko, Georges-Olivier Reymond, Nicolas Schlosser, and Philippe Grangier. Operation of quantum phase gate using neutral atoms in microscopic dipol traps. *Physical Review A*, 52(5) :052301, April 2002. E-print quant - ph/0206007.
- [63] Juan Ignacio Cirac and Peter Zoller. A scalable quantum computer with ions in an array of microtaps. *Nature*, 404(6778), April 2000.
- [64] Jean-Paul Delahaye. Cryptographie quantique. *Pour la Science*, 178 :101–106, August 1992.
- [65] Stephen Wiesner. Conjugate coding. *Sigact News*, 15(1) :78–88, 1983. Rédigé vers 1969–1970, cet article ne fut publié qu’en 1983.
- [66] Charles H. Bennett and Gilles Brassard. Quantum cryptography : public-key distribution and coin tossing. In *Proceedings of the IEEE International Conference on Computers, Systems and Signal Processing*, pages 175–179. IEEE, 1984.
- [67] Charles H. Bennett, Gilles Brassard, and Artur Ekert. La cryptographie quantique. [58], pages 114–117.
- [68] Nicolas Gisin, Grégoire Ribordy, Wolfgang Tittel, and Hugo Zbinden. Quantum cryptography. *Review of Modern Physics*, 74 :174, 2002. E-print quant - ph/0101098v2.

- [69] Philippe Grangier, John Rarity, and Anders Karlsson. *The European Physical Journal D*, 18(2 (Special Issue on Quantum interference and cryptographic keys : novel physics and advancing technologies (QUICK))), February 2002.
- [70] Claude Crépeau. Réconciliation et distillation publiques de secret. Disponible à <http://www.cs.mcgill.ca/~crepeau/theses.html>, 1995.
- [71] Ueli M. Maurer. Secret key agreement by public discussion from common information. *IEEE Transactions on Information Theory*, 39 :733–742, 1993. Disponible à <http://www.crypto.ethz.ch/~maurer/publications.html>.
- [72] Gilles Brassard and Louis Salvail. Secret-key reconciliation by public discussion. In *Advances in cryptology – Eurocrypt’93*, number 765 in Lecture Notes in Computer Science, pages 411–423, New-York, 1993. Springer Verlag.
- [73] Charles H. Bennett, Gilles Brassard, Claude Crépeau, and Ueli M. Maurer. Generalized privacy amplification. *IEEE Transactions on Information Theory*, 41(6) :1915–1935, November 1995. Disponible à <http://www.crypto.ethz.ch/~maurer/publications.html>.
- [74] Kim-Chi Nguyen. Extension des protocoles de réconciliation en cryptographie quantique. Travail de fin d’études, Université Libre de Bruxelles, 2002.
- [75] Charles H. Bennett and Stephen Wiesner. Communication via one- and two-particle operators on Einstein-Podolsky-Rosen states. *Physical Review Letters*, 69(20) :2881–2884, November 1992.
- [76] Charles H. Bennett, Gilles Brassard, Claude Crépeau, Richard Jozsa, Asher Peres, and William K. Wootters. Teleporting an unknown quantum state via dual classical and Einstein-Podolsky-Rosen channels. *Physical Review Letters*, 70(13) :1895–1899, 1993.
- [77] Natalia Korolkova, Christine Silberhorn, Oliver Glöckl, Stefan Lorenz, Ch. Marquadt, and Gerd Leuchs. Direct experimental test of non-separability and other quantum techniques using continuous variables of light. In *The European Physical Journal D* [69], pages 229–235. *E-print quant-ph/0109011v2*.
- [78] Natalia Korolkova, Gerd Leuchs, Rodney Loudon, Timothy C. Ralph, and Christine Silberhorn. Polarisation squeezing and continuous variable polarization entanglement. *Physical Review A*, 65 :052306, 2002. *E-print quant-ph/0108098*.
