



# Intrication de champs quantiques mesoscopiques pour les communications quantiques

Rémi Blandino

## ► To cite this version:

Rémi Blandino. Intrication de champs quantiques mesoscopiques pour les communications quantiques. Autre [cond-mat.other]. Université Paris Sud - Paris XI, 2013. Français. NNT : 2013PA112045 . tel-00827393

**HAL Id: tel-00827393**

**<https://pastel.hal.science/tel-00827393>**

Submitted on 29 May 2013

**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.

INSTITUT D'OPTIQUE  
LABORATOIRE CHARLES FABRY

UNIVERSITÉ PARIS XI  
ÉCOLE DOCTORALE ONDES ET MATIÈRE

## THÈSE

SPÉCIALITÉ : PHYSIQUE QUANTIQUE

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

par  
**RÉMI BLANDINO**

Sujet :

---

### INTRICATION DE CHAMPS QUANTIQUES MESOSCOPIQUES POUR LES COMMUNICATIONS QUANTIQUES

---

Soutenue le 25 mars 2013  
devant la commission d'examen :

|      |                     |                     |
|------|---------------------|---------------------|
| M.   | Jean-François ROCH, | Président           |
| M.   | Antonio ACÍN,       | Rapporteur          |
| M.   | Claude FABRE,       | Rapporteur          |
| M.   | Nicolas CERF,       | Examineur           |
| Mme. | Rosa TUALLE-BROURI, | Directrice de thèse |
| M.   | Philippe GRANGIER,  | Membre invité       |



## Remerciements

Je tiens d’abord à remercier Rosa Tualle-Brouri, ma directrice de thèse, pour son encadrement, l’autonomie, et la confiance qu’elle a su m’accorder. Elle a été à l’écoute de mes centres d’intérêt, et je lui suis très reconnaissant de m’avoir laissé la liberté de m’orienter davantage vers la théorie au cours de ma troisième année. Merci également à Philippe Grangier pour avoir participé à l’encadrement de cette thèse, et avec qui les discussions sont toujours riches d’enseignements.

Je ne saurais trop remercier Marco Barbieri et Franck Ferreyrol, qui étaient respectivement post-doc et thésard sur la « manip femto » avec moi, pour leurs enseignements et les bons moments que nous avons passé ensemble. J’ai également eu beaucoup de plaisir à travailler avec Jean Etesse, qui a pris la relève sur la manip. Je lui souhaite bon courage, ainsi qu’à Bhaskar Kanseri qui l’a récemment rejoint.

Un grand merci à tous les thésards et post-docs qui ont contribué à l’ambiance du groupe d’optique quantique : Aline, Charles, Lucas, Joseph, Ronan, Alpha, Andreas, Tatiana, Erwan, Imad, Sylvain, et tous les autres. Merci à Antoine Browaeys pour les innombrables discussions de physique que nous avons pu avoir. Merci à Alexei Ourjoumtsev pour ses conseils sur la manip. Merci également aux autres membres du groupe pour leur disponibilité, et à Simon Fossier et Anthony Leverrier, avec qui j’ai eu plaisir à collaborer, et qui m’ont beaucoup appris en cryptographie quantique.

Je tiens à remercier tout le personnel administratif et technique de l’Institut d’Optique pour leurs compétences et leur réactivité. Merci également à tous les enseignants avec qui j’ai eu beaucoup de plaisir à enseigner et à apprendre, en particulier Yvan Sortais, Jean-Marie Feybesse, Nicolas Dubreuil, Alexios Beveratos, Fabienne Bernard et Jean-Michel Jonathan.

Merci à mes amis, mineurs, physiciens, ou d’autres horizons, pour toutes ces belles années passées ensemble.

Enfin, mes plus grands remerciements vont à ma famille pour son soutien inébranlable durant ces nombreuses années d’études.





# Table des matières

|          |                                                                            |          |
|----------|----------------------------------------------------------------------------|----------|
| <b>1</b> | <b>Introduction</b>                                                        | <b>1</b> |
| 1.1      | Concepts . . . . .                                                         | 1        |
| 1.1.1    | La physique quantique . . . . .                                            | 1        |
| 1.1.2    | L'information quantique . . . . .                                          | 1        |
| 1.2      | Travail de thèse . . . . .                                                 | 3        |
| 1.2.1    | Première partie : les outils théoriques et expérimentaux . . . . .         | 3        |
| 1.2.2    | Deuxième partie : résultats expérimentaux . . . . .                        | 3        |
| 1.2.3    | Troisième partie : étude théorique d'un amplificateur sans bruit . . . . . | 4        |
| <b>I</b> | <b>Outils théoriques et expérimentaux</b>                                  | <b>5</b> |
| <b>2</b> | <b>Outils théoriques</b>                                                   | <b>7</b> |
| 2.1      | Champ quantique . . . . .                                                  | 8        |
| 2.1.1    | Quantification du champ . . . . .                                          | 8        |
| 2.1.2    | Variables discrètes & états de Fock . . . . .                              | 11       |
| 2.1.3    | Variables continues & quadratures du champ . . . . .                       | 12       |
| 2.1.4    | Lien entre les descriptions discrètes et continues . . . . .               | 13       |
| 2.2      | Evolution temporelle et modes du champ . . . . .                           | 14       |
| 2.2.1    | Evolution temporelle - représentation d'interaction . . . . .              | 14       |
| 2.2.2    | Modes du champ . . . . .                                                   | 14       |
| 2.3      | Descriptions et propriétés des états quantiques . . . . .                  | 17       |
| 2.3.1    | Matrice densité . . . . .                                                  | 18       |
| 2.3.2    | Fonction de Wigner . . . . .                                               | 20       |
| 2.3.3    | Quelle description choisir ? . . . . .                                     | 23       |
| 2.3.4    | Distance entre états quantiques : la fidélité . . . . .                    | 24       |
| 2.3.5    | Mesures et POVM . . . . .                                                  | 25       |
| 2.4      | Quelques transformations unitaires . . . . .                               | 25       |
| 2.4.1    | Le déphasage . . . . .                                                     | 25       |
| 2.4.2    | La lame séparatrice . . . . .                                              | 26       |
| 2.5      | Les états gaussiens . . . . .                                              | 27       |
| 2.5.1    | Etats d'incertitude minimale . . . . .                                     | 28       |
| 2.5.2    | Le vide quantique . . . . .                                                | 29       |
| 2.5.3    | Les états cohérents . . . . .                                              | 29       |
| 2.5.4    | Le vide comprimé monomode . . . . .                                        | 33       |
| 2.5.5    | Le vide comprimé bimode, ou état EPR . . . . .                             | 35       |
| 2.5.6    | Les états thermiques . . . . .                                             | 37       |

|          |                                                                 |           |
|----------|-----------------------------------------------------------------|-----------|
| 2.6      | Etats non gaussiens                                             | 39        |
| 2.6.1    | Les états de Fock                                               | 40        |
| 2.6.2    | Les superpositions d'états cohérents, ou "chats de Schrödinger" | 41        |
| 2.7      | Application à l'information quantique                           | 42        |
| 2.7.1    | Variables discrètes                                             | 43        |
| 2.7.2    | Variables continues                                             | 44        |
| 2.8      | Conclusion                                                      | 45        |
| <b>3</b> | <b>Le dispositif expérimental</b>                               | <b>47</b> |
| 3.1      | Présentation du dispositif                                      | 47        |
| 3.1.1    | Introduction                                                    | 47        |
| 3.1.2    | Le dispositif expérimental                                      | 48        |
| 3.1.3    | La source laser                                                 | 49        |
| 3.2      | Transformations unitaires : optique linéaire                    | 53        |
| 3.2.1    | Déphasage                                                       | 53        |
| 3.2.2    | Lame séparatrice                                                | 53        |
| 3.2.3    | Lames demi-onde $\lambda/2$ et quart-d'onde $\lambda/4$         | 54        |
| 3.2.4    | Cube séparateur de polarisation (PBS)                           | 54        |
| 3.3      | Transformations unitaires : optique non linéaire                | 55        |
| 3.3.1    | Génération de seconde harmonique (GSH)                          | 56        |
| 3.3.2    | Amplification paramétrique optique (OPA)                        | 56        |
| 3.4      | Détection et mesures projectives                                | 61        |
| 3.4.1    | Détection homodyne                                              | 61        |
| 3.4.2    | Photodiode à avalanche                                          | 67        |
| 3.4.3    | Soustraction de photon                                          | 68        |
| 3.5      | Conclusion                                                      | 69        |
| <b>4</b> | <b>Eléments de théorie de l'information</b>                     | <b>71</b> |
| 4.1      | Introduction                                                    | 71        |
| 4.2      | Information classique                                           | 72        |
| 4.2.1    | Entropie de Shannon                                             | 72        |
| 4.2.2    | Entropie de deux variables aléatoires                           | 72        |
| 4.2.3    | Information mutuelle                                            | 73        |
| 4.2.4    | Quelques propriétés de base de l'entropie                       | 73        |
| 4.3      | Information quantique                                           | 74        |
| 4.3.1    | Entropie de von Neumann                                         | 74        |
| 4.3.2    | Entropie et corrélations pour des systèmes bipartites           | 74        |
| 4.3.3    | Etats classiques-quantiques                                     | 76        |
| 4.3.4    | Borne de Holevo                                                 | 76        |
| 4.4      | Modèle du canal gaussien                                        | 77        |
| 4.4.1    | Canal sans bruit                                                | 77        |
| 4.4.2    | Canal avec bruit thermique                                      | 77        |
| 4.4.3    | Information mutuelle                                            | 79        |
| 4.5      | Conclusion                                                      | 80        |

|           |                                                                                        |            |
|-----------|----------------------------------------------------------------------------------------|------------|
| <b>II</b> | <b>Résultats expérimentaux</b>                                                         | <b>81</b>  |
| <b>5</b>  | <b>Estimation de la discordance quantique pour un état EPR</b>                         | <b>83</b>  |
| 5.1       | Introduction . . . . .                                                                 | 83         |
| 5.2       | La discordance quantique . . . . .                                                     | 84         |
| 5.2.1     | Une mesure des corrélations quantiques . . . . .                                       | 84         |
| 5.2.2     | La discordance en information quantique . . . . .                                      | 88         |
| 5.3       | Protocole expérimental . . . . .                                                       | 89         |
| 5.3.1     | Principe . . . . .                                                                     | 89         |
| 5.3.2     | Tri des quadratures et estimation des variances comprimées et anticomprimées . . . . . | 90         |
| 5.3.3     | Estimation des incertitudes sur les variances comprimées et anticomprimées . . . . .   | 93         |
| 5.4       | Modélisation . . . . .                                                                 | 94         |
| 5.4.1     | Etat produit par l'OPA . . . . .                                                       | 94         |
| 5.4.2     | Estimation de $N_s$ , $N_t$ et de leurs incertitudes, par inversion . . . . .          | 95         |
| 5.5       | Estimation de la discordance et de son incertitude . . . . .                           | 95         |
| 5.5.1     | Discordance gaussienne pour un état thermique comprimé . . . . .                       | 95         |
| 5.5.2     | Estimation par inversion . . . . .                                                     | 97         |
| 5.5.3     | Estimation bayésienne . . . . .                                                        | 97         |
| 5.5.4     | Résultats expérimentaux . . . . .                                                      | 99         |
| 5.6       | Comparaison avec les bornes de Cramér-Rao . . . . .                                    | 100        |
| 5.6.1     | Information de Fisher et borne de Cramér-Rao classique . . . . .                       | 101        |
| 5.6.2     | De l'information de Fisher classique à l'information de Fisher quantique . . . . .     | 102        |
| 5.6.3     | Application à l'évaluation de la discordance . . . . .                                 | 104        |
| 5.6.4     | Résultats expérimentaux . . . . .                                                      | 106        |
| 5.7       | Conclusion . . . . .                                                                   | 107        |
| <b>6</b>  | <b>Caractérisation d'une porte de phase quantique</b>                                  | <b>109</b> |
| 6.1       | Introduction . . . . .                                                                 | 109        |
| 6.2       | Présentation de la porte de phase . . . . .                                            | 111        |
| 6.2.1     | Calcul quantique avec des états cohérents . . . . .                                    | 111        |
| 6.2.2     | La porte de phase . . . . .                                                            | 111        |
| 6.3       | Réalisation expérimentale d'une porte de phase $\pi$ . . . . .                         | 113        |
| 6.3.1     | Méthode et dispositif expérimental . . . . .                                           | 113        |
| 6.3.2     | Modélisation . . . . .                                                                 | 115        |
| 6.3.3     | Extraction des paramètres expérimentaux . . . . .                                      | 116        |
| 6.3.4     | Résultats expérimentaux . . . . .                                                      | 117        |
| 6.3.5     | Test du modèle de la porte . . . . .                                                   | 117        |
| 6.3.6     | Incertainitude sur l'estimation de $\xi$ . . . . .                                     | 119        |
| 6.4       | Comment caractériser la porte ? . . . . .                                              | 120        |
| 6.4.1     | Ressemblance des états expérimentaux avec des chats parfaits . . . . .                 | 120        |
| 6.4.2     | Tomographie de processus quantique . . . . .                                           | 121        |
| 6.4.3     | Utilisation de la modélisation de la porte . . . . .                                   | 123        |
| 6.5       | Fidélité pour un qubit initial parfait . . . . .                                       | 125        |
| 6.5.1     | Porte expérimentale . . . . .                                                          | 125        |
| 6.5.2     | Modèle de porte simplifié . . . . .                                                    | 126        |
| 6.5.3     | Fidélités pour des superpositions de même poids . . . . .                              | 127        |
| 6.5.4     | Quelques calculs de fidélité . . . . .                                                 | 128        |

|       |                                                               |     |
|-------|---------------------------------------------------------------|-----|
| 6.6   | Fidélité avec une porte idéale . . . . .                      | 133 |
| 6.6.1 | Principe . . . . .                                            | 133 |
| 6.6.2 | Simulations pour la porte de phase . . . . .                  | 134 |
| 6.6.3 | Invariance du choix de l'état maximalement intriqué . . . . . | 135 |
| 6.7   | Conclusion . . . . .                                          | 137 |

### III Résultats théoriques : amplificateur sans bruit en cryptographie quantique 139

|          |                                                                                         |            |
|----------|-----------------------------------------------------------------------------------------|------------|
| <b>7</b> | <b>Cryptographie quantique</b>                                                          | <b>141</b> |
| 7.1      | Introduction . . . . .                                                                  | 141        |
| 7.1.1    | La cryptographie classique . . . . .                                                    | 142        |
| 7.1.2    | La cryptographie quantique . . . . .                                                    | 143        |
| 7.2      | Principes généraux de cryptographie quantique . . . . .                                 | 144        |
| 7.2.1    | Hypothèses sur le contrôle de l'environnement . . . . .                                 | 144        |
| 7.2.2    | Types d'attaques . . . . .                                                              | 145        |
| 7.3      | Quelques protocoles . . . . .                                                           | 145        |
| 7.3.1    | Variables discrètes . . . . .                                                           | 145        |
| 7.3.2    | Variables continues & protocoles gaussiens . . . . .                                    | 146        |
| 7.3.3    | Autres protocoles . . . . .                                                             | 146        |
| 7.3.4    | Réseaux et applications commerciales . . . . .                                          | 147        |
| 7.4      | Le protocole GG02 : de la QKD avec des états "classiques" . . . . .                     | 147        |
| 7.4.1    | Principe . . . . .                                                                      | 147        |
| 7.4.2    | Modélisation à intrication virtuelle . . . . .                                          | 148        |
| 7.4.3    | Preuves de sécurité . . . . .                                                           | 148        |
| 7.4.4    | Expression des taux secrets . . . . .                                                   | 150        |
| 7.5      | Conclusion . . . . .                                                                    | 152        |
| <b>8</b> | <b>L'amplificateur sans bruit non déterministe</b>                                      | <b>153</b> |
| 8.1      | Introduction . . . . .                                                                  | 153        |
| 8.2      | Amplificateurs déterministes . . . . .                                                  | 154        |
| 8.2.1    | Bruit minimal ajouté par un amplificateur déterministe . . . . .                        | 154        |
| 8.2.2    | Amplificateur indépendant de la phase . . . . .                                         | 155        |
| 8.2.3    | Amplificateur dépendant de la phase . . . . .                                           | 156        |
| 8.3      | Principe et propriétés de base de l'amplificateur sans bruit non déterministe . . . . . | 157        |
| 8.3.1    | Principe . . . . .                                                                      | 157        |
| 8.3.2    | Transformation de quelques états gaussiens . . . . .                                    | 159        |
| 8.4      | Implémentations théoriques et réalisations expérimentales . . . . .                     | 161        |
| 8.4.1    | Les ciseaux quantiques . . . . .                                                        | 161        |
| 8.4.2    | Autres réalisations . . . . .                                                           | 164        |
| 8.5      | Applications . . . . .                                                                  | 165        |
| 8.5.1    | Préparation d'états quantiques . . . . .                                                | 165        |
| 8.5.2    | Communications quantiques . . . . .                                                     | 166        |
| 8.6      | Conclusion . . . . .                                                                    | 166        |

|           |                                                                                         |            |
|-----------|-----------------------------------------------------------------------------------------|------------|
| <b>9</b>  | <b>Propriétés de l'amplificateur sans bruit non déterministe</b>                        | <b>167</b> |
| 9.1       | Introduction . . . . .                                                                  | 167        |
| 9.2       | Bornes supérieures pour la probabilité de succès . . . . .                              | 168        |
| 9.2.1     | NLA parfait et probabilité de succès non nulle . . . . .                                | 168        |
| 9.2.2     | Bornes obtenues par non diminution de la fidélité . . . . .                             | 169        |
| 9.2.3     | Borne de Vidal en dimension finie . . . . .                                             | 171        |
| 9.3       | Application après un canal quantique : système effectif équivalent . . . . .            | 173        |
| 9.3.1     | Amplificateur sans bruit après un canal quantique . . . . .                             | 174        |
| 9.3.2     | Amplificateur sans bruit en amont d'un canal quantique effectif . . . . .               | 177        |
| 9.3.3     | Paramètres effectifs . . . . .                                                          | 178        |
| 9.3.4     | Trois types de canaux effectifs . . . . .                                               | 179        |
| 9.4       | Application aux protocoles de cryptographie quantique . . . . .                         | 180        |
| 9.4.1     | Simplification du système effectif pour $\eta \leq 1$ . . . . .                         | 181        |
| 9.4.2     | Allure des paramètres effectifs . . . . .                                               | 181        |
| 9.4.3     | Sens physique des paramètres effectifs . . . . .                                        | 185        |
| 9.4.4     | Comportements limites des paramètres effectifs . . . . .                                | 187        |
| 9.4.5     | Vérification numérique . . . . .                                                        | 188        |
| 9.4.6     | Prise en compte d'une troncature . . . . .                                              | 189        |
| 9.5       | Application aux communications quantiques . . . . .                                     | 190        |
| 9.5.1     | Suppression des pertes et atténuateur sans bruit . . . . .                              | 190        |
| 9.5.2     | "Concentration de phase" . . . . .                                                      | 192        |
| 9.6       | Conclusion . . . . .                                                                    | 193        |
| <b>10</b> | <b>L'amplificateur sans bruit non déterministe en cryptographie quantique</b>           | <b>195</b> |
| 10.1      | Introduction . . . . .                                                                  | 195        |
| 10.2      | Calcul des taux secrets . . . . .                                                       | 196        |
| 10.2.1    | Taux secrets sans le NLA . . . . .                                                      | 196        |
| 10.2.2    | Taux secrets avec le NLA . . . . .                                                      | 197        |
| 10.3      | Amélioration des performances - attaques collectives . . . . .                          | 197        |
| 10.3.1    | Considérations préliminaires . . . . .                                                  | 198        |
| 10.3.2    | Une distance de transmission augmentée, et une plus grande tolérance au bruit . . . . . | 200        |
| 10.3.3    | Que se passe t'il quand le gain augmente trop ? . . . . .                               | 206        |
| 10.3.4    | Augmentation arbitraire de la distance maximale de transmission . . . . .               | 207        |
| 10.4      | Attaques individuelles et non amélioration des performances avec le NLA . . . . .       | 211        |
| 10.4.1    | Démonstration pour un canal ajoutant du bruit . . . . .                                 | 211        |
| 10.4.2    | Démonstration exacte pour tout $T$ , pour un canal sans bruit . . . . .                 | 214        |
| 10.5      | Conclusion . . . . .                                                                    | 216        |
| <b>11</b> | <b>Conclusion et perspectives</b>                                                       | <b>217</b> |
| 11.1      | Conclusion . . . . .                                                                    | 217        |
| 11.2      | Perspectives . . . . .                                                                  | 218        |
| <b>IV</b> | <b>Annexes</b>                                                                          | <b>219</b> |
| <b>A</b>  | <b>Quelques propriétés opératorielles bien utiles</b>                                   | <b>221</b> |
| A.1       | Evolution d'une fonction d'opérateurs . . . . .                                         | 221        |

|          |                                                                                                          |            |
|----------|----------------------------------------------------------------------------------------------------------|------------|
| A.2      | Evolution des opérateurs $\hat{a}$ et $\hat{a}^\dagger$ sous l'action d'un hamiltonien quadratique . . . | 221        |
| A.3      | “Désintriquer” une exponentielle . . . . .                                                               | 222        |
| A.3.1    | Formule de Baker-Hausdorff . . . . .                                                                     | 222        |
| A.3.2    | Relations de commutation de $SU(2)$ et lame séparatrice . . . . .                                        | 223        |
| A.3.3    | Relations de commutation de $SU(1,1)$ et opérateurs de squeezing . . . . .                               | 223        |
| <b>B</b> | <b>Modèle multimode simplifié de l'OPA</b>                                                               | <b>227</b> |
| B.1      | Modélisation multimode et lien avec le modèle empirique . . . . .                                        | 227        |
| B.1.1    | Efficacité homodyne parfaite . . . . .                                                                   | 228        |
| B.1.2    | Prise en compte de l'efficacité homodyne . . . . .                                                       | 229        |
| <b>C</b> | <b>Eléments de photodétection</b>                                                                        | <b>231</b> |
| <b>D</b> | <b>Modèle simplifié de l'origine de <math>\xi</math></b>                                                 | <b>233</b> |
| <b>E</b> | <b>Une autre implémentation approchée du NLA</b>                                                         | <b>235</b> |
| E.1      | Résultat préliminaire sur les états de Fock . . . . .                                                    | 235        |
| E.2      | Décomposition de $\hat{T}_N$ sur la base des déplacements . . . . .                                      | 236        |
| E.2.1    | Sans coupure d'intégration . . . . .                                                                     | 236        |
| E.2.2    | Avec une coupure d'intégration . . . . .                                                                 | 237        |
| E.3      | Implémentation “physique” d'une superposition de déplacements . . . . .                                  | 238        |
| <b>F</b> | <b>Compléments sur le NLA</b>                                                                            | <b>239</b> |
| <b>G</b> | <b>Développement perturbatif des taux secrets</b>                                                        | <b>241</b> |
| G.1      | Développement de l'information de Holevo . . . . .                                                       | 241        |
| G.2      | Développement de $I_{ab}$ . . . . .                                                                      | 242        |
| G.3      | Développement du taux secret contre les attaques collectives . . . . .                                   | 242        |
| G.4      | Développement de $I_{BE}$ . . . . .                                                                      | 242        |
| <b>H</b> | <b>Autre dérivation des paramètres effectifs</b>                                                         | <b>243</b> |
| H.1      | Obtention des paramètres effectifs . . . . .                                                             | 243        |
| H.1.1    | Conditions en utilisant un état thermique déplacé . . . . .                                              | 243        |
| H.1.2    | Conditions en utilisant un état thermique non déplacé . . . . .                                          | 244        |
| H.1.3    | Résolution du système . . . . .                                                                          | 245        |
| H.2      | Lien avec les paramètres effectifs dans le cas général . . . . .                                         | 245        |
| <b>I</b> | <b>Mesure hétérodyne d'Alice</b>                                                                         | <b>247</b> |
| I.1      | Mesure hétérodyne d'un mode d'un état EPR . . . . .                                                      | 247        |
| I.2      | Mesure hétérodyne d'un état cohérent . . . . .                                                           | 248        |

# Chapitre 1

## Introduction

### 1.1 Concepts

#### 1.1.1 La physique quantique

La physique quantique a connu de profonds changements au cours de ces dernières années. Elle a d'abord permis une compréhension de la matière qui a donné lieu à de nombreuses applications technologiques, qui sont aujourd'hui devenues indispensables. L'exemple le plus important est sans doute celui du laser, omniprésent, dont le principe de fonctionnement est intrinsèquement quantique. Les semi-conducteurs ont également été fortement développés grâce à la physique quantique, permettant la miniaturisation et le développement des ordinateurs à grande échelle. Toutes ces applications sont basées sur une compréhension des phénomènes d'interactions complexes qui opèrent entre les atomes, et également avec la lumière.

La physique quantique est basée sur des concepts pour lesquels nous n'avons que peu d'intuition. À une particule, on ne peut associer qu'une probabilité de présence à un certain endroit de l'espace. Certaines grandeurs physiques, telles que la position et l'impulsion, ne peuvent pas être déterminées en même temps. Parmi les nombreuses propriétés propres aux systèmes quantiques, deux peuvent être considérées comme les plus représentatives des différences avec le monde macroscopique : la première est le principe de superposition. Un système quantique peut être dans une superposition de plusieurs états quantiques à la fois, cette superposition n'étant levée que lorsque qu'une mesure est effectuée sur le système. Plusieurs expériences ont confirmé ce phénomène surprenant, dont la célèbre expérience des fentes d'Young, réalisée en envoyant une seule particule, ou même un seul photon. La seconde propriété est la possibilité d'avoir des corrélations quantiques non locales entre deux systèmes, *l'intrication*, à la base de nombreux protocoles d'information quantique.

#### 1.1.2 L'information quantique

La théorie de l'information remonte aux années cinquante, avec la contribution majeure de Claude Shannon [Shannon48] qui introduisit les outils nécessaires pour quantifier la notion d'information, et de ce fait les ressources nécessaires à sa manipulation et à son stockage. Il a également posé les bases théoriques de la compression des données, et des codes correcteurs d'erreurs. L'information est encodée en utilisant les symboles d'un alphabet, pouvant être composé d'un ensemble d'éléments discret, c'est le cas du codage numérique prenant les valeurs 0 ou 1, ou d'un ensemble d'éléments continu, comme pour le codage analogique. Dans les deux cas, les



éléments de l'alphabet peuvent correspondre à différents états d'un système physique. Les états 0 et 1 peuvent par exemple être stockés dans une mémoire en utilisant deux états d'aimantation d'un matériau, et le codage analogique peut correspondre à l'amplitude d'une tension ou d'un champ électromagnétique. Tous ces systèmes ont la particularité d'être décrits par les lois de la physique classique, du fait de leur taille macroscopique.

En revanche, depuis quelques décennies il a été réalisé que les propriétés de base des systèmes quantiques pouvaient être utilisées comme des ressources fondamentales pour traiter et communiquer de l'information, donnant naissance à un domaine en plein développement : *l'information quantique*. En codant un bit sur deux états  $|0\rangle$  et  $|1\rangle$  d'un système physique obéissant aux lois de la physique quantique – niveaux d'énergie d'un atome, états de vibration dans un piège, polarisation d'un photon, etc... – le bit quantique, ou *qubit*, peut se trouver dans une superposition arbitraire  $c_0|0\rangle + c_1|1\rangle$  alors que le bit classique ne peut être que dans l'état 0 ou 1. Cette superposition est à la base du *calcul quantique*, qui l'exploite afin de pouvoir résoudre certains problèmes exponentiellement plus rapidement qu'avec un ordinateur classique. De nombreux systèmes physiques font l'objet de recherches actives : atomes, ions, photons, supraconducteurs, etc... Chacun de ces systèmes possède ses propres avantages et ses limites pour une mise en œuvre efficace. Exploiter ces propriétés quantiques reste un formidable défi technologique, tant elles sont fragiles et tendent à être estompées par l'environnement externe et les imperfections expérimentales. Beaucoup de recherches sont effectuées en vue de trouver des implémentations robustes et utilisables à grande échelle.

Les corrélations quantiques, ainsi que le principe d'incertitude de Heisenberg – selon lequel certaines grandeurs quantiques ne peuvent pas être déterminées en même temps – ont permis de développer les *communications quantiques*, dont l'exemple le plus avancé est la cryptographie quantique, qui commence à atteindre un stade commercial. La cryptographie quantique, plus précisément la distribution quantique de clé, consiste à transmettre une chaîne de bits entre deux partenaires, à travers un environnement sous le contrôle d'un espion tentant d'intercepter la communication. L'information est le plus souvent transmise en utilisant de la lumière envoyée dans une fibre optique, ou se propageant à l'air libre. Sa grande force est de garantir que la clé extraite en effectuant un certain nombre de manipulations sur les données sera parfaitement secrète, quelles que soient les ressources physiques dont dispose l'espion. Elle peut ensuite être utilisée pour coder un message, soit avec des algorithmes utilisant des clés plus courtes que le message à transmettre, soit en utilisant une clé de la même longueur que le message, ce qui rend ce dernier parfaitement indéchiffrable. Ce secret est cependant obtenu au prix de débits pouvant être très faibles comparés à ceux atteints aujourd'hui par les communications classiques standards. Ce taux de clé secrète est d'autant plus faible que le canal quantique introduit des pertes et du bruit, et peut même devenir nul si les conditions de la transmission sont trop défavorables. Les recherches sur les différents protocoles de cryptographie quantique visent à améliorer ces performances, tout en nécessitant un dispositif expérimental le plus simple possible, voire afin de pouvoir utiliser les technologies telecom très avancées.

L'information quantique ne se limite toutefois pas au calcul quantique et à la cryptographie quantique. De nombreux autres protocoles permettent d'utiliser les ressources quantiques pour le traitement de l'information.

Enfin, deux domaines, plus ou moins disjoints, peuvent être différenciés : celui où l'information est codée sur un ensemble d'états discret. On parle alors de *variables discrètes* ; et celui où l'information est codée sur un ensemble d'états continu, comme l'amplitude du champ électrique. On parle alors de *variables continues*. Certains systèmes physiques sont naturellement décrits par des variables discrètes ou continues : les niveaux d'énergie d'un atome sont discrets, tout comme les états d'excitation d'un oscillateur harmonique, alors que la position ou l'impulsion

d'un électron sont continues. D'autres systèmes, comme le champ électromagnétique quantique, peuvent être aussi bien décrits par des variables discrètes que par des variables continues. Le photon, quantum d'énergie par nature indivisible, semble être naturellement décrit par des variables discrètes, toutefois, le champ électromagnétique qui lui correspond prend des valeurs continues.

## 1.2 Travail de thèse

L'information quantique est une discipline donnant lieu à de nombreuses thématiques. Sur le plan théorique, il est nécessaire de développer de nouveaux outils et de nouveaux protocoles utilisant les ressources de la physique quantique. Sur le plan expérimental, le principal enjeu est la mise en œuvre de ces protocoles, compte tenu des imperfections limitant grandement les performances. Il est nécessaire de pouvoir caractériser les imperfections, et de disposer d'outils de diagnostic utilisant au mieux les ressources expérimentales disponibles.

Ce travail de thèse s'inscrit dans le cadre de l'information quantique avec des variables continues du champ électromagnétique, et présente des contributions aux problématiques théoriques et expérimentales. Il est axé autour des deux thématiques que nous avons présenté dans les paragraphes précédents : le calcul quantique, et les communications quantiques.

### 1.2.1 Première partie : les outils théoriques et expérimentaux

Ce manuscrit est organisé de la manière suivante : dans une première partie, nous présenterons les outils de base, à la fois théoriques et expérimentaux, utilisés pour décrire, produire, et mesurer les états quantiques. Nous détaillerons le formalisme de base et les états quantiques typiques dans le chapitre 2, puis le dispositif expérimental utilisé pour les produire et les analyser sera présenté dans le chapitre 3. Enfin, le chapitre 4 nous permettra d'introduire les notions fondamentales de la théorie de l'information. Les outils propres aux autres chapitres seront ensuite présentés lorsqu'ils seront nécessaires.

### 1.2.2 Deuxième partie : résultats expérimentaux

La deuxième partie regroupe les résultats expérimentaux obtenus au cours de cette thèse. Le chapitre 5 est à la frontière des communications quantiques et du calcul quantique. Il est consacré à l'estimation expérimentale de la *discord quantique* pour un état bipartite de type EPR. La discord est une mesure des corrélations d'origine purement quantique. Elle peut être non nulle pour des états qui ne sont pas intriqués, mais qui possèdent néanmoins des corrélations supérieures à ce qu'il est possible d'avoir classiquement. Elle fait actuellement l'objet de recherches très actives, en particulier car de nombreuses études laissent penser qu'elle pourrait être une ressource fondamentale et que l'intrication ne serait pas toujours nécessaire pour surpasser les protocoles classiques. Afin de poursuivre ces investigations, son estimation expérimentale est donc un enjeu particulièrement important.

Nous présenterons une méthode d'estimation de la discord quantique pour un état EPR. Nous montrerons qu'elle permet d'obtenir une précision proche des valeurs minimales données par les bornes de Cramér-Rao classiques et quantiques, tout en utilisant des outils expérimentaux maintenant standards en optique quantique. Nous verrons également l'amélioration que peut apporter une estimation bayésienne.

Le chapitre 6 est consacré au calcul quantique avec des variables continues. Nous y présenterons une méthode de caractérisation de la qualité d’une porte quantique, afin de la comparer à l’opération théorique que l’on cherche à implémenter. Le qubit est encodé dans deux états cohérents  $|\alpha\rangle$  et  $|\alpha\rangle$  du champ électrique, de même amplitude et de phases opposées. Il est ensuite possible, grâce à des opérations non déterministes, de réaliser des portes quantiques. Un nombre restreint de portes différentes permet de réaliser n’importe quelle opération sur un ou plusieurs qubits. La plus simple d’entre elle est une porte introduisant un facteur de phase  $\varphi$  dans une superposition quantique :  $c_1|\alpha\rangle + c_2|\alpha\rangle \rightarrow c_1|\alpha\rangle + e^{i\varphi}c_2|\alpha\rangle$ . En particulier, une phase  $\phi=\pi$  peut être introduite en appliquant l’opérateur de destruction d’un photon  $\hat{a}$ .

La caractérisation de la qualité de la réalisation expérimentale est une étape nécessaire en vue d’une intégration à plus grande échelle. Nous présenterons une méthode semi-empirique, basée sur une modélisation analytique et sur des résultats expérimentaux : en utilisant un modèle détaillé du dispositif expérimental dont les paramètres sont extraits à partir d’un jeu de mesures, les imperfections du qubit sont séparées des imperfections de la porte de phase elle-même. La porte de phase “expérimentale” est ensuite numériquement appliquée à un qubit parfait, et l’état obtenu est comparé à celui produit par une porte parfaite.

### 1.2.3 Troisième partie : étude théorique d’un amplificateur sans bruit

Enfin, la troisième partie regroupe les résultats théoriques obtenus dans le cadre des communications quantiques. Nous y présenterons une application théorique d’un amplificateur sans bruit non déterministe en cryptographie quantique, dont les principes seront introduits dans le chapitre 7.

Un tel amplificateur possède la propriété surprenante de pouvoir amplifier uniquement le signal, sans en amplifier le bruit quantique, comme nous le verrons dans le chapitre 8. Une telle opération ne peut toutefois pas être réalisée de manière déterministe, mais avec une certaine probabilité de succès. Présenté il y a seulement quelques années, cet amplificateur a fait l’objet de recherches intenses par la communauté scientifique. Plusieurs versions approchées ont été expérimentalement réalisées afin d’apporter des démonstrations de principe de sa faisabilité.

En revanche, les applications potentielles font toujours l’objet de nombreuses recherches théoriques. Dans le chapitre 9, nous montrerons comment borner la probabilité de succès de cet amplificateur, puis nous nous intéresserons à une des applications les plus prometteuses : la possibilité de compenser les imperfections d’un canal de transmission.

Nous utiliserons ces résultats dans le chapitre 10, afin de montrer que cette propriété permet d’améliorer les performances d’une distribution quantique de clé avec des variables continues. Nous détaillerons en particulier comment la distance maximale de transmission peut être augmentée, ainsi que la résistance au bruit ajouté par le canal de transmission.

Première partie

Outils théoriques et expérimentaux



# Chapitre 2

## Outils théoriques

### Sommaire

---

|            |                                                                 |           |
|------------|-----------------------------------------------------------------|-----------|
| <b>2.1</b> | <b>Champ quantique</b>                                          | <b>8</b>  |
| 2.1.1      | Quantification du champ                                         | 8         |
| 2.1.2      | Variables discrètes & états de Fock                             | 11        |
| 2.1.3      | Variables continues & quadratures du champ                      | 12        |
| 2.1.4      | Lien entre les descriptions discrètes et continues              | 13        |
| <b>2.2</b> | <b>Evolution temporelle et modes du champ</b>                   | <b>14</b> |
| 2.2.1      | Evolution temporelle - représentation d'interaction             | 14        |
| 2.2.2      | Modes du champ                                                  | 14        |
| <b>2.3</b> | <b>Descriptions et propriétés des états quantiques</b>          | <b>17</b> |
| 2.3.1      | Matrice densité                                                 | 18        |
| 2.3.2      | Fonction de Wigner                                              | 20        |
| 2.3.3      | Quelle description choisir ?                                    | 23        |
| 2.3.4      | Distance entre états quantiques : la fidélité                   | 24        |
| 2.3.5      | Mesures et POVM                                                 | 25        |
| <b>2.4</b> | <b>Quelques transformations unitaires</b>                       | <b>25</b> |
| 2.4.1      | Le déphasage                                                    | 25        |
| 2.4.2      | La lame séparatrice                                             | 26        |
| <b>2.5</b> | <b>Les états gaussiens</b>                                      | <b>27</b> |
| 2.5.1      | Etats d'incertitude minimale                                    | 28        |
| 2.5.2      | Le vide quantique                                               | 29        |
| 2.5.3      | Les états cohérents                                             | 29        |
| 2.5.4      | Le vide comprimé monomode                                       | 33        |
| 2.5.5      | Le vide comprimé bimode, ou état EPR                            | 35        |
| 2.5.6      | Les états thermiques                                            | 37        |
| <b>2.6</b> | <b>Etats non gaussiens</b>                                      | <b>39</b> |
| 2.6.1      | Les états de Fock                                               | 40        |
| 2.6.2      | Les superpositions d'états cohérents, ou "chats de Schrödinger" | 41        |
| <b>2.7</b> | <b>Application à l'information quantique</b>                    | <b>42</b> |
| 2.7.1      | Variables discrètes                                             | 43        |
| 2.7.2      | Variables continues                                             | 44        |
| <b>2.8</b> | <b>Conclusion</b>                                               | <b>45</b> |

---

Ce chapitre introduit les outils théoriques nécessaires à la description et à la manipulation des états quantiques du champ lumineux. Nous développerons en particulier des concepts particulièrement utiles en optique quantique avec des variables continues, qui seront utilisés dans tout le reste de ce manuscrit.

## 2.1 Champ quantique

Nous allons commencer par présenter les grandes lignes de la quantification du champ électromagnétique. Après un bref rappel des solutions classiques des équations de Maxwell, nous introduirons les opérateurs quantiques associés au champ, en soulignant l'analogie entre le champ et une collection d'oscillateurs harmoniques. Nous verrons que les états du champ peuvent être décrits de manière équivalente par des états discrets, les photons, ou continus, les quadratures.

### 2.1.1 Quantification du champ

#### Equations classiques

La théorie de l'électromagnétisme classique est donnée par les équations de Maxwell, que l'on peut reformuler en introduisant le potentiel électrique et le potentiel vecteur<sup>1</sup>  $\mathbf{A}(\mathbf{r}, t)$  [Jackson98]. Parmi les différents choix de jauge possible, la jauge de Coulomb  $\nabla \cdot \mathbf{A}(\mathbf{r}, t) = 0$  est particulièrement bien adaptée à l'électrodynamique quantique non relativiste, puisqu'elle rend le potentiel vecteur transverse, séparant ainsi clairement les composantes longitudinales et transverses du champ électromagnétique. En l'absence de charge, le potentiel électrique est nul dans cette jauge, et l'on a

$$\mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{A}(\mathbf{r}, t), \quad (2.1a)$$

$$\mathbf{B}(\mathbf{r}, t) = \nabla \wedge \mathbf{A}(\mathbf{r}, t), \quad (2.1b)$$

où  $\mathbf{A}$  vérifie l'équation de d'Alembert. L'approche usuelle consiste à considérer que le système physique peut être englobé dans un volume  $V=L^3$  de taille arbitraire, afin d'avoir un ensemble discret de modes du champ électromagnétique [Grynberg10] :

$$\mathbf{A}(\mathbf{r}, t) = \sum_l \frac{\mathfrak{E}_l}{\omega_l} \boldsymbol{\epsilon}_l \left( \alpha_l(t) e^{i\mathbf{k}_l \cdot \mathbf{r}} + \alpha_l^*(t) e^{-i\mathbf{k}_l \cdot \mathbf{r}} \right) \quad (2.2)$$

où  $l=(n_x, n_y, n_z, \lambda)$  regroupe tous les indices d'un mode,  $\mathbf{k}_l$  a pour composantes  $(\mathbf{k}_l)_j = 2\pi n_j/L$ ,  $\boldsymbol{\epsilon}_l$  est le vecteur de polarisation orthogonal à  $\mathbf{k}_l$  et indexé par  $\lambda \in \{1, 2\}$ ,  $\omega_l = c\|\mathbf{k}_l\|$ , et  $\mathfrak{E}_l$  est une constante réelle que l'on pose égale à  $\sqrt{\frac{\hbar\omega_l}{2\epsilon_0 L^3}}$ .

On montre ensuite à partir des équations de Maxwell que

$$\alpha_l(t) = \alpha_l e^{-i\omega_l t}, \quad \text{avec } \alpha_l \in \mathbb{C}. \quad (2.3)$$

Les différents modes  $l$  ne sont donc pas couplés entre eux lors d'une évolution libre, et évoluent ainsi de manière indépendante. En utilisant (2.3) et (2.2) dans (2.1a), on obtient la décomposition

---

1. Dans tout ce manuscrit, les grandeurs vectorielles sont écrites en gras.

du champ  $\mathbf{E}(\mathbf{r}, t)$  en *modes normaux* :

$$\mathbf{E}(\mathbf{r}, t) = \sum_l \boldsymbol{\epsilon}_l \mathfrak{E}_l [i\alpha_l e^{i\mathbf{k}_l \cdot \mathbf{r} - i\omega_l t} - i\alpha_l^* e^{-i\mathbf{k}_l \cdot \mathbf{r} + i\omega_l t}] \quad (2.4)$$

On peut alors montrer que l'hamiltonien total  $\mathcal{H} = \frac{\epsilon_0}{2} \int_V d^3r (\|\mathbf{E}^2(\mathbf{r}, t)\| + c^2 \|\mathbf{B}^2(\mathbf{r}, t)\|)$  peut se mettre sous la forme  $\mathcal{H} = \sum_l \mathcal{H}_l$ , avec

$$\mathcal{H}_l = \hbar\omega_l |\alpha_l|^2. \quad (2.5)$$

L'énergie totale est donc la somme des énergies  $\mathcal{H}_l$  associées à chaque mode  $l$ . L'hamiltonien peut s'interpréter comme une somme d'oscillateurs harmoniques indépendants, décrits par les variables conjuguées  $q_l$  et  $p_l$ , définies pour chaque mode par :

$$q_l(t) = \sqrt{2\hbar} \Re\{\alpha_l(t)\} \quad (2.6a)$$

$$p_l(t) = \sqrt{2\hbar} \Im\{\alpha_l(t)\} \quad (2.6b)$$

En effet, on peut alors écrire les hamiltoniens partiels sous la forme

$$\mathcal{H}_l = \frac{\omega_l}{2} (q_l^2 + p_l^2), \quad (2.7)$$

et il apparaît clairement que ces variables obéissent aux équations de Hamilton [Landau12] :

$$\frac{d}{dt} q_l = \frac{\partial}{\partial p_l} \mathcal{H} \quad (2.8a)$$

$$\frac{d}{dt} p_l = -\frac{\partial}{\partial q_l} \mathcal{H} \quad (2.8b)$$

Ainsi, les parties réelles et imaginaires des amplitudes  $\alpha_l$  de chaque mode du champ définissent des *variables conjuguées*, similaires aux variables position et impulsion en mécanique, auxquelles on va pouvoir appliquer la procédure de quantification canonique.

### Opérateur champ quantique

La *quantification* du champ électromagnétique est un problème particulièrement ardu lorsqu'elle est faite rigoureusement, et dépend en particulier de la jauge utilisée [Cohen-Tannoudji01]. Parmi les différentes méthodes, la quantification canonique est celle usuellement utilisée en physique quantique. Aux grandeurs  $q_l(t)$  et  $p_l(t)$  d'un mode  $l$ , on associe des opérateurs hermitiens  $\hat{q}_l$  et  $\hat{p}_l$ , indépendants du temps, vérifiant :

$$[\hat{q}_l, \hat{p}_l] = i\hbar \quad (2.9)$$

avec  $[\hat{q}_l, \hat{p}_l] := \hat{q}_l \hat{p}_l - \hat{p}_l \hat{q}_l$ . Ces opérateurs sont analogues aux opérateurs position  $\hat{x}$  et impulsion  $\hat{p}$  d'une particule matérielle, qui vérifient la même relation de commutation. Puisque deux modes  $l$  et  $m$  différents sont indépendants l'un de l'autre, les opérateurs associés aux grandeurs conjuguées de ces deux modes doivent commuter :

$$[\hat{q}_l, \hat{p}_m] = [\hat{p}_l, \hat{p}_m] = [\hat{q}_l, \hat{q}_m] = 0 \quad (2.10)$$

Par définition (2.6), l'amplitude classique  $\alpha_l(t)$  est égale à  $(q_l(t) + ip_l(t))/\sqrt{2\hbar}$ . Sa "transcription" quantique est donc également un opérateur, noté  $\hat{a}_l$ , et défini par

$$\hat{a}_l = \frac{1}{\sqrt{2\hbar}} (\hat{q}_l + i\hat{p}_l). \quad (2.11)$$



Le complexe conjugué  $\alpha_l^*(t)$  est quand à lui égal à  $(q_l(t) - iq_l(t))/\sqrt{2\hbar}$ . L'opérateur quantique associé est donc  $(\hat{q}_l - i\hat{p}_l)/\sqrt{2\hbar}$ , qui n'est autre que  $\hat{a}_l^\dagger$ , puisque  $\hat{q}_l$  et  $\hat{p}_l$  sont hermitiens. Le processus de quantification revient donc en fin de compte à remplacer les amplitudes  $\alpha_l(t)$  et  $\alpha_l^*(t)$  respectivement par les opérateurs  $\hat{a}_l$  et  $\hat{a}_l^\dagger$ . En procédant de la sorte dans (2.4), on obtient l'opérateur champ électrique  $\hat{E}(\mathbf{r})$  :

$$\boxed{\hat{E}(\mathbf{r}) = \sum_l i\epsilon_l \mathfrak{E}_l (\hat{a}_l e^{i\mathbf{k}_l \cdot \mathbf{r}} - \hat{a}_l^\dagger e^{-i\mathbf{k}_l \cdot \mathbf{r}})} \quad (2.12)$$

Les opérateurs  $\hat{a}_l$  et  $\hat{a}_m^\dagger$  de deux modes  $l$  et  $m$  quelconques sont entièrement définis par leurs commutateurs, calculables à partir des relations de commutation (2.9)<sup>2</sup> :

$$[\hat{a}_l, \hat{a}_m^\dagger] = \delta_{lm}, \quad [\hat{a}_l, \hat{a}_m] = 0. \quad (2.13)$$

Les opérateurs  $\hat{a}_l$  vérifient donc des relations de commutations bosoniques, desquelles découlent tout un ensemble de propriétés fondamentales pour ces opérateurs [Cohen-Tannoudji97c]. La première est que les opérateurs  $\hat{n}_l = \hat{a}_l^\dagger \hat{a}_l$  possèdent des valeurs propres entières positives  $n_l$ , associées à leurs vecteurs propres  $|n_l\rangle$ <sup>3</sup>. Ensuite, les opérateurs  $\hat{a}_l$  et  $\hat{a}_l^\dagger$  transforment un état  $|n_l\rangle$  en

$$\hat{a}_l^\dagger |n_l\rangle = \sqrt{n_l + 1} |n_l + 1\rangle, \quad (2.14)$$

$$\hat{a}_l |n_l\rangle = \sqrt{n_l} |n_l - 1\rangle. \quad (2.15)$$

Pour cette raison, les opérateurs  $\hat{a}_l^\dagger$  et  $\hat{a}_l$  sont respectivement appelés opérateurs de *création* et de *destruction*.

L'hamiltonien donné par l'équation (2.7) s'écrit quantiquement

$$\hat{\mathcal{H}}_l = \frac{\omega_l}{2} (\hat{q}_l^2 + \hat{p}_l^2). \quad (2.16)$$

En injectant dans cette expression les opérateurs  $\hat{q}_l = \sqrt{\hbar}(\hat{a}_l + \hat{a}_l^\dagger)/\sqrt{2}$  et  $\hat{p}_l = i\sqrt{\hbar}(\hat{a}_l^\dagger - \hat{a}_l)/\sqrt{2}$  donnés par (2.11), on obtient finalement

$$\hat{\mathcal{H}}_l = \hbar\omega_l \left( \hat{n}_l + \frac{1}{2} \right). \quad (2.17)$$

Un mode du champ est donc bien analogue avec un oscillateur harmonique matériel. La relation (2.17) montre que les états propres de  $\hat{\mathcal{H}}_l$  sont les mêmes que ceux de  $\hat{n}_l$  : il s'agit des états  $|n_l\rangle$ , appelés états de Fock pour le mode  $l$ . On voit également que deux états  $|n_l\rangle$  et  $|n_l - 1\rangle$  ont une différence d'énergie de  $\hbar\omega_l$ . Un état  $|n_l\rangle$  est donc interprété comme contenant  $n_l$  quanta d'énergie  $\hbar\omega_l$ , les *photons*. Un état du rayonnement peut dans une certaine mesure être interprété comme un flux de particules ayant chacune une énergie  $\hbar\omega_l$ , ce qui permet par exemple d'expliquer l'effet photo-électrique [Einstein05].

2. En posant  $l = (n_x, n_y, n_z, \lambda_1)$  et  $m = (m_x, m_y, m_z, \lambda_2)$ , le symbole  $\delta_{l,m}$  est à interpréter comme  $\delta_{n_x m_x} \delta_{n_y m_y} \delta_{n_z m_z} \delta_{\lambda_1 \lambda_2}$ .

3. La positivité des valeurs propres de  $\hat{n}_l$  est en fait une simple conséquence de la positivité de la norme d'un état, puisque  $\langle n_l | \hat{n}_l | n_l \rangle = n_l \langle n_l | n_l \rangle = n_l = \langle n_l | \hat{a}_l^\dagger \hat{a}_l | n_l \rangle = \| \hat{a}_l | n_l \rangle \|^2 \geq 0$ . Les relations de commutations (2.13) imposent ensuite que ces valeurs propres soient entières.

L'action répétée de  $\hat{a}_l$  décroît la valeur de  $n_l$ , qui finit inévitablement par devenir nulle. L'état propre  $|0_l\rangle$  associé à  $n_l=0$  est appelé le *vide quantique* pour le mode  $l$ . Il peut être défini par

$$\hat{a}_l|0_l\rangle = 0. \quad (2.18)$$

Le vide est un état quantique à part entière, qui occupe une place importante en optique quantique. Nous étudierons plus précisément ses propriétés dans la suite de ce chapitre. L'état vide  $|0\rangle$  pour tous les modes est simplement le produit tensoriel du vide pour chaque mode :  $|0\rangle = \bigotimes_l |0_l\rangle$ .

A l'énergie des photons s'ajoute le terme constant  $\hbar\omega_l/2$ , correspondant à l'énergie du vide. Ce terme provient du fait qu'il s'agit de l'état fondamental de l'hamiltonien  $\hat{\mathcal{H}}_l$ , ne contenant certes aucune excitation, mais pour lequel nous verrons que la variance de l'amplitude du champ électrique n'est pas nulle. Ce sont également les fluctuations du vide qui sont responsables de la désexcitation des niveaux atomiques par émission spontanée [Cohen-Tannoudji96].

Nous n'avons jusqu'à présent considéré qu'un seul mode du champ  $l$ . Le traitement *multimode* se généralise sans difficulté par produit tensoriel, puisque le hamiltonien quantique total est toujours  $\hat{\mathcal{H}} = \sum_l \hat{\mathcal{H}}_l$ . Par exemple, pour des modes  $l_1, \dots, l_N$  excités de respectivement  $n_{l_1}, \dots, n_{l_N}$  photons, l'état correspondant est  $|n_{l_1}\rangle \otimes \dots \otimes |n_{l_N}\rangle$ . Notons que lorsque tous les modes sont considérés, le terme  $\sum_l \hbar\omega_l/2$  tend bien sur vers l'infini. En revanche, puisqu'il est constant et n'influe pas sur la dynamique des états quantiques, il est souvent simplement "oublié" en optique quantique. Le lecteur désirant davantage de précisions pourra se reporter à un ouvrage de théorie quantique des champs [Peskin95, Greiner96].

Cette section sera uniquement consacrée aux deux descriptions du champ quantifié. L'application à l'information quantique, avec en particulier l'utilisation de différents états pour former un qubit, sera approfondie à la fin de ce chapitre, lorsque nous aurons présenté tous les outils nécessaires, ainsi que les états "typiques" du champ.