- [79] M. J. Holland, M. J. Colett, D. F. Walls, and M. D. Levenson. Nonideal quantum non-demolition measurements. *Physical Review A*, 42(5) :2995–3005, September 1990.
- [80] Philippe Grangier, Jean-Michel Courty, and Serge Reynaud. Characterization of nonideal quantum non-demolition measurements. *Optics Communications*, 89(1) :99–106, April 1992.
- [81] Juan-Ariel Levenson, Izo Abram, T. Rivera, P. Fayolle, J. C. Garreau, and Philippe Grangier. Quantum optical cloning amplifier. *Physical Review Letters*, 70(3) :267–270, January 1993.
- [82] Karine Vigneron. *Contrôle du bruit quantique de la lumière et mesures quantiques non destructives utilisant des atomes piégés et refroidis*. Thèse de doctorat, Université Paris XI, 1998.
- [83] Jean-François Roch, Karine Vigneron, Philippe Grelu, Alice Sinatra, Jean-Philippe Poizat, and Philippe Grangier. Quantum nondemolition measurement using cold trapped atoms. *Physical Review Letters*, 78 :634–637, 1997.

- [84] Jean-Philippe Poizat, Jean-François Roch, and Philippe Grangier. Characterization of quantum non-demolition measurements in optics. *Annales de Physique* (Paris), 19 :365, 1994.
- [85] Philippe Grangier, Juan-Ariel Levenson, and Jean-Philippe Poizat. Quantum non-demolition measurement in optics. *Nature*, 396(6711) :537–542, 1998.
- [86] Elisabeth Giacobino and Claude Fabre. Quantum noise reduction in parametric oscillators and laser. *Annales de Physique* (Paris), 20(5–6) :509–516, October–December 1995.
- [87] Jean-Philippe Poizat. *Bruit quantique des diodes lasers*. Habilitation à diriger des recherches, Université Paris XI, December 1999.
- [88] Jean-Philippe Poizat, Tiejun Chang, and Philippe Grangier. Quantum intensity noise of laserdiodes and non-orthogonal spatial eigenmodes. *Physical Review A*, 61(043807), 2000.
- [89] Lev Vaidman. Teleportation of quantum states. *Physical Review A*, 49(2) :1473–1476, 1994.
- [90] Samuel L. Braunstein and H. Jeff Kimble. Teleportation of continuous quantum variables. *Physical Review Letters*, 80(4) :869–872, 1998.
- [91] A. Furusawa, J. Sørensen, Samuel L. Braunstein, Christopher A. Fuchs, H. Jeff Kimble, and Eugene S. Polzik. Unconditional quantum teleportation. *Science*, 282(5389) :706–709, 1998.
- [92] B. Lounis and W. E. Moerner. Single photons on demand from a single molecule at room temperature. *Nature*, 407 :491, 2000.
- [93] Charles Santori, David Fattal, Jelena Vučković, Glenn S. Solomon, and Yoshihisa Yamamoto. Indistinguishable photons from a single-photon device. *Nature*, 419 :594, 2002.
- [94] Alexios Beveratos. *Réalisation expérimentale d’une source de photons uniques par fluorescence de centres colorés individuels dans le diamant ; application à la cryptographie quantique*. Thèse de doctorat, Université Paris XI, December 2002.
- [95] Alexios Beveratos, Sergei Kühn, Rosa Brouri, Thierry Gacoin, Jean-Philippe Poizat, and Philippe Grangier. Room temperature stable single photon source. In *The European Physical Journal D* [69], page 191. *E-print quant - ph/0110176*.
- [96] Alexios Beveratos, Rosa Brouri, Thierry Gacoin, André Villing, Jean-Philippe Poizat, and Philippe Grangier. Single photon quantum cryptography. *Physical Review Letters*, 89 :187901, 2002. *E-print quant - ph/0206136*.