### 2.1.2 Variables discrètes & états de Fock

La quantification du champ nous a déjà amené à introduire les états de Fock, états propres de l'opérateur nombre. Cet opérateur étant une observable, les états de Fock forment donc une base de l'espace de Hilbert, qui est de plus non dégénérée [Cohen-Tannoudji97c]. Un état quelconque du champ  $|\psi\rangle$  dans un mode donné<sup>4</sup> se décompose selon

$$|\psi\rangle = \sum_n c_n |n\rangle, \quad \text{avec } c_n = \langle n|\psi\rangle. \quad (2.19)$$

Le champ est ici décrit en termes de photons, correspondant à l'aspect corpusculaire. A partir de la relation (2.14), on voit qu'un état  $|n\rangle$  peut être créé en appliquant  $n$  fois l'opérateur de création sur le vide :

$$|n\rangle = \frac{1}{\sqrt{n!}} (\hat{a}^\dagger)^n |0\rangle \quad (2.20)$$

Notons que compte tenu des précédents commentaires, les états de Fock vérifient bien sûr une relation de fermeture discrète :

$$\sum_n |n\rangle \langle n| = \mathbb{I} \quad (2.21)$$

---

4. Nous ne précisons plus les modes lorsqu'un seul d'entre eux est utilisé.

### 2.1.3 Variables continues & quadratures du champ

Afin de faire ressortir l'aspect ondulatoire du champ, introduisons les opérateurs de quadratures  $\hat{\mathbf{X}}_l$  et  $\hat{\mathbf{P}}_l$ , définis par

$$\hat{\mathbf{X}}_l = \hat{\mathbf{q}}_l \sqrt{\frac{2N_0}{\hbar}}, \quad (2.22a)$$

$$\hat{\mathbf{P}}_l = \hat{\mathbf{p}}_l \sqrt{\frac{2N_0}{\hbar}}, \quad (2.22b)$$

avec  $N_0$  une constante réelle et positive, généralement égale à 1 ou 1/2. Ils vérifient une relation de commutation similaire à (2.9) :

$$[\hat{\mathbf{X}}_l, \hat{\mathbf{P}}_m] = 2iN_0\delta_{lm} \quad (2.23)$$

$N_0$  peut donc être vu comme une “redéfinition” de  $\hbar$  : on aurait pu s'en passer en posant directement  $\hbar=2$  ou  $\hbar=1$  dans (2.9). En utilisant les notations (2.22a) et (2.22b), les opérateurs création et destruction s'écrivent maintenant :

$$\hat{a}_l = \frac{1}{2\sqrt{N_0}}(\hat{\mathbf{X}}_l + i\hat{\mathbf{P}}_l), \quad (2.24a)$$

$$\hat{a}_l^\dagger = \frac{1}{2\sqrt{N_0}}(\hat{\mathbf{X}}_l - i\hat{\mathbf{P}}_l), \quad (2.24b)$$

et réciproquement,

$$\hat{\mathbf{X}}_l = (\hat{a}_l + \hat{a}_l^\dagger)\sqrt{N_0}, \quad (2.25a)$$

$$\hat{\mathbf{P}}_l = (\hat{a}_l^\dagger - \hat{a}_l)i\sqrt{N_0}. \quad (2.25b)$$

L'opérateur champ  $\hat{\mathbf{E}}(\mathbf{r})$  donné par (2.12) se réécrit alors

$$\boxed{\hat{\mathbf{E}}(\mathbf{r}) = - \sum_l \epsilon_l \frac{\mathfrak{E}_l}{\sqrt{N_0}} \left( \hat{\mathbf{X}}_l \sin \mathbf{k}_l \cdot \mathbf{r} + \hat{\mathbf{P}}_l \cos \mathbf{k}_l \cdot \mathbf{r} \right).} \quad (2.26)$$

Les spectres des opérateurs de quadratures sont nécessairement continus du fait de la relation de commutation (2.23) [Cohen-Tannoudji97a]. On notera  $|x\rangle$  (resp.  $|p\rangle$ ) l'état propre de  $\hat{\mathbf{X}}$  (resp.  $\hat{\mathbf{P}}$ ) associé à la valeur propre  $x$  (resp.  $p$ ), pour un mode donné. Ces états propres forment chacun une base de l'espace de Hilbert associée à ce mode :

$$\int dx |x\rangle\langle x| = \int dp |p\rangle\langle p| = \mathbb{I} \quad (2.27)$$

Un état quantique quelconque  $|\psi\rangle$  d'un mode du champ pourra donc être décomposé sur ces états propres selon

$$|\psi\rangle = \int dx \psi(x)|x\rangle = \int dp \psi(p)|p\rangle, \quad (2.28)$$

où les coefficients  $\psi(x)=\langle x|\psi\rangle$  et  $\psi(p)=\langle p|\psi\rangle$  sont les fonctions d'ondes de l'état quantique exprimées respectivement dans les bases  $\{|x\rangle\}$  et  $\{|p\rangle\}$ . Notons que le terme de “fonction d'onde” peut être trompeur, car bien qu'elles possèdent les mêmes propriétés que celles d'une particule

matérielle, ces fonctions d'ondes ne correspondent pas à une position ou à une impulsion du photon. Remarquons également que les états propres des quadratures sont des états non physiques car ils correspondent à une énergie infinie, au même titre que les états propres de la position ou de l'impulsion pour une particule matérielle. Seules des superpositions de ces états donnent des états physiques.

Remarquons enfin qu'un état propre de quadrature peut également se décomposer sur les états de Fock,

$$|x\rangle = \sum_n c_n |n\rangle, \quad \text{avec } c_n = \langle n|x\rangle, \quad (2.29)$$

et qu'un état de Fock se décompose sur les quadratures selon  $|n\rangle = \int dx [\langle x|n\rangle] |x\rangle$ .

Puisque les opérateurs de quadratures ne commutent pas, il n'est pas possible de leur trouver une base commune qui les diagonalise, et ils vérifient une relation d'incertitude de Heisenberg :

$$\Delta X_l \Delta P_m \geq \frac{1}{2} |\langle [\hat{X}_l, \hat{P}_m] \rangle| \quad (2.30)$$

où  $\langle . \rangle$  désigne la valeur moyenne, et  $\Delta A = \sqrt{\langle \hat{A}^2 \rangle - \langle \hat{A} \rangle^2}$  l'écart type. Pour un mode  $l$ , cette relation devient :

$$\boxed{\Delta X_l \Delta P_l \geq N_0} \quad (2.31)$$

Cette relation fondamentale est à la base de nombreux phénomènes d'optique quantique, et d'applications en information quantique. Nous verrons plusieurs états pour lesquels elle est une égalité. C'est en particulier le cas du vide  $|0\rangle$ , et plus généralement pour les états purs ayant une fonction d'onde gaussienne<sup>5</sup>. La convention  $N_0=1$  permet donc d'avoir une variance du vide égale à 1, ce qui facilite les calibrations lors des mesures expérimentales des opérateurs quadratures : il suffit de diviser toutes les mesures par celle correspondant à l'écart-type mesuré pour le vide.

On peut enfin définir des quadratures déphasées d'un angle  $\theta$ , qui vérifient bien entendu les mêmes relations de commutation que  $\hat{X}$  et  $\hat{P}$  :

$$\hat{X}_\theta = \hat{X} \cos \theta + \hat{P} \sin \theta \quad (2.32a)$$

$$\hat{P}_\theta = -\hat{X} \sin \theta + \hat{P} \cos \theta \quad (2.32b)$$

#### 2.1.4 Lien entre les descriptions discrètes et continues

En conclusion, nous désignerons par variables discrètes des états propres associés à un opérateur à spectre discret, par opposition aux variables continues, états propres associés à un opérateur à spectre continu. Chacune des deux descriptions contient la même information sur l'état quantique considéré : on peut former des variables discrètes en formant un "paquet d'onde" avec des variables continues, et inversement obtenir un continuum d'états en superposant de manière adéquat des états discrets.

Le mélange des descriptions discrètes et continues est au cœur de ce travail de thèse. Nous verrons que les outils naturellement adaptés aux variables discrètes peuvent être utilisés pour des états décrits par des variables continues, et vice versa, produisant ainsi des états particulièrement intéressants pour l'information quantique.

5. Puisque les deux représentations sont liées par une transformée de Fourier, être gaussien pour une représentation assure que l'autre l'est également.

## 2.2 Evolution temporelle et modes du champ

### 2.2.1 Evolution temporelle - représentation d'interaction

En optique quantique, de nombreuses transformations sont obtenues en échangeant des photons entre différents modes : une lame séparatrice échange des photons de même énergie entre deux modes spatiaux différents ; les milieux non linéaires peuvent transformer un photon pompe en plusieurs photons d'énergie différente. Nous reviendrons en détail sur ces transformations dans la suite de ce manuscrit. Remarquons simplement pour l'instant que pour une interaction résonante (c'est-à-dire assurant la conservation de l'énergie), l'hamiltonien de couplage est de la forme

$$\hat{\mathcal{W}} = \sum_{\substack{p \neq q \\ \omega_p = \omega_q}} \gamma_{p,q} \hat{\mathbf{a}}_p^\dagger \hat{\mathbf{a}}_q + \sum_{\substack{p,q,m \\ \omega_p + \omega_q = \omega_m}} \zeta_{p,q,m} \hat{\mathbf{a}}_p^\dagger \hat{\mathbf{a}}_q^\dagger \hat{\mathbf{a}}_m + \text{h.c.} + \dots, \quad (2.33)$$

où les indices  $p, q, m, \dots$  des sommes contiennent tous les indices décrivant un mode. On vérifie que la résonance de l'interaction conduit à la relation <sup>6</sup>

$$[\hat{\mathcal{H}}_0, \hat{\mathcal{W}}] = 0, \quad (2.34)$$

où  $\hat{\mathcal{H}}_0 = \sum_m \hbar \omega_m \hat{\mathbf{a}}_m^\dagger \hat{\mathbf{a}}_m$  est l'hamiltonien libre. On pourra donc séparer l'évolution totale  $\hat{U}(t) = \exp[-\frac{i}{\hbar} t (\hat{\mathcal{H}}_0 + \hat{\mathcal{W}})]$  contenant un couplage (2.33) en une évolution libre  $\hat{U}_0(t) = \exp[-\frac{i}{\hbar} t \hat{\mathcal{H}}_0]$ , et une évolution  $\hat{U}_0^\dagger(t) \hat{U}(t) = \exp[-\frac{i}{\hbar} t \hat{\mathcal{W}}]$  due au couplage seul.

Dans ce manuscrit, nous serons de ce fait toujours en *représentation d'interaction*, en appliquant l'évolution libre à l'opérateur champ :

$$\hat{\mathbf{E}}(\mathbf{r}, t) = \hat{U}_0^\dagger(t) \hat{\mathbf{E}}(\mathbf{r}) \hat{U}_0(t) \quad (2.35a)$$

$$= \sum_l i \epsilon_l \mathfrak{E}_l \left( \hat{\mathbf{a}}_l e^{i \mathbf{k}_l \cdot \mathbf{r} - i \omega_l t} - \hat{\mathbf{a}}_l^\dagger e^{-i \mathbf{k}_l \cdot \mathbf{r} + i \omega_l t} \right) \quad (2.35b)$$

$$= - \sum_l \epsilon_l \frac{\mathfrak{E}_l}{\sqrt{N_0}} \left( \hat{\mathbf{X}}_l \sin(\mathbf{k}_l \cdot \mathbf{r} - \omega_l t) + \hat{\mathbf{P}}_l \cos(\mathbf{k}_l \cdot \mathbf{r} - \omega_l t) \right) \quad (2.35c)$$

Les états quantiques évoluent alors selon :

$$|\psi(t_1)\rangle = e^{-\frac{i}{\hbar} \hat{\mathcal{W}}(t_1 - t_0)} |\psi(t_0)\rangle \quad (2.36)$$

En fonction des situations, on pourra plutôt faire évoluer les opérateurs création ou destruction, en gardant les états quantiques indépendants du temps. Dans ce cas, on aura par exemple

$$\hat{\mathbf{a}}(t_1) = e^{+\frac{i}{\hbar} \hat{\mathcal{W}}(t_1 - t_0)} \hat{\mathbf{a}}(t_0) e^{-\frac{i}{\hbar} \hat{\mathcal{W}}(t_1 - t_0)}. \quad (2.37)$$

### 2.2.2 Modes du champ

#### Limite continue

La boîte dans laquelle l'expérience est "enfermée" n'étant que virtuelle, les résultats physiques ne doivent bien entendu pas en dépendre. Nous pourrions nous contenter de garder une

---

6. Grâce à la condition de résonance, l'évolution libre de  $\hat{\mathcal{W}}$  introduit des facteurs de phases qui sont tous compensés :  $\hat{U}_0^\dagger(t) \hat{\mathcal{W}} \hat{U}_0(t) = \hat{\mathcal{W}}$ . On en déduit donc que  $[\hat{\mathcal{H}}_0, \hat{\mathcal{W}}] = 0$ .

sommation discrète, toutefois il est intéressant, dans le cas d'un champ en régime impulsionnel, de s'affranchir de cette boîte en faisant tendre son volume vers l'infini. Repartons pour cela de l'expression de l'opérateur  $\hat{\mathbf{E}}(\mathbf{r})$  obtenue en (2.12) :

$$\hat{\mathbf{E}}(\mathbf{r}) = \sum_l i\epsilon_l \sqrt{\frac{\hbar\omega_l}{2\epsilon_0}} \frac{1}{\sqrt{L^3}} \left( \hat{\mathbf{a}}_l e^{i\mathbf{k}_l \cdot \mathbf{r}} - \hat{\mathbf{a}}_l^\dagger e^{-i\mathbf{k}_l \cdot \mathbf{r}} \right) \quad (2.38)$$

Le volume élémentaire  $\Delta k_j$  dans l'espace des  $\mathbf{k}$  est de  $2\pi/L$  par degré de liberté  $j$ . En posant  $\Delta^3 \mathbf{k} = \Delta k_x \Delta k_y \Delta k_z$ , on a donc  $1/\sqrt{L^3} = [\Delta^3 \mathbf{k}/(2\pi)]^{3/2}$ . On peut alors écrire

$$\hat{\mathbf{E}}(\mathbf{r}) = \sum_l i\epsilon_l \sqrt{\frac{\hbar\omega_l}{16\epsilon_0\pi^3}} \left( \frac{\hat{\mathbf{a}}_l}{\sqrt{\Delta^3 \mathbf{k}}} e^{i\mathbf{k}_l \cdot \mathbf{r}} - \frac{\hat{\mathbf{a}}_l^\dagger}{\sqrt{\Delta^3 \mathbf{k}}} e^{-i\mathbf{k}_l \cdot \mathbf{r}} \right) \Delta^3 \mathbf{k}. \quad (2.39)$$

Lorsque  $L \rightarrow \infty$ ,  $\Delta^3 \mathbf{k} \rightarrow 0$ , et on peut donc remplacer la somme par une intégrale sur  $\mathbf{k}$  :

$$\hat{\mathbf{E}}(\mathbf{r}) = \sum_{\lambda=1,2} \int d^3 \mathbf{k} i\epsilon_\lambda(\mathbf{k}) \sqrt{\frac{\hbar c \|\mathbf{k}\|}{16\epsilon_0\pi^3}} \left( \hat{\mathbf{a}}_\lambda(\mathbf{k}) e^{i\mathbf{k} \cdot \mathbf{r}} - \hat{\mathbf{a}}_\lambda^\dagger(\mathbf{k}) e^{-i\mathbf{k} \cdot \mathbf{r}} \right) \quad (2.40)$$

où  $\epsilon_\lambda(\mathbf{k})$  est le vecteur polarisation associé au vecteur  $\mathbf{k}$  et à la polarisation  $\lambda$ , et où l'on a posé

$$\hat{\mathbf{a}}_\lambda(\mathbf{k}) = \lim_{\Delta^3 \mathbf{k} \rightarrow 0} \frac{\hat{\mathbf{a}}_l}{\sqrt{\Delta^3 \mathbf{k}}}. \quad (2.41)$$

Les commutateurs du champ (2.13) sont également transformés lors du passage à la limite continue [Cohen-Tannoudji01] :

$$[\hat{\mathbf{a}}_{\lambda_1}(\mathbf{k}), \hat{\mathbf{a}}_{\lambda_2}^\dagger(\mathbf{q})] = \delta_{\lambda_1 \lambda_2} \delta(\mathbf{k} - \mathbf{q}) \quad (2.42)$$

Les commutateurs associés à deux opérateurs de destruction ou de création ne sont quant à eux pas changés, ils sont toujours nuls.

### Modes impulsionnels

L'expression (2.40), décompose le champ  $\hat{\mathbf{E}}(\mathbf{r})$  dans la "base" des ondes planes  $\exp(i\mathbf{k} \cdot \mathbf{r})$ . Un état de Fock associé à ces modes correspond donc à un nombre de quanta d'énergie présents dans une onde d'extensions spatiale et temporelle infinies. La base des ondes planes n'est pas forcément la plus pertinente, et il peut être intéressant de décomposer le champ sur d'autres bases de modes pour décrire nos expériences, qui utilisent des impulsions temporelles associées à un profil spatial d'extension fini.

Considérons donc une famille de fonctions  $\{\phi_m(\mathbf{k}, \lambda)\}$ , avec  $m$  un indice discret, vérifiant les deux propriétés suivantes :

$$\sum_\lambda \int d^3 \mathbf{k} \phi_m^*(\mathbf{k}, \lambda) \phi_n(\mathbf{k}, \lambda) = \delta_{mn} \quad (2.43)$$

$$\sum_m \phi_m^*(\mathbf{k}, \lambda) \phi_m(\mathbf{q}, \sigma) = \delta_{\lambda\sigma} \delta(\mathbf{k} - \mathbf{q}) \quad (2.44)$$

Ces deux relations, respectivement dites d'orthogonalité et de fermeture, assurent que les fonctions  $\{\phi_m(\mathbf{k}, \lambda)\}$  constituent bien une base orthonormée sur laquelle on peut décomposer les

modes du champ. Dans cette nouvelle base, on peut ensuite définir des opérateurs  $\hat{\mathbf{c}}_m$  associés aux modes  $m$  :

$$\hat{\mathbf{c}}_m = \sum_{\lambda} \int d^3\mathbf{k} \phi_m^*(\mathbf{k}, \lambda) \hat{\mathbf{a}}_{\lambda}(\mathbf{k}) \quad (2.45)$$

Ces opérateurs vérifient des relations de commutation bosonique, compte tenu de (2.43) :

$$[\hat{\mathbf{c}}_m, \hat{\mathbf{c}}_n^{\dagger}] = \sum_{\lambda} \sum_{\sigma} \int d^3\mathbf{k} d^3\mathbf{q} \phi_m^*(\mathbf{k}, \lambda) \phi_n(\mathbf{q}, \sigma) [\hat{\mathbf{a}}_{\lambda}(\mathbf{k}), \hat{\mathbf{a}}_{\sigma}^{\dagger}(\mathbf{q})] \quad (2.46a)$$

$$= \sum_{\lambda} \int d^3\mathbf{k} \phi_m^*(\mathbf{k}, \lambda) \phi_n(\mathbf{k}, \lambda) \quad (2.46b)$$

$$= \delta_{mn} \quad (2.46c)$$

On pourra donc définir des états nombres  $|n_m\rangle$  associés aux modes  $m$ , même si dans le cas général ces états ne sont plus états propres de l'hamiltonien.

Les opérateurs  $\hat{\mathbf{a}}_{\lambda}(\mathbf{k})$  peuvent facilement s'exprimer en fonction des nouveaux modes, compte tenu de (2.44) :

$$\sum_m \hat{\mathbf{c}}_m \phi_m(\mathbf{k}, \lambda) = \sum_m \left( \sum_{\sigma} \int d^3\mathbf{q} \phi_m^*(\mathbf{q}, \sigma) \hat{\mathbf{a}}_{\sigma}(\mathbf{q}) \right) \phi_m(\mathbf{k}, \lambda) \quad (2.47a)$$

$$= \sum_{\sigma} \int d^3\mathbf{q} \left( \sum_m \phi_m^*(\mathbf{q}, \sigma) \phi_m(\mathbf{k}, \lambda) \right) \hat{\mathbf{a}}_{\sigma}(\mathbf{q}) \quad (2.47b)$$

$$= \hat{\mathbf{a}}_{\lambda}(\mathbf{k}) \quad (2.47c)$$

Insérons maintenant cette décomposition dans l'expression du champ  $\hat{\mathbf{E}}(\mathbf{r}, t)$  (2.35), en posant  $\sqrt{\hbar c \|\mathbf{k}\|} / 16\epsilon_0 \pi^3 = N(k)$  :

$$\hat{\mathbf{E}}(\mathbf{r}, t) = \sum_{\lambda} \int d^3\mathbf{k} N(k) i\epsilon_{\lambda}(\mathbf{k}) \left( \sum_m \hat{\mathbf{c}}_m \phi_m(\mathbf{k}, \lambda) \right) e^{i\mathbf{k} \cdot \mathbf{r} - i\omega_k t} + \text{h.c} \quad (2.48a)$$

$$= \sum_m \hat{\mathbf{c}}_m \underbrace{\left( \sum_{\lambda} \int d^3\mathbf{k} N(k) i\epsilon_{\lambda}(\mathbf{k}) \phi_m(\mathbf{k}, \lambda) e^{i\mathbf{k} \cdot \mathbf{r} - i\omega_k t} \right)}_{\mathbf{u}_m(\mathbf{r}, t)} + \text{h.c} \quad (2.48b)$$

$$= \sum_m \hat{\mathbf{c}}_m \mathbf{u}_m(\mathbf{r}, t) + \text{h.c} \quad (2.48c)$$

Notons que puisque  $\omega \geq 0$ , la fonction vectorielle  $\mathbf{u}_m$  ne contient que des fréquences positives, et correspond donc à une représentation analytique. Ainsi, même lorsque le champ n'est pas quantifié dans une boîte, on peut le décomposer sur un ensemble discret de modes spatio-temporels, que l'on peut choisir en fonction du problème que l'on considère. Cette décomposition est particulièrement utile lorsque seuls quelques modes sont peuplés, ce qui permet de simplifier l'expression du champ en laissant de côté les autres modes. Lorsque le champ est décrit par un mode spatio-temporel donné, on pourra en particulier toujours l'inclure dans une base, en vertu du procédé d'orthonormalisation de Gram-Schmidt [Warusfel04].

On peut naturellement définir des opérateurs de quadratures pour un mode impulsif. Considérons par exemple un mode spatio-temporel <sup>7</sup>  $\mathbf{u}_0$  dont le spectre temporel est centré autour

7. Nous avons montré avec (2.48) qu'il faut travailler avec la représentation analytique du mode. Ceci justifie que nous prenions directement  $e^{-i\omega_0 t}$  et non  $\cos \omega_0 t$ .

d'une fréquence  $\omega_0$  :  $\mathbf{u}_0(\mathbf{r}, t) = \mathcal{A}(\mathbf{r}, t)e^{-i\omega_0 t}$ . Supposons maintenant que l'enveloppe  $\mathcal{A}$  prenne des valeurs réelles. L'opérateur champ correspondant à ce mode s'écrit alors :

$$\hat{\mathbf{E}}_0(\mathbf{r}, t) = \mathcal{A}(\mathbf{r}, t) \left( \hat{\mathbf{c}}_0 e^{-i\omega_0 t} + \hat{\mathbf{c}}_0^\dagger e^{+i\omega_0 t} \right) \quad (2.49)$$

En posant comme précédemment

$$\hat{\mathbf{X}}_0 = (\hat{\mathbf{c}}_0 + \hat{\mathbf{c}}_0^\dagger) \sqrt{N_0}, \quad (2.50a)$$

$$\hat{\mathbf{P}}_0 = (\hat{\mathbf{c}}_0^\dagger - \hat{\mathbf{c}}_0) i \sqrt{N_0}, \quad (2.50b)$$

on retrouve une description du mode  $\mathbf{u}_0$  en termes de quadratures :

$$\hat{\mathbf{E}}_0(\mathbf{r}, t) = \frac{1}{\sqrt{N_0}} \mathcal{A}(\mathbf{r}, t) \left( \hat{\mathbf{X}}_0 \cos \omega_0 t + \hat{\mathbf{P}}_0 \sin \omega_0 t \right) \quad (2.51)$$

Puisque les opérateurs  $\hat{\mathbf{X}}_0$  et  $\hat{\mathbf{P}}_0$  ont également un commutateur égal à  $2iN_0$ , toutes les conclusions précédentes s'appliquent.

L'expression (2.48) a été obtenu sans hypothèse particulière sur la forme du champ. En conséquence, elle s'avère peu maniable pour décrire de manière quantitative un mode du champ. Plusieurs simplifications sont fréquemment trouvées dans la littérature (voir par exemple [Chiao08]) : la première est celle d'un champ scalaire, permettant de sortir le vecteur polarisation  $\boldsymbol{\epsilon}_\lambda(\mathbf{k})$  de l'intégrale sur  $\mathbf{k}$ . Elle est valable dans le cadre de l'approximation paraxiale, selon laquelle le faisceau diverge peu par rapport à l'axe optique. On peut également conserver le caractère vectoriel du champ, en gardant l'approximation paraxiale. L'expression détaillée du champ dans ce cadre peut être trouvée dans [Aiello05, Calvo05]. Enfin, lorsque le mode est centré sur une fréquence centrale, on peut faire l'hypothèse de l'enveloppe lentement variable. Toutes ces hypothèses sont bien sûr reliées, et le lecteur trouvera davantage de détails dans le cadre de l'optique classique dans [Saleh91].

En conclusion, les modes discrets  $\hat{\mathbf{c}}_m$  définis par (2.45) ont des propriétés analogues aux modes normaux associés aux ondes planes. A chaque mode est associé un opérateur de création et de destruction bosoniques correspondant à un nombre d'excitations dans ce mode. Dans la suite de ce manuscrit, nous utiliserons des opérateurs similaires associés à des modes impulsionnels, sauf mention contraire explicite.

## 2.3 Descriptions et propriétés des états quantiques

Nous avons déjà présenté au cours de la section précédente deux représentations possibles du champ : la description en termes de variables discrètes, qui est une suite infinie de coefficients  $(c_0, \dots, c_n, \dots)$  correspondant à la décomposition du champ dans la base de Fock ; et la description en termes de variables continues, qui est une fonction  $\psi(x)$  (ou  $\psi(p)$ ) correspondant à la décomposition du champ dans une base de quadratures. Ces deux descriptions équivalentes sont toutefois limitées : elles ne permettent de décrire qu'un état quantique qui est pur.

Dans de très nombreux cas expérimentaux, il manque cependant une information de nature *classique* sur l'état quantique. Si l'appareil de mesure ne nous donne qu'une information partielle sur le résultat, s'il y a un couplage incontrôlé avec l'environnement, ou encore si une partie du système n'est pas observée, on ne pourra associer qu'une probabilité  $p_k$  d'être dans un état pur appartenant à un ensemble d'états  $\{|\psi_k\rangle\}$ . L'état est alors dans un mélange statistique.

Les descriptions que nous allons rappeler maintenant permettent de tenir compte de cette incertitude classique en décrivant des mélanges statistiques. La matrice densité est une extension des variables discrètes, alors que la fonction de Wigner est une extension des variables continues.



### 2.3.1 Matrice densité

#### Définition et propriétés

Rappelons le formalisme de la matrice densité en considérant un système quantique dans un mélange statistique : son état est un état pur  $|\psi_k\rangle$ , avec une certaine probabilité  $p_k$  (avec  $\sum_k p_k=1$ ). Dans ce cas, il est facile de se convaincre que la valeur moyenne d'une observable  $\hat{A}$  dans cet état doit être égale à la somme pondérée des valeurs moyennes de  $\hat{A}$  dans chacun des états  $|\psi_k\rangle$  :

$$\langle \hat{A} \rangle = \sum_k p_k \langle \psi_k | \hat{A} | \psi_k \rangle \quad (2.52)$$

Puisque  $\langle \psi_k | \hat{A} | \psi_k \rangle = \text{Tr}\{\hat{A} |\psi_k\rangle \langle \psi_k|\}$ , on obtient :

$$\langle \hat{A} \rangle = \sum_k p_k \text{Tr}\{\hat{A} |\psi_k\rangle \langle \psi_k|\} \quad (2.53a)$$

$$= \text{Tr}\left\{\hat{A} \left( \sum_k p_k |\psi_k\rangle \langle \psi_k| \right)\right\} \quad (2.53b)$$

La valeur moyenne de  $\hat{A}$  est ainsi entièrement définie par l'opérateur  $\hat{\rho} = \sum_k p_k |\psi_k\rangle \langle \psi_k|$ , qui est par définition la matrice densité du système. On montre ensuite que la matrice densité est hermitienne et de trace unité :

$$\hat{\rho}^\dagger = \hat{\rho} \quad \text{et} \quad \text{Tr}\{\hat{\rho}\} = 1 \quad (2.54)$$

Elle est également positive, c'est-à-dire que ses termes diagonaux sont positifs pour un état  $|u\rangle$  quelconque :

$$\langle u | \hat{\rho} | u \rangle \geq 0 \quad (2.55)$$

Elle peut donc toujours être diagonalisée, avec des valeurs propres positives. D'une manière générale, les termes diagonaux – les *populations* – correspondent aux probabilités de mesurer les différents états de la base dans laquelle la matrice densité est écrite. Les termes hors diagonale – les *cohérences* – témoignent d'une cohérence quantique due à une superposition des différents états de base. Lorsque les cohérences sont toutes nulles, la matrice densité est sous forme diagonale : l'état est dans un mélange statistique des états de base. Dans le cas où l'état est pur, un seul terme diagonal doit donc être non nul lorsque la matrice densité est diagonalisée. La matrice densité est dans ce cas un projecteur sur le sous-espace engendré par l'état propre  $|\psi\rangle$  :  $\hat{\rho} = |\psi\rangle \langle \psi|$ . Cette propriété n'est en revanche plus vraie pour un mélange statistique :  $\hat{\rho}^2 \neq \hat{\rho}$ . La grandeur

$$\mathcal{P} = \text{Tr}\{\hat{\rho}^2\} \quad (2.56)$$

définit alors la *pureté* de l'état. Elle vaut 1 si et seulement si l'état est pur, et est inférieure à 1 dans le cas d'un mélange statistique. C'est donc un moyen très rapide afin de vérifier un calcul numérique, lorsque l'on sait par exemple que l'état final doit rester pur.

### Trace partielle et purification

**Trace partielle** Supposons que le système soit constitué de deux modes 1 et 2 ayant pour espaces de Hilbert  $\mathcal{H}_1$  et  $\mathcal{H}_2$ , de bases respectives  $\{|u_m\rangle\}$  et  $\{|v_k\rangle\}$ . Ces deux modes peuvent correspondre par exemple à deux directions de propagation d'une impulsion lumineuse. L'espace de Hilbert total est donc  $\mathcal{H}=\mathcal{H}_1\otimes\mathcal{H}_2$ , de base  $\{|u_m\rangle\otimes|v_k\rangle\}$ . Un état quelconque de ce système sera décrit par une matrice densité

$$\hat{\rho} = \sum_{\substack{m,m' \\ k,k'}} \rho_{m,m',k,k'} |u_m\rangle\langle u_{m'}| \otimes |v_k\rangle\langle v_{k'}|. \quad (2.57)$$

Considérons maintenant un opérateur  $\hat{\Pi}_1$  agissant sur  $\mathcal{H}_1$ , et son extension  $\hat{\Pi}=\hat{\Pi}_1\otimes\mathbb{I}_2$  agissant sur l'espace total  $\mathcal{H}$ . La valeur moyenne de  $\hat{\Pi}$  est comme précédemment donnée par :

$$\langle \hat{\Pi} \rangle = \text{Tr}\{\hat{\Pi}\hat{\rho}\} \quad (2.58a)$$

$$= \sum_{p,q} \langle u_p|\otimes\langle v_q| \left( \sum_{m,m',k,k'} \rho_{m,m',k,k'} (\hat{\Pi}_1|u_m\rangle)\langle u_{m'}|\otimes|v_k\rangle\langle v_{k'}| \right) |u_p\rangle\otimes|v_q\rangle \quad (2.58b)$$

$$= \sum_p \langle u_p|\hat{\Pi}_1 \left( \sum_{m,m'} \left( \sum_q \rho_{m,m',q,q} \right) |u_m\rangle\langle u_{m'}| \right) |u_p\rangle \quad (2.58c)$$

$$= \text{Tr}_1\{\hat{\Pi}_1\hat{\rho}_1\} \quad (2.58d)$$

où l'on a posé  $\hat{\rho}_1 = \sum_{m,m'} \left( \sum_q \rho_{m,m',q,q} \right) |u_m\rangle\langle u_{m'}|$ , et où  $\text{Tr}_1\{.\}$  désigne la trace sur  $\mathcal{H}_1$ .

La matrice densité  $\hat{\rho}_1$  permet donc de caractériser l'état du système restreint à  $\mathcal{H}_1$ , en "oubliant" le mode 2. Elle est obtenue en prenant la trace partielle de la matrice densité totale sur  $\mathcal{H}_2$ ,

$$\hat{\rho}_1 = \text{Tr}_2\{\hat{\rho}\}, \quad (2.59)$$

et il est facile de se convaincre qu'elle est également hermitienne, et de trace unité.

Notons enfin que les matrices densité réduites d'un état pur bipartite intriqué seront forcément des mélanges statistiques, car sinon il serait possible d'assigner un vecteur d'état à chaque mode, et l'état ne serait donc pas intriqué.

**Purification** La purification d'un système [Nielsen00] est en quelque sorte l'opération inverse de la trace partielle. Un état mélange statistique peut en effet toujours être interprété comme la trace partielle d'un état pur contenant d'autres modes, pouvant être "fictifs". Prenons l'exemple d'un spin 1/2 dépolarisé, de matrice densité

$$\hat{\rho} = \frac{1}{2} \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}. \quad (2.60)$$

Son espace de Hilbert  $\mathcal{H}_s$  est de dimension 2. Etendons le avec un autre espace  $\mathcal{H}_p$  de même dimension, et de base  $\{|A\rangle, |B\rangle\}$ , afin de former un espace total  $\mathcal{H}=\mathcal{H}_s\otimes\mathcal{H}_p$ . La matrice densité donnée par (2.60) peut alors être obtenue en prenant la trace partielle sur  $\mathcal{H}_p$  de l'état pur  $|\phi\rangle = \frac{1}{\sqrt{2}}(|+z\rangle\otimes|A\rangle + |-z\rangle\otimes|B\rangle)$  :

$$\hat{\rho} = \text{Tr}_{\mathcal{H}_p}\{|\phi\rangle\langle\phi|\} \quad (2.61)$$

L'état  $|\phi\rangle$  est une *purification* de  $\hat{\rho}$ . Il s'agit d'une purification et non de *la* purification, car on aurait tout aussi bien pu utiliser par exemple

$$|\tilde{\phi}\rangle = \frac{1}{\sqrt{2}}|+_z\rangle \otimes \left(\frac{|A\rangle+|B\rangle}{\sqrt{2}}\right) + \frac{1}{\sqrt{2}}|-_z\rangle \otimes \left(\frac{|A\rangle-|B\rangle}{\sqrt{2}}\right), \quad (2.62)$$

qui redonne la même matrice densité  $\hat{\rho}$  après trace partielle. En fait, toutes les purifications d'un mélange statistique donné ne diffèrent que d'une transformation unitaire appliquée sur l'espace ajouté pour la purification [Nielsen00].

### 2.3.2 Fonction de Wigner

#### Définition

La matrice densité est une description bien adaptée pour les états naturellement décrits en base de Fock, et pour lesquels on peut se restreindre à un sous-espace ne contenant que quelques photons. Elle devient en revanche assez peu pratique lorsque les états sont plutôt décrits en termes de variables continues. On peut alors être tenté d'introduire une distribution dans un espace des phases comme on le fait classiquement. Seulement, puisque l'on ne peut pas définir en même temps les deux quadratures avec une précision arbitraire, contrairement à une description classique, on pressent que l'on risque de se heurter à quelques difficultés afin de définir et interpréter physiquement un tel objet.

La fonction de Wigner  $W(x, p)$ , pour un mode donné, n'est donc pas définie comme étant une distribution de probabilité, mais selon un critère moins fort. La seule propriété qu'il lui est demandée est de pouvoir obtenir la distribution de probabilité pour une quadrature quelconque  $\hat{X}_\theta$ , en intégrant selon l'autre quadrature [Leonhardt97] :

$$\text{pr}(x_\theta, \theta) = \langle x_\theta | \hat{\rho} | x_\theta \rangle \quad (2.63a)$$

$$= \int dp_\theta W(x_\theta \cos \theta - p_\theta \sin \theta, x_\theta \sin \theta + p_\theta \cos \theta) \quad (2.63b)$$

où les états  $|x_\theta\rangle$  sont états propres des quadratures (2.32) Pour  $\theta=0$  cette formule se simplifie en

$$\text{pr}(x_\theta=x) = \int dp_\theta W(x, p_\theta) = \int dp W(x, p) = \langle x | \hat{\rho} | x \rangle, \quad (2.64)$$

et pour  $\theta=\pi/2$ ,

$$\text{pr}(x_\theta=p) = \int dp_\theta W(-p_\theta, p) = \int dx W(x, p) = \langle p | \hat{\rho} | p \rangle. \quad (2.65)$$

Il est remarquable que le lien avec la matrice densité puisse ensuite être établi uniquement grâce à la définition (2.63). On peut en effet montrer [Leonhardt97] que la fonction caractéristique  $\mathcal{W}[u, v]$ , définie comme la transformée de Fourier de la fonction de Wigner,

$$\mathcal{W}[u, v] = \int dx dp W(x, p) e^{-iu x - iv p}, \quad (2.66)$$

est reliée à la matrice densité par la relation :

$$\mathcal{W}[u, v] = \text{Tr}\{\hat{\rho} \exp(-iu\hat{X} - iv\hat{P})\} \quad (2.67a)$$

$$= \langle \exp(-iu\hat{X} - iv\hat{P}) \rangle \quad (2.67b)$$

La transformée de Fourier de la fonction de Wigner est donc l'équivalent quantique de la fonction caractéristique d'une distribution de probabilité, présentée par exemple dans [Appel08]. Après quelques étapes de calculs, on en déduit enfin la formule de Wigner, donnant l'expression explicite de la fonction de Wigner à partir de la matrice densité :

$$W(x, p) = \frac{1}{2N_0} \frac{1}{2\pi} \int d\nu e^{i\frac{\nu p}{2N_0}} \langle x - \frac{\nu}{2} | \hat{\rho} | x + \frac{\nu}{2} \rangle \quad (2.68)$$

Pour un état multimode  $\hat{\rho} = \hat{\rho}_1 \otimes \dots \otimes \hat{\rho}_N$ , la fonction de Wigner totale est simplement le produit des fonctions de chaque mode :

$$W(x_1, p_1; \dots; x_N, p_N) = W_1(x_1, p_1) \dots W_N(x_N, p_N) \quad (2.69)$$

Enfin, si les quadratures de  $N$  modes sont transformées linéairement en  $\mathbf{x}' = M\mathbf{x}$ , où  $\mathbf{x} := (x_1, p_1, \dots, x_N, p_N)^T$  est un vecteur regroupant les quadratures avant la transformation et  $\mathbf{x}'$  est défini de la même manière pour les quadratures après la transformation, la nouvelle fonction de Wigner  $W'$  vérifie

$$W'(\mathbf{x}') = W(M^{-1}\mathbf{x}). \quad (2.70)$$

### Propriétés

La fonction de Wigner possède de nombreuses propriétés. Nous n'allons présenter ici que les plus importantes sans les démontrer, les démonstrations pouvant être trouvées dans [Leonhardt97]. On montre tout d'abord que puisque les distributions de probabilité des quadratures sont normalisées, la fonction de Wigner doit l'être également :

$$\int dx dp W(x, p) = 1 \quad (2.71)$$

La généralisation de la fonction de Wigner pour un opérateur monomode  $\hat{A}$  quelconque se fait de manière équivalente à la définition (2.68) [Leonhardt97] :

$$W_A(x, p) = \frac{1}{2N_0} \frac{1}{2\pi} \int d\nu e^{i\frac{\nu p}{2N_0}} \langle x - \frac{\nu}{2} | \hat{A} | x + \frac{\nu}{2} \rangle \quad (2.72)$$

Ceci permet de calculer la trace d'un produit d'opérateurs à partir de leurs fonctions de Wigner :

$$\text{Tr}\{\hat{A}\hat{B}\} = 2\pi 2N_0 \int dx dp W_A(x, p) W_B(x, p) \quad (2.73)$$

Cette formule est sans doute une des plus importantes, car elle permet d'obtenir plusieurs grandeurs. On peut calculer la valeur moyenne d'un opérateur,

$$\langle \hat{A} \rangle = \text{Tr}\{\hat{\rho}\hat{A}\} = 2\pi 2N_0 \int dx dp W(x, p) W_A(x, p), \quad (2.74)$$

et la fidélité entre un état  $\hat{\rho}$  quelconque et un état pur  $|\psi\rangle$  quelconque :

$$\langle \psi | \hat{\rho} | \psi \rangle = 2\pi 2N_0 \int dx dp W(x, p) W_{|\psi\rangle\langle\psi|}(x, p) \quad (2.75)$$

Pour des états ayant une fidélité égale à zéro, l'intégrale ne peut être nulle que s'il existe des zones où une des deux fonctions de Wigner est négative. Nous avons donc un argument simple pour justifier que la fonction de Wigner ne peut pas être positive pour tous les états, et que de ce fait elle ne peut pas être interprétée comme une distribution de probabilité. La négativité de la fonction de Wigner est d'ailleurs un outils extrêmement important pour quantifier le degré "quantique" d'un état. Selon le théorème de Hudson-Piquet, la fonction de Wigner d'un état pur est partout positive si et seulement si elle est gaussienne [Hudson74].

La purété d'un état s'obtient avec :

$$\mathcal{P} = 2\pi 2N_0 \int dx dp W^2(x, p) \quad (2.76)$$

Enfin, on peut obtenir les éléments de matrice dans n'importe quelle base, ce qui montre bien que la fonction de Wigner détermine entièrement l'état quantique et que sa donnée est complètement équivalente à celle de la matrice densité :

$$\langle a | \hat{\rho} | b \rangle = 2\pi 2N_0 \int dx dp W(x, p) W_{|a\rangle\langle b|}(x, p) \quad (2.77)$$

## Fonction P

Si l'on cherche à définir une représentation en prenant un analogue quantique de la fonction caractéristique tel que (2.67), la fonction de Wigner n'est pas la seule possibilité. En fait, puisque  $\hat{X}$  et  $\hat{P}$  ne commutent pas, l'exponentielle classique  $e^{-iux-iup}$  peut être "quantifiée" de plusieurs façons. La fonction  $P$  est la distribution de fonction caractéristique égale à [Leonhardt97] :

$$\mathcal{P}[u, v] = \text{Tr}\{\hat{\rho} \exp(-i\alpha \hat{a}^\dagger) \exp(-i\alpha^* \hat{a})\} \quad (2.78)$$

avec la convention  $N_0=1/2$  et  $\alpha=(u+iv)/\sqrt{2}$ . Cette fonction caractéristique est simplement reliée à celle de la fonction de Wigner par

$$\mathcal{P}[u, v] = \mathcal{W}[u, v] e^{+\frac{1}{4}(u^2+v^2)}. \quad (2.79)$$

En prenant la transformée de Fourier inverse, nous voyons que la fonction de Wigner est égale à la convolution de la fonction  $P$  avec une gaussienne.

La fonction  $P$  est particulièrement intéressante d'un point de vue théorique car elle permet de "diagonaliser" la matrice densité sur un ensemble d'états cohérents<sup>8 9</sup> :

$$\hat{\rho} = \int d^2\alpha P(\alpha) |\alpha\rangle\langle\alpha| \quad (2.80)$$

avec

$$P(\alpha) = \frac{e^{|\alpha|^2}}{\pi^2} \int d^2\beta e^{|\beta|^2} \langle -\beta | \hat{\rho} | \beta \rangle e^{\beta^* \alpha - \beta \alpha^*}. \quad (2.81)$$

8. Nous introduirons les états cohérents dans la section 2.5.3. Ces états sont le pendant quantique des états classiques des modes du champ électromagnétique d'amplitudes normales  $\alpha$ .

9. Cette définition est en fait légèrement différente de (2.78). Dans (2.81), la fonction  $P$  prend comme argument l'amplitude  $\alpha$  d'un état cohérent, alors que la définition (2.78) conduit à prendre comme arguments les quadratures  $x=2\sqrt{N_0}\Re(\alpha)$  et  $p=2\sqrt{N_0}\Im(\alpha)$ . Dans tous les calculs où nous utiliserons la fonction  $P$ , c'est la convention (2.81) qui sera utilisée, qui ne dépend pas du choix de  $N_0$ .

Cette formule est valable même si l'état est pur, pourtant elle n'utilise que des états  $\{|\alpha\rangle\langle\alpha|\}$  et fait penser à un mélange statistique. Cette décomposition contre intuitive s'explique en fait en se souvenant que d'une part les états cohérents ne sont pas orthogonaux entre eux, et en remarquant d'autre part que la fonction  $P$  peut être un objet mathématique plus délicat qu'une simple fonction. Par exemple, pour un état cohérent  $\hat{\rho}=|\gamma\rangle\langle\gamma|$ , qui est pourtant un état tout à fait physique,  $\langle-\beta|\hat{\rho}|\beta\rangle = e^{-|\gamma|^2-|\beta|^2} e^{\gamma^*\beta-\beta^*\gamma}$ , et donc :

$$P(\alpha) = e^{|\alpha|^2-|\gamma|^2} \frac{1}{\pi^2} \int d^2\beta e^{\beta^*(\alpha-\gamma)-\beta(\alpha^*-\gamma^*)} = \delta^2(\alpha-\gamma) \quad (2.82)$$

Pour des états non classiques, elle peut même ne pas exister en tant que distribution tempérée [Leonhardt97]. C'est d'ailleurs une raison pour laquelle elle est difficile à reconstruire expérimentalement, contrairement à la fonction de Wigner qui peut l'être efficacement par tomographie quantique [Leonhardt97].

La fonction  $P$  nous sera très utile lorsque nous étudierons l'utilisation d'un amplificateur sans bruit en cryptographie quantique, aux chapitres 9 et 10.

### Autres représentations

La fonction  $P$  n'est qu'une des autres représentations possibles, mais il en existe en fait une infinité, en généralisant la formule (2.79) en :

$$\mathcal{W}[u, v, s] = \mathcal{W}[u, v] e^{\frac{s}{4}(u^2+v^2)} \quad (2.83)$$

Une telle représentation  $s$  possède des propriétés très proches de celles de la fonction de Wigner. La fonction de Wigner correspond donc à  $s=0$ , et la fonction  $P$  à  $s=1$ . Nous renvoyons le lecteur à la référence [Leonhardt97] pour une présentation plus complète.

Enfin, une autre fonction couramment utilisée est la fonction  $Q$ , pour laquelle  $s=-1$ . Elle est définie par :

$$Q(q, p) = \frac{1}{2\pi} \text{Tr}\{\hat{\rho}|\alpha\rangle\langle\alpha|\} = \frac{1}{2\pi} \langle\alpha|\hat{\rho}|\alpha\rangle \quad (2.84)$$

avec  $\alpha=(q+ip)/\sqrt{2}$  et la convention  $N_0=1/2$ . Ses détails sont fortement estompés par rapport à la fonction de Wigner, car le terme exponentiel dans sa transformée de Fourier agit comme un filtre coupant les hautes fréquences.

En fonction des situations, on pourra privilégier une des trois représentations présentées dans cette section. Pour certains problèmes théoriques, la fonction  $P$  ou la fonction  $Q$  peuvent être plus adaptées. En revanche d'un point de vue expérimental, la fonction de Wigner apparaît comme le meilleur compromis entre un comportement proche d'une distribution de probabilité, et la capacité à pouvoir facilement distinguer les états quantiques.

### 2.3.3 Quelle description choisir ?

Nous venons de voir qu'un état quantique peut être décrit de plusieurs manières équivalentes. Y en a-t-il une à privilégier ? D'un point de vue pratique, la matrice densité possède un certain nombre d'inconvénients : il est tout d'abord assez difficile d'identifier les états quantiques "à l'œil", à part pour les états de Fock ou quelques états bien connus. Il faut ensuite pouvoir décrire correctement l'état du champ : si la dimension du sous-espace n'est pas suffisamment grande, des termes seront oubliés et les autres coefficients seront mal normalisés. Ce problème de dimensions

devient critique lorsque l'on considère plusieurs modes du champ : pour  $M$  modes représentés dans un espace de dimension  $N$ , il faut une matrice densité de dimension  $(NM)^2$ . Lorsque l'on connaît une purification de l'état, une astuce consiste à garder les opérations de trace partielle pour la fin du calcul et à travailler avec un vecteur d'état correspondant à la purification, plutôt qu'avec une matrice densité. Ceci permet de diminuer la dimension des objets à manipuler numériquement.

En revanche, le gros avantage d'une représentation en matrice densité est de pouvoir simuler directement diverses opérations quantiques appliquées sur un état, avec un formalisme matriciel. Il existe un package pour Matlab, appelé *Quantum Optics Toolbox*, permettant d'utiliser un tel formalisme très efficacement, en définissant par exemple des opérateurs de création et de destruction, ou des opérateurs associés aux différents outils que nous présenterons dans ce chapitre : lame séparatrice, squeezing, déplacement... C'est donc tout le dispositif expérimental qui peut être simulé simplement avec ces outils, qu'il convient néanmoins de manier avec précaution pour s'assurer de la pertinence physique des résultats. Nous utiliserons souvent cette description pour les diverses simulations présentées dans ce manuscrit.

La fonction de Wigner permet de résoudre certains des problèmes posés par le formalisme de la matrice densité, au prix toutefois de calculs bien plus lourds. Elle est à privilégier pour obtenir des résultats analytiques, car dans ce cas un travail en base de Fock conduit très vite à une quantité de séries difficiles à manipuler et à interpréter.

Nous combinerons souvent les deux approches pour les simulations numériques de ce manuscrit, en faisant des calculs à l'aide de la matrice densité, puis en représentant l'état final à l'aide de sa fonction de Wigner.

### 2.3.4 Distance entre états quantiques : la fidélité

Lorsqu'un des états est pur, par exemple  $\hat{\sigma} = |\phi\rangle\langle\phi|$ , la fidélité s'écrit simplement :

$$\mathcal{F}(\hat{\rho}, \hat{\sigma}) = \langle\phi|\hat{\rho}|\phi\rangle = \text{Tr}\{\hat{\rho}\hat{\sigma}\} \quad (2.85)$$

Lorsque  $\hat{\rho}$  et  $\hat{\sigma}$  sont tous les deux des mélanges statistiques, l'expression précédente n'est plus valable car elle ne permet pas de satisfaire tous les axiomes requis pour une bonne mesure de fidélité [Jozsa94]. En particulier, si  $\hat{\rho} = \mathbb{1}/2$ , on a  $\text{Tr}\{\hat{\rho}\hat{\rho}\} = 1/2$  alors que l'on souhaite que la fidélité soit égale à 1 pour deux états identiques.

On peut néanmoins trouver une formule valable pour des mélanges statistiques, et qui satisfait tous les axiomes [Jozsa94, Nielsen00] :

$$\mathcal{F}(\hat{\rho}, \hat{\sigma}) = \left( \text{Tr} \left\{ \sqrt{\sqrt{\hat{\rho}} \hat{\sigma} \sqrt{\hat{\rho}}} \right\} \right)^2 \quad (2.86)$$

Cette définition de la fidélité redonne bien sûr la formule (2.85) lorsqu'un des deux états est pur. Le problème qui se pose maintenant est que le calcul analytique des racines carrées de matrices n'est pas facile (voire impossible?), et qu'il faut souvent avoir recourt à des résultats numériques. Cette mesure de la fidélité pour des mélanges statistiques est égale à la fidélité maximale que l'on peut obtenir avec des purifications des états [Nielsen00] :

$$\mathcal{F}(\hat{\rho}, \hat{\sigma}) = \max_{|\psi_1\rangle, |\psi_2\rangle} |\langle\psi_1|\psi_2\rangle|^2 \quad (2.87)$$

où  $|\psi_1\rangle$  et  $|\psi_2\rangle$  sont des purifications respectives de  $\hat{\rho}_1$  et  $\hat{\rho}_2$ .

### 2.3.5 Mesures et POVM

En physique quantique, les mesures sont généralement associées à des observables, et décrites par un ensemble de projecteurs orthogonaux. Il existe cependant des mesures plus générales, pour lesquelles la condition d'orthogonalité peut être relâchée. Dans le cas général, une mesure peut être décrite par un ensemble d'opérateurs  $\{\hat{M}_n\}$ , vérifiant simplement

$$\sum_n \hat{M}_n^\dagger \hat{M}_n = \mathbb{I}. \quad (2.88)$$

Cette condition assure que la somme des probabilités des différentes mesures soit égale à 1. Lorsque le résultat d'une mesure sur  $|\psi\rangle$  est  $m$ , l'état est projeté sur  $\frac{1}{\sqrt{p_m}} \hat{M}_m |\psi\rangle$ , où la quantité

$$p_m = \langle \psi | \hat{M}_m^\dagger \hat{M}_m | \psi \rangle \quad (2.89)$$

est égale à la probabilité de succès d'obtenir  $m$ . Lorsque seules les probabilités de succès interviennent dans un problème, ou lorsque la mesure est effectuée sur un mode auxiliaire qui est ensuite tracé, seuls les éléments  $\{\hat{M}_n^\dagger \hat{M}_n := \hat{E}_m\}$  interviennent : ils définissent alors un POVM (*Positive Operator-Valued Measure*). Le lecteur pourra consulter les références [Paris12], [Nielsen00], et [Preskill98] pour davantage de détails sur ces mesures généralisées.