- [97] Ulf Leonhardt. *Measuring the Quantum State of Light*. Cambridge Studies in Modern Optics. Cambridge University Press, 1997.
- [98] Pierre Bérest. *Calcul des variations. Application à la mécanique et à la physique*. École Polytechnique, Palaiseau (France), 1996.
- [99] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Lagrangien et hamiltonien en mécanique classique*, appendice III, pages 1473–1489. Volume II of [3], 1977.
- [100] Jean-Louis Basdevant and Jean Dalibard. *Lagrangien et hamiltonien ; force de Lorentz en mécanique quantique.*, chapter XIV, pages 261–271. Volume I of [4], 1995.

- [101] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Points de vue de Schrödinger et de Heisenberg*, complément B_{III}, pages 311–313. Volume I of [3], 1977.
- [102] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *L'oscillateur harmonique à une dimension*, chapter V, pages 479–643. Volume I of [3], 1977.
- [103] Gilbert Grynberg and Claude Fabre. *Optique Quantique*. École Polytechnique, 1998. Est également édité chez Ellipse.
- [104] Wolfgang P. Schleich. *Field States*, chapter 9, pages 291–319. In [5], 2001.
- [105] Wolfgang P. Schleich. *Field Quantization*, chapter 10, pages 255–290. In [5], 2001.
- [106] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Étude des états stationnaires en représentation $\{|x\rangle\}$. Polynômes d'Hermite*, complément B_V. Volume I of [3], 1977.
- [107] J. Bertrand and P. Bertrand. A tomographic approach to wigner's function. *Foundations of Physics*, 17(4) :397–405, April 1987.
- [108] Jacques Neveu. *Probabilités*. École Polytechnique, Palaiseau (France), 1996.
- [109] Benjamin Schumacher. Sending entanglement through noisy quantum channels. *Physical Review A*, 54(4), October 1996.
- [110] R. L. Hudson. When is the Wigner quasi-probability density non-negative ? *Rep. Math. Phys.*, 6 :249–252, 1974.
- [111] C. Piquet. Fonctions de type positif associées à deux opérateurs hermitiens. *Comptes-Rendus de l'Académie des Sciences (Paris)*, 279 A :107–109, 1974.
- [112] Patrice Bertet, Alexia Auffeves, Paolo Maioli, Stefano Osnaghi, Tristan Meunier, Michel Brune, Jean-Michel Raimond, and Serge Haroche. Direct measurement of the Wigner function of a one-photon Fock state in a cavity. *Physical Review Letters*, 89(20) :200402, October 2002.
- [113] Hauke Hansen. *Generation and Characterization of New Quantum States of the Light Field*. Dissertation zur erlangung des akademischen grades doctor rerum naturalium (dr. rer. nat.), Fachbereich Physik der Universität Konstanz, April 2000.
- [114] Alexander I. Lvovsky, Hauke Hansen, T. Aichele, O. Benson, J. Mlynek, and S. Schiller. Quantum state reconstruction of the single photon Fock-state. *Physical Review Letters*, 87(5) :050402, July 2001. *E-print quant - ph/0101051*.
- [115] T. Aichele, Alexander I. Lvovsky, and J. Mlynek. Optical mode characterization of single photons prepared by means of conditional measurements on a biphoton state. In *The European Physical Journal D* [69], pages 237–245. *E-print quant - ph/0107080*.
- [116] Lu-Ming Duan, Géza Giedke, Juan Ignacio Cirac, and Peter Zoller. Inseparability criterion for continuous variable systems. *Physical Review Letters*, 84(12) :2722–2725, 2000. *E-print quant - ph/9908056*.
- [117] R. Simon. Peres-Horodecki separability criterion for continuous variable systems. *Physical Review Letters*, 84(12) :2726–2729, 2000. *E-print quant - ph/9909044*.
- [118] Jean-Michel Courty, Philippe Grangier, L. Hilico, and Serge Reynaud. Quantum fluctuations in optical bistability : calculations from linear response theory. *Optics Communications*, 83(3–4) :251–256, June 1991.
- [119] L. Hilico, Claude Fabre, Serge Reynaud, and Elisabeth Giacobino. Linear input-output method for quantum fluctuations in optical bistability with two-level atoms. *Physical Review A*, 46(7) :4397–405, October 1992.

- [120] Carlton M. Caves. Quantum limits on noise in linear amplifiers. *Physical Review D*, 26(8) :1817–1839, October 1982.
- [121] Alexander I. Lvovsky and Jeffrey H. Shapiro. Nonclassical character of statistical mixtures of the single-photon and vacuum optical states. *Physical Review A*, 65(033830), 2002. *E-print quant - ph/0109057*.
- [122] B. Yurke and D. Stoler. Measurement of amplitude probability distributions for photon-number-operator eigenstates. *Physical Review A Rapid Communications*, 36(4) :1955–1959, August 1987.
- [123] Hauke Hansen, C. Hettich, P. Lodahl, Alexander I. Lvovsky, J. Mlynek, and S. Schiller. An ultra-sensitive pulsed balanced homodyne detector : Application to time-domain quantum measurements. *Optics Letters*, 26 :1430, 2001. *E-print quant - ph/0101084*.
- [124] Xiaoying Li, Qing Pan, Jietai Jing, Jing Zhang, Changde Xie, and Kunchi Peng. Quantum dense coding exploiting a bright Einstein-Podolsky-Rosen beam. *Physical Review Letters*, 88(4) :047904, January 2002. *E-print quant - ph/0107068*.
- [125] Mohammad Hafezi. Non-locality tests with continuous variables. Rapport de stage d’option scientifique, École Polytechnique, July 2002.
- [126] Jérôme Wenger. Tomographie quantique impulsionnelle. Rapport de stage de fin d’études, École Supérieure d’Optique, October 2001.
- [127] Z. Y. Ou. Quantum multi-particle interference due to a single-photon state. *Quantum and Semiclassical Optics*, 8(2) :315–322, 1996.
- [128] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Fast Fourier Transform*, chapter 12, pages 496–536. In [6], seconde edition, 1992. Disponible à <http://www.nr.com>.
- [129] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Computing Fourier Integrals Using the FFT*, section 13.9, pages 584–591. In [6], seconde edition, 1992. Disponible à <http://www.nr.com>.
- [130] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Statistical Description of Data*, chapter 14, pages 609–655. In [6], seconde edition, 1992. Disponible à <http://www.nr.com>.
- [131] Jérôme Wenger. Thèse de doctorat, Université Paris XI, 2004.
- [132] Samuel L. Braunstein and H. Jeff Kimble. Dense coding for continuous variables. *Physical Review Letters*, 61 :042302, March 1999.
- [133] Nicolas J. Cerf, M. Lévy, and Gilles Van Assche. Quantum distribution of gaussian keys using squeezed states. *Physical Review A*, 63(5) :052311, 2001. *E-print quant - ph/0005044*.
- [134] Nicolas J. Cerf, Sofyan Iblisdir, and Gilles Van Assche. Cloning and cryptography with quantum continuous variables. In *The European Physical Journal D* [69], pages 211–218. *E-print quant - ph/0107077*.
- [135] William K. Wootters and W. K. Zurek. A single quantum cannot be cloned. *Nature*, 299 :802, 1982.
- [136] D. Dieks. Communications by epr devices. *Physics Letters*, 92A :271–272, 1982.
- [137] Nicolas J. Cerf, A. Ipe, and X. Rottenberg. Cloning of continuous quantum variables. *Physical Review Letters*, 85(8) :1754–1757, 2000. *E-print quant - ph/9909037*.

- [138] Nicolas J. Cerf and Sofyan Iblisdir. Optimal N -to- M cloning of conjugate quantum variables. *Physical Review A* (Rapid Communication), 62(4) :040301(R), 2000. *E-print quant - ph/0005044*.