## 2.4 Quelques transformations unitaires

### 2.4.1 Le déphasage

Un déphasage  $|n\rangle \rightarrow e^{-i\phi} |n\rangle$  est obtenu à l'aide d'un opérateur

$$\boxed{\hat{U}(\theta) = \exp[-i\theta \hat{a}^\dagger \hat{a}]} \quad (2.90)$$

qui est équivalent à une évolution libre pendant un temps  $t = \hbar\theta$ . Cet opérateur transforme l'opérateur destruction en

$$e^{i\phi \hat{a}^\dagger \hat{a}} \hat{a} e^{-i\phi \hat{a}^\dagger \hat{a}} = \hat{a} e^{-i\phi}, \quad (2.91)$$

et les quadratures en

$$\hat{X}' = \hat{X} \cos \theta + \hat{P} \sin \theta, \quad (2.92a)$$

$$\hat{P}' = -\hat{X} \sin \theta + \hat{P} \cos \theta. \quad (2.92b)$$

**Remarque** Une rotation d'un angle  $\theta > 0$ , fait “tourner” l'état quantique dans le sens horaire. Une manière de s'en convaincre est de regarder l'évolution d'un état cohérent :  $\hat{U}(\theta)|\alpha\rangle = |\alpha e^{-i\theta}\rangle$ . En revanche, pour une même rotation, les quadratures tournent dans le sens trigonométrique, comme le montre (2.92).

L'explication est simple : afin de donner les mêmes valeurs moyennes, l'angle (algébrique) doit être le même entre l'état et les quadratures après la rotation (Fig. 2.1).



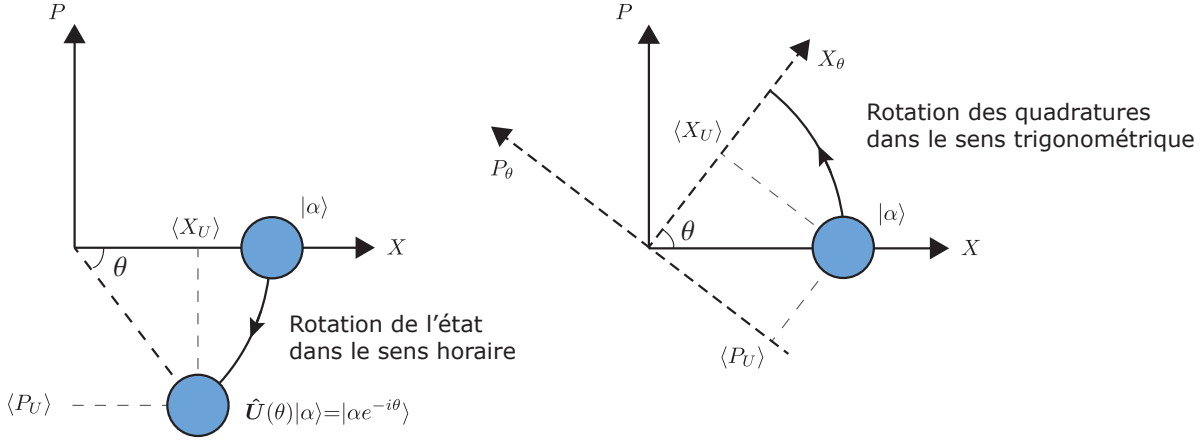


FIGURE 2.1 – Action d’une rotation sur un état cohérent (gauche), et sur les opérateurs de quadratures (droite).

### 2.4.2 La lame séparatrice

Une lame séparatrice est décrite par l’opérateur  $\hat{U}_{BS}(\theta)$  :

$$\hat{U}_{BS}(\theta) = \exp \left[ \theta (\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger) \right] \quad (2.93)$$

où  $\cos^2 \theta = T$ , et où les opérateurs  $\hat{a}$  et  $\hat{b}$  correspondent aux deux modes de la lame séparatrice. Sous l’action de  $\hat{U}_{BS}$ , les opérateurs  $\hat{a}(\theta) = \hat{U}_{BS}^\dagger(\theta) \hat{a} \hat{U}_{BS}(\theta)$  et  $\hat{b}(\theta) = \hat{U}_{BS}^\dagger(\theta) \hat{b} \hat{U}_{BS}(\theta)$  se transforment selon

$$\frac{d}{d\theta} \hat{a}(\theta) = \hat{U}_{BS}^\dagger(\theta) [\hat{a}, \hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger] \hat{U}_{BS}(\theta) = \hat{b}(\theta), \quad (2.94a)$$

$$\frac{d}{d\theta} \hat{b}(\theta) = \hat{U}_{BS}^\dagger(\theta) [\hat{b}, \hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger] \hat{U}_{BS}(\theta) = -\hat{a}(\theta), \quad (2.94b)$$

dont les solutions sont :

$$\begin{aligned} \hat{a}(\theta) &= \hat{a} \cos \theta + \hat{b} \sin \theta \\ \hat{b}(\theta) &= -\hat{a} \sin \theta + \hat{b} \cos \theta \end{aligned} \quad (2.95)$$

Les quadratures évoluent selon :

$$\hat{U}_{BS}^\dagger(\theta) \hat{Q}_a \hat{U}_{BS}(\theta) = \hat{Q}_a \cos \theta + \hat{Q}_b \sin \theta \quad (2.96a)$$

$$\hat{U}_{BS}^\dagger(\theta) \hat{Q}_b \hat{U}_{BS}(\theta) = -\hat{Q}_a \sin \theta + \hat{Q}_b \cos \theta \quad (2.96b)$$

avec  $\hat{Q} = \hat{X}$  ou  $\hat{P}$ .

**Transformation de la fonction de Wigner** En utilisant (2.70) et (2.96), on voit que la fonction de Wigner  $W'$  de deux modes couplés par une lame séparatrice vaut

$$W'(x_a, p_a, x_b, p_b) = W(tx_a - rx_b, tp_a - rp_b, tx_b + rx_a, tp_b + rp_a), \quad (2.97)$$

avec  $t = \sqrt{T}$  et  $r = \sqrt{1-T}$ .

## 2.5 Les états gaussiens

Les états gaussiens sont des états naturellement décrits par des variables continues, et pour lesquels la fonction de Wigner est gaussienne. Ils sont entièrement caractérisés par un vecteur regroupant leurs valeurs moyennes, et par leurs matrices de covariances.

### Matrice de covariance

Les quadratures d'un état à  $N$  modes peuvent être regroupées dans un vecteur

$$\hat{\mathbf{x}} := \left( \hat{\mathbf{X}}_1, \hat{\mathbf{P}}_1, \dots, \hat{\mathbf{X}}_N, \hat{\mathbf{P}}_N \right)^T, \quad (2.98)$$

dont les  $2N$  éléments vérifient

$$[\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_j] = 2iN_0\Omega_{ij}, \quad (2.99)$$

avec  $\Omega_{ij} = \bigoplus_{k=1}^N \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$ . En pratique, nous nous intéresserons principalement aux quadratures d'états à un ou deux modes.

Les valeurs moyennes sont regroupées dans le vecteur

$$\mathbf{d} := \langle \hat{\mathbf{x}} \rangle = \text{Tr}\{\hat{\rho} \hat{\mathbf{x}}\}, \quad (2.100)$$

qui sera nul pour la plupart de nos états. La matrice de covariance  $\mathbf{\Gamma}$  est définie par<sup>10</sup> :

$$\mathbf{\Gamma}_{ij} = \frac{1}{2N_0} \langle \{\hat{\mathbf{x}}_i - \langle \hat{\mathbf{x}}_i \rangle, \hat{\mathbf{x}}_j - \langle \hat{\mathbf{x}}_j \rangle\} \rangle \quad (2.101)$$

$$= \frac{1}{2N_0} \langle \{\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_j\} \rangle - \langle \hat{\mathbf{x}}_i \rangle \langle \hat{\mathbf{x}}_j \rangle \quad (2.102)$$

avec  $\{\hat{\mathbf{C}}, \hat{\mathbf{D}}\} = \hat{\mathbf{C}}\hat{\mathbf{D}} + \hat{\mathbf{D}}\hat{\mathbf{C}}$ . Les relations de Heisenberg entre les différents modes se traduisent alors par [Simon94]

$$\mathbf{\Gamma} + i\Omega \geq 0. \quad (2.103)$$

Un état gaussien quelconque à  $N$  modes a une fonction de Wigner

$$W(x_1, p_1, \dots, x_N, p_N) = \frac{1}{(2\pi N_0)^N \sqrt{\det \mathbf{\Gamma}}} e^{-\frac{1}{2N_0} (\mathbf{x} - \mathbf{d})^T \mathbf{\Gamma}^{-1} (\mathbf{x} - \mathbf{d})}. \quad (2.104)$$

De nombreuses autres propriétés des états gaussiens pourront être trouvées dans les références [Ferraro05, Weedbrook12]. La première est particulièrement complète, alors que la deuxième est relativement plus facile d'accès.

10. On peut aussi trouver une convention différente, avec  $\mathbf{\Gamma}_{ij} = \frac{1}{2} \langle \{\hat{\mathbf{x}}_i - \langle \hat{\mathbf{x}}_i \rangle, \hat{\mathbf{x}}_j - \langle \hat{\mathbf{x}}_j \rangle\} \rangle$ , ce qui conduit à  $\mathbf{\Gamma} + iN_0\Omega \geq 0$ .

### 2.5.1 Etats d'incertitude minimale

#### Relation d'incertitude pour des mélanges statistiques

Revenons sur la relation d'incertitude (2.31) : quels sont les états qui saturent cette inégalité ? Nous allons voir que de tels états doivent être purs et gaussiens. La plupart des ouvrages dérivent cette relation d'incertitude en raisonnant sur un état pur [Leonhardt97, Cohen-Tannoudji97c]. Nous allons plutôt suivre une autre approche, tirée de [Stoler72], permettant de traiter le cas plus général d'un mélange statistique.

Considérons pour commencer un état  $\hat{\rho}$  quelconque, que l'on pourra écrire sous forme diagonale  $\hat{\rho} = \sum_m p_m |\phi_m\rangle\langle\phi_m|$ . On forme les opérateurs

$$\hat{\mathcal{X}} = \hat{X} - \langle\hat{X}\rangle, \quad \text{et} \quad \hat{\mathcal{P}} = \hat{P} - \langle\hat{P}\rangle, \quad (2.105)$$

avec comme d'habitude  $\langle\hat{X}\rangle = \text{Tr}\{\hat{\rho}\hat{X}\}$  et  $\langle\hat{P}\rangle = \text{Tr}\{\hat{\rho}\hat{P}\}$ . Les variances sont obtenues naturellement grâce à ces opérateurs :

$$\Delta^2 X = \text{Tr}\{\hat{\rho}\hat{\mathcal{X}}^2\} \quad \text{et} \quad \Delta^2 P = \text{Tr}\{\hat{\rho}\hat{\mathcal{P}}^2\} \quad (2.106)$$

Définissons ensuite l'opérateur  $\hat{T}_\delta = \hat{\mathcal{X}} - i\delta\hat{\mathcal{P}}$ , avec  $\delta \in \mathbb{R}$ . La positivité de  $\hat{\rho}$  assure que :

$$\text{Tr}\{\hat{\rho}\hat{T}_\delta\hat{T}_\delta^\dagger\} = \text{Tr}\{\hat{T}_\delta^\dagger\hat{\rho}\hat{T}_\delta\} = \sum_k \langle k|\hat{T}_\delta^\dagger\hat{\rho}\hat{T}_\delta|k\rangle = \sum_k \langle\tilde{k}|\hat{\rho}|\tilde{k}\rangle \geq 0 \quad (2.107)$$

avec  $|\tilde{k}\rangle = \hat{T}_\delta|k\rangle$ , puisque chaque terme  $\langle\tilde{k}|\hat{\rho}|\tilde{k}\rangle$  est positif. On a donc :

$$\text{Tr}\{\hat{\rho}\hat{T}_\delta\hat{T}_\delta^\dagger\} = \text{Tr}\{\hat{\rho}(\hat{\mathcal{X}}^2 + \delta^2\hat{\mathcal{P}}^2 - 2N_0\delta)\} \quad (2.108a)$$

$$= \Delta X^2 + \delta^2 \Delta P^2 - 2N_0\delta \quad (2.108b)$$

$$\geq 0 \quad (2.108c)$$

C'est une équation du second degré en  $\delta$  en tout point similaire à celle que l'on peut habituellement obtenir en traitant un état pur et en utilisant la positivité de la norme. Pour pouvoir satisfaire l'inégalité, le discriminant  $4N_0^2 - 4\Delta^2 X \Delta^2 P$  doit être positif ou nul, ce qui nous donne la relation d'incertitude (2.31) :

$$\Delta X \Delta P \geq N_0 \quad (2.109)$$

#### Etats d'incertitude minimale

L'égalité  $\Delta X \Delta P = N_0$  n'est possible que si le discriminant est nul. La solution  $\delta = s$  correspondante qui annule (2.108) est alors  $s = \frac{2N_0}{2\Delta^2 P} = N_0 / \Delta^2 P$ . On a alors

$$\text{Tr}\{\hat{\rho}\hat{T}_s\hat{T}_s^\dagger\} = \sum_m p_m \|\hat{T}_s^\dagger|\phi_m\rangle\|^2 = 0, \quad (2.110)$$

ce qui, vu que les  $p_m$  sont par définition strictement positifs, implique que l'on a pour tout  $m$  :

$$\hat{T}_m^\dagger|\phi_m\rangle = 0 \quad (2.111)$$

Cette dernière équation forme une équation différentielle dont la solution est unique [Leonhardt97, Cohen-Tannoudji97c] : il s'agit d'un état gaussien dont la fonction de Wigner est

$$W(x, p) = \frac{1}{2\pi N_0} e^{-\frac{x-\langle X \rangle}{2\Delta X^2}} e^{-\frac{p-\langle P \rangle}{2\Delta P^2}}, \quad (2.112)$$

avec  $\Delta^2 X = s N_0$  et  $\Delta^2 P = \frac{1}{s} N_0$ .

Les  $|\phi_m\rangle$  sont donc tous identiques, et un état saturant la relation d'incertitude est donc nécessairement un état pur. En revanche, les deux quadratures ne sont pas obligées d'avoir la même variance pour saturer l'inégalité de Heisenberg. On parlera d'*états comprimés* lorsque les fluctuations d'une quadrature sont réduites d'un facteur  $s$ , alors que celles de l'autre quadrature sont augmentées d'un facteur  $1/s$ .

### 2.5.2 Le vide quantique

Le vide quantique  $|0\rangle$  est l'état fondamental du champ, lorsqu'il n'y a aucune excitation ( $n=0$ ). Comme nous le verrons, c'est le seul état de Fock qui soit gaussien. Il joue un rôle fondamental en optique quantique, car il est justement présent chaque fois qu'un mode d'une transformation est "vide".

La valeur moyenne des quadratures est nulle, car  $\langle 0|\hat{\mathbf{a}}|0\rangle = \langle 0|\hat{\mathbf{a}}^\dagger|0\rangle = 0$ . En revanche, ce n'est pas le cas de leur variance, puisque

$$\langle 0|\hat{\mathbf{X}}^2|0\rangle = N_0 \langle 0|(\hat{\mathbf{a}} + \hat{\mathbf{a}}^\dagger)^2|0\rangle = N_0 \quad (2.113)$$

On montre de même que  $\langle 0|\hat{\mathbf{P}}^2|0\rangle = N_0$ . Ainsi la variance des quadratures du vide est égale à  $N_0$ . Elle est connue sous de nombreuses appellations : "shot noise", "bruit de photon", "bruit quantique standard" ou encore "bruit de grenaille". En conclusion, la matrice de covariance et le vecteur déplacement sont donnés par :

$$\mathbf{d} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \mathbf{\Gamma} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (2.114)$$

La fonction de Wigner du vide est (Fig. 2.2) :

$$W_{|0\rangle}(x, p) = \frac{1}{2\pi N_0} e^{-\frac{1}{2N_0}(x^2 + p^2)} \quad (2.115)$$

### 2.5.3 Les états cohérents

#### Définition

La nature étant intrinsèquement quantique, la description classique du champ (2.4) doit pouvoir être retrouvée à partir des valeurs moyennes de l'opérateur champ (2.12). Les *états cohérents* [Loudon00, Leonhardt97], ou états quasi-classiques, sont les états  $|\alpha\rangle$  pour lesquels, en valeur moyenne et pour un mode normal donné, l'amplitude et l'énergie du champ quantique sont égales à celles du champ classique. Ces conditions se traduisent par :

$$\langle \alpha | \left( \hat{\mathbf{a}} e^{i\mathbf{k}\cdot\mathbf{r} - i\omega t} - \hat{\mathbf{a}}^\dagger e^{-i\mathbf{k}\cdot\mathbf{r} + i\omega t} \right) | \alpha \rangle = \alpha e^{i\mathbf{k}\cdot\mathbf{r} - i\omega t} - \alpha^* e^{-i\mathbf{k}\cdot\mathbf{r} + i\omega t} \quad (2.116)$$

$$\hbar\omega \langle \alpha | \hat{\mathbf{n}} | \alpha \rangle = \hbar\omega |\alpha|^2 \quad (2.117)$$

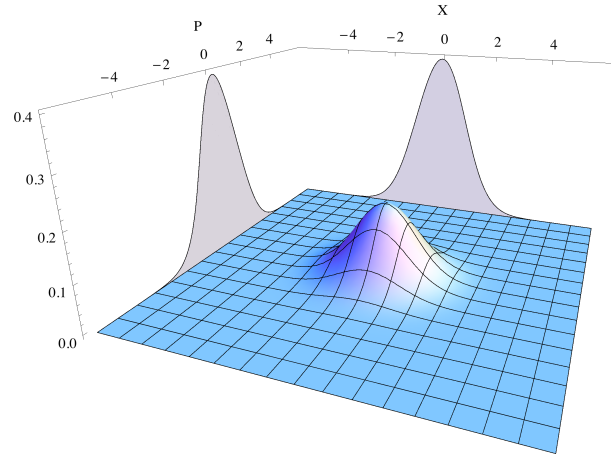


FIGURE 2.2 – *Fonction de Wigner du vide, et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .*

Elle sont suffisantes pour définir les états  $|\alpha\rangle$ , qui doivent être états propres de l'opérateur  $\hat{a}$  [Cohen-Tannoudji97b] :

$$\hat{a}|\alpha\rangle = \alpha|\alpha\rangle, \quad \alpha \in \mathbb{C} \quad (2.118)$$

Rappelons nous que l'opérateur  $\hat{a}$  a été introduit à la place des amplitudes normales classiques. On peut donc l'interpréter comme un opérateur donnant l'amplitude complexe du champ, en gardant à l'esprit qu'il n'est pas une observable car il n'est pas hermitien, et que ses valeurs propres  $\alpha$  peuvent être complexes.

Afin de démontrer la relation (2.118), calculons la norme de l'opérateur  $\hat{b} = \hat{a} - \alpha$  appliqué sur  $\alpha$  [Cohen-Tannoudji97b] :

$$\langle \alpha | \hat{b}^\dagger \hat{b} | \alpha \rangle = \langle \alpha | \hat{a}^\dagger \hat{a} | \alpha \rangle - \alpha \langle \alpha | \hat{a}^\dagger | \alpha \rangle - \alpha^* \langle \alpha | \hat{a} | \alpha \rangle + |\alpha|^2 \quad (2.119)$$

Les conditions (2.116) et (2.117) impliquent que

$$\langle \alpha | \hat{b}^\dagger \hat{b} | \alpha \rangle = 0, \quad (2.120)$$

et donc que  $\hat{a}|\alpha\rangle = \alpha|\alpha\rangle$ .

Les états cohérents sont les états qui se rapprochent le plus de l'image classique d'une onde plane monochromatique. Ils peuvent être produits par un laser très au dessus du seuil [Chiao08, Loudon00].

## Opérateur déplacement

**Définition** L'opérateur déplacement  $\hat{D}(\alpha)$  est défini par

$$\hat{D}(\alpha) = \exp \left[ \alpha \hat{a}^\dagger - \alpha^* \hat{a} \right] \quad (2.121a)$$

$$= \exp \left[ \frac{i}{\sqrt{N_0}} [\Im(\alpha) \hat{X} - \Re(\alpha) \hat{P}] \right] \quad (2.121b)$$

On voit immédiatement que  $\hat{D}^\dagger(\alpha) = \hat{D}(-\alpha)$ , et que  $\hat{D}^\dagger(\alpha)\hat{D}(\alpha) = \mathbb{I}$ . Cet opérateur se décompose grâce à la formule de Baker-Hausdorff (annexe A) :

$$\hat{D}(\alpha) = e^{-\frac{1}{2}|\alpha|^2} e^{\alpha\hat{a}^\dagger} e^{-\alpha^*\hat{a}} \quad (2.122a)$$

$$= e^{+\frac{1}{2}|\alpha|^2} e^{-\alpha^*\hat{a}} e^{\alpha\hat{a}^\dagger} \quad (2.122b)$$

**Transformation des opérateurs** Sous l'action de l'opérateur déplacement, les opérateurs destruction et création sont transformés en (annexe A.2)

$$\hat{D}^\dagger(\alpha) \hat{a} \hat{D}(\alpha) = \hat{a} + \alpha, \quad (2.123a)$$

$$\hat{D}^\dagger(\alpha) \hat{a}^\dagger \hat{D}(\alpha) = \hat{a}^\dagger + \alpha^*, \quad (2.123b)$$

ce qui implique que les quadratures sont transformées en

$$\hat{D}^\dagger(\alpha) \hat{X} \hat{D}(\alpha) = \hat{X} + 2\sqrt{N_0}\Re(\alpha), \quad (2.124a)$$

$$\hat{D}^\dagger(\alpha) \hat{P} \hat{D}(\alpha) = \hat{P} + 2\sqrt{N_0}\Im(\alpha). \quad (2.124b)$$

Elles sont donc simplement déplacées d'une valeur constante par rapport aux quadratures du vide. La variance d'une variable aléatoire ne changeant pas avec l'ajout d'une constante, les quadratures déplacées ont toujours une variance égale à  $N_0$ . La définition (2.118) implique également que les quadratures  $\hat{X}$  et  $\hat{P}$  ont des valeurs moyennes respectivement égales à  $2\sqrt{N_0}\Re(\alpha)$  et  $2\sqrt{N_0}\Im(\alpha)$ , et une variance égale à  $N_0$ . On peut alors en conclure qu'un état cohérent est un *vide déplacé* :

$$|\alpha\rangle = \hat{D}(\alpha)|0\rangle = \exp\left[\alpha\hat{a}^\dagger - \alpha^*\hat{a}\right]|0\rangle \quad (2.125)$$

De plus,  $\hat{a} \left( \hat{D}(\alpha)|0\rangle \right) = \hat{D}(\alpha)\hat{D}^\dagger(\alpha)\hat{a}\hat{D}(\alpha)|0\rangle = \hat{D}(\alpha)[\hat{a} + \alpha]|0\rangle = \alpha \left( \hat{D}(\alpha)|0\rangle \right)$ . L'état  $\hat{D}(\alpha)|0\rangle$  vérifie donc bien la définition (2.118).

### Moments et fonction de Wigner

Les conclusions précédentes nous montrent que la matrice de covariance d'un état cohérent est la même que celle du vide :

$$\mathbf{d} = 2\sqrt{N_0} \begin{pmatrix} \Re(\alpha) \\ \Im(\alpha) \end{pmatrix} \quad \mathbf{\Gamma} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (2.126)$$

Sa fonction de Wigner est donc identique à celle du vide centrée en  $2\sqrt{N_0}\alpha$ . Remarquons que le vide est d'ailleurs un état cohérent d'amplitude  $\alpha=0$ .

La phase de  $\alpha$  correspond simplement à une rotation dans l'espace des phases. Ainsi, lorsqu'un seul mode entre en jeu, on peut la plupart du temps supposer  $\alpha$  réel sans perte de généralité. En revanche lorsque plusieurs modes sont concernés, la phase relative est importante et doit être conservée.

La fonction de Wigner d'un état cohérent est (Fig. 2.3) :

$$W_\alpha(x, p) = \frac{1}{2\pi N_0} e^{-\frac{1}{2N_0}((x-2\sqrt{N_0}\alpha_x)^2 + (p-2\sqrt{N_0}\alpha_p)^2)} \quad (2.127)$$

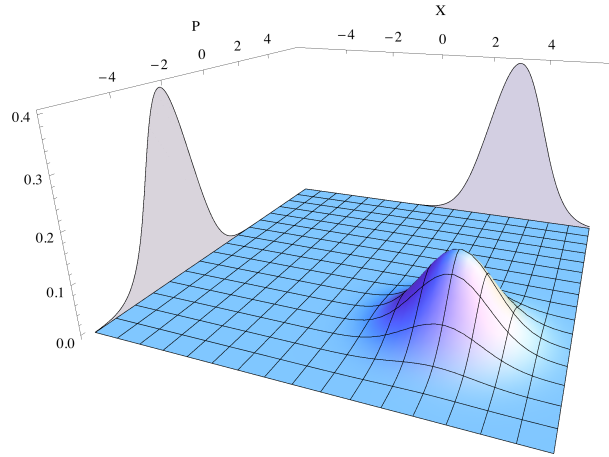


FIGURE 2.3 – Fonction de Wigner d'un état cohérent  $\alpha=1.5-1.2i$ , et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .

### Décomposition en base de Fock

Les coefficients de la décomposition d'un état cohérent en base de Fock sont la plupart du temps trouvés par récurrence, en utilisant les propriétés des opérateurs destruction et création. On peut également les trouver facilement en utilisant (2.122a) et le fait que  $\exp[-\alpha\hat{a}]\lvert 0\rangle=\lvert 0\rangle$  :

$$\lvert \alpha \rangle = e^{-\frac{1}{2}|\alpha|^2} e^{\alpha\hat{a}^\dagger} \lvert 0 \rangle \quad (2.128a)$$

$$= e^{-\frac{1}{2}|\alpha|^2} \sum_n \frac{(\alpha\hat{a}^\dagger)^n}{n!} \lvert 0 \rangle \quad (2.128b)$$

$$= e^{-\frac{1}{2}|\alpha|^2} \sum_n \frac{\alpha^n}{\sqrt{n!}} \lvert n \rangle \quad (2.128c)$$

### Propriétés élémentaires

Les états cohérents ne sont pas orthogonaux entre eux, car leur produit scalaire

$$\langle \alpha | \beta \rangle = e^{-\frac{1}{2}|\alpha|^2} e^{-\frac{1}{2}|\beta|^2} e^{\alpha^*\beta} = e^{-\frac{1}{2}|\alpha-\beta|^2 + \frac{1}{2}(\alpha^*\beta - \alpha\beta^*)} \quad (2.129)$$

n'est jamais strictement nul. En revanche, lorsque  $|\alpha-\beta|^2$  est grand, on pourra faire l'approximation qu'ils le sont. Ils forment également une résolution de l'unité

$$\frac{1}{\pi} \int d^2\alpha \lvert \alpha \rangle \langle \alpha \rvert = \mathbb{I}, \quad (2.130)$$

et constituent pour cette raison une base sur-complète. Cette formule peut être utilisée pour calculer la trace d'un opérateur [Chiao08] :

$$\text{Tr}\{\hat{A}\} = \frac{1}{\pi} \int d^2\alpha \langle \alpha | \hat{A} | \alpha \rangle \quad (2.131)$$

Ces propriétés peuvent être généralisées pour des états cohérents multimodes [Chiao08, Blow90, Loudon00]. Enfin, comme indiqué dans la section 2.3.2 et avec l'équation (2.80), les états cohérents peuvent être utilisés pour décomposer un état quantique  $\hat{\rho}$  à l'aide de la fonction  $P$  :

$$\hat{\rho} = \int d^2\alpha P(\alpha) \lvert \alpha \rangle \langle \alpha \rvert \quad (2.132)$$

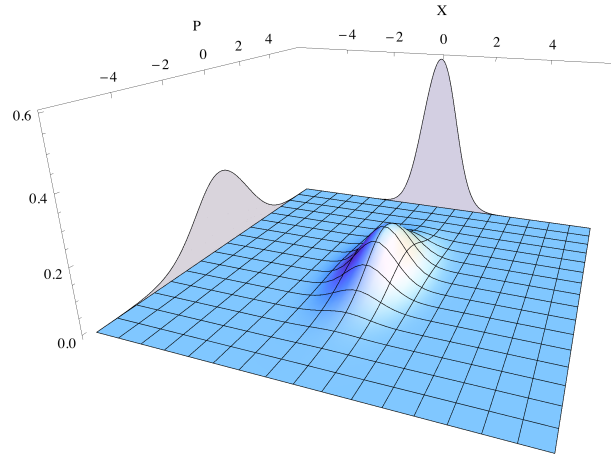


FIGURE 2.4 – Fonction de Wigner d'un état comprimé selon  $\hat{X}$ , de paramètre  $s=0.45$ , et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .

## 2.5.4 Le vide comprimé monomode

### Définition

Le vide comprimé monomode  $|\psi_{\text{sqs}}\rangle$  est un état de moyenne nulle et d'incertitude minimale, dont une des quadratures est comprimée d'un facteur  $s$ , alors que l'autre est amplifiée d'un facteur  $1/s$ . Si la compression est faite selon  $\hat{X}$ , on a :

$$\Delta^2 X = s N_0 \quad (2.133a)$$

$$\Delta^2 P = \frac{1}{s} N_0 \quad (2.133b)$$

Nous verrons que contrairement au vide,  $|\psi_{\text{sqs}}\rangle$  contient des photons. Il peut être obtenu par amplification paramétrique du vide, avec un amplificateur dégénéré, comme nous le détaillerons dans le chapitre 3.

### Moments et fonction de Wigner

La définition du vide comprimé (2.133) suffit à obtenir ses moments

$$\mathbf{d} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \mathbf{\Gamma} = \begin{pmatrix} s & 0 \\ 0 & 1/s \end{pmatrix} \quad (2.134)$$

et sa fonction de Wigner (Fig. 2.4) :

$$W_{\text{sqs}}(x, p) = \frac{1}{2\pi N_0} e^{-\frac{1}{2N_0} \left( \frac{x^2}{s} + \frac{p^2}{1/s} \right)} \quad (2.135)$$

Elle a bien l'allure de celle du vide que l'on aurait comprimé dans une direction et étendue dans l'autre.

### Opérateur de squeezing

**Définition** L'opérateur de squeezing monomode est défini par :

$$\hat{S}(r) = \exp \left[ \frac{r}{2} \hat{a}^2 - \frac{r}{2} (\hat{a}^\dagger)^2 \right] \quad (2.136)$$



C'est un opérateur qui vérifie  $\hat{S}^\dagger(r) = \hat{S}(-r)$  et  $\hat{S}^\dagger(r)\hat{S}(r) = \mathbb{I}$ . On montre en annexe A.3 que  $\hat{S}(r)$  peut s'écrire comme :

$$\hat{S}(r) = \exp \left[ -\frac{1}{2}(\hat{a}^\dagger)^2 \tanh r \right] \exp \left[ -\frac{1}{2}(2\hat{a}^\dagger \hat{a} + 1) \ln(\cosh r) \right] \exp \left[ +\frac{1}{2}\hat{a}^2 \tanh r \right] \quad (2.137)$$

**Transformation des opérateurs** L'opérateur  $\hat{a}(r) = \hat{S}^\dagger(r)\hat{a}\hat{S}(r)$  évolue selon l'équation

$$\frac{d}{dr}\hat{a}(r) = \frac{1}{2}\hat{S}(r)^\dagger \left[ \hat{a}, \hat{a}^2 - \hat{a}^{\dagger 2} \right] \hat{S}(r) = -\hat{a}^\dagger, \quad (2.138)$$

dont la solution est

$$\boxed{\hat{a}(r) = \hat{a} \cosh r - \hat{a}^\dagger \sinh r.} \quad (2.139)$$

Les quadratures sont donc transformées en :

$$\hat{X}(r) = e^{-r} \hat{X} \quad (2.140)$$

$$\hat{P}(r) = e^{+r} \hat{P} \quad (2.141)$$

Ces transformations correspondent à une amplification dépendante de la phase, ce qui conduit bien à la définition (2.133), en posant  $s = e^{-2r}$ . Le vide comprimé est donc obtenu en appliquant  $\hat{S}(r)$  sur le vide :

$$\boxed{|\psi_{\text{sqs}}\rangle = \hat{S}(r)|0\rangle = \exp \left[ \frac{r}{2}\hat{a}^2 - \frac{r}{2}(\hat{a}^\dagger)^2 \right] |0\rangle} \quad (2.142)$$

### Décomposition en base de Fock

Elle est facilement obtenue en utilisant la décomposition (2.137), comme montré en annexe A.3 :

$$\boxed{|\psi_{\text{sqs}}\rangle = \frac{1}{\sqrt{\cosh r}} \sum_{n=0}^{\infty} \frac{1}{n!} \sqrt{(2n)!} \left( -\frac{1}{2} \tanh r \right)^n |2n\rangle} \quad (2.143)$$

On remarque que le vide comprimé ne contient que des nombres pairs de photons.

**Remarque :** Le vide comprimé n'est d'incertitude minimale que pour les quadratures selon lesquelles il est comprimé. Si la compression est selon une quadrature  $\hat{X}_\theta$ , telle que

$$\hat{X}_\theta(r) = e^{-r} \hat{X}_\theta, \quad (2.144a)$$

$$\hat{P}_\theta(r) = e^{+r} \hat{P}_\theta, \quad (2.144b)$$

on peut obtenir les variances  $\hat{X}$  et  $\hat{P}$  en utilisant les expressions (2.32) :

$$\Delta X^2 = \cosh 2r - \sinh 2r \cos 2\theta \quad (2.145a)$$

$$\Delta P^2 = \cosh 2r + \sinh 2r \cos 2\theta \quad (2.145b)$$

### 2.5.5 Le vide comprimé bimode, ou état EPR

#### Définition

Le vide comprimé bimode  $|\psi_{\text{EPR}}\rangle$  est un état de moyenne nulle, dont les quadratures sont *intriquées* entre elles. Deux modes optiques  $\hat{\mathbf{a}}$  et  $\hat{\mathbf{b}}$  entrent en jeux, avec des opérateurs de quadratures  $\hat{\mathbf{X}}_1$  et  $\hat{\mathbf{P}}_1$  pour le premier mode, et  $\hat{\mathbf{X}}_2$  et  $\hat{\mathbf{P}}_2$  pour le second. Regardé indépendamment, chaque mode ne comporte pas de compression. En revanche, les fluctuations de  $\hat{\mathbf{X}}_1 + \hat{\mathbf{X}}_2$  et  $\hat{\mathbf{P}}_1 - \hat{\mathbf{P}}_2$  sont comprimées, alors que celles de  $\hat{\mathbf{X}}_1 - \hat{\mathbf{X}}_2$  et  $\hat{\mathbf{P}}_1 + \hat{\mathbf{P}}_2$  sont amplifiées :

$$\Delta^2 X_+ = sN_0 \quad \Delta^2 P_- = sN_0 \quad (2.146a)$$

$$\Delta^2 X_- = \frac{1}{s}N_0 \quad \Delta^2 P_+ = \frac{1}{s}N_0 \quad (2.146b)$$

où  $\Delta^2 Q_{\pm}$  est la variance de la quadrature  $\hat{Q}_{\pm} = \frac{1}{\sqrt{2}}(\hat{Q}_1 \pm \hat{Q}_2)$ , avec  $\hat{Q} = \hat{\mathbf{X}}$  ou  $\hat{\mathbf{P}}$ . Le vide comprimé bimode est une réalisation physique de l'état utilisé par Einstein, Podolsky et Rosen pour leur fameux paradoxe [Einstein35], qui est obtenu avec  $s \rightarrow 0$ . Pour cette raison et par abus de langage, il est appelé état EPR, même lorsque  $s > 0$ .

Il peut être obtenu par amplification paramétrique du vide, avec un amplificateur non dégénéré [Ou92, Zhang00], ou en combinant deux vides comprimés dans des directions orthogonales sur une lame séparatrice 50/50 [Bowen03, Bowen04, Mizuno05, Furusawa98].

Bien que le vide comprimé bimode fasse partie des états intriqués les “plus simples”, il n'en reste pas moins l'état maximalement intriqué pour une énergie donnée [Barnett89, Barnett91], car c'est un état pur dont chaque mode réduit est dans un état thermique.

#### Moments et fonction de Wigner

Comme précédemment, on peut obtenir les moments du vide comprimé bimode

$$\mathbf{d} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \mathbf{\Gamma} = \begin{pmatrix} \cosh 2r\mathbb{I} & -\sinh 2r\mathbb{Z} \\ -\sinh 2r\mathbb{Z} & \cosh 2r\mathbb{I} \end{pmatrix} \quad (2.147)$$

avec  $\mathbb{I} = \text{diag}(1, 1)$  et  $\mathbb{Z} = \text{diag}(1, -1)$ , et sa fonction de Wigner à partir de sa définition (2.146) :

$$W_{\text{EPR}}(x_1, p_1, x_2, p_2) = \frac{1}{(2\pi N_0)^2} e^{-\frac{1}{2N_0} \left( \frac{1}{2s}(x_1+x_2)^2 + \frac{1}{2s}(p_1-p_2)^2 + \frac{1}{2/s}(x_1-x_2)^2 + \frac{1}{2/s}(p_1+p_2)^2 \right)} \quad (2.148)$$

#### Opérateur de squeezing bimode

**Définition** L'opérateur de squeezing bimode est défini par :

$$\hat{\mathbf{S}}_2(r) = \exp \left[ r\hat{\mathbf{a}}\hat{\mathbf{b}} - r\hat{\mathbf{a}}^\dagger\hat{\mathbf{b}}^\dagger \right] \quad (2.149)$$

Il vérifie  $\hat{\mathbf{S}}_2^\dagger(r) = \hat{\mathbf{S}}_2(-r)$  et  $\hat{\mathbf{S}}_2^\dagger(r)\hat{\mathbf{S}}_2(r) = \mathbb{I}$ , et on montre en annexe A.3 qu'il peut s'écrire comme :

$$\hat{\mathbf{S}}_2(r) = \exp \left[ -\hat{\mathbf{a}}^\dagger\hat{\mathbf{b}}^\dagger \tanh r \right] \exp \left[ -(\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}} + \hat{\mathbf{b}}^\dagger\hat{\mathbf{b}} + 1) \ln(\cosh r) \right] \exp \left[ +\hat{\mathbf{a}}\hat{\mathbf{b}} \tanh r \right] \quad (2.150)$$

**Transformation des opérateurs** Les opérateurs  $\hat{a}(r)=\hat{S}_2^\dagger(r)\hat{a}\hat{S}_2(r)$  et  $\hat{b}(r)=\hat{S}_2^\dagger(r)\hat{b}\hat{S}_2(r)$  évoluent selon

$$\frac{d}{dr}\hat{a}(r) = \hat{S}_2(r) \left[ \hat{a}, \hat{a}\hat{b} - \hat{a}^\dagger\hat{b}^\dagger \right] \hat{S}_2(r) = -\hat{b}^\dagger, \quad (2.151a)$$

$$\frac{d}{dr}\hat{b}(r) = \hat{S}_2(r) \left[ \hat{b}, \hat{a}\hat{b} - \hat{a}^\dagger\hat{b}^\dagger \right] \hat{S}_2(r) = -\hat{a}^\dagger, \quad (2.151b)$$

dont les solutions sont :

$$\boxed{\begin{aligned} \hat{a}(r) &= \hat{a} \cosh r - \hat{b}^\dagger \sinh r \\ \hat{b}(r) &= \hat{b} \cosh r - \hat{a}^\dagger \sinh r \end{aligned}} \quad (2.152)$$

Les quadratures des modes 1 et 2 sont donc mélangées entre elles sous l'action de la compression :

$$\hat{X}_1(r) = \hat{X}_1 \cosh r - \hat{X}_2 \sinh r \quad \hat{X}_2(r) = \hat{X}_2 \cosh r - \hat{X}_1 \sinh r \quad (2.153a)$$

$$\hat{P}_1(r) = \hat{P}_1 \cosh r + \hat{P}_2 \sinh r \quad \hat{P}_2(r) = \hat{P}_2 \cosh r + \hat{P}_1 \sinh r \quad (2.153b)$$

En revanche, les différences et les sommes des quadratures sont bien comprimées ou amplifiées :

$$\hat{X}_1(r) + \hat{X}_2(r) = e^{-r} (\hat{X}_1 + \hat{X}_2) \quad \hat{P}_1(r) - \hat{P}_2(r) = e^{-r} (\hat{P}_1 - \hat{P}_2) \quad (2.154a)$$

$$\hat{X}_1(r) - \hat{X}_2(r) = e^{+r} (\hat{X}_1 - \hat{X}_2) \quad \hat{P}_1(r) + \hat{P}_2(r) = e^{+r} (\hat{P}_1 + \hat{P}_2) \quad (2.154b)$$

Ces transformations correspondent à la définition du vide comprimé bimode (2.146), ce qui montre que :

$$\boxed{|\psi_{\text{EPR}}\rangle = \hat{S}_2(r)|0_{ab}\rangle = \exp \left[ r\hat{a}\hat{b} - r\hat{a}^\dagger\hat{b}^\dagger \right] |0_{ab}\rangle} \quad (2.155)$$

### Décomposition en base de Fock

Puisque le vide comprimé bimode est intriqué en quadratures, il doit l'être également en base de Fock. Comme montré en annexe A.3 en utilisant l'expansion de l'opérateur de squeezing (2.150), sa décomposition en base de Fock est

$$\boxed{|\psi_{\text{EPR}}\rangle = \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} (-\lambda)^n |n\rangle |n\rangle,} \quad (2.156)$$

avec  $\lambda = \tanh r$ . Les deux modes contiennent donc exactement le même nombre de photons. Cette propriété est couramment utilisée pour préparer des états de Fock de manière probabiliste : une mesure de  $n$  photons dans un mode projette l'autre mode dans l'état  $|n\rangle$ .

Afin de simplifier les notations et de s'affranchir du terme  $(-1)^n$  dans (2.156), nous noterons souvent un état EPR  $|\psi_{\text{EPR}}\rangle = \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} \lambda^n |n\rangle |n\rangle$ , ce qui correspond simplement à une rotation dans l'espace des phases.

### Obtention avec deux vides comprimés monomodes

On peut obtenir un vide comprimé bimode à partir d'un vide comprimé monomode selon  $\hat{P}$  (donc de paramètre  $-r$ ), et d'un vide comprimé monomode selon  $\hat{X}$  (donc de paramètre  $r$ ). Ecrivons pour cela l'état résultant du mélange de ces deux vides comprimés avec une lame séparatrice 50/50, décrite par l'opérateur  $\hat{U}_{BS}$  :

$$\hat{U}_{BS}\hat{S}_a(-r)\otimes\hat{S}_b(r)|0_{ab}\rangle = \hat{U}_{BS}\hat{S}_a(-r)\hat{U}_{BS}^\dagger\hat{U}_{BS}\hat{S}_b(r)\hat{U}_{BS}^\dagger|0_{ab}\rangle \quad (2.157)$$

On utilise ensuite la relation (A.1) pour obtenir :

$$\hat{U}_{BS}\hat{S}_a(-r)\hat{U}_{BS}^\dagger = \exp\left[-\frac{r}{2}\left(\frac{\hat{a}-\hat{b}}{\sqrt{2}}\right)^2 + \frac{r}{2}\left(\frac{\hat{a}^\dagger-\hat{b}^\dagger}{\sqrt{2}}\right)^2\right] \quad (2.158a)$$

$$\hat{U}_{BS}\hat{S}_b(+r)\hat{U}_{BS}^\dagger = \exp\left[+\frac{r}{2}\left(\frac{\hat{a}+\hat{b}}{\sqrt{2}}\right)^2 - \frac{r}{2}\left(\frac{\hat{a}^\dagger+\hat{b}^\dagger}{\sqrt{2}}\right)^2\right] \quad (2.158b)$$

Enfin, on vérifie que les termes des deux exponentielles commutent entre eux, ce qui permet de les regrouper dans une seule exponentielle et d'obtenir :

$$\hat{U}_{BS}\hat{S}_a(-r)\otimes\hat{S}_b(r)|0\rangle = \exp\left[r\hat{a}\hat{b}-r\hat{a}^\dagger\hat{b}^\dagger\right]|0\rangle = \hat{S}_2(r)|0\rangle \quad (2.159a)$$

Notons que l'on peut également retrouver ce résultat en utilisant les fonctions de Wigner et la relation (2.97).

### 2.5.6 Les états thermiques

#### Définition

Comme son nom l'indique, un état thermique correspond à l'état du rayonnement lorsqu'il est en équilibre avec un réservoir de température  $T$  [Chiao08, Loudon00] :

$$\hat{\rho}_{th} = \frac{e^{-\beta\hat{H}}}{\text{Tr}\{e^{-\beta\hat{H}}\}} \quad (2.160)$$

Avec  $\beta=1/k_B T$  et  $\hat{H}$  l'hamiltonien libre. Le nombre moyen de photons  $\bar{n}=\text{Tr}\{\hat{\rho}_{th}\hat{a}^\dagger\hat{a}\}$  est égal à  $(e^{\hbar\omega/k_B T}-1)^{-1}$ . En posant  $\bar{n}=\sinh^2 r$  et  $\lambda=\tanh r$ , l'état (2.160) s'écrit :

$$\boxed{\hat{\rho}_{th} = (1-\lambda^2) \sum_{n=0}^{\infty} \lambda^{2n} |n\rangle\langle n|} \quad (2.161)$$

On reconnait la trace partielle sur un mode d'un état EPR de paramètre  $r$  :

$$\hat{\rho}_{th} = \text{Tr}_a\{|\psi_{EPR}\rangle\langle\psi_{EPR}|\} = \text{Tr}_b\{|\psi_{EPR}\rangle\langle\psi_{EPR}|\} \quad (2.162)$$

Un état thermique a donc pour purification un état EPR.

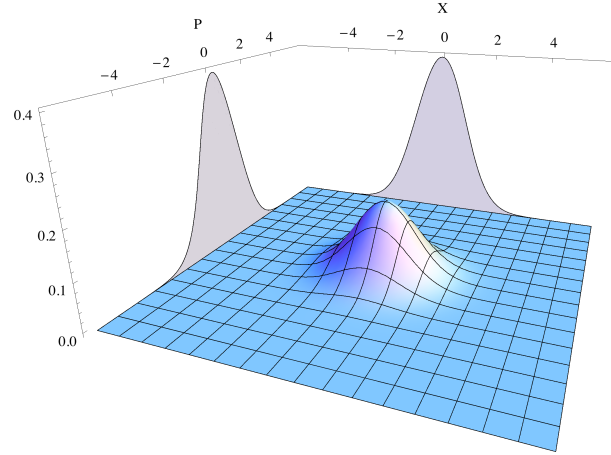


FIGURE 2.5 – Fonction de Wigner d'un état thermique de variance  $V=2$ , et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .

### Quelques propriétés

Les états thermiques maximisent l'entropie de von Neumann  $\mathcal{S}(\hat{\rho}) = -\text{Tr}\{\hat{\rho} \log \hat{\rho}\}$  pour une énergie (donc  $\bar{n}$  et  $\lambda$ ) donnée [Wehrl78]. Cette propriété, associée à (2.162), fait que les états EPR sont des états maximalelement intriqués pour une énergie donnée [Barnett89, Barnett91], comme mentionné précédemment.

Pour toutes les quadratures  $\hat{X}_\theta$ , on montre que leur variance est égale à

$$V = \frac{1+\lambda^2}{1-\lambda^2} N_0 = \cosh(2r) N_0 = (2\bar{n}+1)N_0, \quad (2.163)$$

et que leur valeur moyenne est nulle. Les états thermiques ressemblent donc au vide mais avec une variance élargie. Utilisés dans un mode d'une lame séparatrice, ces fluctuations s'ajoutent à celles du signal injecté dans l'autre mode et correspondent à un ajout de bruit gaussien.

### Moments et fonction de Wigner

Les moments d'un état thermique sont

$$\mathbf{d} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \mathbf{\Gamma} = \begin{pmatrix} V & 0 \\ 0 & V \end{pmatrix} \quad (2.164)$$

ce qui correspond à une fonction de Wigner égale à (Fig. 2.5)

$$W_{\text{th}}(x, p) = \frac{1}{2\pi N_0 V} e^{-\frac{1}{2N_0} \frac{x^2 + p^2}{V}}. \quad (2.165)$$

### Décomposition sur des états cohérents

Nous aurons besoin de la fonction  $P$  d'un état thermique lorsque nous étudierons l'utilisation d'un amplificateur sans bruit en cryptographie quantique. On peut l'obtenir en utilisant la définition (2.81), mais cette approche est très calculatoire. Une méthode plus simple est de

remarquer qu'un état thermique se décompose sur des états cohérents, en utilisant (2.155) et (2.150) (voir aussi (A.28a)) :

$$\hat{\rho}_{\text{th}}(\lambda) = (1-\lambda^2) \text{Tr}_b \{ e^{-\lambda \hat{a}^\dagger \hat{b}^\dagger} |0_{ab}\rangle \langle 0_{ab}| e^{\lambda \hat{a} \hat{b}} \} \quad (2.166a)$$

$$= \frac{1}{\pi} (1-\lambda^2) \int d^2\beta \langle \beta_b | e^{-\lambda \hat{a}^\dagger \hat{b}^\dagger} |0_{ab}\rangle \langle 0_{ab}| e^{\lambda \hat{a} \hat{b}} | \beta_b \rangle \quad (2.166b)$$

$$= \frac{1}{\pi} (1-\lambda^2) \int d^2\beta |\langle \beta_b | 0_b \rangle|^2 e^{-\lambda \beta^* \hat{a}^\dagger} |0_a\rangle \langle 0_a| e^{\lambda \beta \hat{a}} \quad (2.166c)$$

$$= \frac{1}{\pi} (1-\lambda^2) \int d^2\beta e^{-|\beta|^2} e^{-\lambda \beta^* \hat{a}^\dagger} |0_a\rangle \langle 0_a| e^{\lambda \beta \hat{a}} \quad (2.166d)$$

Compte tenu de (2.122a),  $e^{-\lambda \beta^* \hat{a}^\dagger} |0_a\rangle = e^{\frac{1}{2}|\lambda\beta|^2} |-\lambda\beta^*\rangle$ , et on a donc

$$\hat{\rho}_{\text{th}}(\lambda) = \frac{1}{\pi} (1-\lambda^2) \int d^2\beta e^{|\beta|^2(\lambda^2-1)} |-\lambda\beta^*\rangle \langle -\lambda\beta^*|. \quad (2.167)$$

Ensuite, on fait le changement de variable  $-\lambda\beta^* = \alpha$ , ce qui implique  $d^2\beta = -d^2\alpha/\lambda^2$ , et on obtient finalement :

$$\hat{\rho}_{\text{th}}(\lambda) = \frac{1}{\pi \bar{n}} \int d^2\alpha e^{-\frac{1}{\bar{n}}|\alpha|^2} |\alpha\rangle \langle \alpha| \quad (2.168)$$

où  $\bar{n} = \frac{\lambda^2}{1-\lambda^2}$  est le nombre moyen de photons. On reconnaît immédiatement une décomposition qui permet d'identifier la fonction  $P$  :

$$P(\alpha) = \frac{1}{\pi \bar{n}} e^{-\frac{1}{\bar{n}}|\alpha|^2} \quad (2.169)$$

Ainsi, contrairement au cas général, pour un état thermique la fonction  $P$  est bien égale à une distribution de probabilité. Un état thermique est un mélange d'états cohérents de phases aléatoires, dont les amplitudes suivent une distribution de moyenne nulle, et de variance dépendant de  $\bar{n}$ .

On montre par un calcul similaire que la fonction  $P$  d'un état thermique déplacé de  $\gamma$  est égale à  $P(\alpha-\gamma)$  :

$$\hat{D}(\gamma) \hat{\rho}_{\text{th}}(\lambda) \hat{D}^\dagger(\gamma) = \frac{1}{\pi \bar{n}} \int d^2\alpha e^{-\frac{1}{\bar{n}}|\alpha-\gamma|^2} |\alpha\rangle \langle \alpha| \quad (2.170)$$

## 2.6 Etats non gaussiens

Les états présentés jusqu'à maintenant possédaient une fonction de Wigner gaussienne. Même si l'on peut combiner les différentes opérations de déplacement, squeezing, rotation, ou ajout de bruit thermique pour former toute une classe d'états gaussiens, cela ne suffit pas à décrire tous les états quantiques possibles. Comme indiqué dans la section présentant la fonction de Wigner, il existe nécessairement des états non gaussiens, qui devront donc avoir des fonctions de Wigner prenant des valeurs négatives [Hudson74].

Ces états sont extrêmement importants en information quantique, car un très grand nombre de protocoles en ont besoin pour s'avérer efficaces. Par exemple, le calcul quantique avec des variables continues nécessite une opération non gaussienne pour être plus efficace qu'un ordinateur classique [Bartlett02]. La présence d'un élément non gaussien est indispensable, que ce soit

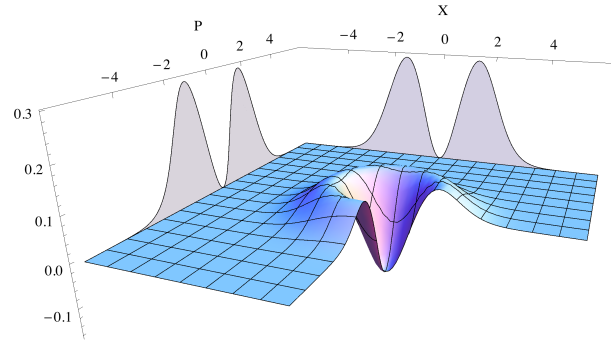


FIGURE 2.6 – *Fonction de Wigner d'un état de Fock  $n=1$ , et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .*

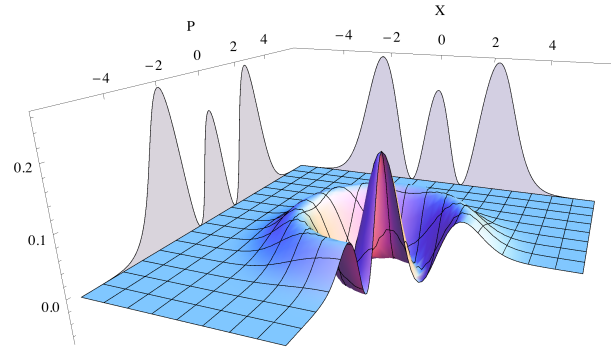


FIGURE 2.7 – *Fonction de Wigner d'un état de Fock  $n=2$ , et densité de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .*

par une évolution non linéaire de type Kerr, l'utilisation de mesures non gaussiennes, ou d'états non gaussiens [Marek09, Braunstein05]. La distillation d'intrication pour des états gaussiens est également impossible sans éléments non gaussiens [Eisert02, Fiurášek02, Giedke02], tout comme la correction d'erreurs [Niset09].