- [139] Giacomo M. D'Ariano, Francesco De Martini, and Massimiliano F. Sacchi. Continuous variable cloning via network of parametric gates. *Physical Review Letters*, 86 :914–917, 2001. *E-print quant - ph/0012025*.
- [140] Samuel L. Braunstein, Nicolas J. Cerf, Sofyan Iblisdir, vanLoock, and Massar. Optimal cloning of coherent states with a linear amplifier and beam splitters. *Physical Review Letters*, 86(21) :4938–4941, May 2001. *E-print quant - ph/0012046*.
- [141] Jaromír Fiurášek. Optical implementation of continuous-variable quantum cloning machine. *Physical Review Letters*, 86(21) :4942–4945, May 2001. *E-print quant - ph/0012048*.
- [142] Timothy C. Ralph. All optical quantum teleportation. *Optics Letters*, 24 :348, 1999. *E-print quant - ph/9812021*.
- [143] Jean-Paul Delahaye. Logique et calculs de la téléportation. *Pour la Science*, 272 :28–34, June 2000.
- [144] Samuel L. Braunstein. A fun talk on teleportation, February 1995. <http://www.resarch.ibm.com/quantuminfo/teleportation/braunstein.html>.
- [145] Anton Zeilinger. La téléportation quantique. *Pour la Science*, 272 :36–44, June 2000.
- [146] Samuel L. Braunstein, Christopher A. Fuchs, and H. Jeff Kimble. Criteria for continuous-variable quantum teleportation. *Journal of Modern Optics*, 47(2–3 (Special issue : Physics of quantum information)) :267–278, 2000. *E-print quant - ph/9910030*.
- [147] Dik Bouwmeester, Jian-Wei Pan, Klaus Mattle, Manfred Eibl, Harald Weinfurter, and Anton Zeilinger. Experimental quantum teleportation. *Nature*, 390(6660) :575–579, 1997. Commentaire et réponse en [150] et [151].
- [148] D. Boschi, S. Branca, Francesco De Martini, L. Hardy, and Sandu Popescu. Experimental realization of teleporting an unknown pure quantum state via dual classical and Einstein-Podolsky-Rosen channels. *Physical Review Letters*, 80(6) :1121–1125, 1998. *E-print quant - ph/9710013*.
- [149] Brian Julsgaard, Alexander Kozhekin, and Eugene S. Polzik. Experimental long-lived entanglement of two macroscopic objects. *Nature*, 413 :400–403, September 2001.
- [150] Samuel L. Braunstein and H. Jeff Kimble. *A posteriori* teleportation. *Nature*, 394 :840–841, 1998. *E-print quant - ph/9810001*. Commentaire de [147], réponse des auteurs en [151].
- [151] Dik Bouwmeester, Jian-Wei Pan, Klaus Mattle, M. Daniell, Harald Weinfurter, M. Żukowski, and Anton Zeilinger. Bouwmeester *et al.* reply. *Nature*, 394(6696) :841, 1998. Réponse à [150]. Voir aussi [147].
- [152] Philippe Grangier and H. Jeff Kimble. Discussion par e-mail sur les critères de téléportation quantique. Cette discussion a été largement diffusée dans la communauté de l'optique quantique, September–December 2000.
- [153] Samuel L. Braunstein, Christopher A. Fuchs, H. Jeff Kimble, and Peter van Loock. *E-print quant - ph/0012001*, December 2000. Version antérieure de [154].

- [154] Samuel L. Braunstein, Christopher A. Fuchs, H. Jeff Kimble, and Peter van Loock. Quantum versus classical domains for teleportation with continuous variable. *Physical Review A*, 2001. Article similaire à [153].
- [155] Tian Cai Zhang, K. W. Goh, C. W. Chou, P. Lodahl, and H. Jeff Kimble. Quantum teleportation of light beams. *E-print quant - ph/0207076*.