Parmi tous les états non gaussiens imaginables, certains se révèlent particulièrement utiles. C'est le cas des états de Fock, et des superpositions d'états cohérents, que nous détaillons maintenant.

### 2.6.1 Les états de Fock

Nous les avons déjà rencontré lors de la quantification du champ. Ces états peuvent être utilisés en tant que tels pour coder de l'information (voir section 2.7.1), mais ils sont aussi couramment utilisés en tant que ressources pour implémenter des opérations quantiques.

Étudions maintenant quelques-unes de leurs propriétés plus en détail. La fonction de Wigner

d'un état de Fock  $|n\rangle$  est

$$W_{|n\rangle}(x, p) = \frac{(-1)^n}{2\pi N_0} e^{-\frac{1}{2N_0}(x^2+p^2)} L_n \left( \frac{x^2+p^2}{N_0} \right), \quad (2.171)$$

avec  $L_n$  le nième polynôme de Laguerre :

$$L_n(x) = \sum_{k=0}^n C_n^k \frac{(-x)^k}{k!} \quad (2.172)$$

A part pour le vide  $|0\rangle$  qui est gaussien, elle prend toujours des valeurs négatives. L'état à un photon, qui est sans doute l'état de Fock le plus utilisé comme ressource, a une fonction de Wigner égale à

$$W_{|1\rangle}(x, p) = \frac{1}{2\pi N_0} \left( \frac{x^2+p^2}{N_0} - 1 \right) e^{-\frac{1}{2N_0}(x^2+p^2)}. \quad (2.173)$$

Elle est négative pour un rayon  $\sqrt{x^2+p^2} < \sqrt{N_0}$ . Cette négativité disparaît pour des pertes supérieures à 3 dB (voir par exemple [Grosshans02a] pour un calcul détaillé).

Les fonctions de Wigner des états  $|1\rangle$  et  $|2\rangle$  sont représentées sur les figures 2.6 et 2.7. Les états de Fock n'ont pas de phase définie : ils sont symétriques et leurs fonctions de Wigner sont invariantes par rotation. Une mesure d'une quadrature quelconque donnera toujours la même distribution, quelle que soit la quadrature choisie.

### 2.6.2 Les superpositions d'états cohérents, ou “chats de Schrödinger”

Dans sa célèbre expérience de pensée [Schrödinger35], E. Schrödinger s'interrogeait sur la possibilité de mettre un chat dans une superposition quantique d'états “mort” et “vivant”, en l'enfermant dans une boîte contenant un atome instable, relié à une fiole de poison. En cas de désintégration de l'atome, la fiole se brise et fait trépasser le chat. Puisque l'atome peut être dans une superposition quantique d'un état non désintégré et désintégré, le chat devrait pouvoir également se retrouver à la fois mort et vivant. Le chat est-il dans cet état tant que l'on n'a pas ouvert la boîte ?

Utiliser un chat n'est bien sûr pas essentiel. Le point crucial est d'avoir un système suffisamment macroscopique pour que nous puissions en avoir une expérience dans notre monde classique. Le paradoxe soulevé par Schrödinger concerne davantage la possibilité de mettre des états macroscopiques “classiques” dans des superpositions quantiques. Nous avons déjà rencontré de tels états, il s'agit des états cohérents. C'est donc naturellement qu'une superposition d'états cohérents est appelée, par analogie, un “chat de Schrödinger”.

Une telle superposition est définie par

$$|\psi_{\text{chat}}\rangle = \mathcal{N} \left( |\alpha\rangle + e^{i\phi} |-\alpha\rangle \right), \quad (2.174)$$

avec  $\mathcal{N} = 1/\sqrt{2(1+e^{-2|\alpha|^2} \cos \phi)}$ . Les états sont pris de même amplitude, car on peut toujours s'y ramener pour des états d'amplitudes quelconques avec un déplacement. La fonction de Wigner associée à  $|\psi_{\text{chat}}\rangle$  s'écrit

$$W_{\text{chat}}(x, p) = \frac{e^{-\frac{1}{2N_0}(x^2+p^2)}}{2\pi N_0 (1+e^{-2|\alpha|^2} \cos \phi)} \left( e^{-2|\alpha|^2} \cosh \left( \frac{2\alpha x}{\sqrt{N_0}} \right) + \cos \left( \frac{2\alpha p}{\sqrt{N_0}} - \phi \right) \right). \quad (2.175)$$



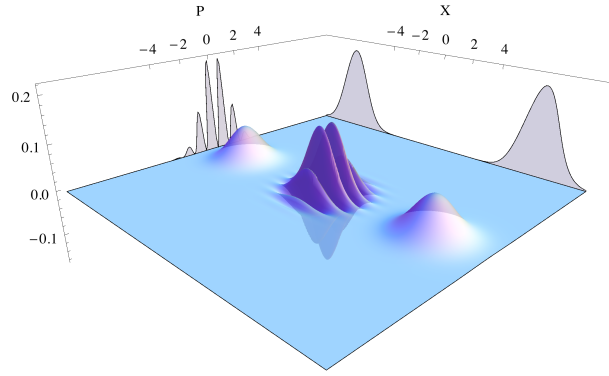


FIGURE 2.8 – Fonction de Wigner d’un chat impair,  $\alpha=3.5$ , et distribution de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ .

Ces états n’ont pas qu’un intérêt fondamental, ils sont aussi particulièrement utiles en information quantique [Gilchrist04, Lund08]. Deux superpositions particulières sont souvent utilisées, les chats pairs  $|\psi_{\text{chat}+}\rangle$  et impairs  $|\psi_{\text{chat}-}\rangle$ , appelés ainsi en raison de leurs distributions de photons :

$$|\psi_{\text{chat}+}\rangle = \mathcal{N}_+ \left( |\alpha\rangle + |-\alpha\rangle \right) = \frac{\sqrt{2}e^{-|\alpha|^2/2}}{\sqrt{1+e^{-2|\alpha|^2}}} \sum_{k=0}^{\infty} \frac{\alpha^{2k}}{\sqrt{(2k)!}} |2k\rangle \quad (2.176)$$

$$|\psi_{\text{chat}-}\rangle = \mathcal{N}_- \left( |\alpha\rangle - |-\alpha\rangle \right) = \frac{\sqrt{2}e^{-|\alpha|^2/2}}{\sqrt{1-e^{-2|\alpha|^2}}} \sum_{k=0}^{\infty} \frac{\alpha^{2k+1}}{\sqrt{(2k+1)!}} |2k+1\rangle \quad (2.177)$$

Un chat pair (resp. impair) ne contient donc qu’un nombre pair (resp. impair) de photons. Leurs fonctions de Wigner sont représentées sur les figures 2.8 et 2.9. Ces deux états sont de plus strictement orthogonaux entre eux quel que soit  $\alpha$ , ce qui n’est pas le cas de  $|\alpha\rangle$  et  $|-\alpha\rangle$ .

## 2.7 Application à l’information quantique

L’utilisation des propriétés quantiques pour le traitement de l’information est un domaine relativement récent, qui a commencé à réellement se développer vers la fin des années 1990. On peut distinguer deux grands domaines : le calcul quantique, et les communications quantiques.

Le calcul quantique [DiVincenzo95] exploite certaines propriétés quantiques fondamentales, tel que le principe de superposition, afin de résoudre certaines tâches exponentiellement plus rapidement qu’avec un ordinateur classique. Parmi les protocoles les plus courants, citons par exemple l’algorithme de Shor [Shor97] pour la factorisation en nombres premiers, l’algorithme de Grover [Grover96] pour la recherche dans une base de données non triée, ou encore l’algorithme de Deutsch-Jozsa [Deutsch92, Cleve98].

Les communications quantiques, de manière générale, sont “l’art de transférer un état quantique d’un endroit à l’autre” [Gisin07]. La téléportation quantique [Bennett93, Bouwmeester97, Braunstein98, Furusawa98] en est un parfait exemple. Le principe est de transférer l’état quantique d’une particule (atome, photon, ...) à une autre particule similaire située à un autre endroit, en utilisant un état intriqué comme ressource et en faisant des mesures de Bell. La

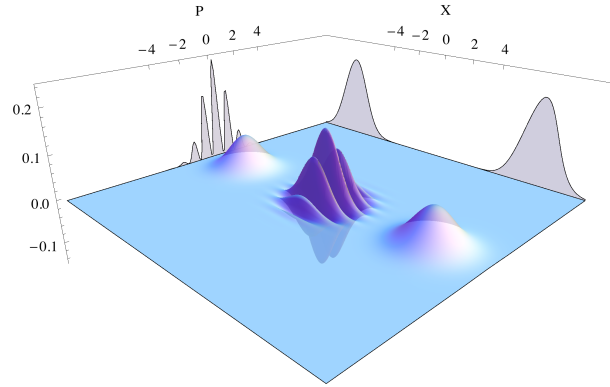


FIGURE 2.9 – Fonction de Wigner d'un chat pair,  $\alpha=3.5$ , et distribution de probabilité des quadratures  $\hat{X}$  et  $\hat{P}$ . On note que les oscillations ont une phase différente de celles d'un chat impair.

cryptographie quantique est également une branche importante des communications quantiques [Bennett84, Scarani09]. Le principe commun à tous les protocoles est d'échanger des états quantiques entre deux partenaires, traditionnellement appelés Alice et Bob, afin de pouvoir établir une clé secrète grâce aux mesures effectuées sur ces états. Nous aurons l'occasion de revenir plus en détail sur les principes de cryptographie quantique au cours du chapitre 7.

Dans tous ces différents domaines de l'information quantique, il existe néanmoins un dénominateur commun : la manière d'encoder l'information, qui peut être discrète, ou continue.

### 2.7.1 Variables discrètes

On parle généralement de variables discrètes lorsque l'information est encodée dans un sous espace de dimension finie, associé à un spectre discret. Le codage utilise deux états quantiques orthogonaux pour former un qubit, généralisant ainsi la notion de bit classique. Ces deux états, traditionnellement dénommés  $|0\rangle$  et  $|1\rangle$ , correspondent aux deux valeurs 0 et 1 d'un bit classique. Alors que classiquement un bit ne peut prendre qu'une seule de ces deux valeurs à la fois, la physique quantique autorise une superposition quelconque  $\alpha|0\rangle + \beta|1\rangle$ , avec  $|\alpha|^2 + |\beta|^2 = 1$ . Cette possibilité de superposition est au cœur des algorithmes de calculs quantiques surpassant leurs homologues classiques. Les qubits peuvent être formés à l'aide d'une multitude de systèmes physiques : niveaux atomiques dans des ions ou des atomes neutres, matériaux supraconducteurs, molécules, ... et bien sûr avec les états du champ lumineux.

Les états de Fock, discrets par nature, peuvent être utilisés pour former un qubit en utilisant simplement les états  $|0\rangle$  et  $|1\rangle$ , contenant respectivement zéro et un photon. Cet encodage pose toutefois de nombreux problèmes technologiques (efficacité des détecteurs pour détecter le vide, effet des pertes, ...) qui peuvent être en partie contournés en utilisant un encodage avec un photon dans deux modes différents, par exemple deux états de polarisation  $|H\rangle$  et  $|V\rangle$ .

Pour le calcul quantique, cet encodage permet de réaliser très simplement les portes quantiques à un qubit (par exemple les rotations, la porte de Hadamard) à l'aide de lames séparatrices et de lames à retards. Nous reviendrons plus en détail sur ces portes dans le chapitre 6 concernant la caractérisation d'une porte de phase. Les photons uniques ont de plus l'avantage d'être

relativement résistants au bruit, et de permettre des transformations donnant une très bonne fidélité [Looock11]. L'inconvénient majeur est par contre la difficulté à faire interagir les photons pour former des portes à deux qubits. Les matériaux non linéaires ne présentent pas une efficacité suffisante pour assurer un couplage entre deux photons uniques, si bien que d'autres alternatives doivent être utilisées. L'une d'entre elles est par exemple l'utilisation de non linéarités géantes dans les atomes de Rydberg [Peyronel12, Saffman10]. Une autre approche très prometteuse est d'utiliser des non linéarités induites par des mesures [Knill01, O'Brien03, O'Brien07], associées avec le principe de la "téléportation de portes quantiques" [Gottesman99]. Une opération quantique arbitraire peut alors être réalisée avec de l'optique linéaire et des compteurs de photons, mais de manière probabiliste. Un fonctionnement "quasi déterministe" est en théorie possible, mais nécessite une quantité de ressources difficilement intégrable à grande échelle, même si des améliorations du protocole existent [Kok07].

De nombreux protocoles de communications quantiques utilisent un encodage discret sous forme de qubit. Outre les protocoles de téléportation quantique [Bennett93, Bouwmeester97], citons par exemple le fameux protocole de cryptographie quantique BB84 [Bennett84]. Le principe est relativement simple : Alice choisit aléatoirement un des quatre états  $|0\rangle$ ,  $|1\rangle$ ,  $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ , ou  $|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$ , qu'elle envoie ensuite à Bob. Ce dernier choisit aléatoirement une base de mesure  $\{|0\rangle, |1\rangle\}$  ou  $\{|-\rangle, |+\rangle\}$ , qu'il révèle ensuite à Alice, en gardant le résultat de sa mesure secret. Tout un ensemble d'algorithmes classiques permettent ensuite d'extraire une clé parfaitement secrète d'un éventuel espion, dont la taille dépend bien sûr des conditions expérimentales.

En plus de la faible interaction entre photons uniques, les variables discrètes posent un certain nombre de défis technologiques qui ne sont pour l'instant pas complètement surmontés. Un des plus importants est que les photons uniques sont difficiles à produire de manière déterministe dans un mode précis. De nombreuses sources de photons existent [Grangier04], mais en général le photon est soit émis de manière déterministe mais dans un mode aléatoire (par exemple avec des atomes piégés [McKeever04, Darquié05, Hijlkema07], ou les centres NV du diamant [Kurtsiefer00, Brouri00]), soit émis dans un mode précis mais de manière non déterministe (par exemple avec la fluorescence paramétrique).

La détection est un autre facteur limitant des variables discrètes. Le détecteur le plus courant est la photodiode à avalanche (APD), qui indique plutôt la présence "d'au moins un photon" sans en préciser le nombre, avec une efficacité d'environ 50 %. D'autres systèmes permettent de résoudre le nombre de photons avec une meilleure efficacité, mais au prix d'une mise en œuvre technique contraignante et/ou d'une plus grande lenteur (par exemple les détecteurs de types VLPC [Waks03], ou supraconducteurs [Rosenberg05]).

### 2.7.2 Variables continues

Les variables continues comprennent les états naturellement associés à un spectre continu. Nous en avons présenté plusieurs, pour la plupart gaussiens, et qui présentent l'avantage de pouvoir être facilement produits expérimentalement. Les mesures associées sont également beaucoup plus simples que pour les variables discrètes, puisque de simples photodiodes suffisent pour mesurer leurs quadratures avec une détection homodyne (dont nous présenterons le principe dans le chapitre suivant). Le codage de l'information est ici analogue au codage analogique classique. Comme ce dernier en revanche, le bruit et les pertes sont plus difficiles à contrer, car les états y sont beaucoup plus sensibles.

Il existe de nombreux protocoles basés sur les variables continues : la téléportation quantique [Braunstein98, Furusawa98], mais aussi la cryptographie quantique, dont un des protocoles ne nécessite que des états cohérents [Grosshans03b, Scarani09], et bien sûr, le calcul quantique

[Lloyd99, Braunstein05].

Le calcul quantique avec des variables continues est une généralisation d’un “ordinateur analogique” classique. La réalisation d’un hamiltonien arbitraire, nécessite au moins une opération non linéaire [Lloyd99], par exemple de type Kerr, qui ne préserve pas la structure gaussienne des états. En contrepartie, un ordinateur “gaussien”, comportant des états gaussiens et des transformations gaussiennes n’apporte pas d’amélioration sur le temps de calcul car peut être simulé sur un ordinateur classique [Bartlett02]. Cet élément peut être un hamiltonien d’ordre supérieur à un quadratique, l’utilisation de mesures non gaussiennes, ou encore d’états non gaussiens [Marek09]. C’est dans une certaine mesure le pendant en variable continue du théorème de Gottesman-Knill [Nielsen00], selon lequel un calcul limité à des transformations du groupe de Clifford sur des qubits peut être simulé classiquement.

Afin d’utiliser les avantages des variables continues, tout en bénéficiant des avantages d’un codage avec des variables discrètes, des approches hybrides se sont développées [Looock11]. Le principe est d’encoder un qubit dans deux états décrits par des variables continues, tel que des superpositions d’états comprimés [Gottesman01], ou des superpositions d’états cohérents [Ralph02, Ralph03, Lund08], avec  $|0\rangle = |-\alpha\rangle$  et  $|1\rangle = |\alpha\rangle$ . Cette dernière méthode sera étudiée plus en détail dans le chapitre sur la caractérisation de la porte de phase. Ces approches hybrides ne sont toutefois pas sans poser quelques difficultés, notamment pour réaliser des portes à plusieurs qubits ou pour préparer les ressources nécessaires.

## 2.8 Conclusion

Dans ce chapitre, nous avons présenté les outils théoriques de base nécessaires à la description du champ, à la fois en termes de variables discrètes, qu’en termes de variables continues. Nous verrons dans le reste de ce manuscrit comment les deux approches peuvent être combinées, afin d’ouvrir de nouvelles perspectives en information quantique.



# Chapitre 3

## Le dispositif expérimental

### Sommaire

|            |                                                         |           |
|------------|---------------------------------------------------------|-----------|
| <b>3.1</b> | <b>Présentation du dispositif</b>                       | <b>47</b> |
| 3.1.1      | Introduction                                            | 47        |
| 3.1.2      | Le dispositif expérimental                              | 48        |
| 3.1.3      | La source laser                                         | 49        |
| <b>3.2</b> | <b>Transformations unitaires : optique linéaire</b>     | <b>53</b> |
| 3.2.1      | Déphasage                                               | 53        |
| 3.2.2      | Lame séparatrice                                        | 53        |
| 3.2.3      | Lames demi-onde $\lambda/2$ et quart-d'onde $\lambda/4$ | 54        |
| 3.2.4      | Cube séparateur de polarisation (PBS)                   | 54        |
| <b>3.3</b> | <b>Transformations unitaires : optique non linéaire</b> | <b>55</b> |
| 3.3.1      | Génération de seconde harmonique (GSH)                  | 56        |
| 3.3.2      | Amplification paramétrique optique (OPA)                | 56        |
| <b>3.4</b> | <b>Détection et mesures projectives</b>                 | <b>61</b> |
| 3.4.1      | Détection homodyne                                      | 61        |
| 3.4.2      | Photodiode à avalanche                                  | 67        |
| 3.4.3      | Soustraction de photon                                  | 68        |
| <b>3.5</b> | <b>Conclusion</b>                                       | <b>69</b> |

### 3.1 Présentation du dispositif

#### 3.1.1 Introduction

Nous avons vu dans le chapitre précédent plusieurs exemples d'états quantiques, parmi lesquels les états non gaussiens qui sont essentiels en information quantique. Voyons maintenant comment les produire expérimentalement, et les caractériser.

La production d'états non gaussiens peut être réalisée de plusieurs manières. La première serait d'utiliser une source qui les produit directement. Si cela peut être envisageable pour des états de Fock, cette méthode reste toutefois limitée pour la production d'états plus "exotiques", tels que les chats de Schrödinger, ou d'autres superpositions. La seconde manière, plus couramment utilisée, est d'utiliser un élément non gaussien induisant une évolution non gaussienne. On peut alors partir d'un état gaussien, facile à produire, et le faire évoluer jusqu'à l'état recherché.

Une évolution non gaussienne déterministe requiert de fortes non linéarités. Un couplage d'ordre deux préservant la structure gaussienne, il est nécessaire d'utiliser au moins un couplage d'ordre trois tel que l'effet Kerr. Malheureusement, il n'existe pas pour l'instant de technique permettant d'avoir un tel couplage de manière suffisante. L'autre solution consiste alors à utiliser des mesures non gaussiennes pour induire les transformations recherchées. Cette technique, non déterministe, offre de plus une grande versatilité puisque l'on peut appliquer différentes transformations en changeant simplement les mesures effectuées.

Il reste alors à partir d'un état gaussien. Si les états cohérents sont les états les plus simples à produire, ils n'ont pas non plus un potentiel très important pour pouvoir y appliquer des transformations induites par des mesures non gaussiennes. Si l'on prélève une partie du faisceau à l'aide d'une lame séparatrice, les deux modes de sortie ne sont pas intriqués, et un conditionnement sur un des modes n'aura aucun effet sur le second. Après les états cohérents, les états purs gaussiens les plus simples sont les états comprimés, monomodes ou bimodes, et sont utilisés dans la plupart des expériences d'optique quantique.

La production de ces états se fait le plus simplement par amplification paramétrique optique [Ou92, Zhang00, Loudon00], en utilisant un couplage non linéaire du second ordre. Afin d'obtenir des non linéarités suffisantes, deux méthodes sont utilisées. La première est d'utiliser un laser continu, et de placer les cristaux dans des cavités pour multiplier le nombre de passages. Les états produits sont de très bonnes qualités, mais cette méthode présente plusieurs inconvénients. Les différentes cavités doivent être verrouillées entre elles, et les modes quantiques sont fréquentiels, ce qui nécessite d'appliquer un filtrage sur les données.

La seconde méthode est d'utiliser des impulsions laser très courtes (femto ou picoseconde) avec de très fortes puissances instantanées, suffisantes pour n'effectuer qu'un seul passage. Outre le fait de ne plus avoir besoin de cavités, cette méthode procure également un échantillonnage temporel naturel, bien adapté pour les communications quantiques.

Notre dispositif expérimental est basé sur la production d'états comprimés monomodes et bimodes en régime impulsif. Les états quantiques sont ensuite produits grâce à des mesures non gaussiennes, appliquées par exemple sur un des modes du vide comprimé, ou en prélevant une partie du faisceau. Nous détaillerons chacun de ces éléments dans la suite de ce chapitre.

### 3.1.2 Le dispositif expérimental

L'outil de base de notre dispositif est le laser impulsif. Les impulsions femtosecondes (de longueur d'onde 850 nm) sont créées dans une cavité contenant un cristal de titane-saphir, puis extraites à l'aide d'une cellule de Bragg. Elles ont ensuite plusieurs utilités. La majeure partie de la puissance sert à produire un faisceau bleu à 425 nm par doublage de fréquence (SHG) dans un premier cristal non linéaire d'ordre deux. Ce faisceau bleu joue ensuite le rôle de pompe pour le deuxième cristal non linéaire servant d'amplificateur paramétrique optique (OPA), utilisé dans toutes nos expériences jusqu'à présent.

L'OPA est l'outil de base pour la production des états. En configuration dégénérée, il produit un vide comprimé monomode, et en configuration non dégénérée, un état EPR. Le dispositif après l'OPA est adapté en fonction des expériences, en utilisant principalement des miroirs, des lames à retard, et des cubes séparateurs de polarisation pour séparer ou recombiner les faisceaux. Nous disposons de deux systèmes de photodiodes à avalanche (APD), qui se comportent comme des "détecteurs" de photons, et permettent d'appliquer les transformations non gaussiennes. Ces transformations sont destructives, puisque les photons sont absorbés.

Une petite partie du faisceau prélevée en amont sert à la caractérisation du laser : puissance, spectre, profil temporel, et éventuellement profil spatial. Une autre partie sert de sonde pour les

réglages et les alignements, et peut, selon les expériences, servir d'état cohérent (passant dans l'OPA, ou utilisé après). Enfin, environ 5% sert d'oscillateur local pour mesurer les états avec les détections homodynes. Selon les expériences, une seule ou deux détections peuvent être utilisées.

Nous disposons donc d'un dispositif complet pour produire des états non classiques, et les analyser.

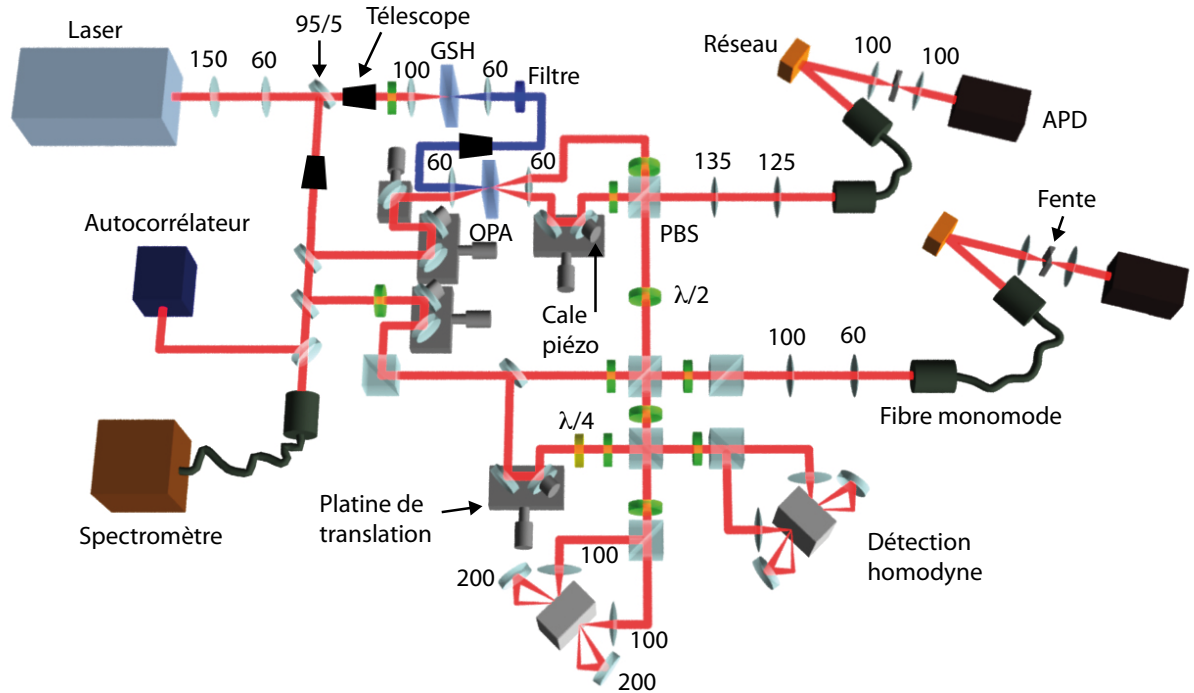


FIGURE 3.1 – Schéma du dispositif expérimental.

### 3.1.3 La source laser

#### Le laser impulsionnel

Nous utilisons un laser impulsionnel *Coherent Mira 900*, en configuration femtoseconde, pompé par un laser *Coherent Verdi V5*, à une puissance d'environ 4.2 W (Fig. 3.2).

Le laser fonctionne normalement à une cadence de 76 MHz. Pour augmenter la puissance des impulsions, la cavité est étendue avec un cavity dumper *Pulse Switch* de la société *APE GmbH*, qui permet aux impulsions de faire plusieurs passages avant d'être extraites avec une cellule de Bragg. Un boîtier externe permet de contrôler l'extraction, en jouant principalement sur la puissance, la durée, et la phase de l'onde RF venant moduler l'indice de réfraction de la cellule de Bragg par effet acousto-optique.

L'impulsion se forme en verrouillant les différents modes longitudinaux de la cavité par effet Kerr : lorsque tous les modes sont en phase, l'intensité crête est plus importante et le passage dans le cristal de titane-saphir induit une auto-focalisation du faisceau. En plaçant une fente de largeur réglable dans la cavité, on favorise ainsi la formation d'une impulsion, qui est ensuite ré-amplifiée à chaque passage, et davantage focalisée. Un mécanisme de démarrage permet d'initier automatiquement ce verrouillage de modes : en changeant l'orientation d'une petite plaque de verre placée dans la cavité, on change sa longueur effective, et on favorise l'apparition de plusieurs modes longitudinaux, qui permettent ensuite d'avoir suffisamment de puissance pour l'effet Kerr.



### Mira Optima 900-F / Mira 900-S

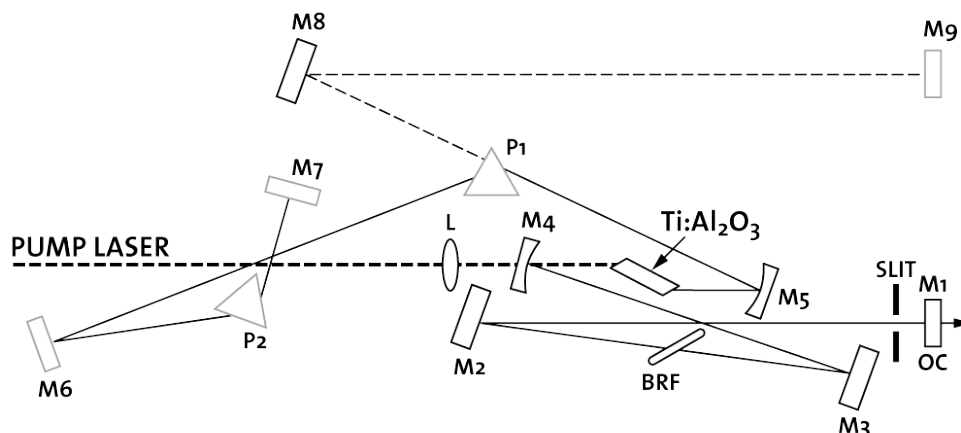


FIGURE 3.2 – Schéma de la cavité du Mira 900, tiré de sa documentation technique.

|                       |                                   |
|-----------------------|-----------------------------------|
| Longueur d'onde       | 850 nm (réglable)                 |
| Largeur spectrale     | entre 3 et 4 nm                   |
| Durée des impulsions  | environ 200 fs (auto-corrélation) |
| Puissance moyenne     | entre 33 et 38 mW                 |
| Cadence de répétition | 800 kHz (réglable)                |
| Puissance crête       | ~ 250 kW                          |
| Energie par impulsion | ~ 40 nJ                           |

TABLE 3.1 – Caractéristiques du laser MIRA 900.

Un système de filtrage permet de sélectionner la longueur d'onde centrale des impulsions, réglée pour nous à 850 nm. La dispersion est compensée par un train de prismes réglables.

Les principales caractéristiques du laser sont données dans la table 3.1.

### Spectre des impulsions

Le spectre est contrôlé à l'aide d'un spectromètre optique *Advantest Q8384* avec une résolution de l'ordre de 50 pm. La longueur d'onde centrale est réglée sur 850 nm en utilisant le système de filtrage du laser. Elle peut légèrement dévier de temps en temps, mais elle est en général très stable. La forme du spectre dépend en revanche beaucoup du réglage de l'onde RF pour l'extraction de l'impulsion. La figure 3.3 montre un profil typique que l'on peut obtenir. Lorsque la puissance de sortie du laser est trop importante, la forme gaussienne peut devenir plus aplatie sur sa partie supérieure avec un profil plus en dents de scie. Dans ce cas, il faut ajuster la phase de l'onde, et/ou jouer sur l'ouverture de la fente.

### Auto-corrélateur

Nous utilisons un auto-corrélateur *PulseCheck 15* de la société *APE GmbH* pour contrôler la durée des impulsions. Le principe de fonctionnement est schématisé sur la figure 3.5. Le faisceau est d'abord séparé à l'aide d'une lame semi-réfléchissante. Un des bras introduit un

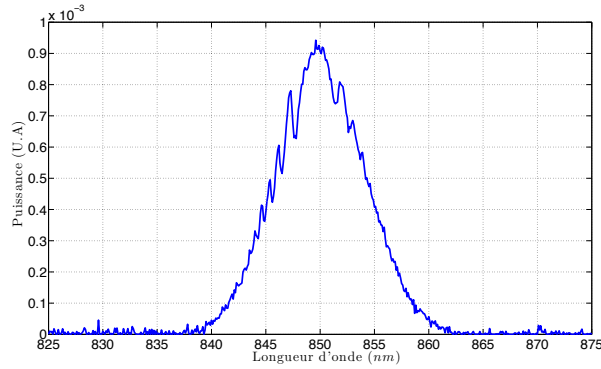


FIGURE 3.3 – Spectre du laser femtoseconde.

décalage temporel  $\tau$  et un décalage spatial. Les deux bras sont ensuite superposés de façon non-colinéaire dans un cristal non linéaire du second ordre, qui produit un faisceau doublé d'intensité  $I_{2\omega}(\tau) = \int dt I_{\omega}(t) I_{\omega}(t+\tau)$ , mesuré par un photo-multiplicateur. On reconstruit le signal en balayant le décalage  $\tau$ .

L'auto-corrélateur ne donne pas directement accès au profil temporel, mais permet quand même d'obtenir une information sur la durée de l'impulsion. En faisant l'hypothèse d'un profil en sécante hyperbolique  $\text{sech}^2 = \cosh^{-2}$ , on montre que la largeur à mi-hauteur mesurée par l'auto-corrélateur  $\Delta\tau_{\text{FWHM,AC}}$  et celle de l'impulsion  $\Delta\tau_{\text{FWHM}}$  sont reliées par  $\Delta\tau_{\text{FWHM}} = \Delta\tau_{\text{FWHM,AC}}/1.543$ .

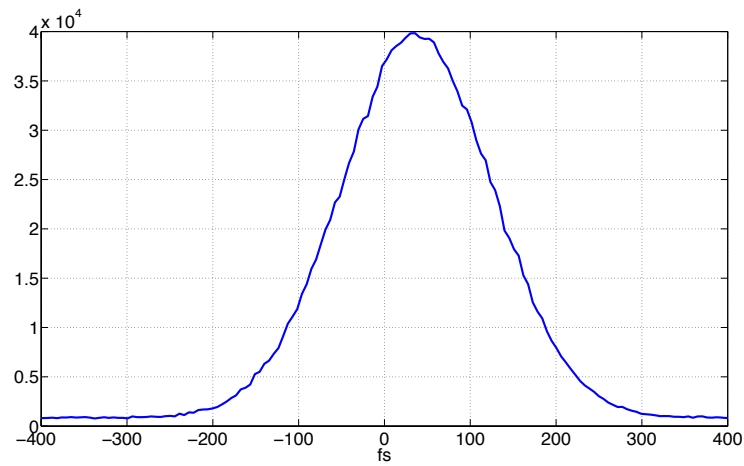
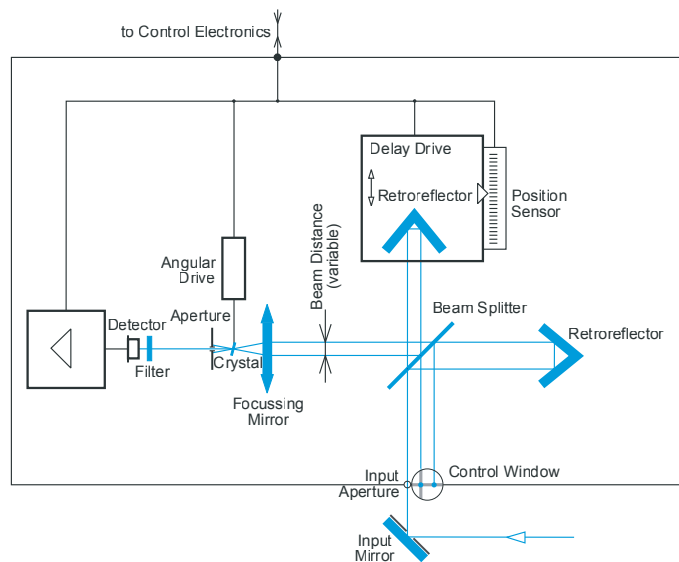
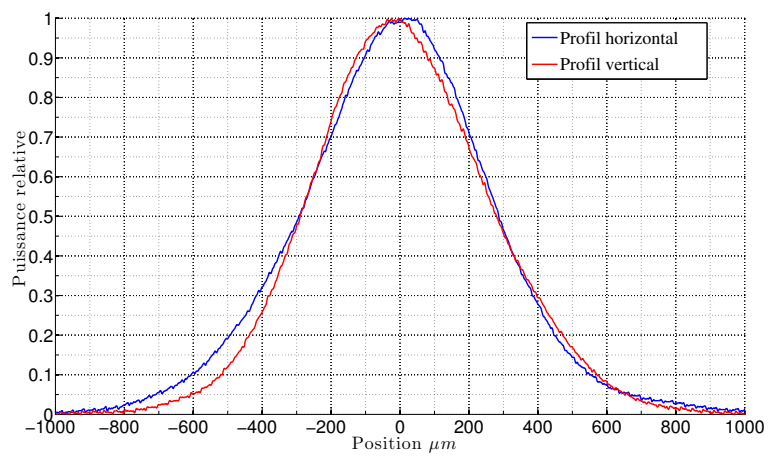
L'auto-corrélateur permet également de détecter d'éventuelles doubles impulsions, résultant d'un mauvais réglage du laser. Pour cela, on peut balayer manuellement le délai  $\tau$  sur une plage de 100 ps.

### Profil spatial

Le profil spatial est contrôlé à l'aide d'un profilomètre *BP109-UV* de la société *Thorlabs*, qui donne le profil selon deux axes orthogonaux. Un profil typique est montré sur la figure 3.6. Une dégradation peut être due à un mauvais réglage de la cellule de Bragg. En particulier, lorsque qu'une double bosse apparaît, c'est en général parce que les impulsions ne sont pas bien superposées dans la cavité du laser. Il peut aussi arriver que le faisceau soit légèrement coupé car trop proche du bord d'un miroir. On peut alors jouer sur la hauteur et l'angle du cristal pour améliorer le profil, voire sur les miroirs de fin de cavité.

Un mauvais alignement sur les optiques peut aussi être la cause d'un mauvais profil spatial. Il faut en particulier contrôler la sortie de chaque télescope utilisé pour modifier la taille du faisceau.

On peut aussi contrôler le faisceau à l'œil avec un viseur infrarouge, lorsque sa puissance n'est pas suffisante pour utiliser le profilomètre. C'est le cas par exemple pour le faisceau sonde servant à l'alignement.

FIGURE 3.4 – *Signal de l'auto-corrélateur.*FIGURE 3.5 – *Schéma du fonctionnement de l'auto-corrélateur.*FIGURE 3.6 – *Profil spatial des impulsions.*

## 3.2 Transformations unitaires : optique linéaire

Nous présentons ici les principales transformations unitaires utilisées dans notre expérience. Certaines ont déjà été abordées dans le chapitre décrivant les outils théoriques. Nous donnons ici une description plus expérimentale.

### 3.2.1 Déphasage

Un déphasage est obtenu très simplement en modifiant la longueur parcourue par le faisceau. Plusieurs platines de translation permettent d'introduire des déphasages dus à plusieurs centimètres, ce qui nous donne une plage suffisante pour pouvoir réussir à superposer temporellement les différents faisceaux utilisés dans notre dispositif. Pour une longueur  $L$ , le déphasage introduit est  $\theta = \frac{2\pi}{\lambda}L$ , et les quadratures se transforment selon l'opérateur  $\hat{U}(\theta)$  défini en (2.90) :

$$\hat{X}' = \hat{X} \cos \theta + \hat{P} \sin \theta \quad (3.1a)$$

$$\hat{P}' = -\hat{X} \sin \theta + \hat{P} \cos \theta \quad (3.1b)$$

Plusieurs cales piézo-électriques permettent de balayer la phase du faisceau, pour les réglages des interférences, des cristaux non linéaires, ou pour les mesures homodynes.

### 3.2.2 lame séparatrice

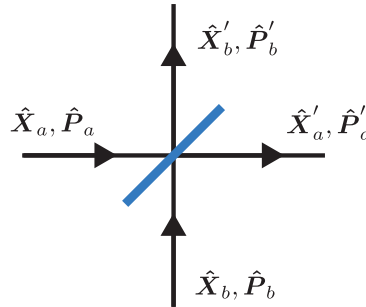


FIGURE 3.7 – *Lame séparatrice.*

Elle est constituée d'une lame de verre traitée pour réfléchir une fraction  $R$  de l'intensité du faisceau. Son action est décrite par l'opérateur  $\hat{U}_{BS}(\theta) = \exp[\theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger)]$  (2.93), avec  $\theta = \arccos \sqrt{T}$ , qui transforme les quadratures selon :

$$\hat{Q}'_a = \hat{Q}_a \sqrt{T} + \hat{Q}_b \sqrt{R} \quad (3.2a)$$

$$\hat{Q}'_b = -\hat{Q}_a \sqrt{R} + \hat{Q}_b \sqrt{T} \quad (3.2b)$$

avec  $\hat{Q} = \hat{X}$  ou  $\hat{P}$ . Les quadratures  $\hat{X}$  et  $\hat{P}$  ne sont pas couplées entre elles. En l'absence de pertes, on a  $T+R=1$ . En pratique, les pertes sont faibles ( $< 0.25\%$ ), et les lames sont traitées afin de minimiser la dispersion des impulsions.

Lorsque qu'un des modes d'entrée de la séparatrice est vide et que l'on oublie la partie réfléchie du faisceau, on modélise des pertes pour le faisceau transmis. Cette modélisation est

extrêmement générale, et peut être utilisée chaque fois que des pertes affectent un état ou une mesure. Pour un signal décrit par un opérateur  $\hat{a}$ , les pertes se traduisent donc par

$$\hat{a}' = \sqrt{T}\hat{a} + \sqrt{1-T}\hat{a}_0, \quad (3.3)$$

où  $\hat{a}_0$  est un mode vide entrant dans la deuxième voie de la séparatrice.

Pour justifier cette modélisation des pertes, l'approche suivie par la référence [Leonhardt97] est particulièrement simple et élégante. Il suffit de postuler qu'un état cohérent  $|\alpha\rangle$  est transformé en un état cohérent  $|\sqrt{T}\alpha\rangle$  comme on peut s'y attendre classiquement. La linéarité des pertes et la décomposition d'un état quelconque avec la fonction  $P$  suffisent ensuite à montrer l'analogie avec une lame séparatrice.

### 3.2.3 Lames demi-onde $\lambda/2$ et quart-d'onde $\lambda/4$

Ce sont des lames biréfringentes, qui déphasent les composantes de polarisation selon l'axe rapide et l'axe lent. Pour un angle  $\phi$  entre la verticale et l'axe lent, la lame demi-onde ( $\lambda/2$ ) ajoute une phase de  $\pi$ , et réalise la transformation

$$\hat{Q}'_H = \hat{Q}_H \cos 2\phi + \hat{Q}_V \sin 2\phi \quad (3.4a)$$

$$\hat{Q}'_V = -\hat{Q}_H \sin 2\phi + \hat{Q}_V \cos 2\phi \quad (3.4b)$$

avec  $\hat{Q} = \hat{X}$  ou  $\hat{P}$ . Cette transformation est très similaire à celle réalisée par une lame séparatrice (3.2). La seule différence est que les modes sont séparés en polarisation, au lieu de l'être spatialement.

Une lame quart-d'onde ( $\lambda/4$ ) ajoute une phase de  $\pi/2$  entre l'axe rapide et l'axe lent, transformant une polarisation linéaire en une polarisation elliptique, et vice versa. La transformation des quadratures est

$$\begin{pmatrix} \hat{X}'_H \\ \hat{P}'_H \\ \hat{X}'_V \\ \hat{P}'_V \end{pmatrix} = \begin{pmatrix} 1 & \cos 2\phi & 0 & \sin 2\phi \\ -\cos 2\phi & 1 & -\sin 2\phi & 0 \\ 0 & \sin 2\phi & 1 & -\cos 2\phi \\ -\sin 2\phi & 0 & \cos 2\phi & 1 \end{pmatrix} \begin{pmatrix} \hat{X}_H \\ \hat{P}_H \\ \hat{X}_V \\ \hat{P}_V \end{pmatrix} \quad (3.5)$$

### 3.2.4 Cube séparateur de polarisation (PBS)

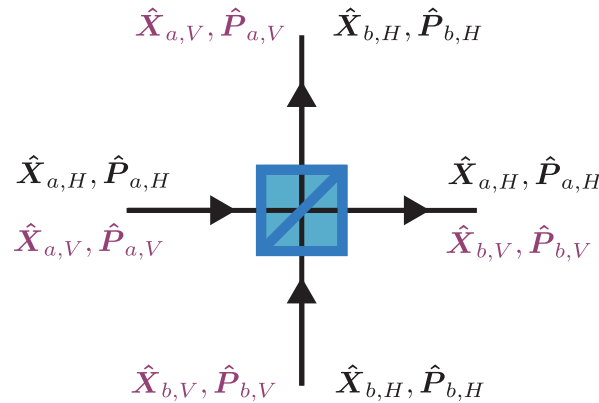


FIGURE 3.8 – Cube séparateur de polarisation.

Il s'agit d'un cube formé de deux prismes accolés, qui sont séparés par un traitement diélectrique. La polarisation verticale est réfléchiée, alors que la polarisation horizontale, qui arrive sous une incidence proche de Brewster, est transmise. Le principe même de son fonctionnement empêche d'avoir à la fois une bonne transmission et un bon taux de réjection des polarisations. La transmission est typiquement de 95-98%, et les taux de réjection de 1/100 à 1/1000.

Utilisé avec une lame  $\lambda/2$ , l'ensemble se comporte comme une lame séparatrice, dont on peut régler la transmission en tournant la  $\lambda/2$ . Les faisceaux peuvent également suivre le même chemin optique avant d'interférer, ce qui leur permet de subir les mêmes fluctuations de phase. Pour les calculs, un ensemble  $\lambda/2 + \text{PBS}$  sera traité comme une lame séparatrice, et décrit par un opérateur similaire à (2.93).

### 3.3 Transformations unitaires : optique non linéaire

Afin de produire nos états gaussiens comprimés, nous avons besoin d'un cristal non linéaire du deuxième ordre. Ce cristal est lui-même pompé par un faisceau à 425 nm, obtenu par doublement de fréquence du faisceau laser infrarouge à l'aide d'un premier cristal non linéaire d'ordre deux. A chaque fois, ces cristaux servent de médiateurs pour échanger de l'énergie entre le faisceau pompe, et deux modes, appelés *signal* et *idler* (complémentaire). Pour assurer une bonne conversion, il faut pouvoir satisfaire deux conditions élémentaires : la conservation de l'énergie  $\omega_p = \omega_s + \omega_i$ , et l'accord de phase  $\mathbf{k}_p = \mathbf{k}_s + \mathbf{k}_i$ . Pour cela, on utilise des cristaux biréfringents, pour lesquels on modifie l'orientation et / ou la température. En fonction des conditions permettant l'accord de phase, on favorisera ainsi un fonctionnement dégénéré, où les modes signal et idler sont identiques, ou non dégénéré, où ils sont séparés spatialement.

Nous utilisons deux cristaux de niobate de potassium  $\text{KNbO}_3$  de 100  $\mu\text{m}$  d'épaisseur, de la société *FFEGmbH*. Ils présentent des caractéristiques intéressantes pour nos applications (table 3.2), avec un fort coefficient non linéaire, et une absence totale de *walk off* (accord de phase non critique). Ces cristaux sont biréfringents biaxiaux et traités antireflets pour nos longueurs d'ondes. Ils sont coupés perpendiculairement à l'axe  $a$ . Leur accord de phase est colinéaire de type I (ooe) : les faisceaux signal et idler sont polarisés selon l'axe ordinaire  $b$ , et le faisceau pompe selon l'axe extraordinaire  $c$ .

L'accord de phase est obtenu en réglant la température du cristal (Fig. 3.9 (b)). Les cristaux sont placés dans des enceintes à vide pour éviter la condensation (Fig. 3.9 (a)), et sont refroidis par effet Peltier.

|                          |                                        |
|--------------------------|----------------------------------------|
| Coefficient non linéaire | $d_{\text{eff}} = 12 \text{ pm/V}$     |
| Indice de réfraction     | $n_{b,\omega} = n_{c,2\omega} = 2.281$ |
| GVD                      | $0.38 \text{ fs}^2/\mu\text{m}$        |
| GVM                      | $1.2 \text{ ps/mm}$                    |
| Seuil de dommage         | $2 \text{ GW/cm}^2$                    |

TABLE 3.2 – Caractéristiques des cristaux de  $\text{KNbO}_3$ . La GVD (Group Velocity Dispersion) décrit l'étalement de l'impulsion dû à la dispersion ; la GVM (Group Velocity Mismatch) rend compte de la différence des vitesses de propagation des ondes pompe et signal/idler de par leurs fréquences différentes.

Les tailles des faisceaux peuvent être ajustées à l'aide de plusieurs télescopes de la société *Thorlabs*, modèle BE02-05-A pour le bleu et BE02-05-B pour le rouge (zoom  $\times 2$  à  $\times 5$ ).

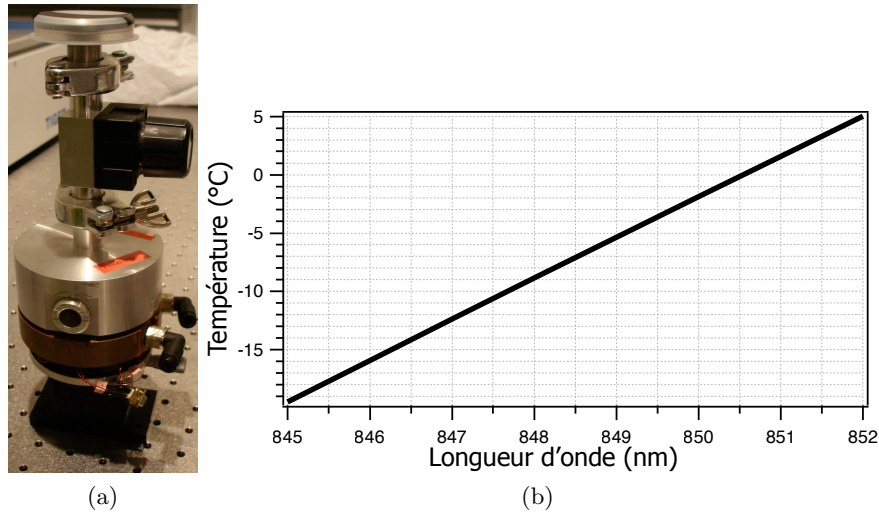


FIGURE 3.9 – (a) Enceinte à vide contenant le cristal de  $\text{KNbO}_3$  ; (b) Température d'accord de phase en fonction de la longueur d'onde [Zysset92].

Une étude détaillée des montages présentés ci dessous peut être trouvée dans les thèses de Jérôme Wenger [Wenger04a] et Alexei Ourjountsev [Ourjountsev07a].

### 3.3.1 Génération de seconde harmonique (GSH)

La génération de seconde harmonique permet de créer le faisceau bleu à 425 nm. Le cristal est dans une configuration d'accord de phase dégénéré, afin de convertir deux photons infrarouge en un photon bleu. Théoriquement, l'efficacité de cette conversion devrait atteindre 90% [Wenger04a]. En pratique, l'efficacité maximale observée a été de 32% pour un waist de 16  $\mu\text{m}$ . Cette limitation est due à plusieurs effets parasites, tel que l'absorption à deux photons, qui diminue l'efficacité de conversion lorsque la focalisation dans le cristal est trop importante (Fig. 3.10) [Wenger04a, Ourjountsev07a, Ferreyrol11a]. Les effets photo-refractifs [Reeves91] sont également responsables d'une dégradation du mode spatial due à une modification locale de l'indice de réfraction. Pour limiter ces effets, il peut être nécessaire de défocaliser le faisceau, et de translater transversalement le cristal afin d'utiliser une autre zone.

Après avoir effectué plusieurs essais, un waist  $\omega_0=20 \mu\text{m}$  dans le cristal, correspondant à une efficacité de conversion d'environ 20%, donne un bon compromis entre la puissance et la qualité du profil spatial obtenues pour le faisceau bleu. Afin d'éliminer les photons infrarouges résiduels qui pourraient aller vers l'OPA, on utilise une série de filtres laissant passer le bleu avec une efficacité de 80%, et bloquant l'infrarouge avec une transmission de  $10^{-16}$ . Au final, la puissance de bleu disponible se situe aux alentours de 4 mW.

### 3.3.2 Amplification paramétrique optique (OPA)

L'amplification paramétrique optique correspond à la conversion d'un photon pompe en deux photons infrarouges. Comme pour la GSH, en fonction des conditions d'accord de phase, on privilégie soit un régime dégénéré où les photons sont émis par paires dans le même mode, soit un régime non dégénéré où les photons sont émis par paires dans deux modes différents.

Lorsque le mode signal ou idler n'est pas vide, l'OPA produit une amplification dépendante de la phase en configuration dégénérée, et indépendante de la phase en configuration non dégénérée.

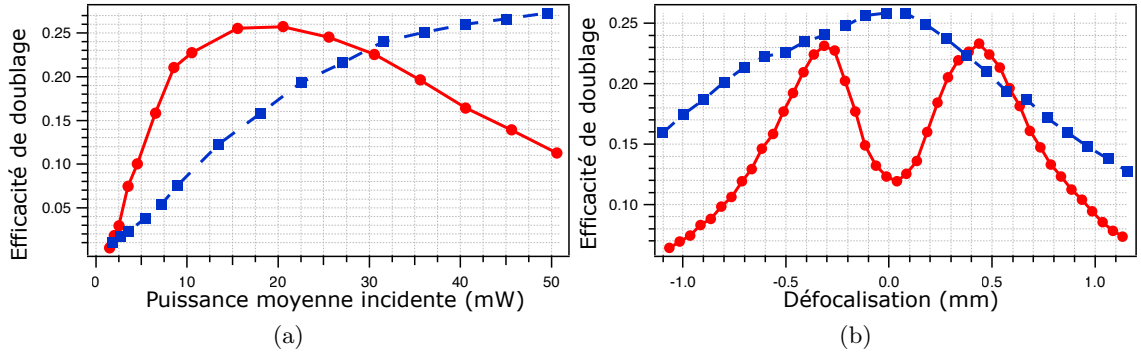


FIGURE 3.10 – (a) Efficacité de conversion en fonction de la puissance moyenne incidente ; (b) Efficacité de conversion en fonction de la défocalisation. Les ronds et les carrés correspondent respectivement à des waists  $8\mu\text{m}$  et  $16\mu\text{m}$  dans le cristal. Données tirées de [Wenger04a].

En sortie du cristal, le faisceau amplifié a la même fréquence  $\omega_s$  que le signal, et l'idler a une fréquence  $\omega_i = \omega_p - \omega_s$ . L'utilisation d'un faisceau sonde classique permet de régler le dispositif et l'alignement avec la pompe.

Le cristal est placé dans un montage identique à celui du GSH. Le rapport des waists entre la pompe et le signal est un compromis entre la qualité du mode de sortie, et le gain de l'amplification. Si le rapport est trop grand, le gain sera trop faible. En revanche, s'il est trop petit, les bords du faisceau ne voient pas la même intensité de pompe, et sont donc amplifiés différemment (phénomène de *gain induced diffraction* [La Porta91]). Un rapport empirique d'environ  $\omega_p = \sqrt{2}\omega_s$  donne un bon compromis [Wenger04c]. Pour nos expériences, nous serons dans des conditions similaires, avec un waist de pompe  $\omega_p = 15\mu\text{m}$  et un waist signal  $\omega_s = 12\mu\text{m}$ .

On peut raisonnablement faire l'hypothèse d'une pompe classique et non déplétée, et l'hamiltonien d'interaction s'écrit dans ce cas, avec un traitement multimode [Hong85, Ghosh86, Ou97] :

$$\hat{\mathcal{W}} = -i\hbar \int d^3\mathbf{k}_1 d^3\mathbf{k}_2 \Phi(\mathbf{k}_1, \mathbf{k}_2) \hat{\mathbf{a}}^\dagger(\mathbf{k}_1) \hat{\mathbf{a}}^\dagger(\mathbf{k}_2) + \text{h.c} \quad (3.6)$$

Le terme  $\Phi(\mathbf{k}_1, \mathbf{k}_1)$  tient compte de tous les paramètres de l'interaction : puissance de la pompe, caractéristiques du cristal, accord de phase, ... Un traitement multimode complet permet une modélisation très fine des imperfections expérimentales, mais devient vite fort complexe. La plupart du temps cependant on pourra avec une bonne approximation faire un traitement "monomode", en couplant seulement deux modes impulsionnels  $\hat{\mathbf{a}}_s$  et  $\hat{\mathbf{a}}_i$ . On a alors l'hamiltonien

$$\hat{\mathcal{W}}_{\text{ND}} = i\hbar\tilde{\zeta}_{\text{ND}} \left( \hat{\mathbf{a}}_s \hat{\mathbf{a}}_i - \hat{\mathbf{a}}_s^\dagger \hat{\mathbf{a}}_i^\dagger \right) \quad (3.7)$$

pour la configuration non dégénérée, et

$$\hat{\mathcal{W}}_{\text{D}} = i\hbar\tilde{\zeta}_{\text{D}} \left( \hat{\mathbf{a}}_s^2 - \hat{\mathbf{a}}_s^{\dagger 2} \right) \quad (3.8)$$

pour la configuration dégénérée. Pour une longueur d'interaction  $L$ , en posant  $r = \tilde{\zeta}_{\text{ND}}L$  ou  $r/2 = \tilde{\zeta}_{\text{D}}L$ , les opérateurs d'évolution en représentation d'interaction prennent respectivement la forme des opérateurs de squeezing monomode (2.136), ou bimode (2.149). On rappelle que pour un squeezing monomode, les quadratures se transforment en

$$\begin{pmatrix} \hat{X}'_s \\ \hat{P}'_s \end{pmatrix} = \begin{pmatrix} e^{-r} & 0 \\ 0 & e^{+r} \end{pmatrix} \begin{pmatrix} \hat{X}_s \\ \hat{P}_s \end{pmatrix}, \quad (3.9)$$



et pour un squeezing bimode, elles se transforment en

$$\begin{pmatrix} \hat{X}'_s - \hat{X}'_i \\ \hat{P}'_s - \hat{P}'_i \\ \hat{X}'_s + \hat{X}'_i \\ \hat{P}'_s + \hat{P}'_i \end{pmatrix} = \begin{pmatrix} e^{+r} & 0 & 0 & 0 \\ 0 & e^{-r} & 0 & 0 \\ 0 & 0 & e^{-r} & 0 \\ 0 & 0 & 0 & e^{+r} \end{pmatrix} \begin{pmatrix} \hat{X}_s - \hat{X}_i \\ \hat{P}_s - \hat{P}_i \\ \hat{X}_s + \hat{X}_i \\ \hat{P}_s + \hat{P}_i \end{pmatrix} \quad (3.10)$$

Lorsque les modes signal et idler sont vides, les états produits par fluorescence paramétrique correspondent donc à des vides comprimés dans le cas dégénéré, et à des états EPR dans le cas non dégénéré. L'émission est ici fortement multimode, puisqu'il n'y a pas de fréquence particulière du mode signal stimulant l'émission. Nous verrons que l'on peut malgré tout se ramener à un traitement monomode modélisant efficacement les imperfections.

### Modélisation des imperfections

Plusieurs imperfections expérimentales font que l'état produit par l'OPA ne correspond pas tout fait à (3.9) ou (3.10). Une des contributions majeures provient d'un mauvais recouvrement de la pompe et du signal, produisant des photons "solitaires" au lieu d'être par paires, qui se comportent comme du bruit. Une autre contribution importante vient également du caractère fortement multimode de la fluorescence paramétrique, qui fait que la compression n'est pas définie pour un seul mode, mais pour tout une série de modes différents du mode sélectionné par la détection homodyne.

Un modèle empirique donnant de très bons résultats consiste à modéliser l'OPA réel par un OPA parfait "monomode" de paramètre  $r$ , suivi d'une amplification parfaite de gain  $h = \cosh^2 \gamma r$  indépendante de la phase. Les modèles associés aux configurations dégénérées et non dégénérées sont schématisés sur les figures 3.11 et 3.12. Nous verrons au chapitre 8 qu'une amplification indépendante de la phase ajoute nécessairement du bruit, même si elle est parfaite. Dans une certaine mesure, elle permet donc de modéliser les photons solitaires dont nous avons parlé.

On peut donc se ramener à un traitement monomode en modélisant donc l'état produit par l'OPA parfait en configuration dégénérée par un vide comprimé  $\exp\left[\frac{r}{2}\hat{\mathbf{a}}^2 - \frac{r}{2}\hat{\mathbf{a}}^{\dagger 2}\right] |0\rangle$ , et l'état produit en configuration non dégénérée par un état EPR  $\exp\left[r\hat{\mathbf{a}}\hat{\mathbf{b}} - r\hat{\mathbf{a}}^{\dagger}\hat{\mathbf{b}}^{\dagger}\right] |0\rangle$ . Le mode  $\hat{\mathbf{a}}$  n'est pas une propriété de l'OPA en soit, mais dépend du mode homodyne. L'amplification parasite se fait ensuite en appliquant un OPA non dégénéré  $\exp[\gamma r(\hat{\mathbf{a}}\hat{\mathbf{c}} - \hat{\mathbf{a}}^{\dagger}\hat{\mathbf{c}}^{\dagger})]$  sur chacun des modes de sortie (un ou deux) de l'OPA parfait, et en traçant sur le mode fictif  $\hat{\mathbf{c}}$ .

Nous présentons un modèle simplifié de traitement multimode en annexe B.1, qui permet de comprendre l'origine de  $r$ , et  $\gamma$ . La référence [Tualle-Broui09] fournit un modèle multimode plus complet, permettant de justifier rigoureusement la validité de ce modèle empirique.

### Configuration dégénérée

La configuration dégénérée correspond aux mêmes conditions d'accord de phase que pour la GSH. On peut accéder aux paramètres de l'amplification  $r$  et  $\gamma$  en utilisant une sonde classique, qui permet d'optimiser les réglages et l'alignement avec la pompe. En balayant la phase de la sonde avec une cale piézoélectrique modulée à environ 80 Hz, on mesure le gain maximal  $g_{\min}$  et minimal  $g_{\max}$  de l'amplification à l'aide d'un oscilloscope. En l'absence d'amplification parasite, on devrait avoir  $g_{\min} = e^{-2r} = s$  et  $g_{\max} = e^{2r} = 1/s$ , et donc un produit  $g_{\min}g_{\max} = 1$ . En pratique,

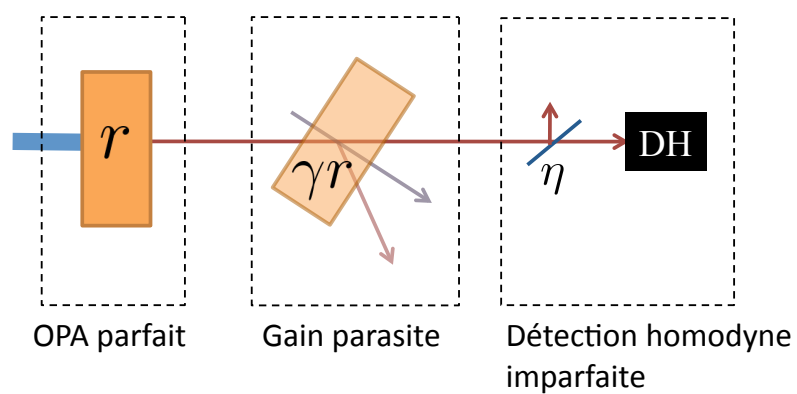


FIGURE 3.11 – Modélisation de l'OPA imparfait en configuration dégénérée.

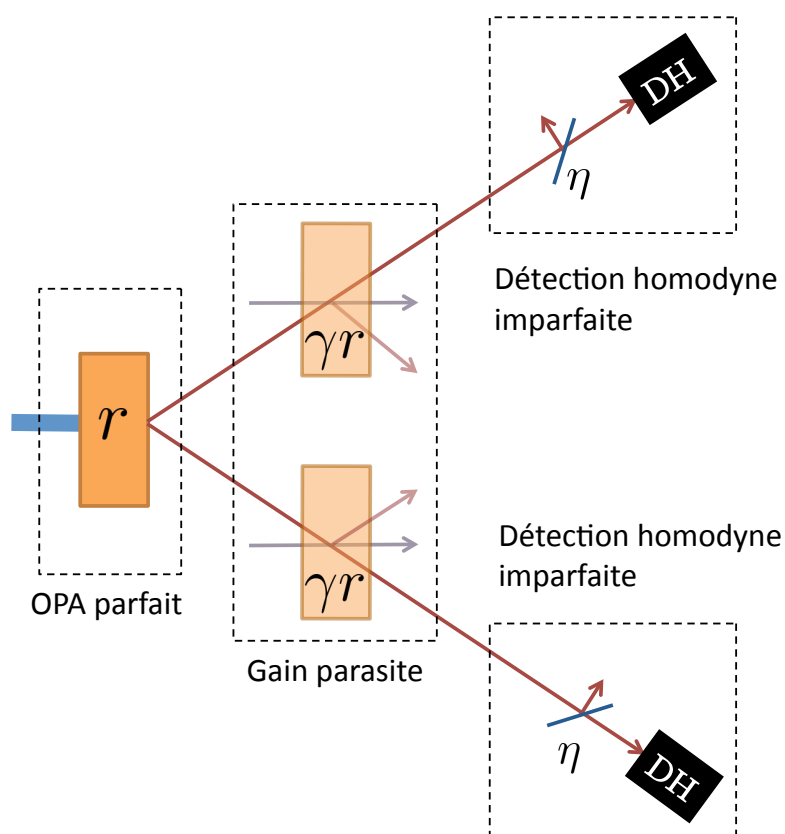


FIGURE 3.12 – Modélisation de l'OPA imparfait en configuration non dégénérée.

on a plutôt  $g_{\min}g_{\max} \simeq 1.03$ . En utilisant le modèle de la figure 3.11, ces gains valent :

$$g_{\min} = e^{-2r} \cosh^2 \gamma r = hs \quad (3.11a)$$

$$g_{\max} = e^{+2r} \cosh^2 \gamma r = \frac{h}{s} \quad (3.11b)$$

En inversant ces équations, on obtient  $r$  et  $\gamma$  :

$$r = \frac{1}{4} \ln \frac{g_{\max}}{g_{\min}} \quad (3.12a)$$

$$\gamma = \frac{1}{r} \operatorname{arcosh} \left( (g_{\min}g_{\max})^{\frac{1}{4}} \right) \quad (3.12b)$$

Ces gains  $g_{\min}$  et  $g_{\max}$  sont montrés sur la figure 3.13 (a) en fonction de la puissance de pompe. Ils donnent ici un paramètre  $\gamma=0.6$  relativement constant, traduisant le fait que le recouvrement entre la pompe et le signal dépend peu de la puissance de pompe. La figure 3.13 (b) montre le paramètre  $r$  en fonction de la racine carrée de la puissance de pompe. Le fait qu'il passe en dessous de la droite théorique est la signature d'effets parasites, principalement une absorption multiphotonique [Ourjountsev07a].

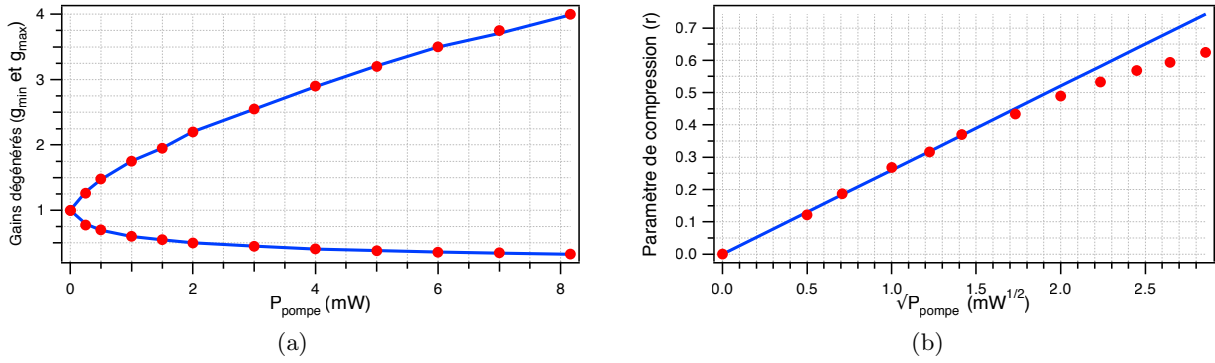


FIGURE 3.13 – (a) Gains  $g_{\min}$  et  $g_{\max}$  en fonction de la puissance de pompe [Ourjountsev07a]. La pente de  $0.26 \text{ mW}^{-1/2}$  est obtenue à partir des 6 premiers points ; (b) Paramètre  $r$  en fonction de la racine carrée de la puissance de pompe [Ourjountsev07a].

### Configuration non dégénérée

La configuration non dégénérée s'obtient en modifiant l'accord de phase, en abaissant la température d'environ  $6^\circ\text{C}$  et en translatant le faisceau pour avoir un angle d'environ  $5^\circ$  avec la pompe. En pratique, on part de la configuration dégénérée, puis on translate le faisceau jusqu'à ce que les oscillations de l'amplification dépendante de la phase disparaissent.

Le gain paramétrique vaut  $G=gh=\cosh^2 r \cosh^2 \gamma r$ . Sa variation est linéaire avec la puissance de pompe :  $G \simeq 1 + r^2(1 + \gamma^2) = 1 + \alpha^2 P_{\text{pompe}}$ , avec  $\alpha = 0.21 \text{ mW}^{-1/2}$ , comme le montre la figure 3.14.

Théoriquement, les variances des quadratures  $\hat{X}_0 + \hat{X}_1$  et  $\hat{P}_0 - \hat{P}_1$  sont comprimées avec un gain  $g_{\min}=s$ , alors que les variances des quadratures  $\hat{X}_0 - \hat{X}_1$  et  $\hat{P}_0 + \hat{P}_1$  sont anticomprimées avec un gain  $g_{\max}=1/s$ . En pratique, on montre que l'amplification parasite  $h$  modifie ces gains selon [Ourjountsev07a] :

$$g_{\min} = hs + h - 1 \quad (3.13a)$$

$$g_{\max} = \frac{h}{s} + h - 1 \quad (3.13b)$$

Enfin, la table 3.3 montre des valeurs typiques obtenues pour les caractéristiques de l'OPA.

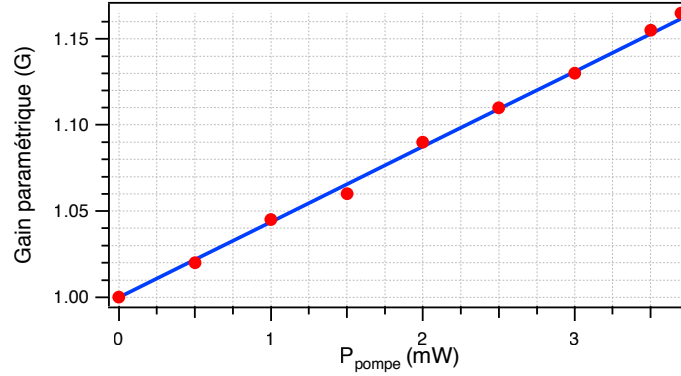


FIGURE 3.14 – Gain paramétrique en fonction de la puissance de pompe [Ourjountsev07a].

| Paramètre                        | Valeur typique                  |
|----------------------------------|---------------------------------|
| Puissance de pompe               | $\simeq 4$ mW                   |
| Gain paramétrique (non dégénéré) | $G \simeq 1.05 - 1.20$          |
| Gain minimum (dégénéré)          | $g_{\min} \simeq 0.51 - 0.59$   |
| Gain maximum (dégénéré)          | $g_{\max} \simeq 1.96 - 2.07$   |
| Produit des gains                | $g_{\min} g_{\max} \simeq 1.03$ |
| Compression                      | $r \simeq 0.3 - 0.35$           |
| Gain parasite                    | $\gamma = 0.5 - 0.9$            |

TABLE 3.3 – Valeurs typiques des caractéristiques de l'OPA.

### 3.4 Détection et mesures projectives

Nous présentons maintenant les mesures que nous pouvons effectuer sur nos états quantiques. Elles sont de deux natures : les mesures non gaussiennes sont utilisées pour préparer des états quantiques à partir des états de l'OPA. Nous disposons pour cela de deux APD, servant de détecteurs de photons. Les mesures gaussiennes servent ensuite à caractériser les états produits. Elles sont réalisées avec une ou deux détecteurs homodynes.

#### 3.4.1 Détection homodyne

##### Principe

A l'heure actuelle, il n'existe pas de technologie permettant de mesurer l'amplitude d'un champ électromagnétique aux fréquences optiques. En revanche, des méthodes interférométriques permettent de mesurer l'amplitude d'une enveloppe évoluant moins rapidement. Pour nos états quantiques, cette enveloppe correspond aux quadratures, et on peut les mesurer en utilisant une *détection homodyne*.

Elle permet de mesurer les quadratures du champ grâce à une interférence avec un faisceau intense de référence, appelé oscillateur local. Le signal à mesurer et l'oscillateur local interfèrent sur une lame séparatrice 50/50 mélangeant leurs quadratures, qui sont mesurées à l'aide de

photodiodes. En mesurant les deux intensités dans les deux voies de sortie de la séparatrice, on obtient deux photocourants

$$\hat{i}_{\pm} \propto \frac{1}{2}I_{\text{OL}} + \frac{1}{2}\hat{\mathbf{a}}_s^{\dagger}\hat{\mathbf{a}}_s \pm \frac{1}{2}\sqrt{\frac{I_{\text{OL}}}{N_0}}\hat{\mathbf{X}}_{s|\text{OL}}, \quad (3.14)$$

où  $\hat{\mathbf{a}}_s$  correspond au mode du signal. La différence de ces photocourants est donc proportionnelle à la quadrature  $\hat{\mathbf{X}}_{s|\text{OL}}$  du signal dans le mode de la détection homodyne, définie par la phase de l'oscillateur local. En balayant celle-ci, on peut mesurer n'importe quelle quadrature  $\hat{\mathbf{X}}_{s,\theta}$  du signal.

Ainsi, les quadratures peuvent être mesurées avec des composants optiques standards, et de simples photodiodes. C'est un intérêt majeur par rapport à la mesure du nombre de photons, qui a beaucoup contribué au développement des variables continues. Contrairement à un fonctionnement en régime fréquentiel [Yuen83, Slusher87], on ne peut pas s'affranchir des bruits techniques à basses fréquences en filtrant les mesures. Notre fonctionnement impulsionnel requiert de travailler dans le domaine temporel avec des détecteurs larges bandes, et bien sûr de minimiser le bruit et les pertes.

### Modélisation théorique

Commençons par donner quelques justifications permettant d'aboutir à l'équation (3.14) en régime impulsionnel. Nous mettrons ainsi en lumière une propriété fondamentale d'une détection homodyne : elle n'est sensible qu'aux composantes du signal dans le même mode que l'oscillateur local, en agissant comme un filtre.

L'oscillateur local est d'abord mélangé avec le signal sur une lame séparatrice 50/50, ce qui transforme les opérateurs de leurs modes selon

$$\begin{cases} \hat{\mathbf{a}}_s(\mathbf{k}) \\ \hat{\mathbf{a}}_{\text{OL}}(\mathbf{k}) \end{cases} \rightarrow \begin{cases} \hat{\mathbf{a}}_+(\mathbf{k}) = \frac{1}{\sqrt{2}}(\hat{\mathbf{a}}_s(\mathbf{k}) + \hat{\mathbf{a}}_{\text{OL}}(\mathbf{k})) \\ \hat{\mathbf{a}}_-(\mathbf{k}) = \frac{1}{\sqrt{2}}(\hat{\mathbf{a}}_s(\mathbf{k}) - \hat{\mathbf{a}}_{\text{OL}}(\mathbf{k})) \end{cases} \quad (3.15)$$

On mesure ensuite l'intensité dans les deux voies de sortie. Sous certaines approximations, on montre en annexe C qu'un photocourant est proportionnel au nombre total de photons dans les modes excités. Les photocourants dans les deux voies de sorties sont donc donnés par

$$\hat{i}_{\pm} \propto \int d^3\mathbf{k} \hat{\mathbf{a}}_{\pm}^{\dagger}(\mathbf{k})\hat{\mathbf{a}}_{\pm}(\mathbf{k}), \quad (3.16)$$

dont la différence correspond au signal de la détection homodyne :

$$\hat{i}_{\text{DH}} = \hat{i}_+ - \hat{i}_- \quad (3.17a)$$

$$\propto \int d^3\mathbf{k} \hat{\mathbf{a}}_+^{\dagger}(\mathbf{k})\hat{\mathbf{a}}_+(\mathbf{k}) - \int d^3\mathbf{k} \hat{\mathbf{a}}_-^{\dagger}(\mathbf{k})\hat{\mathbf{a}}_-(\mathbf{k}) \quad (3.17b)$$

$$= \int d^3\mathbf{k} \left( \hat{\mathbf{a}}_s^{\dagger}(\mathbf{k})\hat{\mathbf{a}}_{\text{OL}}(\mathbf{k}) + \hat{\mathbf{a}}_{\text{OL}}^{\dagger}(\mathbf{k})\hat{\mathbf{a}}_s(\mathbf{k}) \right) \quad (3.17c)$$

**Un filtre pour le mode du signal** Puisque nous sommes en régime impulsionnel, on peut définir un opérateur destruction  $\hat{\mathbf{a}}_{\text{OL}}$  lié au mode de l'oscillateur local,

$$\hat{\mathbf{a}}_{\text{OL}} = \int d^3\mathbf{k} \alpha_{\text{OL}}^*(\mathbf{k})\hat{\mathbf{a}}_{\text{OL}}(\mathbf{k}), \quad (3.18)$$

où  $\alpha_{OL}$  correspond à l'enveloppe normalisée du mode. Si l'oscillateur local est suffisamment intense, d'intensité <sup>1</sup>  $I_{OL}$ , on peut remplacer  $\hat{\mathbf{a}}_{OL}(\mathbf{k})$  par sa valeur moyenne  $\langle \hat{\mathbf{a}}_{OL}(\mathbf{k}) \rangle = \sqrt{I_{OL}} \alpha_{OL}(\mathbf{k})$ . La différence des photocourants (3.17c) devient alors égale à

$$\hat{i}_{DH} = \sqrt{I_{OL}} \int d^3\mathbf{k} \left( \alpha_{OL}(\mathbf{k}) \hat{\mathbf{a}}_s^\dagger(\mathbf{k}) + \alpha_{OL}^*(\mathbf{k}) \hat{\mathbf{a}}_s(\mathbf{k}) \right). \quad (3.19)$$

C'est donc un opérateur proportionnel à l'opérateur de quadrature  $\hat{\mathbf{X}} = (\hat{\mathbf{d}}^\dagger + \hat{\mathbf{d}}) \sqrt{N_0}$  du signal dans le mode de l'oscillateur local défini, par

$$\hat{\mathbf{d}} = \int d^3\mathbf{k} \alpha_{OL}^*(\mathbf{k}) \hat{\mathbf{a}}_s(\mathbf{k}). \quad (3.20)$$

La détection homodyne agit donc comme un filtre qui ne sélectionne que la composante du signal qui est dans le mode spatiotemporel de l'oscillateur local. En déphasant l'oscillateur local de  $-\theta$ ,  $\langle \hat{\mathbf{a}}_{OL}(\mathbf{k}) \rangle$  est transformé en  $\sqrt{I_{OL}} \alpha_{OL}(\mathbf{k}) e^{+i\theta}$ , et la différence des photocourants devient proportionnelle à

$$\hat{i}_{DH} \propto \sqrt{\frac{I_{OL}}{N_0}} \hat{\mathbf{X}}_\theta. \quad (3.21)$$

Une détection homodyne permet donc de mesurer n'importe quelle quadrature  $\hat{\mathbf{X}}_\theta$ .

## Imperfections

**Pertes** Les pertes sont modélisées par une lame séparatrice de transmission  $\sqrt{\eta}$ . Au total, elles peuvent provenir de quatre contributions :

- Les pertes optiques ( $\eta_{opt} \simeq 0.87$ ), dues aux imperfections des traitements anti-reflets, aux pertes dans les cubes séparateurs de polarisation, ...
- L'efficacité quantique des photodiodes ( $\eta_{quant} \simeq 0.95$ ), correspondant à la probabilité qu'un photon incident crée un porteur de charge.
- L'adaptation des modes ou *mode matching* ( $\eta_{mod} \simeq 0.83$ ), due à un recouvrement imparfait entre le mode du signal et le mode de l'oscillateur local. Elle peut être mesurée avec les interférences de deux états cohérents de même amplitude : le contraste donne directement  $\sqrt{\eta_{mod}}$  [Grosshans01].

Voyons comment un recouvrement imparfait se traduit par des pertes. Supposons que le mode signal s'écrive

$$\hat{\mathbf{c}} = \int d^3\mathbf{k} \phi_s^*(\mathbf{k}) \hat{\mathbf{a}}_s(\mathbf{k}), \quad (3.22)$$

et utilisons une identité triviale afin de décomposer le mode de l'oscillateur local sur celui ci <sup>2</sup> :

$$\alpha_{OL} = \langle \phi_s | \alpha_{OL} \rangle \phi_s + \left( \alpha_{OL} - \langle \phi_s | \alpha_{OL} \rangle \phi_s \right) \quad (3.23)$$

Puisque c'est le mode oscillateur local qui est mesuré, c'est bien lui qui doit être décomposé sur le mode signal, et non l'inverse. On vérifie facilement que les deux vecteurs

1. En toute rigueur, on devrait plutôt parler de nombre de photons.

2. Le recouvrement entre les modes est noté  $\langle \phi_s | \alpha_{OL} \rangle$  par abus de notation. Il est défini par  $\int d^3\mathbf{k} \phi_s^*(\mathbf{k}) \alpha_{OL}(\mathbf{k})$ .

de la deuxième partie de (3.23) sont orthogonaux. Supposons maintenant que  $\langle \phi_s | \alpha_{OL} \rangle$  soit réel et positif, de manière à pouvoir poser

$$\sqrt{\eta_{\text{mod}}} = \langle \phi_s | \alpha_{OL} \rangle. \quad (3.24)$$

On peut écrire le deuxième membre de (3.23) en fonction d'une fonction  $\phi_s^\perp$  normalisée et orthogonale à  $\phi_s$ ,

$$\alpha_{OL} - \langle \phi_s | \alpha_{OL} \rangle \phi_s = \|\alpha_{OL} - \langle \phi_s | \alpha_{OL} \rangle \phi_s\| \phi_s^\perp, \quad (3.25)$$

dont la norme s'exprime très simplement en fonction de  $\eta_{\text{mod}}$  :

$$\|\alpha_{OL} - \langle \phi_s | \alpha_{OL} \rangle \phi_s\|^2 = 1 - \eta_{\text{mod}} \quad (3.26)$$

Le mode oscillateur local s'exprime donc en fonction du mode signal selon

$$\alpha_{OL} = \sqrt{\eta_{\text{mod}}} \phi_s + \sqrt{1 - \eta_{\text{mod}}} \phi_s^\perp, \quad (3.27)$$

ce qui se traduit par

$$\boxed{\hat{\mathbf{d}} = \sqrt{\eta_{\text{mod}}} \hat{\mathbf{c}} + \sqrt{1 - \eta_{\text{mod}}} \hat{\mathbf{c}}^\perp} \quad (3.28)$$

avec  $\hat{\mathbf{c}}^\perp = \int d^3\mathbf{k} \phi_s^{*\perp}(\mathbf{k}) \hat{\mathbf{a}}(\mathbf{k})$ .

Puisque tout le signal est par définition dans le mode  $\hat{\mathbf{c}}$ , le mode  $\hat{\mathbf{c}}^\perp$  est vide. On reconnaît la transformation associée à une lame séparatrice de transmission  $\eta_{\text{mod}}$ , qui est donc équivalente à un mauvais recouvrement de modes.

- Le bruit électronique ( $\eta_{\text{elec}} \simeq 0.99$ ), que l'on peut modéliser par des pertes lorsque l'on utilise les mesures homodynes du vide pour normaliser les données [Ferreyrol11a, Appel07].

Au final, nous obtenons une efficacité d'environ  $\eta = \eta_{\text{opt}} \eta_{\text{quant}} \eta_{\text{mod}} \eta_{\text{elec}} \simeq 0.68 \pm 0.04$ , qui peut fluctuer selon les expériences, principalement à cause du mode matching, très sensible à la superposition spatiale et temporelle du signal et de l'oscillateur local.

**Bruit de l'oscillateur local** Pour un oscillateur local bruité  $I_{OL} + \delta I_{OL}$ , l'équation (3.14) donne une mesure proportionnelle à  $\sqrt{I_{OL} + \delta I_{OL}} \hat{\mathbf{X}}_{s|OL}$ . On pourra donc négliger les fluctuations si elles sont faibles devant l'intensité moyenne.

**Mauvais équilibrage** Le principe interférométrique de la détection homodyne nécessite que les intensités de l'oscillateur local soient égales dans les deux voies de sortie de la séparatrice pour pouvoir s'annuler. Si celle-ci n'est pas réellement 50/50, mais a une transmission  $T = \frac{1}{2} + \epsilon$ , la différence des photocourants devient proportionnelle à  $\sqrt{\frac{I_{OL}}{N_0}} \hat{\mathbf{X}}_\theta + 2\epsilon(I_{OL} + \hat{\mathbf{a}}_s^\dagger \hat{\mathbf{a}}_s) \simeq \sqrt{\frac{I_{OL}}{N_0}} \hat{\mathbf{X}}_\theta + 2\epsilon I_{OL}$ . Puisque  $\hat{\mathbf{X}}_\theta$  prend des valeurs de l'ordre de  $\sqrt{N_0}$ , l'effet du déséquilibre sera négligeable si  $\epsilon \ll \frac{1}{2\sqrt{I_{OL}}}$ . Dans notre cas, la puissance de l'oscillateur local est d'environ 20  $\mu\text{W}$ , ce qui équivaut à  $\sim 10^8$  photons par impulsion. L'équilibrage doit donc être meilleur que  $10^{-4}$ .

## Montage

**Montage optique** Nous disposons de deux détections homodynes, dont le schéma du montage optique est présenté sur la figure 3.15. Un premier cube séparateur de polarisation permet de diriger le signal et l'oscillateur local vers les deux détections. Il permet également de séparer les composantes du signal dans les polarisations  $H$  et  $V$ , ce qui nous sera utile pour mesurer les quadratures d'un état EPR. L'OPA en configuration non dégénérée produit un état dont les deux

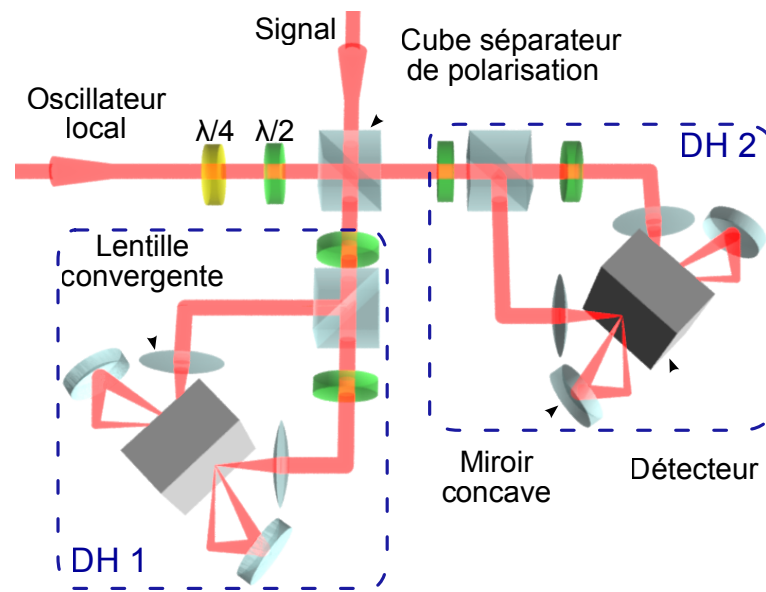


FIGURE 3.15 – Schéma du montage optique des deux détecteurs homodynes. Illustration tirée de [Ferreyrol11a].

modes sont séparés spatialement, mais de même polarisation. En tournant l'une d'entre elles de  $90^\circ$ , on peut superposer les faisceaux sans qu'ils interfèrent, et les diriger vers les détecteurs homodynes avec le même chemin optique, avant lesquels ils sont séparés à nouveau. Un couple de lames  $\lambda/2$  et  $\lambda/4$  permet de régler la proportion et le déphasage de l'oscillateur local allant vers les deux détecteurs, qui peut être réglé en mesurant un état cohérent de faible amplitude.

Pour chaque détection homodyne, les interférences avec l'oscillateur local sont effectuées avec un couple  $\lambda/2$  - cube séparateur de polarisation, afin de pouvoir équilibrer précisément les voies. La polarisation du faisceau transmis est remise à la verticale afin de limiter les pertes, et les réflexions sur les photodiodes sont récupérées par un jeu de deux miroirs qui les renvoient sur celles-ci.

**Acquisition** Les faisceaux sont mesurés par des photodiodes *Hamamatsu S3883*, ayant une bonne efficacité quantique (94.5%), et un faible courant d'obscurité (28 pA). Elles sont polarisées en inverse à  $\pm 6$  V. Ces tensions sont réglables afin d'égaliser les capacités parasites. Les photocourants sont directement soustraits par une loi des nœuds, puis amplifiés par le circuit présenté sur la figure 3.16. Il est conçu pour avoir une bande passante suffisante pour résoudre les impulsions individuelles à 800 kHz, tout en ayant un bruit suffisamment bas pour ne pas perturber les mesures (2.5 mV en écart-type).

L'acquisition des données est effectuée avec une carte *National Instrument PCI-6110*, disposant de 4 voies analogiques, auxquelles les APD sont également branchées. Elle permet d'acquérir jusqu'à 5 millions de points par seconde, avec un bruit négligeable de 0.1 mV en écart-type. Un signal TTL provenant du laser et synchronisé avec les impulsions à l'aide d'une ligne à retard permet de déclencher l'acquisition.

Les mesures sont ensuite contrôlées par des programmes spécifiques à chaque expérience, écrits en C++. Pour les réglages du dispositif, on utilise le logiciel *Igor*.



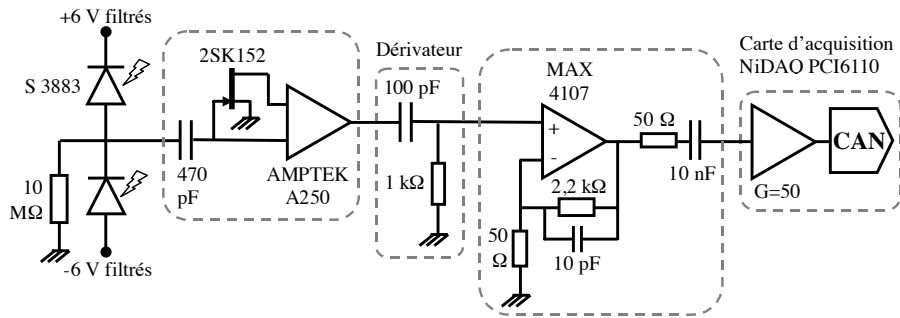


FIGURE 3.16 – Schéma du circuit électrique de la détection homodyne.

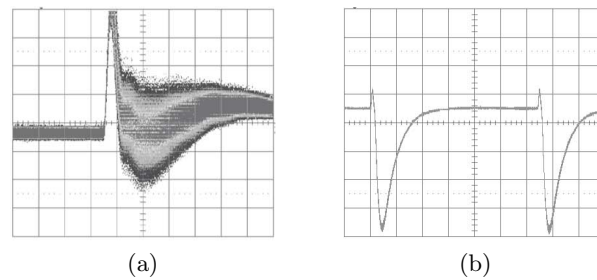


FIGURE 3.17 – (a) Détection homodyne équilibrée (échelle : 50 ns/div et 50 mV/div). (b) Détection homodyne déséquilibrée (échelle : 200 ns/div et 500 mV/div).

**Équilibrage** L'équilibrage de la détection doit être précisément réglé avant les expériences. Pour cela, on égalise l'intensité des deux voies avant les photodiodes avec la  $\lambda/2$ , de manière à ramener à zéro la valeur moyenne du signal observé sur un oscilloscope (Fig. 3.17). On vérifie ensuite que le signal ne présente pas trop de bruit technique provenant du laser en effectuant une transformée de Fourier du signal avec *Igor*. Le cas échéant, ce bruit peut être réduit en modifiant les réglages de la cavité, ou en ajustant la position des faisceaux sur les photodiodes. L'équilibrage est ensuite testé en vérifiant que la variance du vide mesurée augmente bien linéairement avec la puissance de l'oscillateur local (Fig. 3.18).

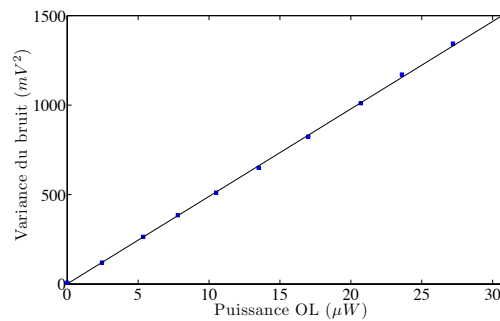


FIGURE 3.18 – Test de calibration : variance du vide en fonction de la puissance de l'oscillateur local.

### 3.4.2 Photodiode à avalanche

Les photodiodes à avalanche (APD) se comportent comme des détecteurs de photons, sans toutefois pouvoir les compter : un clic traduit la présence d'*au moins un photon*, sans en indiquer le nombre. Leur fonctionnement est basé sur un principe d'avalanche : un photon crée une paire électron-trou, qui est ensuite accélérée par une tension de polarisation inverse. Cette paire peut à son tour créer d'autres paires qui sont également accélérées, et qui, par effet d'avalanche, produisent un courant macroscopique qui est ensuite converti en une impulsion TTL. Les avalanches n'étant pas contrôlées, il n'est pas possible de distinguer les contributions de plusieurs photons. Cette incapacité à compter n'a pas été problématique pour les expériences réalisées avec notre dispositif. Au besoin, un multiplexage spatial avec deux APD permet de détecter la présence d'au moins deux photons, ce qui est suffisant pour la plupart des états produits. Les APD sont de plus très simples à mettre en œuvre et ne nécessitent pas de cryostat, contrairement à d'autres dispositifs plus sophistiqués (tels que les VLPC [Waks03], ou utilisant des supraconducteurs [Rosenberg05]).

Nos expériences utilisent deux APD *Perkin Elmer SPCM-AQR-13*, d'efficacité quantique  $\simeq 45\%$ . Leur courant d'obscurité (*dark count*) d'environ 200 coups par seconde, est ramené à environ 10-15 coups par seconde par synchronisation avec un signal TTL provenant du laser. Elles ont un temps mort de 50 ns, ce qui ne nous pose pas de problème puisque nos impulsions sont séparées de 1.25  $\mu\text{s}$ .

Bien que la fluorescence paramétrique soit fortement multimode, nous avons vu que la détection homodyne permet de se ramener à un modèle monomode grâce au filtrage qu'elle réalise. Ce filtrage n'est en revanche pas présent pour les APD, qui peuvent être déclenchées par des photons provenant de tous les modes. Afin de réaliser un conditionnement dans le mode de la détection homodyne, il est nécessaire d'utiliser un dispositif de filtrage, spatial et fréquentiel. Le filtrage spatial est réalisé à l'aide d'une fibre optique monomode de 2 mètres, sélectionnant le mode  $\text{TEM}_{00}$ . Un réseau de diffraction (efficacité 90%) couplé à une fente de largeur réglable placée au foyer d'une lentille de focale 100 mm (transmission globale  $\simeq 30\%$ ) permet ensuite d'appliquer un filtrage fréquentiel d'environ 1 nm de largeur spectrale, centré autour de la fréquence d'une sonde dans le même mode que l'oscillateur local. Au final, l'efficacité globale est d'environ 10%.

Malgré le système de filtrage, le conditionnement peut quand même être déclenché par des photons non corrélés, ou provenant d'autres modes. Cet effet sera modélisé en supposant que le conditionnement provient du bon mode avec une probabilité  $\xi$ , qui peut être interprétée comme une pureté modale, et qu'il provient d'autres modes avec une probabilité  $1-\xi$ . Si un conditionnement APD dans le bon mode induit une transformation  $\mathcal{E}^\vee(\hat{\rho})$ , alors qu'une absence de conditionnement est décrite par  $\mathcal{E}^\chi(\hat{\rho})$ , la transformation totale avec un filtrage imparfait sera

$$\hat{\rho}' = \xi \mathcal{E}^\vee(\hat{\rho}) + (1-\xi) \mathcal{E}^\chi(\hat{\rho}). \quad (3.29)$$

Une APD parfaite est modélisée par un POVM contenant deux éléments  $\hat{\Pi}_1$  et  $\hat{\Pi}_0$ , qui correspondent respectivement à la présence d'au moins un photon, et à une absence de conditionnement :

$$\hat{\Pi}_1 = \mathbb{I} - |0\rangle\langle 0| = \sum_{n=1}^{\infty} |n\rangle\langle n| \quad (3.30a)$$

$$\hat{\Pi}_0 = |0\rangle\langle 0| \quad (3.30b)$$

Pour une efficacité de détection  $\lambda$ , un conditionnement est réussi si au moins un photon parvient jusqu'à l'APD. Pour un état  $|n\rangle$ , la probabilité qu'au moins un photon soit détecté est  $1-(1-\lambda)^n$ . Le POVM associé au conditionnement est alors :

$$\hat{\Pi}_1 = \sum_{n=1}^{\infty} [1-(1-\lambda)^n] |n\rangle\langle n| \quad (3.31)$$

$$\hat{\Pi}_0 = \mathbb{I} - \hat{\Pi}_1 \quad (3.32)$$

L'efficacité des APD à très peu d'influence sur la qualité des états produits par conditionnement : elle influe principalement sur le taux de succès. De toute manière, nos états ne contiennent que très peu de photons, et on s'arrange en général pour que la probabilité que plus d'un photon arrive sur l'APD soit négligeable. Ainsi, pour une faible efficacité, le photon ne sera tout simplement pas détecté. En revanche, il nous est impossible de faire un conditionnement sur l'absence de photon.

### 3.4.3 Soustraction de photon

Les opérateurs destruction  $\hat{a}$  et création  $\hat{a}^\dagger$  ne sont pas que de simples outils théoriques : ils peuvent être appliqués expérimentalement sur un état quantique. Les bases de notre expérience ont été posées au cours de la thèse de J. Wenger, avec l'obtention d'une statistique non gaussienne en soustrayant un photon du vide comprimé monomode [Wenger04b]. Le dispositif a ensuite été optimisé lors de la thèse d'A. Ourjoumtsev qui a obtenu des fonctions de Wigner négatives, produisant ainsi des “chatons de Schrödinger” [Ourjoumtsev06b], et qui a également démontré une augmentation d'intrication en utilisant une soustraction de photon [Ourjoumtsev07b]. J. Neergaard-Nielsen *et al.* ont également préparé des superpositions de vide comprimé et de vide comprimé soustrait d'un photon [Neergaard-Nielsen10], en utilisant une combinaison de soustraction et de déplacement. Citons également V. Parigi *et al.*, qui ont démontré expérimentalement les relations de commutation bosoniques  $[\hat{a}, \hat{a}^\dagger]=1$  en utilisant une combinaison de soustraction et d'addition de photon [Parigi07]. L'addition de photon a également été utilisée dans notre groupe, afin d'étudier la non gaussianité d'un état cohérent auquel on a ajouté un photon [Barbieri10, Ferreyrol11a].

Montrons maintenant comment réaliser une soustraction de photon avec une lame séparatrice et une APD. Nous utiliserons ce principe pour la caractérisation expérimentale de la porte de phase, au chapitre 6. Le principe consiste à prélever une très faible partie du signal à l'aide d'une lame séparatrice ayant une transmission proche de 1, et à détecter un photon dans le mode réfléchi. Si la réflectivité est suffisamment faible, l'état résultant est soustrait d'un photon.

On rappelle que l'opérateur associé à la lame séparatrice est  $\hat{U}_{BS} = \exp[\theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger)]$ , où  $\hat{a}$  est le mode de l'état  $\hat{\rho}$  auquel on veut soustraire un photon, et  $\hat{b}$  est un mode vide. L'application de  $\hat{U}_{BS}$ , suivie de la détection d'au moins un photon sur la voie APD (mode  $\hat{b}$ ) avec l'opérateur  $\hat{M}_b = \mathbb{I}_b - |0_b\rangle\langle 0_b|$ , donne un état

$$\hat{M}_b \hat{U}_{BS} (\hat{\rho} \otimes |0_b\rangle\langle 0_b|) \hat{U}_{BS}^\dagger \hat{M}_b = \left( \hat{M}_b \hat{U}_{BS} |0_b\rangle \right) \hat{\rho} \left( \langle 0_b| \hat{U}_{BS}^\dagger \hat{M}_b \right). \quad (3.33)$$

Pour  $T \simeq 1$ , on peut développer  $\hat{U}_{BS}$  au premier ordre :

$$\hat{U}_{BS} \simeq \mathbb{I} + \theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger) \quad (3.34)$$

On a donc :

$$\hat{M}_b \hat{U}_{BS} |0_b\rangle \simeq \left( \mathbb{I}_b - |0_b\rangle\langle 0_b| \right) \left( \mathbb{I} + \theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger) \right) |0_b\rangle \quad (3.35a)$$

$$= \left( \mathbb{I}_b - |0_b\rangle\langle 0_b| \right) \left( |0_b\rangle - \theta \hat{a} |1_b\rangle \right) \quad (3.35b)$$

$$= -\theta \hat{a} |1_b\rangle \quad (3.35c)$$

Avec ce développement, l'équation(3.33) devient alors égale à

$$\theta^2 \hat{a} \hat{\rho} \hat{a}^\dagger \otimes |1_b\rangle\langle 1_b|. \quad (3.36)$$

Après trace partielle sur le mode  $\hat{b}$ , l'état du mode  $\hat{a}$  est donc finalement proportionnel à  $\hat{a} \hat{\rho} \hat{a}^\dagger$ .

La qualité de la soustraction dépend de la validité du développement : si  $T$  est trop faible, il faut prendre en compte les termes d'ordres supérieurs en  $\theta$ . Nous avons également considéré une APD parfaite. Une efficacité de détection inférieure à un ne change pratiquement que le taux de succès, du moment que  $T$  est suffisamment proche de 1. On pourra consulter la référence [Kim08] pour un calcul plus détaillé.

On peut aussi se convaincre très facilement du principe de la soustraction de photon en regardant l'action de  $\hat{U}_{BS}$  sur un état  $|n\rangle|0\rangle$  [Leonhardt97] :

$$\hat{U}_{BS} |n_a\rangle |0_b\rangle = \sum_{k=0}^n \sqrt{C_n^k} t^k (-r)^{n-k} |k_a\rangle |n-k_b\rangle \quad (3.37)$$

avec  $t=\sqrt{T}$  et  $r=\sqrt{1-T}$ . Lorsque  $t \simeq 1$ , le terme d'ordre un en  $r$  (réflexion d'un photon) correspond à  $k=n-1$ . Si  $r$  est suffisamment petit pour que l'on puisse négliger les termes d'ordres supérieurs, et que l'on conditionne sur la détection d'un photon dans le mode réfléchi, l'état final est

$$\langle 1_b | \hat{U}_{BS} [ |n\rangle |0\rangle ] = -\sqrt{n} t^{n-1} r |n-1_a\rangle, \quad (3.38)$$

puisque  $C_n^{n-1}=n$ . Si on a  $t^{n-1} \simeq 1$ , on retrouve le fait que  $|n\rangle$  est transformé proportionnellement à  $\sqrt{n} |n-1\rangle$ , ce qui correspond bien à l'application de  $\hat{a}$ .

En pratique, une transmission  $T=0.9$  assure un bon compromis entre le taux de succès et la qualité des états produits.

### 3.5 Conclusion

Nous disposons à présent d'un dispositif permettant de préparer et de mesurer une grande variété d'états quantiques, utilisant à la fois les outils des variables discrètes et continues. Nos états de base sont gaussiens : il peut s'agir d'états cohérents, provenant du laser fortement atténué, de vides comprimés monomodes, ou d'états EPR, produits par l'OPA.

De nombreux protocoles d'information quantique font appel à des états non gaussiens. Notre dispositif a permis d'en produire un certain nombre, grâce à des mesures non gaussiennes implémentées par nos APD [Wenger04b, Ourjoumtsev06a, Ourjoumtsev06b, Ourjoumtsev07c, Ourjoumtsev09]. Fort de ces préparations, nous pouvons également utiliser ces états comme ressources, afin de réaliser des protocoles élémentaires d'information quantique [Ourjoumtsev07b, Ferreyrol10], ou encore pour étudier et caractériser certaines propriétés quantiques [Barbieri10, Blandino12a, Blandino12b].

Nous disposons également d'un ensemble d'outils théoriques nous permettant de modéliser nos états et les imperfections de notre dispositif. De ce fait, nous pouvons d'abord simuler une expérience afin d'en estimer la faisabilité, puis déterminer les paramètres clés à optimiser pour le réglage du dispositif, et enfin comprendre l'origine des imperfections.

# Chapitre 4

## Eléments de théorie de l'information

### Sommaire

|            |                                                       |           |
|------------|-------------------------------------------------------|-----------|
| <b>4.1</b> | <b>Introduction</b>                                   | <b>71</b> |
| <b>4.2</b> | <b>Information classique</b>                          | <b>72</b> |
| 4.2.1      | Entropie de Shannon                                   | 72        |
| 4.2.2      | Entropie de deux variables aléatoires                 | 72        |
| 4.2.3      | Information mutuelle                                  | 73        |
| 4.2.4      | Quelques propriétés de base de l'entropie             | 73        |
| <b>4.3</b> | <b>Information quantique</b>                          | <b>74</b> |
| 4.3.1      | Entropie de von Neumann                               | 74        |
| 4.3.2      | Entropie et corrélations pour des systèmes bipartites | 74        |
| 4.3.3      | Etats classiques-quantiques                           | 76        |
| 4.3.4      | Borne de Holevo                                       | 76        |
| <b>4.4</b> | <b>Modèle du canal gaussien</b>                       | <b>77</b> |
| 4.4.1      | Canal sans bruit                                      | 77        |
| 4.4.2      | Canal avec bruit thermique                            | 77        |
| 4.4.3      | Information mutuelle                                  | 79        |
| <b>4.5</b> | <b>Conclusion</b>                                     | <b>80</b> |

### 4.1 Introduction

Considérons une pièce de monnaie un peu spéciale : lorsqu'on la lance, elle retombe sur le côté pile avec une probabilité  $p$ , et sur le côté face avec une probabilité  $1-p$ . Dans le cas extrême où  $p=1$ , la pièce retombera à chaque fois sur le côté pile, et un lancer ne nous apprendra rien puisque nous connaissons son résultat à l'avance. En revanche, lorsque  $p=1/2$ , nous n'avons plus aucune information. A chaque fois, nous obtenons un des deux résultats possibles avec la même probabilité. Si l'on associe la valeur 0 au résultat pile, et 1 au résultat face, chaque lancer nous apprend un bit d'information. Intuitivement, lorsque  $p=0$  ou  $p=1$ , nous n'obtenons aucun bit d'information, et pour une valeur intermédiaire de  $p$ , nous obtiendrons une information comprise entre 0 et 1 bit.