- [156] Warwick P. Bowen, Nicolas Treps, Ben C. Buchler, Roman Schnabel, Timothy C. Ralph, Hans-A. Bachor, Thomas Symul, and Ping Koy Lam. Experimental investigation of continuous variable quantum teleportation. *E-print quant - ph/0207179*.
- [157] Lu-Ming Duan, Géza Giedke, Juan Ignacio Cirac, and Peter Zoller. Entanglement purification of gaussian continuous variable quantum states. *Physical Review Letters*, 84(17) :4002–4005, 2000. *E-print quant - ph/9912017*.
- [158] Timothy C. Ralph and Ping Koy Lam. Teleportation with bright squeezed light. *Physical Review Letters*, 81 :5668–5671, 1998.
- [159] Edgar Allan Poe. *Le scarabé d'or*. Éditions Mille et une nuits, 1993. Traduction de [164] par Charles Baudelaire.
- [160] Auguste Kerckhoffs. La cryptographie militaire. *Journal des sciences militaires*, IX :5–38 et 161–191, January–February 1883. Disponible sur <http://www.cl.cam.ac.uk/~fapp2/kerckhoffs/index.html>.
- [161] Capitaine Roger Baudoin. *Éléments de cryptographie*. Éditions A. Pedone, 1939.
- [162] Simon Singh. *Histoire des codes secrets. De l'Égypte des Pharaons à l'ordinateur quantique*. JC Lattès, 1999. Traduction française de [165].
- [163] Jean-Paul Delahaye. La cryptographie RSA vingt ans après. *Pour la Science*, 267 :104–108, January 2000.
- [164] Edgar Allan Poe. The golden bug, 1843.
- [165] Simon Singh. *The code book*. Fourth Estate Limited, 1999.
- [166] David J. Bernstein. Circuits for integer factorization : a proposal, October 2001. <http://cr.yp.to/papers.html>.
- [167] Arjen K. Lenstra, Adi Shamir, Jim Tomlinson, and Eran Tromer. Analysis of bernstein's factorization circuit, 2002. <http://www.wisdom.weizmann.ac.il/~tromer/Ö>.
- [168] Claude Elwood Shannon. Communication theory of secrecy systems. *Bell System Technical Journal*, 28(4) :656–715, 1949. Réédité à <http://www.cs.ucla.edu/~jkong/research/security/shannon.html>.
- [169] Charles H. Bennett. Quantum cryptography using any two nonorthogonal states. *Physical Review Letters*, 68(21) :3121–3124, May 1992.
- [170] Mark Hillery. Quantum cryptography with squeezed states. *Physical Review A*, 61(2) :022309, 2000. *E-print quant - ph/9909006*.
- [171] Timothy C. Ralph. Continuous variable quantum cryptography. *Physical Review A*(Rapid Communication), 61 :010303(R), December 1999. *E-print quant - ph/9907073*.
- [172] Timothy C. Ralph. Security of continuous-variable quantum cryptography. *Physical Review A*, 62 :062306, November 2000. *E-print quant - ph/0007024*. Cet article est une version longue de [172].

- [173] Timothy C. Ralph. *Quantum Key Distribution with Continuous Variables in Optics*. Kluwer Academic Publisher, 2002. *E-print quant - ph/0109096*. Cet article et une extension de [171, 172].
- [174] Margaret D. Reid. Quantum cryptography with a predetermined key, using continuous-variable einstein-podolsky-rosen correlations. *Physical Review A*, 62(6) :062308, November 2000. *E-print quant - ph/9909030*.
- [175] S. F. Pereira, Z. Y. Ou, and H. Jeff Kimble. Quantum communication with correlated nonclassical states. *Physical Review A*, 62(4) :042311, September 2000. *E-print quant - ph/0003094*.
- [176] D. Gottesmann and John Preskill. Secure quantum key distribution using squeezed states. *Physical Review A*, 63(2) :022309, 2001. *E-print quant - ph/0008046*.