A travers cet exemple, nous avons vu que l'incertitude associée au tirage d'une variable aléatoire peut être quantifiée en terme de bits d'information. La théorie de l'information classique formalise ces notions, et permet aussi de décrire les corrélations entre deux variables aléatoires

et de calculer l'information maximale qu'il est possible d'en extraire. En se basant sur les mêmes concepts, la théorie de l'information quantique généralise la notion d'information à des états quantiques.

Dans ce chapitre, nous présentons les concepts principaux de ces deux théories de l'information. Une très bonne introduction détaillée peut être trouvée dans le livre écrit par Nielsen et Chuang [Nielsen00]. Le lecteur pourra également consulter l'ouvrage de Vedral [Vedral07], qui fournit une approche plus concise, ou encore le cours de Preskill [Preskill98], un peu plus technique.

## 4.2 Information classique

### 4.2.1 Entropie de Shannon

Considérons une variable aléatoire classique  $X$  donnant les résultats  $\{x_i\}$  avec des probabilités  $\{p_i\}$ . On dit aussi que chaque valeur  $x_i$  est un symbole, appartenant à l'alphabet  $\{x_i\}$ . L'entropie caractérise l'incertitude ou l'imprévisibilité avant le tirage, ou de manière équivalente, l'information apportée par le résultat du tirage. Elle est définie par [Shannon48]

$$H(X) = - \sum_i p_i \log_2 p_i. \quad (4.1)$$

L'unité de cette mesure est le *bit*, *binary unit*. Cette définition provient de quelques propriétés que l'on demande à une mesure de l'incertitude. On veut en particulier que la mesure de l'incertitude de deux variables aléatoires indépendantes soit additive, ce qui justifie l'utilisation d'un logarithme.

L'entropie est reliée aux ressources physiques nécessaires pour stocker ou transmettre un message. Asymptotiquement, il est possible de stocker un message de  $N$  symboles en utilisant  $NH(X)$  bits physiques, avec une erreur tendant vers 0 lorsque  $N$  tend vers l'infini [Shannon48]. Chaque symbole de l'alphabet apporte donc en moyenne  $H(X)$  bits d'information.

### 4.2.2 Entropie de deux variables aléatoires

Pour deux variables aléatoires  $X$  et  $Y$ , on définit l'entropie conjointe

$$H(X, Y) = - \sum_{x,y} p(x, y) \log_2 p(x, y). \quad (4.2)$$

C'est une mesure de l'incertitude totale pour le couple de variables aléatoires  $(X, Y)$ . Lorsqu'elles sont indépendantes, on a  $H(X, Y) = H(X) + H(Y)$ .

Lorsque  $X$  et  $Y$  sont corrélées, la connaissance de l'une d'entre elles,  $Y$  par exemple, réduit l'incertitude qu'il reste sur  $X$ . L'incertitude restante sur  $X$ , connaissant l'information apportée par  $Y$ , est donnée par l'entropie conditionnelle

$$H(X|Y) = H(X, Y) - H(Y). \quad (4.3)$$

Elle est nulle quand  $X$  et  $Y$  sont parfaitement corrélées. C'est en moyenne la quantité de bits par symbole manquante avec un code optimal pour parfaitement décrire une chaîne de symboles de  $X$  sachant  $Y$  [Preskill98].

### 4.2.3 Information mutuelle

Considérons un couple de variables aléatoires  $(X, Y)$ . Nous avons vu que l'entropie conditionnelle de  $X$  sachant  $Y$  correspond à l'incertitude restant sur  $X$ , connaissant  $Y$ . Plutôt que de considérer l'incertitude, on peut également s'intéresser à l'information que  $X$  apporte sur  $Y$ . Cette information correspond à l'incertitude sur  $X$ , moins l'incertitude restante sur  $X$  connaissant  $Y$ , et définit *l'information mutuelle* :

$$I(X:Y) = H(X) - H(X|Y) \quad (4.4a)$$

$$= H(Y) - H(Y|X) \quad (4.4b)$$

$$= H(X) + H(Y) - H(X, Y) \quad (4.4c)$$

C'est le nombre de bits par symbole appris sur  $X$  en connaissant  $Y$ . Inversement, c'est aussi ce que l'on sait de  $Y$  connaissant  $X$ . C'est une mesure de leurs corrélations, symétrique en  $X$  et  $Y$ , correspondant à la quantité maximale de bits que l'on peut en extraire.

### 4.2.4 Quelques propriétés de base de l'entropie

Les différentes entropies possèdent de nombreuses propriétés que l'on peut trouver par exemple dans la référence [Nielsen00]. Nous en indiquons quelques unes parmi les plus importantes :

- $H(X|Y), H(Y|X) \geq 0$ , ce qui implique  $I(X:Y) \leq H(X), H(Y)$
- $H(X), H(Y) \leq H(X, Y)$
- $H(X, Y) \leq H(X) + H(Y)$  (sous-additivité)

Toutes ces différentes mesures sont résumées sur la figure 4.1.

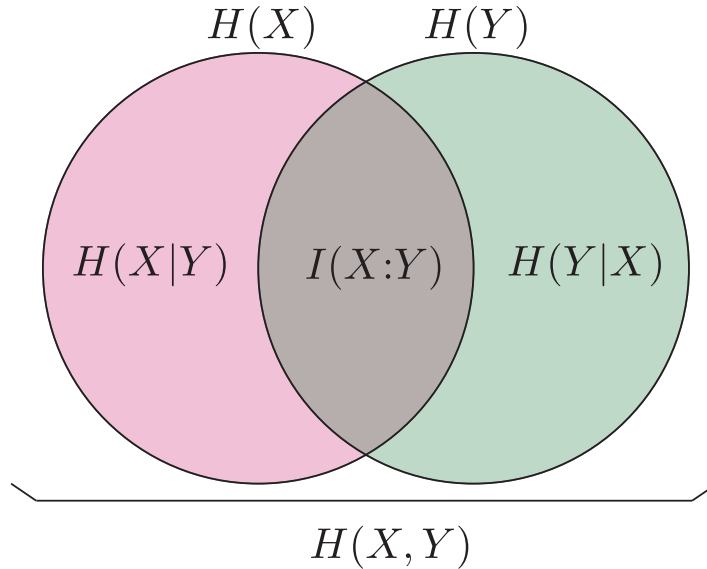


FIGURE 4.1 – Les différentes entropies et mesures des corrélations.

Pour une variable continue gaussienne de variance  $V$ , l'entropie est égale à [Shannon48]

$$H(X) = \frac{1}{2} \log_2 V + C, \quad (4.5)$$

où  $C$  est une constante.



## 4.3 Information quantique

### 4.3.1 Entropie de von Neumann

L'entropie de von Neumann généralise la notion d'entropie à des états quantiques, en remplaçant la distribution de probabilité par une matrice densité. Elle est définie par

$$S(\hat{\rho}) = -\text{Tr}\{\hat{\rho} \log_2 \hat{\rho}\}. \quad (4.6)$$

En écrivant la matrice densité sous forme diagonale,  $\hat{\rho} = \sum_i \lambda_i |i\rangle\langle i|$ , l'entropie de von Neumann devient

$$S(\hat{\rho}) = -\sum_i \lambda_i \log_2 \lambda_i. \quad (4.7)$$

Elle correspond donc à l'entropie de Shannon pour la distribution de probabilité associée à la diagonalisation de  $\hat{\rho}$ .

Comme l'entropie de Shannon, l'entropie de von Neumann a une interprétation en termes de ressources nécessaires pour coder l'information à la limite asymptotique [Nielsen00, Preskill98]. Supposons qu'Alice forme un message en choisissant pour symboles des états  $\{|\phi_x\rangle\}$  avec une probabilité  $p_x$ , parmi un alphabet  $\{|\phi_x\rangle, p_x\}$ . Avant d'être choisi, chaque symbole a pour matrice densité  $\hat{\rho} = \sum_x p_x |\phi_x\rangle\langle\phi_x|$ . Un message formé de  $N$  symboles est donc décrit par  $\hat{\rho}^{\otimes N}$ . L'entropie de von Neumann  $S(\hat{\rho})$  correspond alors au nombre de qubits par symbole nécessaires pour coder le message, à la limite asymptotique. En d'autres termes, un message de  $N$  symboles  $\hat{\rho}^{\otimes N}$  aura un support presque entièrement contenu dans un espace de dimension  $2^{NS(\hat{\rho})}$ , pour  $N$  grand.

L'entropie de von Neumann est également reliée à l'information classique. Elle correspond au nombre maximal de bits classiques que l'on peut obtenir d'un mélange statistique d'états purs en faisant une mesure optimale [Nielsen00, Preskill98].

### 4.3.2 Entropie et corrélations pour des systèmes bipartites

#### Entropie jointe et entropie conditionnelle

Pour un état bipartite  $\hat{\rho}_{AB}$ , on définit l'entropie jointe

$$S(A, B) := -\text{Tr}\{\hat{\rho}_{AB} \log_2 \hat{\rho}_{AB}\}, \quad (4.8)$$

et l'entropie conditionnelle

$$S(A|B) = S(A, B) - S(B), \quad (4.9)$$

où  $S(B) := S(\hat{\rho}_B)$ . Par la suite, nous utiliserons de manière équivalente la notation  $S(A, B)$  ou  $S(\hat{\rho}_{AB})$  pour l'entropie d'un système bipartite.

#### Information mutuelle

Les corrélations totales présentes dans le système sont mesurées par l'information mutuelle quantique, définie d'une manière similaire à l'information mutuelle classique [Modi12, Vedral02]. :

$$S(A:B) = S(A) + S(B) - S(A, B) \quad (4.10)$$

Comme l'information mutuelle classique, cette quantité est symétrique :  $S(A:B) = S(B:A)$ . Ces corrélations peuvent être séparées en une partie contenant des corrélations classiques, et une partie contenant des corrélations purement quantiques [Modi12]. Nous reviendrons sur ces notions dans le chapitre 5 consacré à l'estimation de la discorde quantique.

## Entropie relative

L'entropie relative est définie par

$$S(\hat{\rho}||\hat{\sigma}) = \text{Tr}\{\hat{\rho} \log_2 \hat{\rho}\} - \text{Tr}\{\hat{\rho} \log_2 \hat{\sigma}\} = -S(\hat{\rho}) - \text{Tr}\{\hat{\rho} \log_2 \hat{\sigma}\}. \quad (4.11)$$

Elle est très souvent utilisée comme une mesure de distance entre états quantiques, bien que son asymétrie n'en fasse pas rigoureusement une norme, et elle joue un rôle très important en information quantique. Le lecteur pourra consulter l'article [Vedral02] qui lui est entièrement consacré pour une présentation détaillée.

L'information mutuelle est égale à l'entropie relative entre l'état bipartite et le produit tensoriel des deux états réduits :

$$S(\hat{\rho}_{AB}||\hat{\rho}_A \otimes \hat{\rho}_B) = S(\hat{\rho}_A) + S(\hat{\rho}_B) - S(\hat{\rho}_{AB}). \quad (4.12)$$

De même, pour un état bipartite classique

$$\hat{\rho}_{AB} = \sum_{xy} p_{xy} |x\rangle\langle x| \otimes |y\rangle\langle y|, \quad (4.13)$$

où  $\{|x\rangle\}$  et  $\{|y\rangle\}$  sont des états orthogonaux, l'information mutuelle classique  $I(X, Y)$  correspond à l'entropie relative avec  $\hat{\rho}_A \otimes \hat{\rho}_B$ . Les corrélations ont donc une interprétation géométrique, en terme de distance par rapport à un état  $\hat{\rho}_A \otimes \hat{\rho}_B$  totalement dépourvu de corrélations.

## Entropie d'un état gaussien

Pour un état gaussien  $\hat{\rho}$  à  $N$  modes, l'entropie de von Neumann est égale à [Weedbrook12]

$$S(\hat{\rho}) = \sum_{k=1}^N g\left(\frac{\nu_k - 1}{2}\right), \quad (4.14)$$

où  $\{\nu_k\}$  sont les valeurs propres symplectiques<sup>1</sup> de la matrice de covariance de  $\hat{\rho}$ , et

$$g(x) = (x+1) \log_2(x+1) - x \log_2 x. \quad (4.15)$$

## Quelques propriétés

Encore une fois, nous renvoyons le lecteur aux références [Nielsen00] ou [Preskill98] pour une liste détaillée des propriétés de l'entropie de von Neumann. Nous en présentons quelques unes ci-dessous :

- $S(\hat{\rho})=0$  si  $\hat{\rho}$  est un état pur.
- Pour un état  $\hat{\rho}_{AB}$  pur,  $S(\hat{\rho}_A)=S(\hat{\rho}_B)$ .
- $S(\hat{\rho}_{AB}) \leq S(\hat{\rho}_A) + S(\hat{\rho}_B)$  (sous-additivité) avec égalité si  $\hat{\rho}_{AB}=\hat{\rho}_A \otimes \hat{\rho}_B$ .
- $S(\hat{\rho}_{AB}) \geq |S(\hat{\rho}_A)-S(\hat{\rho}_B)|$  (inégalité du triangle). Ainsi, lorsque  $\hat{\rho}_{AB}$  est pur, l'entropie du système bipartite est nulle, mais ce n'est pas forcément le cas de celle des sous systèmes. Ceci contraste avec le cas classique, où l'entropie de Shannon de deux variables est toujours supérieure ou égale à celle d'une seule des variables.

---

1. Le théorème de Williamson [Weedbrook12] assure que toute matrice de covariance  $\Gamma$  peut être mise sous forme diagonale à l'aide de transformations symplectiques, c'est à dire linéaires en les opérateurs création et destruction, et préservant les relations de commutation. La forme diagonale de  $\Gamma$  s'écrit  $\bigoplus_{k=1}^N \nu_k \mathbb{I}_2$ . Les  $\{\nu_k\}$  sont les valeurs propres symplectiques.

### 4.3.3 Etats classiques-quantiques

En cryptographie quantique, Alice et Bob cherchent à créer des corrélations classiques en utilisant des états quantiques. Pour la plupart des protocoles, Alice prépare un état parmi un ensemble  $\{|\psi_x\rangle, p_x\}$ , qu'elle envoie ensuite à Bob à travers un canal quantique. Les états  $\{|\psi_x\rangle\}$  ne sont pas tous orthogonaux afin d'assurer la sécurité du protocole, comme nous le verrons dans le chapitre 7. Sous l'effet des pertes, du bruit, ou d'autres imperfections, un état pur  $|\psi_x\rangle$  est transformé en un mélange statistique  $\hat{\rho}_x$ . Bob cherche ensuite à déterminer l'état envoyé par Alice.

Cette préparation d'état par Alice, suivie d'un envoi à Bob, est équivalente au partage d'un état classique-quantique

$$\hat{\rho}_{AB} = \sum_x p_x |i_x\rangle\langle i_x| \otimes \hat{\rho}_x, \quad (4.16)$$

classique pour Alice et quantique pour Bob. Les états  $\{|i_x\rangle\}$  sont des états fictifs orthogonaux, témoignant du fait qu'Alice sait quel état elle a préparé. En revanche, les états  $\{\hat{\rho}_x\}$  peuvent être quelconques.

Pour un état classique-quantique de la forme (4.16), l'entropie de von Neumann possède quelques propriétés importantes [Nielsen00] :

$$S(A) = H(p_X) \quad (4.17a)$$

$$S(B) = S(\hat{\rho}_B), \quad \text{avec } \hat{\rho}_B = \sum_x p_x \hat{\rho}_x \quad (4.17b)$$

$$S(A, B) = H(p_X) + \sum_x p_x S(\hat{\rho}_x) \quad (4.17c)$$

En utilisant les relations (4.17), on montre que l'information mutuelle est donnée par

$$S(A:B) = S(\hat{\rho}_B) - \sum_x p_x S(\hat{\rho}_x). \quad (4.18)$$

Lorsque les états  $\{\hat{\rho}_x\}$  sont purs, l'expression (4.18) se réduit donc à  $S(A:B) = S(\hat{\rho}_B)$ .

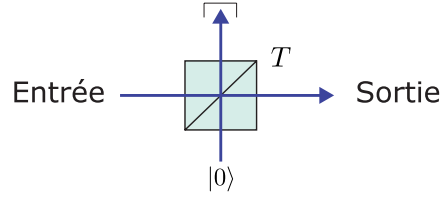
### 4.3.4 Borne de Holevo

Supposons qu'Alice prépare des états parmi un ensemble  $\{\hat{\rho}_x, p_x\}$ , suivant le résultat d'une variable aléatoire  $X$ . Bob réalise des mesures à l'aide d'un POVM  $\{\hat{\mathbf{E}}_y\} = \{\hat{\mathbf{E}}_0, \dots, \hat{\mathbf{E}}_m\}$ , donnant un résultat  $Y$ . Quelle que soit la mesure effectuée par Bob, l'information mutuelle classique  $I(X:Y)$  entre  $X$  et  $Y$  vérifie [Nielsen00]

$$I(X:Y) \leq S(\hat{\rho}) - \sum_x p_x S(\hat{\rho}_x), \quad (4.19)$$

avec  $\hat{\rho} = \sum_x p_x \hat{\rho}_x$ . Cette borne supérieure est la borne de Holevo. Pour être atteinte, elle nécessite en général que Bob puisse effectuer des mesures collectives. Cette borne correspond également à l'information mutuelle quantique (4.18) de l'état classique-quantique équivalent pour Alice et Bob [Cerf96, Vedral07].

Cette borne sera utile, sous une forme légèrement différente, afin de borner l'information qu'Eve peut obtenir durant un protocole de cryptographie quantique.


 FIGURE 4.2 – Modélisation d'un canal quantique de transmission  $T$  par une lame séparatrice.

## 4.4 Modèle du canal gaussien

### 4.4.1 Canal sans bruit

Un canal linéaire et symétrique, de transmission  $T$ , peut être modélisé simplement par une lame séparatrice [Leonhardt97, Weedbrook12], dont le deuxième mode d'entrée  $\hat{b}$  est vide (Fig. 4.2). Nous avons vu avec l'équation (2.96) que dans ce cas, les opérateurs de quadrature sont transformés en

$$\hat{X}_{\text{out}} = \sqrt{T}\hat{X} + \sqrt{1-T}\hat{X}_b, \quad (4.20a)$$

$$\hat{P}_{\text{out}} = \sqrt{T}\hat{P} + \sqrt{1-T}\hat{P}_b. \quad (4.20b)$$

Les quadratures ne sont pas seulement atténuées, elles sont aussi “contaminées” par le vide. Ce bruit n'aura pas d'influence sur les valeurs moyennes. En revanche, il compense d'une certaine manière la diminution du bruit initial afin de préserver le commutateur<sup>2</sup> de  $\hat{X}$  et  $\hat{P}$ . Les variances sont donc transformées en

$$\Delta^2 Q_{\text{out}} = T\Delta^2 Q + (1-T)\Delta^2 Q_b, \quad (4.21a)$$

$$= T\Delta^2 Q + (1-T)N_0, \quad (4.21b)$$

$$= T\left(\Delta^2 Q + \frac{1-T}{T}N_0\right), \quad (4.21c)$$

avec  $Q=X$  ou  $P$ . La deuxième partie de la troisième ligne fait intervenir la notion de *bruit ajouté ramené à l'entrée*, égal à  $\frac{1-T}{T}N_0$  pour les pertes.

### 4.4.2 Canal avec bruit thermique

Le bruit ajouté par un canal quantique peut être caractérisé par son équivalent ramené à l'entrée. Par définition, *l'excès de bruit*  $\epsilon$  sera le bruit ajouté en plus de la contribution due aux pertes, en unité de bruit de photon<sup>3</sup>. Un canal ajoutant un excès de bruit gaussien  $\epsilon$  transforme donc les variances des quadratures en

$$\Delta^2 Q_{\text{out}} = T\left(\Delta^2 Q + \frac{1-T}{T}N_0 + \epsilon N_0\right) = T\Delta^2 Q + (1-T)\left[\frac{T}{1-T}\left(\frac{1-T}{T}N_0 + \epsilon N_0\right)\right]. \quad (4.22)$$

On reconnaît une forme similaire à la deuxième partie de l'expression (4.21a), avec cette fois ci  $\Delta^2 Q_b = \frac{T}{1-T}\left(\frac{1-T}{T}N_0 + \epsilon N_0\right) := V_{\text{th}}$ . Puisque l'on a toujours  $\langle \hat{Q}_b \rangle = 0$ , le canal peut donc être modélisé en injectant un état thermique de variance  $V_{\text{th}}$  dans le mode  $\hat{b}$  (Fig. 4.3).

2. Nous reviendrons sur ce point dans le chapitre présentant l'amplificateur sans bruit.

3. Par abus de langage, on appellera parfois la quantité  $\epsilon$  simplement bruit ajouté ou excès de bruit, au lieu d'excès de bruit ramené à l'entrée.

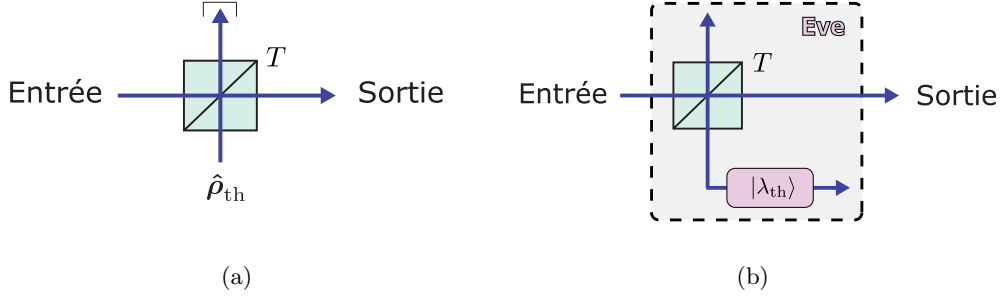


FIGURE 4.3 – Modélisation d’un canal quantique de transmission  $T$  et d’excès de bruit  $\epsilon$  : (a) avec un état thermique  $\hat{\rho}_{th}$  ; (b) avec un état EPR  $|\lambda_{th}\rangle$  purifiant  $\hat{\rho}_{th}$  (voir section 2.3.1 pour un rappel de la purification). Les deux systèmes ne sont pas discernables par Alice ou Bob.

En posant

$$\chi_{line} = \frac{1-T}{T} + \epsilon, \quad (4.23)$$

qui est égal au bruit total ajouté par le canal ramené à l’entrée, en unité de bruit de photon, on définit les paramètres  $\lambda_{th}$  et  $r_{th}$  de l’état thermique tels que

$$\frac{1}{N_0} V_{th} = \cosh 2r_{th} = \frac{1+\lambda_{th}^2}{1-\lambda_{th}^2} = \frac{T}{1-T} \chi_{line}. \quad (4.24)$$

Puisqu’un état thermique peut être obtenu en prenant la trace d’un état EPR, les deux schémas de la figure 4.3 sont équivalents du point de vue d’Alice et de Bob. Si Eve remplace le dispositif de gauche par le dispositif de droite, elle peut acquérir de l’information sur l’état transmis dans le canal en effectuant des mesures sur son état EPR. Ce dispositif est connu sous le nom de “cloneuse intrigante” [Grosshans03b, Grosshans02a].

**Détection homodyne imparfaite** Une détection homodyne d’efficacité  $\nu$  et de bruit électronique  $\kappa$  peut se modéliser comme une détection homodyne parfaite après un canal virtuel de transmission  $\nu$  et de bruit ajouté ramené en entrée de ce canal [Lodewyck07]

$$\chi_{hom} = \frac{1-\nu}{\nu} + \frac{\kappa}{\nu}, \quad (4.25)$$

en suivant une modélisation similaire à la figure 4.3. Une détection homodyne imparfaite après un canal de transmission  $T$  et de bruit ajouté  $\chi_{line}$  est donc équivalente, pour Alice et Bob, à une détection homodyne parfaite après un canal de transmission  $G=\nu T$ , et de bruit ajouté total ramené à l’entrée

$$\chi_{tot} = \chi_{line} + \frac{\chi_{hom}}{T}. \quad (4.26)$$

De manière générale, on peut considérer qu’un bruit ramené à l’entrée  $\chi_{tot}$  correspond à la variance d’une quadrature  $\hat{X}_{tot}$  qui serait ajoutée à celle du signal avant le canal de transmission  $G$ ,

$$\hat{X}_{out} = \sqrt{G}(\hat{X} + \hat{X}_{tot}), \quad (4.27)$$

avec  $\Delta^2 X_{tot} = \chi_{tot} N_0$ . Le bruit est par hypothèse indépendant du signal.

### 4.4.3 Information mutuelle

Anticipons un peu sur le protocole de cryptographie quantique utilisant des états cohérents, que nous présenterons dans le chapitre 7. Supposons qu'Alice choisisse deux variables  $x_A$  et  $p_A$  selon une distribution gaussienne de moyenne nulle et de variance  $V_A N_0$ . Elle prépare ensuite un état cohérent centré en  $(x_A, p_A)$ , qu'elle envoie à Bob à travers un canal quantique gaussien. On suppose que Bob mesure la quadrature  $\hat{X}_B$ . Avant le canal quantique, cette quadrature s'écrit

$$\hat{X}_{\text{in}} = x_A + \hat{X}_0, \quad (4.28)$$

où  $\hat{X}_0$  correspond à la quadrature du vide. Le canal quantique transforme ensuite  $\hat{X}_{\text{in}}$  en

$$\hat{X}_B = \sqrt{G}(\hat{X}_{\text{in}} + \hat{X}_{\text{tot}}) = \sqrt{G}(x_A + \hat{X}_N), \quad (4.29)$$

où l'on a introduit  $\hat{X}_N = \hat{X}_0 + \hat{X}_{\text{tot}}$ . La variance  $\Delta^2 X_N = (1 + \chi_{\text{tot}})N_0$  correspond au bruit total ajouté sur la quadrature d'Alice  $x_A$ . Pour chaque quadrature, on peut adopter une modélisation classique, selon laquelle Alice fait le tirage d'une variable aléatoire<sup>4</sup>  $X_A$  de variance  $V_A N_0$  et de moyenne nulle, alors que Bob reçoit une variable aléatoire  $X_B$  de variance  $G(V_A + 1 + \chi_{\text{tot}})N_0$ , également de moyenne nulle.

L'information qu'ils pourront extraire de leur communication dépend naturellement du bruit ajouté par le canal, et de son amplitude par rapport à la modulation d'Alice. Ces deux grandeurs sont comparées par le rapport signal-à-bruit (SNR), défini par

$$\text{SNR} = \frac{\langle X_A^2 \rangle}{\langle X_N^2 \rangle}. \quad (4.30)$$

L'information mutuelle entre les variables  $X_A$  et  $X_B$  s'obtient en utilisant l'entropie d'une variable gaussienne (4.5), et l'additivité des entropies entre  $X_A$  et  $X_N$ . Puisque le signal et le bruit ne sont pas corrélés, on a  $H(X_A, X_B) = H(X_A) + H(GX_N)$ , et par conséquent :

$$I(A:B) = H(X_A) + H(X_B) - H(X_A, X_B) \quad (4.31a)$$

$$= H(X_B) - H(GX_N) \quad (4.31b)$$

$$= \frac{1}{2} \log_2 \frac{\langle X_B^2 \rangle}{G \langle X_N^2 \rangle} \quad (4.31c)$$

Cette information mutuelle peut également s'écrire en fonction du rapport signal-à-bruit : puisque  $\langle X_A X_N \rangle = 0$ , on a  $\langle X_B^2 \rangle = G(\langle X_A^2 \rangle + \langle X_N^2 \rangle)$ . On obtient alors la *formule de Shannon* :

$$I(A:B) = \frac{1}{2} \log_2 [1 + \text{SNR}] \quad (4.32)$$

Cette formule nous sera utile pour le calcul des taux secrets en cryptographie quantique, dans le chapitre 7.

4. Nous notons  $X$  une variable aléatoire classique, et  $x$  un tirage de cette variable aléatoire.

## 4.5 Conclusion

La théorie de l'information permet de répondre à deux problématiques. La première concerne les ressources nécessaires pour stocker un message constitué de symboles, tirés selon une certaine loi de probabilité. Puisque les symboles n'ont pas tous la même probabilité d'être choisis, certains sont moins probables que d'autres, et apportent ainsi plus d'information. L'entropie permet de quantifier cette notion, ainsi que les ressources nécessaires pour stocker une suite de symboles sans perte d'information à la limite asymptotique.

La seconde problématique concerne les communications, qu'elles soient classiques ou quantiques. Toujours en choisissant des symboles parmi un certain alphabet, quelle est la quantité d'information qui peut être extraite de l'envoi d'un message dans un canal de transmission imparfait ? La réponse à cette deuxième question a permis le développement de la cryptographie quantique, que nous présenterons dans le chapitre [7](#).

Deuxième partie

Résultats expérimentaux





## Chapitre 5

# Estimation de la discorde quantique pour un état EPR

### Sommaire

---

|            |                                                                              |            |
|------------|------------------------------------------------------------------------------|------------|
| <b>5.1</b> | <b>Introduction</b>                                                          | <b>83</b>  |
| <b>5.2</b> | <b>La discorde quantique</b>                                                 | <b>84</b>  |
| 5.2.1      | Une mesure des corrélations quantiques                                       | 84         |
| 5.2.2      | La discorde en information quantique                                         | 88         |
| <b>5.3</b> | <b>Protocole expérimental</b>                                                | <b>89</b>  |
| 5.3.1      | Principe                                                                     | 89         |
| 5.3.2      | Tri des quadratures et estimation des variances comprimées et anticomprimées | 90         |
| 5.3.3      | Estimation des incertitudes sur les variances comprimées et anticomprimées   | 93         |
| <b>5.4</b> | <b>Modélisation</b>                                                          | <b>94</b>  |
| 5.4.1      | Etat produit par l'OPA                                                       | 94         |
| 5.4.2      | Estimation de $N_s$ , $N_t$ et de leurs incertitudes, par inversion          | 95         |
| <b>5.5</b> | <b>Estimation de la discorde et de son incertitude</b>                       | <b>95</b>  |
| 5.5.1      | Discorde gaussienne pour un état thermique comprimé                          | 95         |
| 5.5.2      | Estimation par inversion                                                     | 97         |
| 5.5.3      | Estimation bayésienne                                                        | 97         |
| 5.5.4      | Résultats expérimentaux                                                      | 99         |
| <b>5.6</b> | <b>Comparaison avec les bornes de Cramér-Rao</b>                             | <b>100</b> |
| 5.6.1      | Information de Fisher et borne de Cramér-Rao classique                       | 101        |
| 5.6.2      | De l'information de Fisher classique à l'information de Fisher quantique     | 102        |
| 5.6.3      | Application à l'évaluation de la discorde                                    | 104        |
| 5.6.4      | Résultats expérimentaux                                                      | 106        |
| <b>5.7</b> | <b>Conclusion</b>                                                            | <b>107</b> |

---

### 5.1 Introduction

L'intrication a longtemps été considérée comme un élément essentiel en information quantique, et synonyme de corrélations quantiques. Elle n'est pourtant pas la seule propriété témoignant du caractère quantique des corrélations. De manière surprenante, certains mélanges

d'états non intriqués peuvent quand même présenter des corrélations supérieures aux corrélations classiques.

Plutôt que de mesurer l'intrication, on peut donc mesurer les *corrélations quantiques*, dont la définition est plus large. Une possibilité est de comparer les corrélations quantiques totales d'un état, que l'on peut obtenir en utilisant l'entropie de von Neumann, et les corrélations classiques que l'on a suite à une mesure quelconque d'un sous système. La différence de ces deux quantités correspond alors aux corrélations d'origine purement quantique, et définit *la discorde quantique*<sup>1</sup> [Ollivier01, Henderson01, Modi12].

La discorde inclut les corrélations dues à de l'intrication, mais sa définition est plus large puisque des mélanges d'états non intriqués peuvent quand même avoir une discorde non nulle. Introduite en 2001, elle a récemment fait l'objet de recherches intenses, lorsque la communauté scientifique a réalisé qu'elle pourrait être une ressource utile en information quantique. L'avantage, par rapport à l'intrication, est qu'elle est beaucoup moins "fragile", et beaucoup plus facile à produire expérimentalement. Même s'il n'y a pas encore de consensus sur son rôle précis par rapport à l'intrication, de nombreuses études ont montré un lien étroit entre la présence de discorde et certains protocoles d'information quantique. Nous en présenterons quelques-uns dans la section suivante.

Il n'existe pas d'observable associée à la discorde. Expérimentalement, la capacité à l'estimer précisément est maintenant un enjeu majeur, afin de pouvoir continuer l'investigation de son potentiel en information quantique. Dans ce chapitre, nous présentons l'évaluation expérimentale de la discorde pour des états appartenant à la classe des états thermiques comprimés, tels que ceux produits par notre OPA en configuration non dégénérée pour différentes puissances de pompe. Les deux modes sont mesurés avec deux détecteurs homodynes, nous permettant d'accéder aux caractéristiques des états sans passer par une tomographie quantique complète. Nous comparons ensuite la précision de notre estimation à la borne de Cramér-Rao classique, correspondant aux mesures homodynes, et à la borne de Cramér-Rao quantique [Helstrom76, Paris09, Braunstein94].

Ces bornes sont des limites fondamentales sur la variance minimale que doit avoir un estimateur. Nous montrons que notre estimateur de la discorde est presque optimal pour la borne classique, et qu'il présente un niveau de bruit raisonnable de 10 dB par rapport à la borne quantique. Nous montrons également qu'une méthode d'estimation bayésienne améliore la précision pour de faibles puissances de pompe. Ces travaux ont fait l'objet de la publication [Blandino12b], en collaboration avec Marco Genoni et Matteo Paris.

L'optimalité d'une estimation expérimentale de l'intrication a fait l'objet de travaux théoriques [Genoni08], et d'une démonstration expérimentale pour des variables discrètes [Brida10]. Cependant, les précédentes expériences concernant la discorde se basent sur une tomographie quantique ou une reconstruction indirecte de la matrice densité [Lanyon08, Chiuri11, Gu12, Madsen12, Dakić12, Passante11], sans étude précise sur l'optimalité des méthodes utilisées.

## 5.2 La discorde quantique

### 5.2.1 Une mesure des corrélations quantiques

#### Etats factorisables, classiques, séparables, et intriqués

Commençons par définir les états quantiques associés aux différentes corrélations possibles. Le premier cas est lorsqu'un état  $\hat{\pi}$  ne présente aucune corrélation entre ses  $N$  modes. Il s'agit

1. Nous ne précisons plus le terme "quantique" sauf quand cela sera nécessaire.

d'un *état factorisable*

$$\hat{\pi} = \hat{\pi}_1 \otimes \dots \otimes \hat{\pi}_N, \quad (5.1)$$

pour lequel chaque mode est complètement décorrélié des autres. Une mesure sur un des modes n'apprend donc rien sur les autres modes.

Le premier “degré” de corrélation correspond ensuite aux corrélations classiques, décrites par la théorie classique de l'information de Shannon. Un *état classique* est tel qu'il n'est pas perturbé par une mesure locale [Luo08], et doit donc s'écrire comme

$$\hat{\chi} = \sum_{\{k_n\}} p_{k_1, \dots, k_n} |k_1\rangle \langle k_1| \otimes \dots \otimes |k_N\rangle \langle k_N|, \quad (5.2)$$

où les états  $\{|k_n\rangle\}$  forment une base orthonormale.

Un mélange d'états classiques ne reste pas forcément classique, si les bases ne sont plus orthogonales. On obtient dans le cas général un *état séparable*

$$\hat{\sigma} = \sum_i p_i \hat{\pi}_1^{(i)} \otimes \dots \otimes \hat{\pi}_N^{(i)}, \quad (5.3)$$

qui peut être vu comme la réduction d'un état classique appartenant à un espace plus grand [Li08]. Un état séparable n'est pas intriqué, mais il peut pourtant avoir des corrélations supérieures aux corrélations classiques, et c'est ce que mesure la discordie.

Enfin, un *état intriqué* est un état qui n'appartient pas à l'ensemble des états séparables. La discordie inclut également les corrélations liées à l'intrication.

### Définition de la discordie quantique

Rappelons quelques résultats du chapitre 4. Pour deux variables aléatoires classiques  $A$  et  $B$ , l'information mutuelle

$$I(A:B) = H(A) + H(B) - H(A, B) \quad (5.4)$$

est une mesure de leurs corrélations. La généralisation quantique s'obtient naturellement avec l'entropie de von Neumann (4.10),

$$I(\hat{\rho}_{AB}) = S(\hat{\rho}_A) + S(\hat{\rho}_B) - S(\hat{\rho}_{AB}) \quad (5.5)$$

qui capture toutes les corrélations présentes dans  $\hat{\rho}_{AB}$ , qu'elles soient d'origine classique et/ou quantique [Cerf97, Adami97, Modi12, Vedral02].

Puisqu'un état classique n'est pas perturbé par une mesure d'un des sous-systèmes, (5.4) s'écrit également sous la forme

$$I(A:B) = H(A) - H(A|B) = H(B) - H(B|A), \quad (5.6)$$

qui correspond à l'information gagnée sur  $A$  suite à la mesure de  $B$ , et vice versa. La généralisation quantique de (5.6) n'est en revanche pas immédiate, puisqu'en général une mesure perturbe un état quantique. En conséquence, la symétrie entre les deux sous-systèmes est rompue, la quantité obtenue dépend du type de mesure, et elle peut être différente de (5.5). C'est précisément cette différence que mesure la discordie.

Supposons que Bob fasse une mesure à l'aide d'un POVM  $\{E_k\}$ . Pour un résultat de mesure  $k$ , la matrice densité d'Alice est

$$\hat{\rho}_{A|k} = \frac{\text{Tr}_B\{E_k \hat{\rho}_{AB}\}}{p_B(k)}, \quad (5.7)$$

avec  $p_B(k) = \text{Tr}_{AB}\{E_k \hat{\rho}_{AB}\}$ . On peut alors définir une version quantique de (5.6) comme

$$J(A|\{E_k\}) = S(\hat{\rho}_A) - \sum_k p_B(k) S(\hat{\rho}_{A|k}), \quad (5.8)$$

qui correspond à l'information mutuelle classique obtenue avec le POVM  $\{E_k\}$  [Henderson01]. Les corrélations classiques totales de  $\hat{\rho}_{AB}$  sont obtenues en maximisant  $J(A|\{E_k\})$  sur tous les POVM  $\{E_k\}$  :

$$J(A|B) = \sup_{\{E_k\}} J(A|\{E_k\}) \quad (5.9)$$

La *discorde quantique* est alors définie comme la différence entre les corrélations totales  $I(\hat{\rho}_{AB})$  et les corrélations classiques  $J(A|B)$  :

$$D(A|B) = I(\hat{\rho}_{AB}) - J(A|B) \quad (5.10)$$

Elle correspond donc aux corrélations d'origine purement quantique. Elle est toujours positive ou nulle, et dans le cas général, elle peut être asymétrique. Cette asymétrie reflète une asymétrie dans la répartition des corrélations classiques et quantiques, pour Alice et Bob, bien que les corrélations totales soient symétriques. Notons que presque tous les états quantiques ont une discorde non nulle [Ferraro10].

Pour un état bipartite pur, la discorde est égale à l'intrication, mesurée par l'entropie d'un des états réduits [Modi12], qui est d'ailleurs la seule mesure d'intrication dans ce cas [Horodecki09]. La différence entre l'intrication et la discorde apparaît pour les états impurs. S'ils contiennent de l'intrication, celle-ci est d'une certaine manière contenue dans la discorde, bien qu'une relation directe entre ces deux grandeurs ne soit pas immédiate. La raison est que l'intrication peut être mesurée de plusieurs façons pour un mélange statistique, chacune ayant une interprétation opérationnelle différente [Horodecki09].

Une solution possible afin d'unifier discorde et intrication est par exemple d'adopter une définition différente de la discorde, en utilisant une interprétation géométrique, avec comme mesure de distance l'entropie relative [Modi10], ou la norme de Hilbert Schmidt [Dakić10, Bellomo12]. Chaque propriété, corrélations totales, discorde, intrication, est alors mesurée par la distance entre l'état quantique, et un état ne possédant pas la propriété en question. Cette définition conduit à une mesure des corrélations quantiques différente de (5.10), bien que le concept reste similaire.

Par exemple, la référence [Modi10] définit l'entropie relative de discorde d'un état  $\hat{\rho}$  comme étant la distance minimale entre cet état et un état classique  $\hat{\chi}$  défini par (5.2) :  $D_{\text{rel}} = \min_{\hat{\chi}} S(\hat{\rho} \parallel \hat{\chi})$ , où  $S(\cdot \parallel \cdot)$  est l'entropie relative définie par (4.11). L'intrication est quantifiée avec une mesure similaire, l'entropie relative d'intrication [Plenio05, Horodecki09] :  $E_{\text{rel}} = \min_{\hat{\sigma}} S(\hat{\rho} \parallel \hat{\sigma})$ , où  $\hat{\sigma}$  est un état séparable défini par (5.3). Enfin les corrélations totales  $C$  mesurées par l'entropie de von Neumann peuvent aussi s'exprimer en fonction de l'entropie relative,  $C = S(\hat{\rho} \parallel \hat{\rho}_1 \otimes \dots \otimes \hat{\rho}_N)$ , comme nous l'avons vu avec (4.12).

Il existe encore d'autres façons de définir les corrélations quantiques, et nous renvoyons le lecteur à la référence [Modi12] pour une présentation plus complète. Dans tout ce chapitre, la discordie considérée<sup>2</sup> sera bien celle définie par (5.10).

### Exemple

Illustrons maintenant le concept de discordie quantique et son asymétrie par un exemple simple. L'état

$$\hat{\rho}_{AB} = p|0\rangle\langle 0|\otimes|0\rangle\langle 0| + (1-p)|+\rangle\langle +|\otimes|1\rangle\langle 1| \quad (5.11)$$

est séparable, et donc non intriqué. En revanche,  $D(A|B)=0$  et  $D(B|A)\neq 0$ . Pour le montrer, supposons pour commencer que Bob fasse une mesure dans la base  $\{|0\rangle, |1\rangle\}$ . Avec une probabilité  $p$  il obtient  $|0\rangle$ , et il sait que l'état d'Alice est  $|0\rangle$ , ou bien avec une probabilité  $1-p$  il obtient  $|1\rangle$  et il sait que l'état d'Alice est  $|+\rangle$ . L'information conditionnelle donnée par (5.8), associée aux projecteurs  $E_0=|0\rangle\langle 0|$  et  $E_1=|1\rangle\langle 1|$ , est donc

$$J(A|\{E_k\}) = S(\hat{\rho}_A), \quad (5.12)$$

puisque  $S(\hat{\rho}_{A0}=|0\rangle\langle 0|)=S(\hat{\rho}_{A1}=|+\rangle\langle +|)=0$ . Puisque selon la relation (4.18) on a  $I(\hat{\rho}_{AB})=S(\hat{\rho}_A)$ , la discordie  $D(A|B)$  est donc nulle. Bob peut accéder à toute l'information localement, sans perturber l'état, ce qui est d'ailleurs une conséquence de la relation

$$E_0\hat{\rho}_{AB}E_0 + E_1\hat{\rho}_{AB}E_1 = \hat{\rho}_{AB}. \quad (5.13)$$

Dans le cas où la mesure est faite par Alice, les choses sont différentes, puisqu'elle ne pourra pas discerner avec certitude ses deux états  $|0\rangle$  et  $|+\rangle$ , et obtiendra donc moins d'information que Bob. Supposons par exemple qu'elle fasse une mesure projective dans la base  $\{|0\rangle, |1\rangle\}$ . Si elle obtient  $|1\rangle$ , elle sait que l'état de Bob est également  $|1\rangle$ . Mais si elle obtient  $|0\rangle$ , de son point de vue l'état de Bob est (sans normaliser)  $p|0\rangle\langle 0| + (1-p)/2|1\rangle\langle 1|$ . Elle a donc plus d'incertitude concernant l'état de Bob après sa mesure que Bob n'en avait à propos de son état. Elle pourrait également utiliser un POVM distinguant  $|0\rangle$  et  $|+\rangle$  avec certitude, mais de manière non déterministe [Nielsen00]. Pour une mesure dans la base  $\{|0\rangle, |1\rangle\}$ ,

$$J(B|\{E_k\}) = S(\hat{\rho}_B) - \sum_a p_a \hat{\rho}_{B|a} = H(P) - \frac{1+p}{2} S(\hat{\rho}_{B|0}) - \frac{1-p}{2} S(\hat{\rho}_{B|1}), \quad (5.14)$$

avec  $H(P)$  l'entropie binaire de Shannon,  $\hat{\rho}_{B|0}=\frac{2p}{1+p}|0\rangle\langle 0| + \frac{1-p}{1+p}|1\rangle\langle 1|$  et  $\hat{\rho}_{B|1}=|1\rangle\langle 1|$ . En prenant par exemple  $p=0.3$ , on a  $H(P)\simeq 0.72$ ,  $S(\hat{\rho}_A)\simeq 0.43$ ,  $S(\hat{\rho}_{B|0})\simeq 0.92$  et bien sûr  $S(\hat{\rho}_{B|1})=0$ . On a donc  $J(B|\{E_k\})\simeq 0.72-0.6\times 0.92=0.17$ , qui est inférieur à 0.43. Pour cette mesure, la discordie (non optimisée) vaut donc  $0.43-0.17=0.26$ . En optimisant, la discordie finale sera peut être plus faible, mais non nulle [Modi12].

Dans un cas un peu plus général, on montre que la discordie  $D(A|B)$  est toujours nulle pour un état quantique-classique [Modi12] :

$$\boxed{\hat{\rho}_{AB} = \sum_b p_b \hat{\rho}_{A|b} \otimes |k_b\rangle\langle k_b| \Leftrightarrow D(A|B)=0} \quad (5.15)$$

2. En fait, nous utiliserons la discordie gaussienne [Giorda10, Adesso10], pour laquelle l'optimisation porte uniquement sur des POVM gaussiens. Nous reviendrons sur ce point plus loin dans ce chapitre.

où  $\{|k_b\rangle\}$  forme une base orthonormale. Là encore, la raison est que Bob peut accéder localement à l'information sans perturber le système, contrairement à Alice.

Contrairement à l'intrication, très fragile, la discorde est plus robuste, et peut même augmenter par un couplage avec un environnement, ce qui en fait une ressource particulièrement intéressante d'un point de vue expérimental [Ciccarello12a, Ciccarello12b].

### 5.2.2 La discorde en information quantique

La discorde est activement étudiée en raison de son rôle potentiel dans plusieurs protocoles d'information quantique. Nous en présentons quelques-uns, en renvoyant encore une fois le lecteur à la référence [Modi12] pour une liste plus détaillée.

**Protocole DQC1** En 1998, E. Knill et R. Laflamme réalisèrent que certaines tâches peuvent être réalisées exponentiellement plus rapidement qu'avec un ordinateur classique [Knill98], sans toutefois nécessiter d'intrication. Leur protocole de calcul quantique DQC1 utilise un seul qubit de pureté non nulle, couplé à un ensemble de qubits dans des états complètement dépolarisés. C'est un modèle non universel car il ne peut pas implémenter n'importe quel algorithme, mais il permet par exemple de calculer la trace (normalisée par  $2^n$ ) d'un opérateur unitaire  $\hat{U}_n$  à  $n$  qubits.

A. Datta *et al.* ont montré que l'intrication était exponentiellement petite [Datta05, Datta07], et que la discorde pouvait être à l'origine de l'efficacité du protocole [Datta08, Merali11], bien que certains contre exemples aillent à l'encontre de cette interprétation [Dakić10]. Ce protocole a fait l'objet de plusieurs études expérimentales [Lanyon08, Passante11].

**Intrication** La discorde peut être “convertie” en intrication. M. Piani *et al.* ont montré que la discorde, mesurée par l'entropie relative, peut être utilisée pour la création d'intrication multipartite [Piani11]. A. Streltsov *et al.* ont également montré un lien entre la discorde et la génération d'intrication par une mesure [Streltsov11].

**State merging** Le “state merging” est le transfert de l'état d'Alice à Bob, gardant la cohérence avec un système auxiliaire  $C$  [Horodecki05]. Cette tâche requiert l'utilisation de  $S(A|B)$  ebits (états maximales intriqués) d'information. V. Madhok *et al.* ont montré que la discorde est liée à la différence de quantité d'intrication nécessaire suivant que Bob garde ou non ses corrélations avec le système  $C$  [Madhok11]. Un protocole similaire, “l'extended state merging”, inclut de plus la quantité d'intrication nécessaire pour la formation de l'état  $AB$  : la discorde  $D(A|C)$  est alors reliée à la quantité d'intrication nécessaire pour la formation de l'état et le state merging [Cavalcanti11].

**Complete positivity** La discorde peut être reliée à l'évolution d'un système : l'évolution d'un système couplé à un environnement est décrite par opération complètement positive si le système et l'environnement sont initialement dans un état de discorde nulle [Rodríguez-Rosario08, Shabani09]. Un article récent semble toutefois indiquer que cette condition ne serait pas toujours nécessaire [Brodutch12].

**Autre études expérimentales** Citons pour finir quelques autres études expérimentales. Dans la référence [Chiuri11], A. Chiuri *et al.* ont produit et analysé différents types d'états possédant des corrélations non classiques. M. Gu *et al.* ont montré théoriquement et expérimentalement

que l'on peut relier la discordie d'un état bipartite au gain d'information apporté par des mesures globales, plutôt que par des mesures locales sur chacun des modes [Gu12]. Citons enfin B. Dakić *et al.* [Dakić12], qui interprètent une mesure géométrique de la discordie [Dakić10] comme ressource pour la “remote state preparation” [Bennett01], dont le principe est similaire à une téléportation où Alice connaît l'état à téléporter.

## 5.3 Protocole expérimental

### 5.3.1 Principe

Notre étude a pour but d'estimer expérimentalement la discordie avec des mesures homodynes pour des états gaussiens [Giorda10, Adesso10], et de comparer son incertitude aux limites fondamentales données par les bornes de Cramér-Rao. Nous avons choisi d'utiliser des états comprimés bimodes produits par notre OPA en configuration non dégénérée, à la base de nombreuses expériences d'optique quantique avec des variables continues. En fonction de leur pureté, ces états peuvent présenter ou non de l'intrication, mais en revanche ils présentent toujours de la discordie, sauf s'il n'y a pas de compression et qu'ils sont factorisables [Giorda10].

Le dispositif expérimental est schématisé sur la figure 5.1. En faisant varier l'angle de la lame  $\lambda/2$  devant le premier cristal doubleur de la SHG, on fait varier la puissance de la pompe, et donc le facteur de compression du vide. La puissance maximale de 4.3 mW est obtenue pour un angle de  $60^\circ$ . On modifie ensuite cet angle par pas de  $2^\circ$ , modifiant le facteur de compression de l'OPA selon la relation

$$r(\theta) = r_{60} \cos^2 \left( \frac{\pi}{90} [60 - \theta] \right), \quad (5.16)$$

où  $\theta$  est l'angle de la  $\lambda/2$ , et  $r_{60}$  est le facteur de compression pour  $\theta=60^\circ$  (estimé à environ 0.32, comme expliqué plus loin dans ce chapitre). Lorsque  $r(\theta)$  est trop faible, les caractéristiques de l'état sont plus difficiles à estimer à partir des mesures : on se limite donc à  $\theta=32^\circ$ , correspondant à  $r \simeq 0.1$ .

Nous utilisons deux détecteurs homodynes pour estimer les variances comprimées et anti-comprimées. Prise séparément, chaque détection mesure un état thermique. En revanche, les corrélations apparaissent lorsqu'on s'intéresse à la variance de la différence des quadratures, comme nous l'avons vu dans les sections 3.3.2 et 2.5.5. Plus précisément, les quadratures<sup>3</sup>

$$\hat{Q}^{(1)} = \frac{\hat{X}_0 + \hat{X}_1}{\sqrt{2}}, \quad \hat{Q}^{(4)} = \frac{\hat{P}_0 - \hat{P}_1}{\sqrt{2}}, \quad (5.17)$$

sont comprimées avec une variance  $\sigma^2(Q_{sq})$ , alors que les quadratures

$$\hat{Q}^{(2)} = \frac{\hat{X}_0 - \hat{X}_1}{\sqrt{2}}, \quad \hat{Q}^{(3)} = \frac{\hat{P}_0 + \hat{P}_1}{\sqrt{2}}, \quad (5.18)$$

sont anticomprimées avec une variance  $\sigma^2(Q_{asq})$ .

Ces variances sont mesurées expérimentalement, et constituent le point de départ de notre méthode d'estimation de la discordie.

---

3. La notation  $\hat{Q}$  sera réservée aux quadratures qui sont des combinaisons des quadratures des deux détecteurs homodynes. Lorsque nous parlerons de la mesure d'une quadrature  $\hat{Q}$ , il sera sous entendu que cette mesure est effectuée avec les deux détecteurs homodynes.



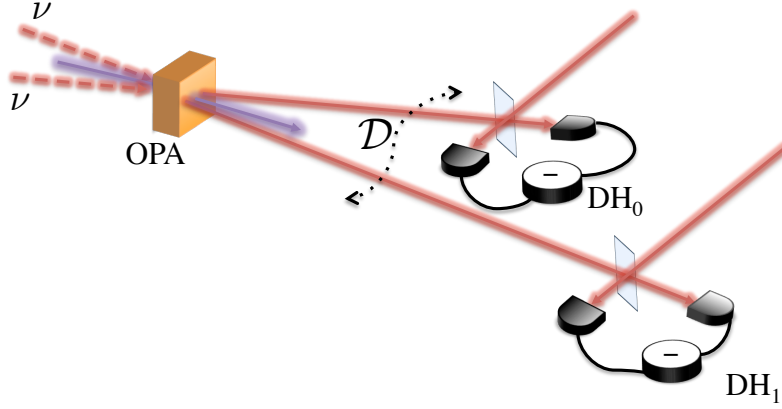


FIGURE 5.1 – Principe de l'expérience.  $\hat{\mathbf{D}}$  correspond à un état thermique fictif utilisé pour modéliser l'OPA, comme expliqué dans la section 5.4.1.

### 5.3.2 Tri des quadratures et estimation des variances comprimées et anti-comprimées

#### Tri des quadratures

Pour chaque angle de la  $\lambda/2$ , on enregistre  $1.6 \times 10^6$  mesures pour chacune des deux détecteurs homodynes (que nous appelons DH0 et DH1), en balayant la phase de l'oscillateur local à une fréquence de l'ordre d'une dizaine de Hz, avant sa séparation vers les deux détecteurs. En utilisant une  $\lambda/2$  après l'OPA, on rend les polarisations des deux modes de l'état EPR orthogonales, de manière à être séparées par un PBS avant les deux détecteurs homodynes.

Les états EPR possèdent une propriété fort intéressante : ils sont invariants vis-à-vis d'un déphasage relatif  $2\varphi$  entre leurs deux modes  $\hat{\mathbf{a}}$  et  $\hat{\mathbf{b}}$  (correspondant aux quadratures indicées 0 et 1). En effet, en se souvenant qu'ils sont obtenus avec l'opérateur  $\hat{\mathbf{S}}_2(r)$  (2.149), on voit qu'un déphasage  $\hat{\mathbf{U}}_a(\varphi)$  appliqué sur le mode  $\hat{\mathbf{a}}$ , et  $\hat{\mathbf{U}}_b(-\varphi)$  appliqué sur le mode  $\hat{\mathbf{b}}$ , se traduit par

$$\hat{\mathbf{U}}_a(\varphi)\hat{\mathbf{U}}_b(-\varphi)\hat{\mathbf{S}}_2(r)|0\rangle = \hat{\mathbf{U}}_a(\varphi)\hat{\mathbf{U}}_b(-\varphi)\hat{\mathbf{S}}_2(r)\hat{\mathbf{U}}_a^\dagger(\varphi)\hat{\mathbf{U}}_b^\dagger(-\varphi)|0\rangle. \quad (5.19)$$

En utilisant la formule (A.1), on peut inclure ces déphasages dans l'expression de  $\hat{\mathbf{S}}_2(r)$ . Puisque  $\hat{\mathbf{a}}$  est transformé en  $\hat{\mathbf{a}}e^{+i\varphi}$  et  $\hat{\mathbf{b}}$  est transformé en  $\hat{\mathbf{b}}e^{-i\varphi}$ , ces termes de phases se compensent et on a finalement

$$\hat{\mathbf{U}}_a(\varphi)\hat{\mathbf{U}}_b(-\varphi)\hat{\mathbf{S}}_2(r)|0\rangle = \hat{\mathbf{S}}_2(r)|0\rangle \quad (5.20)$$

La variance de  $\hat{\mathbf{X}}_0(\varphi) + \hat{\mathbf{X}}_1(-\varphi)$  sera donc égale à celle de  $\hat{\mathbf{X}}_0 + \hat{\mathbf{X}}_1$ , et de même pour les autres quadratures  $\hat{\mathbf{Q}}^{(k)}$  définies en (5.17) et (5.18). Si les deux modes sont déphasés de deux quantités différentes  $\varphi_0$  et  $\varphi_1$ , on peut se ramener au cas précédent en introduisant un déphasage commun de  $(\varphi_0 + \varphi_1)/2$ , et un déphasage relatif de  $\pm(\varphi_0 - \varphi_1)/2$ . Ce déphasage commun correspond à une simple évolution libre. En fin de compte, nous n'avons pas besoin de contrôler la phase relative des deux oscillateurs locaux. On considérera simplement que le minimum de variance pour la différence des mesures correspond à la mesure  $\hat{\mathbf{Q}}^{(4)}$ .

La méthode utilisée pour extraire les mesures des quadratures  $\hat{\mathbf{X}}_0, \hat{\mathbf{P}}_0, \hat{\mathbf{X}}_1, \hat{\mathbf{P}}_1$  des données brutes, pour chaque angle de la  $\lambda/2$ , est décrite ci dessous :

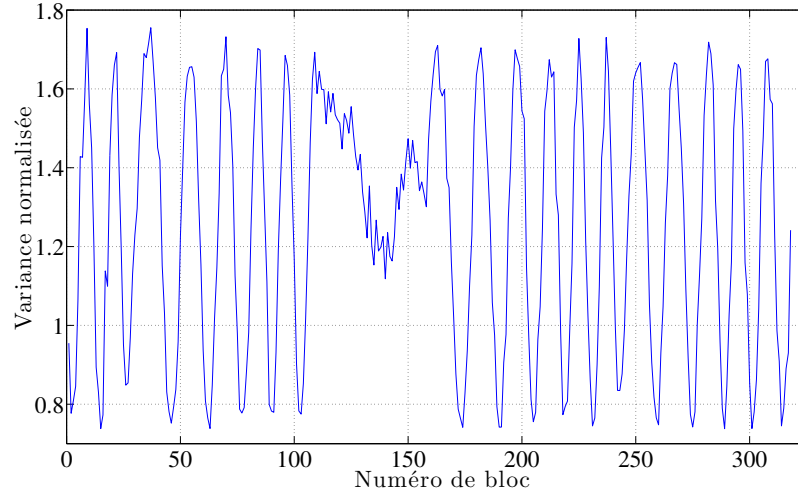


FIGURE 5.2 – Variance normalisée de  $\frac{1}{\sqrt{2}}(\hat{\mathbf{X}}_{\theta,1} - \hat{\mathbf{X}}_{\theta,2})$ , calculée par blocs de 5000 points, pour un angle de la  $\lambda/2$  de  $60^\circ$  donnant la puissance de pompe maximale. Les données correspondant aux blocs 105 à 160 ne sont pas prises en compte pour l'estimation de la discorde.

- Afin de limiter les dérives, on regroupe les données en blocs de 1000 points, et on calcule la moyenne de chaque bloc qui est ensuite soustraite du bloc. Les données sont ensuite dégroupées pour être traitées comme des données brutes.
- Les quadratures des deux détectations homodynes sont normalisées en les divisant par l'écart-type mesuré pour le vide. Ce choix correspond à  $N_0=1$ .
- On calcule ensuite la variance par blocs de 5000 points de la différence des quadratures mesurées par les deux détectations homodynes. En traçant ces variances, des oscillations apparaissent : les minimums correspondent aux valeurs de  $\hat{\mathbf{Q}}^{(4)}$  et les maximums aux valeurs de  $\hat{\mathbf{Q}}^{(2)}$ . Un exemple est donnée sur la figure 5.2.
- On sélectionne ensuite graphiquement et manuellement les zones où les oscillations sont régulières. Les zones “chaotiques” correspondent à des effets parasites dûs à la cale piézoélectrique, et les mesures correspondantes sont supprimées. On supprime également les 10000 premières mesures afin d'éviter les régimes transitoires.
- Un programme sous Matlab repère ensuite les positions des blocs associées aux minimums et aux maximums. Ces positions sont utilisées afin d'extraire les blocs correspondant aux quadratures  $\hat{\mathbf{X}}_0$ ,  $\hat{\mathbf{P}}_0$ ,  $\hat{\mathbf{X}}_1$  et  $\hat{\mathbf{P}}_1$  des données brutes. La taille de bloc de 5000 points est un compromis entre une taille suffisamment grande pour limiter les fluctuations statistiques lors du calcul de la variance, tout en étant suffisamment petite afin que l'on puisse considérer que la phase de l'oscillateur local varie peu, et que toutes les mesures du bloc correspondent à une même quadrature.
- Les différents blocs des quatre quadratures sont regroupés en quatre ensembles, qui correspondent aux mesures de  $\hat{\mathbf{X}}_0$ ,  $\hat{\mathbf{P}}_0$ ,  $\hat{\mathbf{X}}_1$ ,  $\hat{\mathbf{P}}_1$ .

### Estimation des variances comprimées et anticomprimées

Les quadratures triées sont ensuite utilisées pour calculer les variances comprimées et anticomprimées. Nous noterons respectivement  $V_{sq}$  et  $V_{asq}$  les estimateurs de  $\sigma^2(Q_{sq})$  et  $\sigma^2(Q_{asq})$ . La méthode utilisée est la suivante :

| Angle | $V_{sq}$ | $V_{asq}$ |
|-------|----------|-----------|
| 60    | 0.768    | 1.698     |
| 58    | 0.763    | 1.681     |
| 56    | 0.760    | 1.645     |
| 54    | 0.757    | 1.610     |
| 52    | 0.784    | 1.605     |
| 50    | 0.791    | 1.561     |
| 48    | 0.789    | 1.521     |
| 46    | 0.793    | 1.466     |
| 44    | 0.786    | 1.415     |
| 42    | 0.806    | 1.378     |
| 40    | 0.828    | 1.323     |
| 38    | 0.824    | 1.262     |
| 36    | 0.851    | 1.198     |
| 34    | 0.8761   | 1.174     |
| 32    | 0.895    | 1.127     |

TABLE 5.1 – Variances  $V_{sq}$  et  $V_{asq}$  normalisées, avec  $N_0=1$ .

- Pour des raisons de temps de calcul avec la méthode bayésienne que nous détaillerons plus loin, on ne garde que les  $M_q=2 \times 10^4$  premières mesures pour chaque quadrature  $\hat{\mathbf{X}}_0, \hat{\mathbf{P}}_0, \hat{\mathbf{X}}_1, \hat{\mathbf{P}}_1$ , que l'on regroupe respectivement dans des ensembles  $\mathbf{x}_0=\{x_1^0, \dots, x_{M_q}^0\}$ ,  $\mathbf{p}_0=\{p_1^0, \dots, p_{M_q}^0\}$ ,  $\mathbf{x}_1=\{x_1^1, \dots, x_{M_q}^1\}$  et  $\mathbf{p}_1=\{p_1^1, \dots, p_{M_q}^1\}$ .
- On forme ensuite les ensembles  $\mathbf{q}^{(k)}$  contenant les données des quadratures  $\hat{\mathbf{Q}}^{(k)}$ , en utilisant les données précédentes. On note :

$$\mathbf{q}^{(1)} = \frac{\mathbf{x}_0 + \mathbf{x}_1}{\sqrt{2}} \quad \mathbf{q}^{(4)} = \frac{\mathbf{p}_0 - \mathbf{p}_1}{\sqrt{2}} \quad (5.21a)$$

$$\mathbf{q}^{(2)} = \frac{\mathbf{x}_0 - \mathbf{x}_1}{\sqrt{2}} \quad \mathbf{q}^{(3)} = \frac{\mathbf{p}_0 + \mathbf{p}_1}{\sqrt{2}} \quad (5.21b)$$

- On regroupe ensuite les données  $\mathbf{q}^{(1)}$  et  $\mathbf{q}^{(4)}$  des quadratures comprimées dans un ensemble

$$\mathbf{q}_{sq} = \{q_1^{(1)}, \dots, q_{M_q}^{(1)}, q_1^{(4)}, \dots, q_{M_q}^{(4)}\}, \quad (5.22)$$

et les données  $\mathbf{q}^{(2)}$  et  $\mathbf{q}^{(3)}$  des quadratures anticomprimées dans un ensemble

$$\mathbf{q}_{asq} = \{q_1^{(2)}, \dots, q_{M_q}^{(2)}, q_1^{(3)}, \dots, q_{M_q}^{(3)}\}. \quad (5.23)$$

Chacun de ces deux ensembles contient donc  $2 \times M_q = 4 \times 10^4$  mesures.

- On calcule ensuite les variances  $V_{sq}$  et  $V_{asq}$  à partir de  $\mathbf{q}_{sq}$  et  $\mathbf{q}_{asq}$  :

$$V_{sq} = \text{Var}(\mathbf{q}_{sq}) \quad (5.24a)$$

$$V_{sq} = \text{Var}(\mathbf{q}_{asq}) \quad (5.24b)$$

Les variances obtenues sont données dans la table 5.1.

### 5.3.3 Estimation des incertitudes sur les variances comprimées et anticomprimées

#### Fluctuations statistiques d'une variable aléatoire gaussienne

Afin de calculer l'incertitude sur l'estimation de la discorde, nous avons besoin de connaître l'incertitude sur notre estimation de  $V_{sq}$  et  $V_{asq}$ . Pour cela, commençons par rappeler quelques résultats de statistiques. Considérons une variable aléatoire  $X$  suivant une distribution normale, de variance  $\sigma^2$  et de moyenne  $\mu$ . On réalise  $N$  tirages indépendants de  $X$ , donnant les résultats  $\mathbf{x} := \{x_1, \dots, x_N\}$ , et on cherche à estimer  $\sigma^2$  à partir de ces mesures. Un estimateur sans biais est donné par  $\sigma^{*2} = \frac{1}{n-1} \sum_{k=1}^N (x_k - \bar{\mathbf{x}})^2$ , avec  $\bar{\mathbf{x}} = \frac{1}{n} \sum_{k=1}^N x_k$  la moyenne de l'échantillon. On peut alors montrer [Knight99] que la variable aléatoire

$$\frac{\sum_{k=1}^N (X_k - \bar{\mathbf{x}})^2}{\sigma^2} = (n-1) \frac{\sigma^{*2}}{\sigma^2} \quad (5.25)$$

suit une loi du  $\chi^2$  à  $(n-1)$  degrés de liberté<sup>4</sup>. Il en résulte qu'elle a une variance égale à  $2(n-1)$ , et donc que  $\sigma^{*2}$  a une variance égale à

$$\text{Var}(\sigma^{*2}) = \frac{\sigma^4}{(n-1)^2} 2(n-1) = \frac{2}{n-1} \sigma^4. \quad (5.26)$$

De plus, pour un nombre de degrés de liberté  $n \rightarrow \infty$  (ce qui sera largement vérifié pour  $n=2M_q$ ), une distribution  $\chi^2$  tend vers une distribution gaussienne de moyenne  $n$  et de variance  $2n$ . On peut donc écrire

$$n \frac{\sigma^{*2}}{\sigma^2} = n + U\sqrt{2n}, \quad (5.27)$$

où  $U$  est une variable normale centrée réduite, et où l'on a supposé  $n-1 \simeq n$ . Considérons maintenant la valeur  $u_p$ , telle que  $U$  vérifie  $-u_p \leq U \leq u_p$  avec une probabilité  $p$ . On a alors la relation

$$1 - u_p \sqrt{\frac{2}{n}} \leq \frac{\sigma^{*2}}{\sigma^2} \leq 1 + u_p \sqrt{\frac{2}{n}} \quad (5.28)$$

avec une probabilité  $p$ , ce qui donne

$$\frac{\sigma^{*2}}{1 + u_p \sqrt{\frac{2}{n}}} \leq \sigma^2 \leq \frac{\sigma^{*2}}{1 - u_p \sqrt{\frac{2}{n}}}. \quad (5.29)$$

Pour  $n$  grand, un développement des dénominateurs donne

$$\sigma^{*2} \left( 1 - u_p \sqrt{\frac{2}{n}} \right) \leq \sigma^2 \leq \sigma^{*2} \left( 1 + u_p \sqrt{\frac{2}{n}} \right). \quad (5.30)$$

Cela correspond au même intervalle de confiance qu'en faisant l'hypothèse que la variance  $\sigma^2$  suit une loi normale de moyenne  $\sigma^{*2}$  et de variance  $2\sigma^{*4}/n$ . Nous utiliserons donc cette approche pour estimer les incertitudes sur les variances mesurées  $V_{sq}$  et  $V_{asq}$ .

4. Soient  $U_1, \dots, U_k$ ,  $k$  variables aléatoires indépendantes qui suivent chacune une loi normale centrée réduite. On rappelle que par définition, la variable  $\chi_k^2 = U_1^2 + \dots + U_k^2$  suit une loi du  $\chi^2$  à  $k$  degrés de libertés. Sa moyenne et sa variance valent respectivement  $k$  et  $2k$ .

### Application à nos mesures

Les deux ensembles  $\mathbf{q}_{\text{sq}}$  et  $\mathbf{q}_{\text{asq}}$  contenant chacun  $2M_q=4\times 10^4$  mesures, les conditions de validité du paragraphe précédent sont bien vérifiées. L'hypothèse selon laquelle  $V_{\text{sq}}$  et  $V_{\text{asq}}$  suivent des lois gaussiennes est également bien vérifiée pour deux raisons. La première est que les mesures de chaque quadrature, pour chaque puissance de pompe, sont déjà très proches de distributions gaussiennes. La seconde est que le nombre de mesures est suffisamment grand pour utiliser le théorème central limite.

On considère donc que les variances mesurées  $V_{\text{sq}}$  et  $V_{\text{asq}}$  proviennent respectivement de la réalisation de deux variables aléatoires  $\tilde{V}_{\text{sq}}$  et  $\tilde{V}_{\text{asq}}$ , de moyennes  $V_{\text{sq}}$  et  $V_{\text{asq}}$  et de variances

$$\sigma^2(V_{\text{sq}}) = \frac{2V_{\text{sq}}^2}{2M_q}, \quad (5.31a)$$

$$\sigma^2(V_{\text{asq}}) = \frac{2V_{\text{asq}}^2}{2M_q}. \quad (5.31b)$$

On peut donc écrire formellement  $\tilde{V}_{\text{sq}}$  et  $\tilde{V}_{\text{asq}}$  comme

$$\tilde{V}_{\text{sq}} = V_{\text{sq}} + U_{\text{sq}}\sigma(V_{\text{sq}}), \quad (5.32a)$$

$$\tilde{V}_{\text{asq}} = V_{\text{asq}} + U_{\text{asq}}\sigma(V_{\text{asq}}), \quad (5.32b)$$

où  $U_{\text{sq}}$  et  $U_{\text{asq}}$  sont deux variables normales centrées réduites. Cette modélisation nous sera utile pour propager les incertitudes jusqu'à l'évaluation de la discorde, comme nous le montrerons plus loin dans ce chapitre.

## 5.4 Modélisation

### 5.4.1 Etat produit par l'OPA

Nous avons vu au chapitre présentant les outils expérimentaux que l'état produit par l'OPA peut être modélisé avec deux paramètres  $r$  et  $\gamma$ , reliés respectivement à la compression et au gain parasite. De manière plus générale, cet état appartient à la classe des états thermiques comprimés. Il peut être décrit par la compression de deux états thermiques  $\hat{\nu}(N_t)$  fictifs en entrée d'un OPA parfait de paramètre de compression  $\zeta$  :

$$\hat{\rho} = \hat{S}_2(\zeta) \left[ \hat{\nu}(N_t) \otimes \hat{\nu}(N_t) \right] \hat{S}_2^\dagger(\zeta) \quad (5.33)$$

avec  $\hat{S}_2(\zeta) = \exp[\zeta(\hat{a}\hat{b} - \hat{a}^\dagger\hat{b}^\dagger)]$ . Les nombres  $N_s = \sinh^2 \zeta$  et  $N_t$ , sont respectivement le nombre de photons comprimés et de photons thermiques. Les photons thermiques viennent des imperfections expérimentales qui se traduisent par une dégradation des corrélations.

Ces deux paramètres peuvent inclure le gain parasite dû à  $\gamma$ , mais aussi l'efficacité homodyne  $\eta$ . Ils paramétrisent entièrement l'état mesuré, dont la matrice de covariance s'écrit [Ferraro05]

$$\mathbf{\Gamma} = \begin{pmatrix} A\mathbb{I} & -C\mathbb{Z} \\ -C\mathbb{Z} & A\mathbb{I} \end{pmatrix}, \quad (5.34)$$

où l'on rappelle que  $\mathbb{Z} = \text{diag}(1, -1)$ , et avec

$$A = (1 + 2N_s)(1 + 2N_t), \quad (5.35a)$$

$$C = 2(1 + 2N_t)\sqrt{N_s(1 + N_s)}. \quad (5.35b)$$

Les variances des quadratures comprimées et anticomprimées s'expriment ensuite en fonction de  $N_s$  et  $N_t$  :

$$\sigma^2(Q_{sq}) = \left[ 1 + 2N_s - 2\sqrt{N_s(1 + N_s)} \right] (1 + 2N_t) \quad (5.36a)$$

$$\sigma^2(Q_{asq}) = \left[ 1 + 2N_s + 2\sqrt{N_s(1 + N_s)} \right] (1 + 2N_t) \quad (5.36b)$$

En inversant les formules (5.36a) et (5.36b), on obtient les expressions de  $N_s$  et  $N_t$  en fonction de  $\sigma^2(Q_{sq})$  et  $\sigma^2(Q_{asq})$  :

$$N_s = \frac{1}{4} \left( \sqrt{\frac{\sigma^2(Q_{sq})}{\sigma^2(Q_{asq})}} + \sqrt{\frac{\sigma^2(Q_{asq})}{\sigma^2(Q_{sq})}} - 2 \right) \quad (5.37a)$$

$$N_t = \frac{\sqrt{\sigma^2(Q_{sq})\sigma^2(Q_{asq})} - 1}{2} \quad (5.37b)$$

#### 5.4.2 Estimation de $N_s$ , $N_t$ et de leurs incertitudes, par inversion

Ayant obtenu les variances  $V_{sq}$  et  $V_{asq}$ , et leurs incertitudes, nous pouvons maintenant utiliser ces résultats afin d'obtenir une estimation de  $N_s$ ,  $N_t$  et de leurs incertitudes. La méthode d'estimation par inversion utilise directement les valeurs  $V_{sq}$  et  $V_{asq}$  pour calculer  $N_s$  et  $N_t$  à l'aide des formules (5.37a) et (5.37b), et des résultats des sections 5.3.2 et 5.3.3. Nous avons choisi une méthode numérique de type Monte Carlo afin de propager les incertitudes sur  $V_{sq}$  et  $V_{asq}$ , dont le principe est décrit ci-dessous :

- On commence par tirer  $n=10^6$  valeurs de  $\tilde{V}_{sq}$  et  $\tilde{V}_{asq}$  en utilisant les formules (5.32a) et (5.32b). Ce nombre assure un bon compromis entre le temps de calcul nécessaire et la précision obtenue.
- Pour chaque tirage  $k$ , avec  $k=1, \dots, n$ , on calcule ensuite  $N_s^k$  et  $N_t^k$  en utilisant les formules (5.37a) et (5.37b). On obtient ainsi deux ensembles  $\mathbf{N}_s = \{N_s^1, \dots, N_s^n\}$  et  $\mathbf{N}_t = \{N_t^1, \dots, N_t^n\}$  de valeurs de  $N_s$  et  $N_t$ .
- On calcule ensuite les valeurs moyennes de  $\mathbf{N}_s$  et  $\mathbf{N}_t$

$$N_s^{\text{inv}} = \langle \mathbf{N}_s \rangle, \quad (5.38a)$$

$$N_t^{\text{inv}} = \langle \mathbf{N}_t \rangle, \quad (5.38b)$$

et leurs variances

$$\sigma^2(N_s^{\text{inv}}) = \text{Var}(\mathbf{N}_s), \quad (5.39a)$$

$$\sigma^2(N_t^{\text{inv}}) = \text{Var}(\mathbf{N}_t). \quad (5.39b)$$

Cette méthode numérique présente l'avantage de ne pas introduire une linéarisation par rapport aux variables comme cela est souvent le cas pour un calcul d'erreur analytique. Ceci sera particulièrement important dans le cas du calcul de la discordie, dont les dérivées présentent des divergences pour les petites valeurs de  $N_s$  et  $N_t$ . Elle nous fournit ainsi une meilleure propagation des erreurs, tout en étant plus simple à mettre en œuvre.

### 5.5 Estimation de la discordie et de son incertitude

#### 5.5.1 Discordie gaussienne pour un état thermique comprimé

La définition de la discordie nécessite une optimisation sur tous les POVM, qui se révèle être une tâche extrêmement complexe pour des variables continues. On peut toutefois obtenir

des résultats analytiques en se restreignant à des mesures gaussiennes [Giorda10, Adesso10]. La discorde obtenue correspond alors à la *discorde gaussienne*. Elle peut en théorie différer de la discorde optimisée sur des POVM quelconques, bien qu'il semblerait que les mesures gaussiennes soient optimales pour les états gaussiens [Giorda12]. Même si cela n'était pas le cas, il a quand même été montré que les états gaussiens avec une discorde gaussienne non nulle possèdent également une discorde quantique non nulle [Rahimi-Keshari13]. Il n'y a donc pas de risque qu'un état soit classique s'il possède une discorde gaussienne non nulle. Notons enfin que cette restriction à des mesures gaussiennes est aussi justifiée d'un point de vue expérimental, où des mesures non gaussiennes ne sont pas toujours utilisées.

La discorde gaussienne<sup>5</sup> a été calculée explicitement pour un état thermique comprimé tel que (5.33) dans la référence [Giorda10]. Exprimée en fonction de  $N_s$  et  $N_t$ , elle vaut

$$D(N_s, N_t) = 2N_t \log_2(N_t) - 2(N_t + 1) \log_2(N_t + 1) - \mathcal{A} \log_2 \mathcal{A} - \mathcal{B} \log_2 \mathcal{B} + \mathcal{C} \log_2 \mathcal{C} + \mathcal{D} \log_2 \mathcal{D}, \quad (5.40)$$

avec

$$\mathcal{A} = N_s + N_t + 2N_s N_t \quad \mathcal{C} = 1 + N_s + N_t + 2N_s N_t \quad (5.41a)$$

$$\mathcal{B} = \frac{N_t(N_t + 1)}{1 + N_s + N_t + 2N_s N_t} \quad \mathcal{D} = \frac{N_s + 2N_s N_t + (1 + N_t)^2}{1 + N_s + N_t + 2N_s N_t}. \quad (5.41b)$$

Puisqu'un état thermique comprimé peut s'écrire de manière équivalente en fonction des paramètres  $r$ ,  $\gamma$  et  $\eta$  introduits au chapitre 3, on peut obtenir l'expression de  $N_s$  et  $N_t$  en fonction de ces paramètres, en comparant les matrices de covariances obtenues avec les deux paramétrisations :

$$N_s = \frac{1}{2} \left( -1 + \frac{A(r, \gamma, \eta)}{\sqrt{\eta^2 \cosh^4 r \cosh^2(2r\gamma) + B(r, \eta)^2 + 2\eta \cosh^2 r (-2\eta \cosh^4(r\gamma) \sinh^2 r + \cosh(2r\gamma) B(r, \eta))}} \right) \quad (5.42a)$$

$$N_t = \frac{1}{2} \left( -1 + \sqrt{(A(r, \gamma, \eta) - \eta \cosh^2(r\gamma) \sinh 2r)(A(r, \gamma, \eta) + \eta \cosh^2(r\gamma) \sinh 2r)} \right) \quad (5.42b)$$

avec

$$A(r, \gamma, \eta) = 1 - \eta + \eta \cosh^2 r \cosh 2r\gamma + \eta \sinh^2 r, \quad (5.43a)$$

$$B(r, \eta) = 1 - \eta + \eta \sinh^2 r. \quad (5.43b)$$

On pourra ainsi utiliser ces expressions dans la formule (5.40), afin d'obtenir la discorde en fonction de  $r$ ,  $\gamma$  et  $\eta$ . Comme nous le verrons, cela nous sera utile pour calculer les bornes de Cramér-Rao.

Nous avons utilisé deux méthodes pour l'estimation de la discorde à partir des résultats de la section 5.4.2. La première est une méthode directe par inversion. La seconde, plus complexe à mettre en œuvre, utilise une analyse bayésienne. Nous verrons qu'elle améliore l'estimation de la discorde pour de faibles puissances de pompe.

---

5. Dans la suite de ce chapitre, nous ne ferons plus la distinction entre la discorde et la discorde gaussienne, étant entendu que c'est toujours de la seconde qu'il s'agit.

### 5.5.2 Estimation par inversion

Avec cette méthode, on calcule directement la discordie avec la formule (5.40) et les valeurs  $N_s^{\text{inv}}$  et  $N_t^{\text{inv}}$ . Les incertitudes sont propagées par une méthode Monte Carlo, en supposant que  $N_s^{\text{inv}}$  et  $N_t^{\text{inv}}$  sont la réalisation de deux variables normales  $\tilde{N}_s^{\text{inv}}$  et  $\tilde{N}_t^{\text{inv}}$ , respectivement de moyennes  $N_s^{\text{inv}}$  et  $N_t^{\text{inv}}$  et de variances  $\sigma^2(N_s^{\text{inv}})$  et  $\sigma^2(N_t^{\text{inv}})$  :

$$\tilde{N}_s^{\text{inv}} = N_s^{\text{inv}} + U_s \sigma(N_s^{\text{inv}}) \quad (5.44a)$$

$$\tilde{N}_t^{\text{inv}} = N_t^{\text{inv}} + U_t \sigma(N_t^{\text{inv}}) \quad (5.44b)$$

où  $U_s$  et  $U_t$  sont deux variables normales centrées réduites. On réalise  $10^6$  tirages de  $\tilde{N}_s^{\text{inv}}$  et  $\tilde{N}_t^{\text{inv}}$ , en calculant à chaque fois la discordie  $D^{\text{inv}}(\tilde{N}_s^{\text{inv}}, \tilde{N}_t^{\text{inv}})$  correspondante.

Il faut noter ici que, pour les faibles valeurs de  $N_s^{\text{inv}}$  et  $N_t^{\text{inv}}$ , certains tirages de (5.44) peuvent donner des valeurs négatives, non physiques, qui ne permettent pas une estimation de la discordie par (5.40). Ceci est en fait un problème central, et nous allons le gérer ici en ne prenant tout simplement pas en compte ces valeurs négatives. Nous considérons donc, au lieu de (5.44), des distributions gaussiennes tronquées aux valeurs positives. Nous allons cependant voir à la section suivante que l'estimation bayésienne permet de mieux gérer cette information *a priori* qu'est la positivité de paramètres  $N_s$  et  $N_t$ .

On calcule enfin la valeur moyenne et la variance des  $10^6$  valeurs de discordie obtenues. On obtient ainsi l'estimation de la discordie  $D^{\text{inv}}$  et sa variance  $\sigma^2(D^{\text{inv}})$ , avec la méthode par inversion.

### 5.5.3 Estimation bayésienne

La seconde méthode utilisée fait appel à une analyse bayésienne. Elle est plus longue à mettre en œuvre, mais permet une meilleure exploitation des données. Son principe est en théorie très simple : supposons que nous ayons mesuré un ensemble de données  $D$ . Nous savons que ces données ont une distribution de probabilité qui dépend de paramètres  $H$ , qui nous sont inconnus, et que nous cherchons à estimer. Nous avons aussi une certaine information  $I$  *a priori* sur ces paramètres. On peut alors relier, grâce au théorème de Bayes, la probabilité  $P(H|D, I)$  que les paramètres aient une valeur  $H$ , étant données les mesures  $D$  et l'information  $I$ , à la probabilité  $P(D|H, I)$  d'obtenir les mesures  $D$  à partir de paramètres  $H$  et de l'information  $I$ , et à la probabilité que les paramètres aient une valeur  $H$  compte tenu de l'information  $I$  :

$$p(H|D, I) = \frac{p(D|H, I) \times p(H|I)}{p(D|I)} \quad (5.45)$$

On obtient ensuite la valeur moyenne de  $H$  avec  $\bar{H} = \int dH H p(H|D, I)$ , et sa variance avec  $\sigma^2(H) = \int dH (H - \bar{H})^2 p(H|D, I)$ .

Le lecteur pourra trouver une très bonne introduction à l'estimation bayésienne et à l'estimation de paramètres dans les ouvrages [Sivia06] et [Knight99], et dans l'article [Toussaint11] pour des exemples d'applications concrètes en physique.

### Application à l'estimation de la discordie

Voyons maintenant comment utiliser cette analyse bayésienne avec nos mesures. Pour des raisons de puissance de calcul qui seront plus claires par la suite, on commence par découper les



mesures de chaque quadrature  $\hat{Q}^{(k)}$  en  $N_b=10^2$  blocs de  $K=200$  points. Puis pour chaque bloc  $m=0, \dots, N_b-1$ , on regroupe les  $4K$  mesures dans un ensemble

$$\mathcal{X}_m = \{q_{mK+1}^{(1)}, \dots, q_{mK+K}^{(1)}, q_{mK+1}^{(2)}, \dots, q_{mK+K}^{(2)}, q_{mK+1}^{(3)}, \dots, q_{mK+K}^{(3)}, q_{mK+1}^{(4)}, \dots, q_{mK+K}^{(4)}\}. \quad (5.46)$$

Considérons maintenant un bloc  $m$ . Pour des valeurs  $N_s$  et  $N_t$  données, la densité de probabilité d'obtenir une mesure homodyne  $q^{(k)}$  pour une quadrature  $\hat{Q}^{(k)}$  est donnée par

$$p_k(q^{(k)}|N_s, N_t) = \frac{1}{\sqrt{2\pi\sigma_k^2}} \exp\left(-\frac{(q^{(k)})^2}{2\sigma_k^2}\right), \quad (5.47)$$

où  $\sigma_k^2$  est donné par (5.36a) pour  $k=\{1, 4\}$ , et par (5.36b) pour  $k=\{2, 3\}$ . Puisque les mesures homodynes sont indépendantes entre elles, la probabilité d'obtenir l'ensemble de résultats  $\mathcal{X}_m$  est égale à

$$p(\mathcal{X}_m|N_s, N_t) = \prod_{k=1}^4 \prod_{j=1}^K p_k(q_{mK+j}^{(k)}|N_s, N_t). \quad (5.48)$$

Le calcul de cette probabilité nécessite la multiplication d'un grand nombre de valeurs proches de zéro, et reste possible jusqu'à  $4K \simeq 800$ . Au delà, un ordinateur standard n'a pas la puissance de calcul nécessaire pour manipuler le résultat correctement. C'est la raison pour laquelle les données de chaque quadrature sont découpées en blocs de 200 points.

En utilisant ensuite le théorème de Bayes, on obtient la probabilité *a posteriori* que l'ensemble  $\mathcal{X}_m$  soit dû à des valeurs  $N_s$  et  $N_t$ ,

$$p(N_s, N_t|\mathcal{X}_m) = \frac{1}{\mathcal{N}} p(\mathcal{X}_m|N_s, N_t) p_0(N_s) p_0(N_t), \quad (5.49)$$

où  $p_0(N_s)$  and  $p_0(N_t)$  sont les probabilités dites *a priori*, et

$$\mathcal{N} = \int dN_s dN_t p(\mathcal{X}_m|N_s, N_t) p_0(N_s) p_0(N_t). \quad (5.50)$$

est un terme de normalisation. Les probabilités *a priori* reflètent une connaissance préalable sur les distributions de  $N_s$  et  $N_t$  avant d'effectuer les mesures du bloc  $m$ . Dans notre cas, une telle connaissance nous est donnée par les résultats de l'analyse par inversion. On peut en effet considérer que  $p_0(N_s)$  et  $p_0(N_t)$  suivent des lois normales, de moyennes respectivement égales à  $N_s^{\text{inv}}$  et  $N_t^{\text{inv}}$ , et de variances  $\sigma^2(N_s^{\text{inv}})$  et  $\sigma^2(N_t^{\text{inv}})$ , tronquées aux valeurs positives :

$$p_0(N_j) = \begin{cases} \frac{1}{\sqrt{2\pi\sigma^2(N_j^{\text{inv}})}} \exp\left(-\frac{(N_j - N_j^{\text{inv}})^2}{2\sigma^2(N_j^{\text{inv}})}\right) & \text{si } N_j > 0, \\ = 0 & \text{si } N_j \leq 0 \end{cases} \quad (5.51)$$

avec  $j=s, t$ .

On pourrait arguer qu'un tel choix de distribution ne reflète pas vraiment un savoir *avant* les mesures, puisque celle-ci a été établie avec toutes les mesures, dont celles du bloc  $m$ . Toutefois, nous ferons l'hypothèse que les contributions de chaque bloc restent faibles, compte tenu du nombre total de mesures. Autrement dit, nous supposons que les valeurs  $N_j^{\text{inv}}$  et  $\sigma^2(N_j^{\text{inv}})$  dépendent peu des mesures du bloc  $m$ , et que les distributions  $p_0$  reflètent bien un savoir *a priori* pour chaque bloc. La distribution (5.51) apporte également une information *a priori* sur

la positivité des paramètres, qui est prise en compte de manière plus rigoureuse avec l'approche bayésienne. Cette dernière information contribue également à affiner notre connaissance des distributions de  $N_s$  et  $N_t$ , puisque la positivité des  $N_j$  dans les distributions (5.51) implique notamment celle des  $N_j$  dans (5.49). Nous verrons d'ailleurs que c'est pour les petites valeurs de  $N_s$  et de  $N_t$ , lorsque se pose ce problème de positivité, que l'approche bayésienne diffère de l'approche par inversion plus classique.

On utilise ensuite la distribution *a posteriori* (5.49) afin d'obtenir une estimation des deux paramètres  $N_s^{\text{bay},m}$  et  $N_t^{\text{bay},m}$  et de leurs variances pour le bloc  $m$  :

$$N_j^{\text{bay},m} = \int dN_s dN_t N_j p(N_s, N_t | \mathcal{X}_m) \quad (5.52)$$

$$\sigma^2(N_j^{\text{bay},m}) = \int dN_s dN_t (N_j - N_j^{\text{bay},m})^2 p(N_s, N_t | \mathcal{X}_m) \quad (5.53)$$

Finalement, on moyenne les quantités obtenues pour chaque bloc selon

$$N_j^{\text{bay}} = \frac{\left[ \sum_{m=0}^{N_b-1} N_j^{\text{bay},m} / \sigma^2(N_j^{\text{bay},m}) \right]}{\left[ \sum_{m=0}^{N_b-1} 1 / \sigma^2(N_j^{\text{bay},m}) \right]}, \quad (5.54)$$

Ce type de pondération est optimal afin d'estimer une quantité à partir de plusieurs échantillons ayant des variances différentes [Sivia06]. La variance de cette estimation est donnée par

$$\sigma^2(N_j^{\text{bay}}) = \left[ \sum_{m=0}^{N_b-1} 1 / \sigma^2(N_j^{\text{bay},m}) \right]^{-1}. \quad (5.55)$$

La discordie est ensuite calculée à partir de ces valeurs, en utilisant une méthode Monte Carlo similaire aux précédentes. On note  $D^{\text{bay}}$  et  $\sigma^2(D^{\text{bay}})$  les valeurs obtenues par cette méthode bayésienne.

#### 5.5.4 Résultats expérimentaux

Les valeurs de discordie  $D^{\text{bay}}$  obtenues avec la méthode bayésienne sont présentées sur la figure 5.3, en fonction du paramètre de compression  $r$ . Les deux méthodes d'estimation donnent des valeurs très proches, contenues dans la largeur des points du graphique comme les incertitudes statistiques qui leur sont liées. Le plus gros écart obtenu entre les deux méthodes est inférieur à  $0.5 \times \sigma(D^{\text{inv}})$  pour les faibles compressions, et il est de l'ordre de grandeur de  $0.05 \times \sigma(D^{\text{inv}})$  pour les plus grandes valeurs de  $r$  (figure 5.4). Nous verrons en revanche que les incertitudes sont différentes, et que la méthode bayésienne s'avère plus précise pour une faible compression.

L'efficacité homodyne est estimée à  $\eta=0.62$ . En inversant les équations (5.42), on peut obtenir les valeurs de  $r$  et  $\gamma$  correspondant à chaque puissance de pompe. Ceci nous permet de tracer les valeurs de discordes estimées en fonction de  $r$ . On peut ensuite raisonnablement considérer que seul le paramètre  $r$  varie avec la puissance de la pompe, alors que  $\gamma$  et  $\eta$  gardent une valeur constante, comme expliqué dans la section 3.3.2. Afin de vérifier cette hypothèse, et pour obtenir une valeur moyenne de  $\gamma$  pour toutes les puissances de pompe, on fait un fit de la courbe théorique de la discordie (5.40), calculée à partir des valeurs de  $r$  obtenues, et en prenant  $\eta=0.62$ , ce qui nous donne  $\gamma=0.73$ .

On voit que l'accord entre les valeurs expérimentales et la courbe théorique est très bon, ce qui confirme la validité de notre modélisation. On pourra donc considérer que les bornes de Cramér-Rao théoriques décriront bien les limites fondamentales pour la précision de nos estimateurs de la discordie.

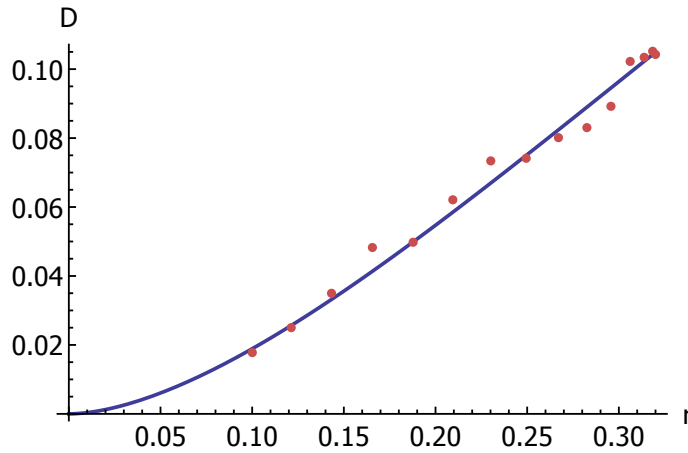


FIGURE 5.3 – Discorde  $D^{\text{bay}}$  estimée par la méthode bayésienne. Les points sont les valeurs estimées à partir des données expérimentales, et la ligne correspond au modèle théorique, pour  $\eta=0.62$ , et  $\gamma=0.73$  obtenu par un fit. Les incertitudes expérimentales sont contenues dans les points : l'écart par rapport à la courbe théorique s'explique par le fait que l'on cherche ici une valeur de  $\gamma$  "moyenne" pour tous les points.

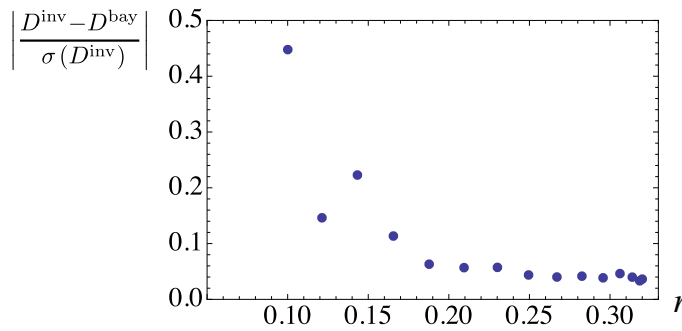


FIGURE 5.4 – Comparaison des valeurs de discorde obtenues avec les deux méthodes : nombre d'écart-types  $\sigma(D^{\text{inv}})$  séparant les estimations  $D^{\text{inv}}$  et  $D^{\text{bay}}$ , en fonction de la compression  $r$ .

## 5.6 Comparaison avec les bornes de Cramér-Rao

L'incertitude sur la valeur d'une grandeur estimée statistiquement dépend naturellement du nombre de mesures utilisées. Il est par exemple bien connu que l'incertitude sur la moyenne d'un échantillon de  $N$  mesures varie en  $1/\sqrt{N}$ , et on comprend intuitivement que dans la plupart des cas on ne peut pas estimer une grandeur avec une précision parfaite, si le nombre de mesures est fini.

Etant donné un nombre de mesures donné, quelle est donc la limite fondamentale sur la précision d'un estimateur d'un paramètre ? Dans notre cas, notre estimation de la discorde est-elle une "bonne" estimation ? Afin de répondre à ces questions, faisons maintenant un petit détour par la théorie de l'estimation de paramètres, afin d'introduire les bornes de Cramér-Rao, pour une variable classique, et pour un état quantique. Le lecteur pourra consulter [Helstrom76, Knight99, Braunstein94, Hayashi06] et [Paris09] pour une présentation plus détaillée.

### 5.6.1 Information de Fisher et borne de Cramér-Rao classique

#### Information de Fisher

Considérons une variable aléatoire classique  $X$  prenant des valeurs  $x$ , avec une densité de probabilité  $p(x|\boldsymbol{\lambda})$  dépendant d'un ensemble de  $n$  paramètres  $\boldsymbol{\lambda}=\{\lambda_1, \dots, \lambda_n\}$ , que l'on cherche à estimer. La forme de la distribution est en générale connue, grâce à une modélisation de la grandeur mesurée, ou grâce à certaines propriétés statistiques (par exemple le théorème central limite). En mesurant  $M$  fois  $X$ , on obtient un ensemble de résultats  $\mathbf{x}=\{x_1, \dots, x_M\}$  à partir duquel on construit un estimateur  $\bar{\boldsymbol{\lambda}}(\mathbf{x})$  de  $\boldsymbol{\lambda}$ .

La matrice d'information de Fisher<sup>6</sup>  $\mathbf{F}(\boldsymbol{\lambda})$  est une mesure de l'information que la loi de probabilité de  $X$  apporte sur  $\boldsymbol{\lambda}$  [Hayashi06]. Ses termes sont définis par

$$\mathbf{F}(\boldsymbol{\lambda})_{\mu\nu} = \int d\mathbf{x} p(\mathbf{x}|\boldsymbol{\lambda}) \frac{\partial \ln p(\mathbf{x}|\boldsymbol{\lambda})}{\partial \lambda_\mu} \frac{\partial \ln p(\mathbf{x}|\boldsymbol{\lambda})}{\partial \lambda_\nu} \quad (5.56)$$

avec  $\mu, \nu=1, \dots, n$ . Elle dépend donc uniquement de la loi de probabilité, et non d'un tirage particulier. On montre sans difficulté que pour  $M$  réalisations indépendantes de  $X$ , l'information de Fisher est additive. Cette propriété est, comme pour l'entropie, assez intuitive : on s'attend à ce que l'information apportée par  $M$  variables indépendantes soit bien la somme de l'information des  $M$  variables.

Ce sont les termes diagonaux  $\mathbf{F}(\boldsymbol{\lambda})_{\mu\mu}$  qui vont nous intéresser pour notre étude de la discordie. Comme nous allons maintenant le montrer, ils permettent d'établir la *borne de Cramér-Rao*.

#### Borne de Cramér-Rao

Considérons un estimateur  $\bar{\lambda}_\mu(\mathbf{x})$  de l'un des paramètres  $\lambda_\mu$ , que l'on suppose non biaisé (sa valeur moyenne  $\mathbb{E}[\bar{\lambda}_\mu(\mathbf{x})]$  est égale à  $\lambda_\mu$ ), et fonction des  $M$  réalisations indépendantes de  $X$ . Puisque le nombre de mesures est fini, on s'attend, au moins à cause des fluctuations statistiques, à ce que la variance de  $\bar{\lambda}_\mu(\mathbf{x})$  soit non nulle.

La borne de Cramér-Rao fixe une limite fondamentale sur la variance de  $\bar{\lambda}_\mu(\mathbf{x})$ , proportionnelle à l'inverse de l'information de Fisher [Knight99, Hayashi06]. La raisonnement est le suivant : puisque l'estimateur est sans biais, on a  $\mathbb{E}(\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) = 0$ . On peut donc écrire

$$\int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) p(\mathbf{x}|\boldsymbol{\lambda}) = 0, \quad (5.57)$$

avec  $\mathbf{x}=(x_1, \dots, x_M)$ ,  $d\mathbf{x}=dx_1 \dots dx_M$  et  $p(\mathbf{x}|\boldsymbol{\lambda})=p(x_1|\boldsymbol{\lambda}) \dots p(x_M|\boldsymbol{\lambda})$ . On dérive ensuite l'équation (5.57) par rapport à  $\lambda_\mu$  :

$$\frac{\partial}{\partial \lambda_\mu} \int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) p(\mathbf{x}|\boldsymbol{\lambda}) = - \underbrace{\int d\mathbf{x} p(\mathbf{x}|\boldsymbol{\lambda})}_{=1} + \int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) \frac{\partial}{\partial \lambda_\mu} p(\mathbf{x}|\boldsymbol{\lambda}) \quad (5.58a)$$

$$\Rightarrow \int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) p(\mathbf{x}|\boldsymbol{\lambda}) \frac{\partial}{\partial \lambda_\mu} \ln p(\mathbf{x}|\boldsymbol{\lambda}) = 1 \quad (5.58b)$$

6. Nous parlerons simplement de l'information de Fisher lorsque nous considérerons un terme diagonal de la matrice d'information de Fisher.

En utilisant l'inégalité de Cauchy-Schwarz, on peut ensuite écrire :

$$\begin{aligned} & \left| \int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu) \sqrt{p(\mathbf{x}|\boldsymbol{\lambda})} \sqrt{p(\mathbf{x}|\boldsymbol{\lambda})} \frac{\partial}{\partial \lambda_\mu} \ln p(\mathbf{x}|\boldsymbol{\lambda}) \right|^2 \\ & \leq \underbrace{\left[ \int d\mathbf{x} (\bar{\lambda}_\mu(\mathbf{x}) - \lambda_\mu)^2 p(\mathbf{x}|\boldsymbol{\lambda}) \right]}_{\sigma^2(\bar{\lambda}_\mu)} \underbrace{\left[ \int d\mathbf{x} \left( \frac{\partial}{\partial \lambda_\mu} \ln p(\mathbf{x}|\boldsymbol{\lambda}) \right)^2 p(\mathbf{x}|\boldsymbol{\lambda}) \right]}_{M\mathbf{F}(\boldsymbol{\lambda})_{\mu\mu}} \end{aligned} \quad (5.59)$$

Puisque la première ligne est égale à 1, on obtient finalement la borne de Cramér-Rao :

$$\boxed{\sigma^2(\bar{\lambda}_\mu) \geq \frac{1}{M\mathbf{F}(\boldsymbol{\lambda})_{\mu\mu}}} \quad (5.60)$$

Cette borne fixe une limite sur la précision que peut avoir un estimateur, quelle que soit la manière dont il est construit. Lorsqu'elle est atteinte pour toutes valeurs de  $\boldsymbol{\lambda}$ , l'estimateur est dit efficace. Lorsque cela n'est pas le cas, ce qui arrive la plupart du temps, la comparaison de la variance de l'estimateur avec cette borne permet de juger de la précision de l'estimateur. Notons que l'on peut relier d'une façon similaire les termes non diagonaux  $\mathbf{F}(\boldsymbol{\lambda})_{\mu\nu}$  et les covariances entre les estimations des paramètres  $\lambda_\mu$  et  $\lambda_\nu$ .

### 5.6.2 De l'information de Fisher classique à l'information de Fisher quantique

#### Borne de Cramér-Rao quantique

Jusqu'à maintenant, nous avons considéré l'estimation de paramètres pour une variable aléatoire classique. Autrement dit, une variable ayant une certaine loi de probabilité  $p(x|\boldsymbol{\lambda})$  bien définie, qui dépend des paramètres  $\boldsymbol{\lambda}$  à estimer. C'est ce que l'on obtient en mesurant un état quantique  $\hat{\rho}_\lambda$  dépendant des paramètres  $\boldsymbol{\lambda}$ , avec un POVM  $\{\hat{\Pi}_x\}$  tel que

$$p(x|\boldsymbol{\lambda}) = \text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x\}. \quad (5.61)$$

En introduisant la *dérivée logarithmique symétrique*  $\hat{\mathbf{L}}_\mu$  comme étant un opérateur hermitien vérifiant

$$\frac{\partial \hat{\rho}_\lambda}{\partial \lambda_\mu} = \frac{\hat{\mathbf{L}}_\mu \hat{\rho}_\lambda + \hat{\rho}_\lambda \hat{\mathbf{L}}_\mu}{2}, \quad (5.62)$$

et en remarquant<sup>7</sup> que  $\partial_{\lambda_\mu} p(x|\boldsymbol{\lambda}) = \text{Tr}\{\partial_{\lambda_\mu} \hat{\rho}_\lambda \hat{\Pi}_x\} = \Re(\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x \hat{\mathbf{L}}_\mu\})$ , l'information de Fisher  $\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu}$  définie par (5.56) et relative à  $\{\hat{\Pi}_x\}$  s'écrit [Paris09]

$$\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu} = \int d\mathbf{x} \frac{\Re(\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x \hat{\mathbf{L}}_\mu\})^2}{\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x\}}. \quad (5.63)$$

Lorsque l'on cherche à estimer les paramètres d'un état quantique, la borne de Cramér-Rao (5.60) dépend donc du type de mesure effectuée, ici le POVM  $\{\hat{\Pi}_x\}$ . En optimisant sur tous les

7. On vérifie que  $\text{Tr}\{\hat{\rho} \hat{\mathbf{L}}_\mu \hat{\Pi}_x\} = (\text{Tr}\{\hat{\rho} \hat{\Pi}_x \hat{\mathbf{L}}_\mu\})^*$  en écrivant les décompositions  $\hat{\rho} = \sum_n a_n |\psi_n\rangle\langle\psi_n|$  et  $\hat{\mathbf{L}}_\mu = \sum_m c_m |\phi_\mu^m\rangle\langle\phi_\mu^m|$ .