- [177] Christine Silberhorn, Natalia Korolkova, and Gerd Leuchs. Quantum key distribution with bright entangled beams. *Physical Review Letters*, 88 :167902, 2002. *E-print quant - ph/0109009*.
- [178] Kamel Bencheikh, Thomas Symul, A. Jankovic, and Juan-Ariel Levenson. Quantum key distribution with continuous variables. *Journal of Modern Optics*, 48 :1903, 2001.
- [179] Patrick Navez, A. Gatti, and L. A. Lugiato. Invisible transmission in quantum cryptography using continuous variables : A proof of eve's vulnerability. *Physical Review A*, 65(3) :032307, February 2002. *E-print quant - ph/010113* sous le titre « A "quantum public key" based cryptographic scheme for continuous variables ».
- [180] Patrick Navez. Statistical confidentiality tests for a quantum transmission using continuous variables. In *The European Physical Journal D* [69], pages 219–228.
- [181] Christine Silberhorn, Timothy C. Ralph, Norbert Lütkenhaus, and Gerd Leuchs. Continuous variable quantum cryptography : Beating the 3 dB loss limit. *Physical Review Letters*, 89(16) :167901, 2002. *E-print quant - ph/0204064*.
- [182] D. Gottesmann, A. Kitaev, and John Preskill. Encoding a qudit in an oscillator. *Physical Review A*, 64(1) :012310, 2001. *E-print quant - ph/0008040*.
- [183] Gilles Van Assche, Jean Cardinal, and Nicolas J. Cerf. *E-print cs.CR/0107030*, July 2001. Soumis à IEEE Transactions on Information Theory.
- [184] I. Csiszàr and J. Kröner. Broadcast channel with confidential message. *IEEE Transactions on Information Theory*, 24 :339–348, 1978.
- [185] Jean-Paul Delahaye. Aléas du hasard informatique. *Pour la Science*, 245 :92–97, mar 1998.
- [186] Hoi-Kwong Lo, H. F. Chau, and Mohammed Ardehali. Efficient quantum key distribution scheme and proof of its unconditional security. *E-print quant - ph/0011056v2*, November 2000.
- [187] Gilles Van Assche. PhD thesis, Université Libre de Bruxelles, vers 2004.
- [188] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Random Numbers*, chapter 7, pages 274–328. In [6], seconde edition, 1992. Disponible à <http://www.nr.com>.
- [189] IEEE standard for binary floating point number. ANSI/IEEE Std 754–1985, IEEE, New-York, 1985.

- [190] William H. Press, Saul A. Teutolsky, and William T. Vetterling. *Diagnosing Machine Parameters*, section 20.1, pages 889–893. In [6], seconde edition, 1992. Disponible à <http://www.nr.com>.
- [191] J. L. Carter and M. N. Wegman. Universal classes of hash functions. *J. of Comp. and Syst. Sci.*, 18 :143–154, 1979.
- [192] R. P. Brent, S. Larvala, and P. Zimmermann. A fast algorithm for testing irreducibility of trinomials mod 2. Tech. rep., Oxford University Computing Laboratory, 2000.
- [193] A. Schönhage. Schnelle multiplikation von polynomen über körper de charakteristik 2. *Acta Informatica*, 7 :395–398, 1977.
- [194] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Rappels de quelques propriétés utiles des opérateurs linéaires*, complément B_{II}, pages 166–175. Volume I of [3], 1977.
- [195] Claude Cohen-Tannoudji, Bernard Diu, and Franck Laloë. *Quelques propriétés générales de deux observables Q et P dont le commutateur est égal à $i\hbar$* , complément E_{II}, pages 187–189. Volume I of [3], 1977.
- [196] Wolfgang P. Schleich. *Position and Momentum Eigenstates*, section 2.1, pages 36–40. In [5], 2001.