POVM, on peut ensuite borner (5.63) par une limite quantique fondamentale, indépendante du type de mesure. On a en effet [Paris09, Braunstein94] :

$$\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu} \leq \int dx \left| \frac{\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x \hat{\mathbf{L}}_\mu\}}{\sqrt{\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x\}}} \right|^2 \quad (5.64a)$$

$$= \int dx \left| \text{Tr} \left\{ \frac{\sqrt{\hat{\rho}_\lambda} \sqrt{\hat{\Pi}_x}}{\sqrt{\text{Tr}\{\hat{\rho}_\lambda \hat{\Pi}_x\}}} \sqrt{\hat{\Pi}_x} \hat{\mathbf{L}}_\mu \sqrt{\hat{\rho}_\lambda} \right\} \right|^2 \quad (5.64b)$$

En utilisant ensuite l'inégalité de Cauchy-Schwarz  $|\text{Tr}\{\hat{\mathbf{A}}^\dagger \hat{\mathbf{B}}\}|^2 \leq \text{Tr}\{\hat{\mathbf{A}}^\dagger \hat{\mathbf{A}}\} \text{Tr}\{\hat{\mathbf{B}}^\dagger \hat{\mathbf{B}}\}$  pour (5.64b), on majore à nouveau l'information de Fisher :

$$\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu} \leq \int dx \text{Tr}\{\hat{\Pi}_x \hat{\mathbf{L}}_\mu \hat{\rho}_\lambda \hat{\mathbf{L}}_\mu\} \quad (5.65)$$

Enfin, en utilisant le fait que  $\int dx \hat{\Pi}_x = \mathbb{I}$ , on obtient

$$\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu} \leq \text{Tr}\{\hat{\rho}_\lambda \hat{\mathbf{L}}_\mu^2\} \quad (5.66)$$

Cette dernière quantité est appelée *information de Fisher quantique* :

$$\boxed{\mathbf{H}(\boldsymbol{\lambda})_{\mu\mu} = \text{Tr}\{\hat{\rho}_\lambda \hat{\mathbf{L}}_\mu^2\}} \quad (5.67)$$

Elle ne dépend pas du type de mesure effectuée. C'est une borne supérieure pour l'information de Fisher  $\mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu}$  que l'on peut obtenir avec n'importe quel POVM. Compte tenu de (5.66) et de (5.60), on obtient la borne de Cramér-Rao quantique :

$$\boxed{\sigma^2(\bar{\lambda}_\mu) \geq \frac{1}{M \mathbf{F}_{\{\hat{\Pi}_x\}}(\boldsymbol{\lambda})_{\mu\mu}} \geq \frac{1}{M \mathbf{H}(\boldsymbol{\lambda})_{\mu\mu}}} \quad (5.68)$$

Cette borne définit une limite quantique fondamentale pour la variance de tout estimateur  $\bar{\lambda}_\mu$ .

Dans la démonstration précédente, nous n'avons considéré que les termes diagonaux de la matrice d'information de Fisher quantique  $\mathbf{H}(\boldsymbol{\lambda})$ . Les autres termes sont définis d'une manière similaire [Paris09]

$$\mathbf{H}(\boldsymbol{\lambda})_{\mu\nu} = \text{Tr} \left\{ \hat{\rho}_\lambda \frac{\hat{\mathbf{L}}_\mu \hat{\mathbf{L}}_\nu + \hat{\mathbf{L}}_\nu \hat{\mathbf{L}}_\mu}{2} \right\}. \quad (5.69)$$

En écrivant  $\hat{\rho}_\lambda$  sous sa forme diagonale,  $\hat{\rho}_\lambda = \sum_n a_n |\psi_n\rangle \langle \psi_n|$ , on montre que  $\mathbf{H}(\boldsymbol{\lambda})_{\mu\nu}$  s'écrit explicitement [Paris09]

$$\mathbf{H}(\boldsymbol{\lambda})_{\mu\nu} = \sum_n \frac{(\partial_\mu a_n)(\partial_\nu a_n)}{a_n} + \sum_{n \neq m} \frac{(a_n - a_m)^2}{a_n + a_m} \left( \langle \psi_n | \partial_\mu \psi_m \rangle \langle \partial_\nu \psi_m | \psi_n \rangle + \langle \psi_n | \partial_\nu \psi_m \rangle \langle \partial_\mu \psi_m | \psi_n \rangle \right), \quad (5.70)$$

où on a posé  $\partial_\mu := \partial_{\lambda_\mu}$ . La matrice densité  $\hat{\rho}_\lambda$  peut dépendre de  $\boldsymbol{\lambda}$  à travers ses valeurs propres, mais aussi à travers ses vecteurs propres. Les termes tels que  $|\partial_\mu \psi_n\rangle$  sont donc à interpréter comme  $|\partial_\mu \psi_n\rangle = \sum_k \partial_\mu \psi_{nk} |k\rangle$ , où  $|\psi_n\rangle = \sum_k \psi_{nk} |k\rangle$  est une décomposition de  $|\psi_n\rangle$  dans une base  $\{|k\rangle\}$  ne dépendant pas de  $\boldsymbol{\lambda}$ .

### 5.6.3 Application à l'évaluation de la discorde

#### Calcul de l'information de Fisher quantique

Le calcul de la matrice d'information de Fisher quantique nécessite de connaître la façon dont l'état quantique dépend du paramètre à estimer, ici la discorde. Nous avons vu sur la figure 5.3 que nos données expérimentales sont bien modélisées par un état thermique comprimé (5.33), et nous pourrions donc considérer que la borne de Cramér-Rao calculée à partir de ce modèle correspond bien à la précision maximale qu'il est possible d'obtenir.

L'état (5.33) est paramétré le plus simplement par les nombres  $N_s$  et  $N_t$ . Une application directe de la formule (5.70) permet donc de calculer la matrice d'information de Fisher quantique  $\mathbf{H}^{(1)}$  pour le couple de paramètre  $\boldsymbol{\lambda}_1 = \{N_s, N_t\}$  :

$$\mathbf{H}^{(1)} = \text{diag} \left( \frac{(1 + 2N_t)^2}{N_s(1 + N_s)(1 + 2N_t + 2N_t^2)}, \frac{1}{N_t(1 + N_t)} \right). \quad (5.71)$$

Les termes diagonaux de  $\mathbf{H}^{(1)}$  permettraient ensuite d'établir des bornes de Cramér-Rao pour l'estimation de  $N_s$  et  $N_t$ . Pour calculer la borne de Cramér-Rao relative à la discorde, deux approches sont possibles. La première serait d'écrire explicitement la dépendance de l'état (5.33) en fonction de la discorde et d'un autre paramètre  $\lambda_X$ , par exemple  $N_s$ ,  $N_t$  ou  $\gamma$  (en supposant  $\eta$  constant), et d'utiliser la formule (5.70). Si elle est la plus directe, cette méthode n'en est pas moins fort complexe. La seconde méthode, qui est celle que nous utiliserons, est en revanche beaucoup plus simple. On peut en effet obtenir très facilement la matrice d'information  $\widetilde{\mathbf{H}}$  pour un ensemble de paramètres  $\tilde{\boldsymbol{\lambda}} = \{\tilde{\lambda}_\nu(\boldsymbol{\lambda})\}$  fonctions de paramètres  $\boldsymbol{\lambda}$ , à partir de la matrice d'information  $\mathbf{H}$  pour les paramètres  $\boldsymbol{\lambda}$ . Puisque  $\tilde{\partial}_\mu = \sum_\nu B_{\mu\nu} \partial_\nu$  avec  $B_{\mu\nu} = \partial \lambda_\nu / \partial \tilde{\lambda}_\mu$ , on a

$$\tilde{L}_\mu = \sum_\nu B_{\mu\nu} L_\nu, \quad (5.72)$$

et il est facile de se convaincre que

$$\widetilde{\mathbf{H}} = \mathbf{B} \mathbf{H} \mathbf{B}^T. \quad (5.73)$$

Pour des raisons de simplicité de calcul, on obtient d'abord la matrice d'information  $\mathbf{H}^{(2)}$  correspondant à un premier changement de variable  $\boldsymbol{\lambda}_1 \rightarrow \boldsymbol{\lambda}_2 = \{r, \gamma\}$  :

$$\mathbf{H}^{(2)} = B_{12} \mathbf{H}^{(1)} B_{12}^T, \quad (5.74)$$

avec

$$B_{12} = \begin{pmatrix} \frac{\partial N_s}{\partial r} & \frac{\partial N_t}{\partial r} \\ \frac{\partial N_s}{\partial \gamma} & \frac{\partial N_t}{\partial \gamma} \end{pmatrix}. \quad (5.75)$$

On effectue ensuite un dernier changement de variable  $\boldsymbol{\lambda}_2 \rightarrow \boldsymbol{\lambda}_3 = \{D, \gamma\}$ , nous permettant d'obtenir une matrice d'information contenant la discorde :

$$\mathbf{H}^{(3)} = B_{23} \mathbf{H}^{(2)} B_{23}^T \quad (5.76)$$

avec

$$B_{23} = \begin{pmatrix} \frac{\partial r}{\partial D} & \frac{\partial \gamma}{\partial D} \\ \frac{\partial r}{\partial \gamma} & \frac{\partial \gamma}{\partial \gamma} \end{pmatrix} = \begin{pmatrix} \frac{\partial r}{\partial D} & 0 \\ \frac{\partial r}{\partial \gamma} & 1 \end{pmatrix}. \quad (5.77)$$

L'expression obtenue est une fonction beaucoup trop longue pour être exprimée ici, mais elle ne présente pas de difficulté particulière.

En utilisant (5.68), on obtient ensuite la borne de Cramér-Rao quantique pour un estimateur  $\tilde{D}$  de la discordie :

$$\sigma^2(\tilde{D}) \geq \frac{1}{M\mathbf{H}(\boldsymbol{\lambda}_3)_{DD}} \quad (5.78)$$

Cette borne, comme la borne de Cramér-Rao classique, dépend de la vraie valeur des paramètres  $\boldsymbol{\lambda}_3$ , ici la discordie et  $\gamma$ . Bien sûr, nous n'avons pas accès à ces vraies valeurs, puisque nous cherchons justement à les estimer. Cependant, compte tenu des résultats de la section 5.5.4, nous pouvons considérer que l'expression théorique de la discordie  $D(r, \gamma, \eta)$  en fonction de  $r$ ,  $\gamma$ , et  $\eta$ , permet quand même d'obtenir une borne de Cramér-Rao pertinente à laquelle on pourra comparer notre estimation expérimentale.

### Calcul de l'information de Fisher pour la détection homodyne

La borne de Cramér Rao quantique (5.78) est une borne optimisée sur tous les POVM possibles. Il est fort probable que la mesure que nous utilisons – la mesure homodyne des quadratures comprimées et anticomprimées – ne permette pas de saturer cette borne. Par contre, il est également intéressant de comparer notre estimateur de la discordie avec la borne de Cramér-Rao classique associée aux mesures homodynes.

La matrice d'information de Fisher pour l'état (5.33) et des mesures homodynes se calcule avec un changement de paramètre similaire au paragraphe précédent. On commence avec un premier ensemble de paramètres  $\boldsymbol{\lambda}_1 = \{N_s, N_t\}$ . Pour une quadrature  $\hat{Q}^{(k)}$ , la densité de probabilité de mesurer une valeur  $q^{(k)}$  avec une détection homodyne est simplement une fonction gaussienne de moyenne nulle et de variance  $\sigma^2(Q^{(k)})$ . En utilisant la formule (5.56), on montre alors que

$$F_{\mu\nu} = \frac{1}{2[\sigma^2(Q^{(k)})]^2} \frac{\partial \sigma^2(Q^{(k)})}{\partial \lambda_\mu} \frac{\partial \sigma^2(Q^{(k)})}{\partial \lambda_\nu}, \quad (5.79)$$

avec  $\{\lambda_\mu\} = N_s, N_t$ . On utilise ensuite les expressions (5.36a) et (5.36b) des variances des quadratures en fonction de  $\boldsymbol{\lambda}_1$  pour obtenir explicitement la matrice d'information de Fisher, pour les quadratures comprimées ( $\mathbf{F}^{\text{sq}}$ ), et anticomprimées ( $\mathbf{F}^{\text{asq}}$ ) :

$$\mathbf{F}^{\text{sq/asq}} = \begin{pmatrix} \frac{1}{2N_s + 2N_s^2} & \mp \frac{1}{\sqrt{N_s(1+N_s)(1+2N_t)}} \\ \mp \frac{1}{\sqrt{N_s(1+N_s)(1+2N_t)}} & \frac{2}{(1+2N_t)^2} \end{pmatrix} \quad (5.80)$$

Puisque l'information de Fisher est additive, et que la moitié de nos mesures portent sur des quadratures comprimées, alors que l'autre moitié porte sur les quadratures anticomprimées, la matrice d'information moyenne pour les paramètres  $\boldsymbol{\lambda}_1$  est finalement

$$\mathbf{F}^{(1)} = \frac{\mathbf{F}^{\text{sq}} + \mathbf{F}^{\text{asq}}}{2} = \begin{pmatrix} \frac{1}{2N_s + 2N_s^2} & 0 \\ 0 & \frac{2}{(1+2N_t)^2} \end{pmatrix}. \quad (5.81)$$

Afin d'obtenir la matrice d'information de Fisher pour l'estimation de la discordie, on procède ensuite de la même manière que nous l'avons fait pour  $\mathbf{H}$ . On fait un premier changement de variable  $\boldsymbol{\lambda}_1 \rightarrow \boldsymbol{\lambda}_2 = \{r, \gamma\}$ , pour lequel la nouvelle matrice d'information est  $\mathbf{F}^{(2)} = B_{12}\mathbf{F}^{(1)}B_{12}^T$ ,



avec  $B_{12}$  définie par (5.75). Puis on fait le changement de variable  $\lambda_2 \rightarrow \lambda_3 = \{D, \gamma\}$ , pour lequel la nouvelle matrice d'information est  $\mathbf{F}^{(3)} = B_{23}\mathbf{F}^{(2)}B_{23}^T$ , avec  $B_{23}$  définie par (5.77).

On obtient alors la borne de Cramér-Rao associée à la mesure homodyne, pour un estimateur  $\tilde{D}$  de la discorde :

$$\sigma^2(\tilde{D}) \geq \frac{1}{M\mathbf{F}(\lambda_3)_{DD}} \quad (5.82)$$

#### 5.6.4 Résultats expérimentaux

##### Discussion sur les ressources utilisées

Les bornes de Cramér-Rao (5.78) et (5.82) font intervenir le nombre de mesures  $M$  utilisées pour construire l'estimateur  $\tilde{D}$ . Nous avons introduit deux estimateurs différents dans la section 5.5. Le premier,  $D^{\text{inv}}$ , est basé sur une méthode d'inversion. Pour cette estimation, on utilise les  $M_q = 2 \times 10^4$  mesures homodynes de chacune des quadratures  $\hat{Q}^{(k)}$ . Le nombre total de mesures intervenant dans les bornes de Cramér-Rao est donc

$$M^{\text{inv}} = M_T = 4M_q = 8 \times 10^4. \quad (5.83)$$

En revanche, pour le second estimateur  $D^{\text{bay}}$  basé sur une analyse bayésienne, le nombre de mesures à considérer est différent. En effet, nous avons d'abord regroupé les mesures homodynes en  $N_b$  blocs de  $4 \times 200$  points (200 points par quadrature). Puis, pour chacun de ces blocs, notre analyse bayésienne a nécessité des probabilités *a priori* utilisant les résultats de l'analyse par inversion, pour  $N_j^{\text{inv}}$  et  $\sigma^2(N_j^{\text{inv}})$  (avec  $j=s, t$ ), obtenus avec  $M_T$  mesures. Au final, tout se passe donc comme si nous avions utilisé  $N_b \times M_T$  mesures, au lieu de  $M_T$ . Pour l'estimation bayésienne, il faudra donc prendre

$$M^{\text{bay}} = N_b \times M_T \quad (5.84)$$

pour les bornes de Cramér-Rao.

##### Précision de notre estimateur de la discorde

Les bornes de Cramér-Rao sont calculées en prenant  $\eta=0.62$ , et avec les valeurs de  $r$  et  $\gamma$  obtenues en inversant les formules (5.42), pour chaque puissance de pompe. Notons que la valeur moyenne  $\gamma=0.73$  obtenue pour le fit de la figure 5.3 n'est pas utilisée ici. Nos résultats expérimentaux sont présentés sur la figure 5.5. Pour les deux méthodes d'estimation, on trace le rapport (exprimé en dB) entre la variance obtenue expérimentalement et la borne de Cramér-Rao classique et quantique :

$$K_M^H = \frac{\sigma^2(\tilde{D})}{1/[M\mathbf{H}^{(3)}(\lambda_3)_{DD}]} \quad (5.85)$$

$$K_M^F = \frac{\sigma^2(\tilde{D})}{1/[M\mathbf{F}^{(3)}(\lambda_3)_{DD}]} \quad (5.86)$$

avec  $\tilde{D}=D^{\text{inv}}, D^{\text{bay}}$  et  $M=M^{\text{inv}}, M^{\text{bay}}$ . Un rapport égal à 1 (ou 0 dB) signifie que l'estimation est optimale, et que la borne de Cramér-Rao correspondante est saturée.

Commençons par analyser la comparaison avec la borne de Cramér-Rao pour la détection homodyne, donnée par le rapport  $K_M^F$ . Pour de faibles valeurs de discorde, l'estimation bayésienne

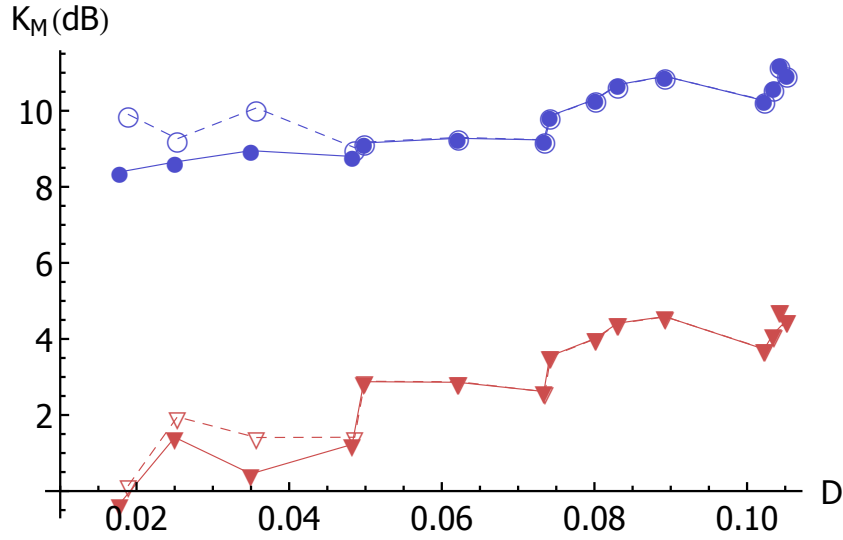


FIGURE 5.5 – Rapports  $K_M^H$  et  $K_M^F$  pour l'estimation par inversion, et l'estimation bayésienne, en fonction de la discorde. Les cercles bleus (resp. triangles rouges) correspondent à la borne de Cramér-Rao quantique  $K_M^H$  (resp. classique  $K_M^F$ ). Les marqueurs pleins correspondent à l'estimation bayésienne, et les marqueurs vides à l'estimation par inversion.

est pratiquement optimale, alors que l'estimation par inversion est un peu plus bruitée, avec un rapport plus important. Pour le point correspondant à la valeur de discorde la plus faible, on remarque que le rapport est légèrement inférieur à 0 dB. Nous attribuons ce comportement à une moins bonne adéquation de notre modélisation pour de faibles valeurs de pompe. On voit ensuite que pour de plus fortes valeurs de pompes, les variances s'écartent un peu plus des valeurs optimales, tout en restant proches. Cela peut s'expliquer par des variations de la discorde au cours de l'expérience due à diverses fluctuations, supérieures aux fluctuations statistiques.

Concernant la borne de Cramér-Rao quantique, comme on pouvait le supposer notre estimation n'est pas optimale, le rapport étant d'environ 10 dB. C'est le prix à payer – assez raisonnable – pour sa simplicité. Là encore, l'estimation bayésienne procure de meilleurs résultats pour de faibles valeurs de discorde, tout en donnant des performances comparables pour des puissances de pompe plus importantes.

## 5.7 Conclusion

La discorde quantique est une mesure des corrélations d'origine purement quantiques, que peuvent posséder des états qui ne sont pas intriqués. De nombreuses études laissent supposer qu'elle pourrait jouer un rôle important en information quantique, et surtout que l'intrication pourrait ne pas toujours être indispensable pour surpasser les protocoles classiques. Un des principaux intérêts de ces travaux est bien sûr expérimental, puisque la discorde est une propriété beaucoup moins fragile que l'intrication, et beaucoup plus facile à produire.

Dans ce chapitre, nous avons estimé la discorde gaussienne pour une classe d'états particulièrement importants en optique quantique avec des variables continues : les états thermiques comprimés, correspondant par exemple aux états produits par un OPA imparfait. En utilisant uniquement des mesures homodynes, parmi les plus courantes, mais également parmi les plus simples pour des états gaussiens, nous avons montré que notre estimation bayésienne est prati-

quement optimale pour la borne de Cramér-Rao associée à ces mesures, pour de faibles valeurs de compression. Nous avons également montré qu'elle ajoute environ 10 dB de bruit par rapport à la limite quantique fondamentale donnée par la borne de Cramér-Rao quantique, ce qui reste raisonnable étant donnée sa simplicité.

La capacité à pouvoir caractériser précisément la discorde d'un état est un élément indispensable pour pouvoir analyser le rôle de la discorde en information quantique. Ces travaux s'inscrivent dans ce cadre, en proposant une méthode simple et efficace.

## Chapitre 6

# Caractérisation d'une porte de phase quantique

### Sommaire

---

|            |                                                                        |            |
|------------|------------------------------------------------------------------------|------------|
| <b>6.1</b> | <b>Introduction</b>                                                    | <b>109</b> |
| <b>6.2</b> | <b>Présentation de la porte de phase</b>                               | <b>111</b> |
| 6.2.1      | Calcul quantique avec des états cohérents                              | 111        |
| 6.2.2      | La porte de phase                                                      | 111        |
| <b>6.3</b> | <b>Réalisation expérimentale d'une porte de phase <math>\pi</math></b> | <b>113</b> |
| 6.3.1      | Méthode et dispositif expérimental                                     | 113        |
| 6.3.2      | Modélisation                                                           | 115        |
| 6.3.3      | Extraction des paramètres expérimentaux                                | 116        |
| 6.3.4      | Résultats expérimentaux                                                | 117        |
| 6.3.5      | Test du modèle de la porte                                             | 117        |
| 6.3.6      | Incertitude sur l'estimation de $\xi$                                  | 119        |
| <b>6.4</b> | <b>Comment caractériser la porte ?</b>                                 | <b>120</b> |
| 6.4.1      | Ressemblance des états expérimentaux avec des chats parfaits           | 120        |
| 6.4.2      | Tomographie de processus quantique                                     | 121        |
| 6.4.3      | Utilisation de la modélisation de la porte                             | 123        |
| <b>6.5</b> | <b>Fidélité pour un qubit initial parfait</b>                          | <b>125</b> |
| 6.5.1      | Porte expérimentale                                                    | 125        |
| 6.5.2      | Modèle de porte simplifié                                              | 126        |
| 6.5.3      | Fidélités pour des superpositions de même poids                        | 127        |
| 6.5.4      | Quelques calculs de fidélité                                           | 128        |
| <b>6.6</b> | <b>Fidélité avec une porte idéale</b>                                  | <b>133</b> |
| 6.6.1      | Principe                                                               | 133        |
| 6.6.2      | Simulations pour la porte de phase                                     | 134        |
| 6.6.3      | Invariance du choix de l'état maximalement intriqué                    | 135        |
| <b>6.7</b> | <b>Conclusion</b>                                                      | <b>137</b> |

---

### 6.1 Introduction

Un des défis à relever afin de pouvoir utiliser les propriétés quantiques fondamentales pour le traitement de l'information est de contrôler la décohérence due au couplage avec l'environ-

nement. Pour cela, les implémentations optiques offrent des solutions intéressantes, puisque la lumière interagit peu avec l'environnement. Historiquement, les recherches se sont principalement axées sur l'utilisation de variables discrètes [Kok07], ou continues [Lloyd99, Braunstein05], ayant chacune leurs avantages et inconvénients d'un point de vue expérimental. Ainsi, les variables discrètes offrent un encodage de type "numérique", plus résistant aux pertes et au bruit, mais nécessitent de produire et de détecter des photons uniques. Les variables continues, en revanche, possèdent un certain nombre d'avantages expérimentaux, tels que l'utilisation d'états plus faciles à produire, mais souffrent des problèmes d'un encodage "analogique". Afin de combiner les forces de ces deux approches, des protocoles hybrides se sont développés au cours de ces dernières années [Loock11].

L'un d'entre eux consiste à encoder un qubit dans une superposition de deux états cohérents  $|\alpha\rangle$  et  $|\beta\rangle$ , formant respectivement les deux états de base  $|0\rangle$  et  $|1\rangle$  [Ralph02, Ralph03]. Bien que ces deux états ne soient pas strictement orthogonaux, une amplitude  $\alpha=1.5$  rend déjà leur recouvrement assez faible,  $|\langle\alpha|\beta\rangle|^2 \simeq 10^{-4}$ , ce qui rend le protocole compatible avec des protocoles de corrections d'erreurs [Lund08]. La comparaison avec d'autres méthodes en termes de ressources pour une intégration à grande échelle semble également favorable [Lund08, Dawson06]. Ce type d'encodage est également avantageux en terme de correction d'erreurs lorsque le qubit est envoyé dans un canal quantique imparfait [Glancy04].

Expérimentalement, les opérations réalisées sont inévitablement des versions approchées des opérations parfaites, dont le degré de ressemblance nécessaire est déterminé par l'efficacité des codes correcteurs d'erreurs [Devitt09]. La caractérisation des opérations effectivement réalisées est donc une étape indispensable afin de pouvoir accroître la complexité des circuits quantiques.

Dans ce chapitre, nous nous intéressons à la caractérisation d'une porte quantique agissant sur des superpositions d'états cohérents, en introduisant une phase  $\varphi$  :

$$x|\alpha\rangle + y|\beta\rangle \rightarrow x|\alpha\rangle + e^{i\varphi}y|\beta\rangle \quad (6.1)$$

Malgré sa simplicité, cette porte joue un rôle important, car elle peut faire partie d'un ensemble universel permettant de décomposer une opération arbitraire [Nielsen00]. La porte de phase  $\pi$  est implémentée très simplement [Marek10b], en se souvenant que les états cohérents sont des états propres de l'opérateur destruction :  $\hat{a}|\pm\alpha\rangle = \pm\alpha|\pm\alpha\rangle$ . Une phase différente de  $\pi$  peut être insérée en utilisant en plus un déplacement, comme nous le détaillerons plus loin.

Nous caractérisons expérimentalement une porte de phase  $\pi$  [Blandino12a], en nous basant sur une modélisation de son fonctionnement, plutôt que sur une approche de type boîte noire [Chuang97]. La porte est d'abord implémentée expérimentalement sur le vide comprimé produit par l'OPA, afin d'extraire les paramètres expérimentaux du modèle. Nous vérifions ensuite la consistance du modèle en appliquant la porte sur un autre état test, obtenu en soustrayant un photon du vide comprimé. Cette modélisation permet de *simuler* l'action de la porte expérimentale sur un état initial parfait, nous permettant ainsi de séparer les imperfections de la porte de celles des états produits expérimentalement. Ce modèle nous permet également d'étendre la caractérisation à d'autres phases  $\varphi$ . Nous obtenons ensuite des critères caractérisant la porte en calculant la fidélité avec des états cibles, produits par une porte parfaite.

Notre méthode permet de caractériser la porte pour la zone où elle est effective, et est avant tout orientée vers une simplicité expérimentale. Elle offre un bon compromis par rapport à une tomographie de processus quantique [O'Brien04, Lobino08], plus coûteuse expérimentalement, dont nous montrerons qu'elle serait de toute façon difficilement utilisable.

## 6.2 Présentation de la porte de phase

Commençons par présenter brièvement quelques particularités du calcul quantique avec des états cohérents. Nous détaillerons ensuite comment implémenter expérimentalement la porte de phase proposée par P. Marek et J. Fiurášek [Marek10b].

### 6.2.1 Calcul quantique avec des états cohérents

Un ordinateur quantique est dit universel s'il peut implémenter une opération unitaire arbitraire sur ses variables [DiVincenzo95]. Pour cela, il suffit de disposer d'un ensemble restreint de portes quantiques agissant sur un ou deux qubits [Nielsen00]. Plusieurs ensembles peuvent être utilisés, composés d'une porte "intriquante" à deux qubits, et de rotations à un qubit. En particulier, un ensemble constitué d'une porte de phase (équivalente à une rotation autour de  $Z$ ), d'une porte C-phase, et d'une porte de Hadamard, est universel.

Si le principe du calcul quantique à état cohérent est séduisant, la manipulation efficace d'un qubit  $x|\alpha\rangle + y|-\alpha\rangle$  fait toujours l'objet de recherches afin de réduire les ressources expérimentales nécessaires. Plusieurs méthodes ont été proposées ces dernières années, toutes basées sur l'utilisation d'états cohérents pour former un qubit, mais implémentant les portes de différentes manières. La première méthode utilisant des états cohérents faisait appel à un encodage dans les états  $|\alpha\rangle + |-\alpha\rangle$  et  $|\alpha\rangle - |-\alpha\rangle$ , et nécessitait des fortes non linéarités [Cochrane99], ce qui la rendait peu exploitable. H. Jeong *et al.* ont ensuite proposé l'utilisation de  $|\alpha\rangle$  et  $|-\alpha\rangle$  comme états de base [Jeong02], mais leur méthode nécessitait toujours des interactions non linéaires. T. Ralph *et al.* ont pu contourner ce problème, en proposant l'utilisation de chats de Shrödinger comme ressource non classique [Ralph02]. Ces états peuvent en théorie être préparés de manière probabiliste avant l'étape de calcul et stockés dans une mémoire quantique. Plusieurs améliorations du protocole ont ensuite permis de diminuer l'amplitude des états à  $\alpha > 2$  [Ralph03], puis  $\alpha > 1.2$  [Lund08], ce qui le rend davantage compatible avec les états chats produits expérimentalement [Ourjoumtsev06b, Neergaard-Nielsen06, Ourjoumtsev07c, Takahashi08, Neergaard-Nielsen10, Gerrits10].

Enfin, une dernière méthode proposée par P. Marek et J. Fiurášek apporte une simplification supplémentaire du dispositif expérimental [Marek10b], en contrepartie d'un fonctionnement davantage non déterministe. En effet, dans les références [Ralph03, Lund08] les portes sont réalisées à l'aide de téléportations - pouvant être rendues quasi-déterministes-, qui nécessitent entre autre l'utilisation d'un ou plusieurs chats pour former les états de Bell. Le protocole de P. Marek et J. Fiurášek ne nécessite un chat que pour implémenter la porte de Hadamard. Les autres portes (la porte de phase, et la porte C-phase), n'utilisent que des soustractions de photon, des déplacements, et des mesures APD, ce qui rend ce protocole plus adapté à une faisabilité expérimentale. Le groupe d'U. Andersen a d'ailleurs réalisé une démonstration de principe de la porte de Hadamard, en la testant pour les deux états logiques  $|\alpha\rangle$  et  $|-\alpha\rangle$  [Tipsmark11].

### 6.2.2 La porte de phase

#### Principe

Une porte de phase  $\varphi$  réalise la transformation :

$$x|\alpha\rangle + y|-\alpha\rangle \rightarrow x|\alpha\rangle + ye^{i\varphi}|-\alpha\rangle \quad (6.2)$$

La porte de phase  $\pi$  est implémentée très simplement, en appliquant l'opérateur  $\hat{a}$

$$\hat{a}(x|\alpha\rangle + y|-\alpha\rangle) = \alpha(x|\alpha\rangle - y|-\alpha\rangle), \quad (6.3)$$

qui correspond bien à l'introduction d'une phase  $\pi$  après normalisation. On vérifie sans difficulté qu'une phase  $\varphi$  peut être introduite par l'action d'un opérateur  $\hat{\mathbf{a}} + \beta$  :

$$[\hat{\mathbf{a}} + \beta](x|\alpha\rangle + y|-\alpha\rangle) \propto x|\alpha\rangle + e^{i\varphi}y|-\alpha\rangle \quad (6.4)$$

avec  $\beta$  qui satisfait :

$$\frac{\beta - \alpha}{\beta + \alpha} = e^{i\varphi} \Rightarrow \beta = i \frac{\alpha}{\tan(\varphi/2)} \quad (6.5)$$

L'opérateur  $\hat{\mathbf{a}} + \beta$  peut être implémenté en utilisant deux déplacements et la soustraction d'un photon du fait de la relation suivante :

$$\hat{\mathbf{D}}^\dagger(\beta)\hat{\mathbf{a}}\hat{\mathbf{D}}(\beta) = \hat{\mathbf{a}} + \beta \quad (6.6)$$

Du fait de la soustraction de photon, le fonctionnement de cette porte est donc probabiliste. En plus de la probabilité de succès de la soustraction, notons qu'il y a également une contrainte fondamentale empêchant un fonctionnement déterministe, liée à la non orthogonalité des états cohérents. Par exemple, une transformation de Hadamard  $|\pm\alpha\rangle \rightarrow |\pm\rangle = \mathcal{N}_\pm (|\alpha\rangle \pm |-\alpha\rangle)$  ne conserve pas le produit scalaire, puisque  $\langle\alpha|-\alpha\rangle \neq 0$ , alors que  $\langle+|- \rangle = 0$ . Elle ne peut donc pas être unitaire, même si elle peut assez rapidement tendre vers une transformation unitaire pour  $\alpha$  suffisamment grand. Les portes à états cohérents sont donc au mieux quasi-déterministes.

### Implémentation expérimentale

On peut en fait obtenir la même transformation que (6.6) en utilisant un seul déplacement appliqué avant l'APD, tel que schématisé sur la figure 6.1. La soustraction de photon est effectuée en prélevant  $R \simeq 10\%$  du faisceau avec une lame séparatrice (combinaison  $\lambda/2$ +PBS). Nous avons montré dans la section 3.4.3 que l'état après la lame séparatrice est approximativement

$$\hat{\mathbf{U}}_{\text{BS}}|\psi_{\text{in}}\rangle \otimes |0\rangle \simeq |\psi_{\text{in}}\rangle \otimes |0\rangle - \theta[\hat{\mathbf{a}}|\psi_{\text{in}}\rangle] \otimes |1\rangle, \quad (6.7)$$

où le deuxième mode est dirigé vers l'APD. En appliquant un déplacement  $\hat{\mathbf{D}}(\zeta)$  juste avant l'APD, l'état total devient

$$|\psi_D\rangle = |\psi_{\text{in}}\rangle \otimes |\zeta\rangle - \theta[\hat{\mathbf{a}}|\psi_{\text{in}}\rangle] \otimes [\hat{\mathbf{D}}(\zeta)|1\rangle]. \quad (6.8)$$

Or, puisque  $\hat{\mathbf{D}}(\zeta)\hat{\mathbf{b}}^\dagger\hat{\mathbf{D}}^\dagger(\zeta) = \hat{\mathbf{b}}^\dagger - \zeta^*$ , on a  $\hat{\mathbf{D}}(\zeta)|1\rangle = \hat{\mathbf{D}}(\zeta)\hat{\mathbf{b}}^\dagger\hat{\mathbf{D}}^\dagger(\zeta)\hat{\mathbf{D}}(\zeta)|0\rangle = (\hat{\mathbf{b}}^\dagger - \zeta^*)|\zeta\rangle$ . L'équation (6.8) s'écrit donc

$$|\psi_D\rangle = |\psi_{\text{in}}\rangle \otimes |\zeta\rangle - \theta[\hat{\mathbf{a}}|\psi_{\text{in}}\rangle] \otimes [(\hat{\mathbf{b}}^\dagger - \zeta^*)|\zeta\rangle]. \quad (6.9)$$

En modélisant le conditionnement réussi par l'APD par un projecteur  $|1\rangle\langle 1|$ , l'état total devient :

$$|\psi_{\text{out}}^{\text{tot}}\rangle = (\mathbb{1} \otimes |1\rangle\langle 1|)|\psi_D\rangle \quad (6.10a)$$

$$= e^{-\frac{|\zeta|^2}{2}} \left[ \zeta|\psi_{\text{in}}\rangle \otimes |1\rangle - \theta[\hat{\mathbf{a}}|\psi_{\text{in}}\rangle] \otimes (1 - |\zeta|^2)|1\rangle \right] \quad (6.10b)$$

$$= e^{-\frac{|\zeta|^2}{2}} \left[ \zeta - \theta(1 - |\zeta|^2)\hat{\mathbf{a}} \right] |\psi_{\text{in}}\rangle \otimes |1\rangle \quad (6.10c)$$

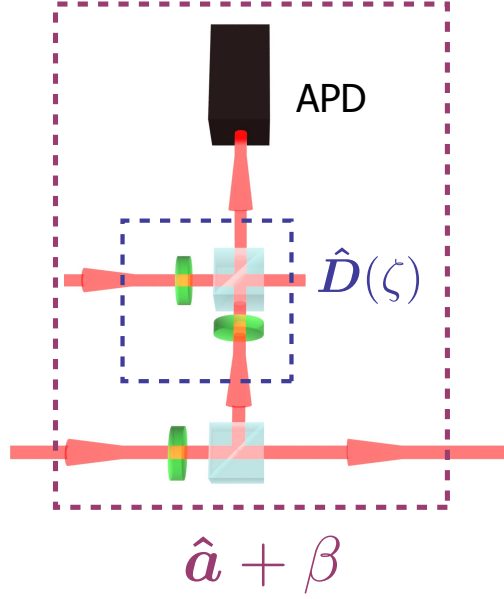


FIGURE 6.1 – Schéma de l'implémentation expérimentale de la porte de phase.

Après normalisation et trace partielle sur le mode de l'APD, l'état final est bien

$$|\psi_{\text{out}}\rangle = (\hat{a} + \beta) |\psi_{\text{in}}\rangle, \quad (6.11)$$

avec  $\beta = \frac{\zeta}{\theta(|\zeta|^2 - 1)}$ .

Le déplacement  $\hat{D}(\zeta)$  peut être implémenté avec une lame séparatrice de transmission  $T_{\text{dep}} \simeq 1$ , et un état cohérent intense d'amplitude  $\alpha_{\text{dep}}$  [Paris96]. En effet, dans ce cas  $|\alpha_{\text{dep}}\rangle$  est état propre de  $\hat{b}$ , mais aussi approximativement de  $\hat{b}^\dagger$  avec la valeur propre  $\alpha_{\text{dep}}^*$ , et donc

$$\exp \left[ \theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger) \right] |\alpha_{\text{dep}}\rangle_b \simeq \exp \left[ \theta(\hat{a}^\dagger \alpha_{\text{dep}} - \hat{a} \alpha_{\text{dep}}^*) \right] |\alpha_{\text{dep}}\rangle_b \quad (6.12)$$

correspond à un déplacement d'amplitude  $\theta \alpha_{\text{dep}}$  pour le mode  $\hat{a}$ , après trace partielle sur le mode  $\hat{b}$ . Pour  $T \simeq 1$ ,  $\theta \simeq \sqrt{1-T}$ , et on réalise donc un déplacement de  $\zeta$  en utilisant un état cohérent d'amplitude  $\alpha_{\text{dep}} = \zeta / \sqrt{1-T}$ .

La soustraction de photon a déjà été réalisée expérimentalement par plusieurs groupes. En revanche, son utilité en tant que porte n'avait pas été réalisée avant la référence [Marek10b], et elle n'avait pas fait l'objet d'une caractérisation en tant que telle.

## 6.3 Réalisation expérimentale d'une porte de phase $\pi$

### 6.3.1 Méthode et dispositif expérimental

Nous avons réalisé expérimentalement la porte de phase  $\pi$  à l'aide du dispositif expérimental présenté sur la figure 6.2. Nous pouvons implémenter deux soustractions de photon. L'une d'entre elle est interprétée comme la porte de phase que nous cherchons à caractériser (indice 1). L'autre nous sert à la préparation d'un état quantique de test (indice 0). Notre méthode est résumée ci-dessous, et détaillée dans les paragraphes suivants.



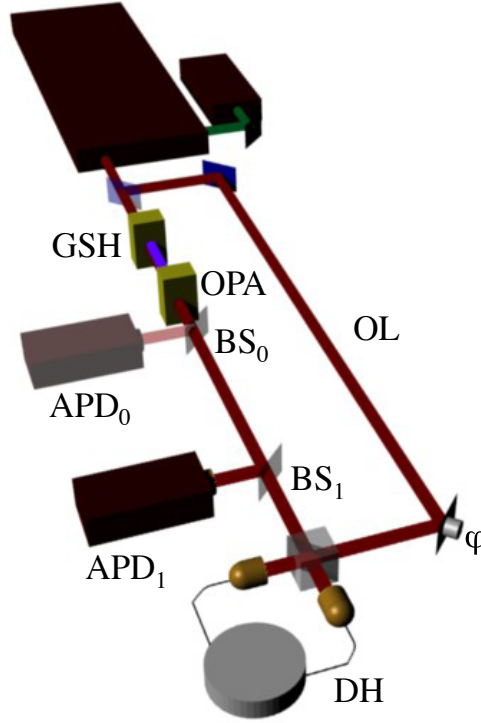


FIGURE 6.2 – Schéma du dispositif expérimental

Nous disposons d'un modèle du fonctionnement de la porte dont on cherche à estimer expérimentalement les paramètres. La première étape est donc d'appliquer la porte sur un état quantique que l'on sait également modéliser, puis d'extraire tous les différents paramètres à partir des mesures homodynes. Parmi les états candidats, les états cohérents n'apparaissent pas très adaptés, puisqu'il ne sont pas modifiés par une soustraction de photon. La porte agissant sur des superpositions d'états cohérents, il est préférable d'utiliser une telle superposition pour la tester. On peut pour cela utiliser le vide comprimé produit par l'OPA en configuration dégénérée, qui est une approximation raisonnable d'un chaton pair pour une faible compression. Cette approximation n'est valable que pour de faibles amplitudes, car le vide comprimé est un état gaussien, contrairement à un chat pair. On peut s'en convaincre en comparant la décomposition d'un chat pair (2.176)  $|+\rangle \propto |0\rangle + \frac{\alpha^2}{\sqrt{2}}|2\rangle$  et d'un vide comprimé (2.143)  $|\psi_{\text{sqz}}\rangle \propto |0\rangle - \sqrt{2}(\frac{1}{2} \tanh r)|2\rangle$ , en ne gardant que les deux premiers termes. On peut s'attendre à une bonne fidélité si  $\alpha^2 = -\tanh r$ , ce qui correspond à  $|\alpha| = 0.54$  pour  $r = 0.3$ . Numériquement, on trouve que l'état sortant de l'OPA avec nos paramètres expérimentaux a une fidélité maximale égale à 0.98 pour une amplitude  $\alpha = 0.55$ . Insistons sur le fait que pour cette estimation des paramètres expérimentaux, l'APD d'indice 0 n'est pas utilisée.

Afin de tester la qualité de notre modélisation, nous appliquons ensuite la porte sur un autre état test, et nous comparons les histogrammes obtenus expérimentalement avec ceux prédits par le modèle utilisant les paramètres expérimentaux. L'état test est obtenu en soustrayant un photon du vide comprimé avec l'APD d'indice 0 (figure 6.2). On peut voir que c'est une relativement bonne approximation d'un chat impair, en comparant  $|-\rangle \propto |1\rangle + \frac{\alpha^2}{\sqrt{6}}|3\rangle$  et  $\hat{a}|\psi_{\text{sqz}}\rangle \propto |1\rangle + (-\frac{1}{2} \tanh r)\sqrt{6}|3\rangle$ . On peut s'attendre à une bonne fidélité si  $\alpha^2 = -3 \tanh r$ , ce qui donne  $|\alpha| = 0.93$  pour  $r = 0.3$ . Après avoir appliqué la porte sur le vide comprimé, nous verrons

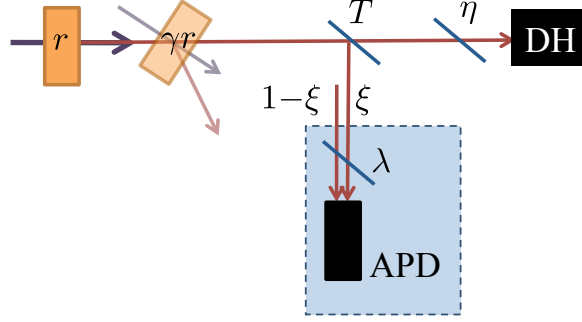


FIGURE 6.3 – Schéma du modèle utilisé pour l'extraction des paramètres expérimentaux.

que la fidélité maximale avec un chat impair est de 0.58 pour  $\alpha=0.92$ .

Nous utilisons donc en fait deux soustractions de photon : une interprétée en tant que porte, que l'on cherche à caractériser, et l'autre servant à préparer un chat impair utilisé pour tester le modèle. Nous verrons que ce dernier est en très bon accord avec les résultats expérimentaux, ce qui nous permettra de l'utiliser pour simuler l'action de la porte sur d'autres états.

### 6.3.2 Modélisation

#### Modélisation de la porte

Détaillons maintenant le modèle du fonctionnement de la porte, associé à un mode  $\hat{a}$  (figure 6.3). Expérimentalement, elle est réalisée en utilisant une lame séparatrice ( $\lambda/2$  + PBS) de transmission  $T=10\%$ , mesurée indépendamment. Le faisceau réfléchi est ensuite dirigé vers l'APD (associée à un mode  $\hat{b}$ ). Nous cherchons à implémenter une porte de phase avec  $\varphi=\pi$ , et n'avons donc pas besoin d'un déplacement. Ce déplacement rajouterait comme principaux paramètres : la transmission de la lame séparatrice supplémentaire, la qualité de l'interférence avec le faisceau à déplacer, l'incertitude sur la phase de l'état cohérent utilisé pour le déplacement et éventuellement les fluctuations d'intensité du laser, se traduisant par des fluctuations d'amplitude.

Le faisceau passe ensuite par le système de filtrage d'efficacité globale  $\lambda \simeq 10\%$ , comprenant l'efficacité quantique de l'APD. A la limite où  $\lambda \rightarrow 0$ , le conditionnement correspond à une soustraction de photon sur le mode  $\hat{b}$ . On peut s'en convaincre en remarquant que cela correspond à la soustraction de photon présentée dans la section 3.4.3, avec le mode de conditionnement dans le mode faiblement transmis, alors que presque tout le faisceau est réfléchi. On peut aussi le voir immédiatement en faisant un développement du POVM de l'APD (3.31) au premier ordre en  $\lambda$ . Dans ce cas,  $\hat{\Pi}_1 \simeq \lambda \sum_n n |n\rangle_b \langle n|_b = \lambda \hat{n}_b$ . Si  $\hat{\rho}$  est l'état quantique total, comprenant le mode du signal  $\hat{a}$  et le mode du conditionnement  $\hat{b}$ , on a bien  $\text{Tr}\{\hat{\Pi}_1 \hat{\rho}\} \simeq \lambda \text{Tr}\{\hat{n}_b \hat{\rho}\} \propto \text{Tr}\{\hat{b} \hat{\rho} \hat{b}^\dagger\}$ . Cette modélisation de l'APD permet de simplifier les calculs analytiques, tout en donnant des résultats très proches de ceux obtenus avec  $\hat{\Pi}_1$ . Elle permet également de s'affranchir de certains problèmes numériques lorsque l'on utilise des fonctions de Wigner, dus à de petites différences de grands nombres [Ourjoumtsev07a].

Le système de filtrage permet de sélectionner un mode adapté à la détection homodyne, mais il n'est pas parfait. Plusieurs sources peuvent déclencher l'APD sans que cela n'induisse de transformation sur l'état : le conditionnement peut provenir de photons décorrélés dus aux imperfections de l'OPA ou à son caractère multimode, de photons dans un mode orthogonal à celui

de la détection homodyne, ou de toute autre source de bruit. On considère qu'une fraction  $(1-\xi)$  des conditionnements provient de ces diverses sources, ce qui n'induit pas de soustraction pour l'état mesuré par la détection homodyne, mais seulement des pertes liées à la lame séparatrice de prélèvement. Pour une fraction  $\xi$ , le conditionnement a bien lieu dans le bon mode.

Les photons décorrélés susceptibles d'induire un mauvais conditionnement dépendent de l'état initial. Pour un état ne contenant que des photons dans le mode de la détection homodyne, et sans autre source de bruit, on aurait  $\xi=1$ . Malgré cela, nous avons choisi d'inclure la pureté modale  $\xi$  comme un paramètre propre à la porte, en l'interprétant comme une capacité à filtrer les photons dans les mauvais modes. De ce fait, nous considérons une caractérisation de la porte de phase pour une technologie expérimentale donnée, pour laquelle les états quantiques sont "entourés" de bruit multimode.

### Modélisation de l'état utilisé pour tester la porte

Comme nous l'avons vu dans la section 3.3.2, l'OPA est modélisée par un OPA dégénéré parfait de compression  $s=e^{-2r}$ , suivi d'un amplificateur indépendant de la phase, introduisant un gain parasite  $h=\cosh^2 \gamma r$ .

### 6.3.3 Extraction des paramètres expérimentaux

L'application de la porte de phase  $\pi$  sur le vide comprimé produit un état proche d'un chaton de Schrödinger, à partir duquel on peut extraire les paramètres du modèle avec des mesures homodynes. Aux 3 paramètres,  $\xi$ ,  $h$ , et  $s$  à estimer, il faut rajouter l'efficacité homodyne  $\eta$ , qui est estimée indépendamment à la valeur de  $\eta=0.68$ . A. Ourjoumtsev a développé un modèle analytique permettant d'extraire ces paramètres à partir des moments des données expérimentales [Ourjoumtsev07a], que nous détaillons brièvement. Le calcul analytique de la fonction de Wigner du chaton produit conduit à la définition de quatre variables  $a$ ,  $b$ ,  $a'$ , et  $b'$  :

$$a = 1 + \eta T(h/s + h - 2) \quad b = 1 + \eta T(hs + h - 2) \quad (6.13a)$$

$$a' = \frac{\eta \xi T(h/s + h - 2)^2}{h(s + 1/s) + 2h - 4} \quad b' = \frac{\eta \xi T(hs + h - 2)^2}{h(s + 1/s) + 2h - 4} \quad (6.13b)$$

qui dépendent des paramètres du modèle à estimer. Le principe est d'obtenir des valeurs  $a^*$ ,  $b^*$ ,  $a'^*$  et  $b'^*$  de ces variables à partir des données expérimentales, ce qui donne ensuite accès aux paramètres recherchés par un fit.

Pour une quadrature  $\hat{X}_\theta$ , on montre que la distribution de probabilité  $P_\theta(x_\theta)$  ne dépend que de deux paramètres  $c$  et  $c'$ , donnés par

$$c = a \cos^2 \theta + b \sin^2 \theta, \quad \text{et} \quad c' = a' \cos^2 \theta + b' \sin^2 \theta. \quad (6.14)$$

Ces paramètres s'expriment également en fonction des moments d'ordres deux et quatre, respectivement notés  $\mu_2$  et  $\mu_4$  :

$$c = 2 \left( \mu_2 - \sqrt{\mu_2^2 - \mu_4/3} \right) \quad c' = \sqrt{\mu_2^2 - \mu_4/3} \quad (6.15)$$

On peut donc en obtenir des estimations  $\bar{c}$  et  $\bar{c}'$  à partir des estimations  $V_2$  et  $V_4$  de  $\mu_2$  et  $\mu_4$ ,

$$\bar{c} = 2 \left( \bar{V}_2 - \sqrt{\bar{V}_2^2 - \bar{V}_4/3} \right), \quad \bar{c}' = \sqrt{\bar{V}_2^2 - \bar{V}_4/3}. \quad (6.16)$$

En utilisant plusieurs phases  $\theta$ , on peut obtenir  $a^*$ ,  $b^*$ ,  $a'^*$  et  $b'^*$  par régression linéaire, en posant  $\bar{c}=a^*y+b^*(1-y)$  et  $\bar{c}'=a'^*y+b'^*(1-y)$ , avec  $y=\cos^2\theta$ . Enfin, on détermine les valeurs de  $\xi$ ,  $h$  et  $s$  en minimisant la fonction

$$\mathcal{L} = (a-a^*)^2 + (b-b^*)^2 + (a'-a'^*)^2 + (b'-b'^*)^2 \quad (6.17)$$

par rapport à ces paramètres.

### 6.3.4 Résultats expérimentaux

Nous avons d'abord appliqué la porte de phase  $\pi$  sur le vide comprimé monomode produit par l'OPA afin d'extraire les paramètres du modèle. Nous avons acquis 400 000 mesures de quadratures, avec un taux d'environ 6000 conditionnements par seconde. La phase de l'oscillateur local n'étant pas contrôlée, on répartit les mesures en 6 intervalles de phase en utilisant la variance du vide comprimé mesurée sur 1000 points juste après chaque conditionnement. La variance la plus faible correspond à la quadrature  $\hat{X}$ , et la variance la plus élevée à la quadrature  $\hat{P}$ . Notons que ce choix est purement conventionnel et dépend de la convention utilisée pour l'opérateur de squeezing modélisant l'OPA. Si on utilise une convention  $\hat{S}=\exp[\frac{r}{2}(\hat{a}^{\dagger 2}-\hat{a}^2)]$ , c'est la quadrature  $\hat{P}$  qui est comprimée. Cette méthode de tri suppose de plus que l'état est symétrique, puisque nous ne pouvons pas distinguer un état comprimé selon  $\hat{X}_\theta$  d'un état comprimé selon  $\hat{X}_{\pi-\theta}$ .

Les données ordonnées sont ensuite divisées en 6 blocs de  $6.6 \times 10^4$  points et regroupées en histogrammes de 64 bins. Chaque bloc  $k$  correspond à une phase  $[k\pi/12, (k+1)\pi/12]$ ,  $k$  allant de 0 à 5. A partir de ces données, on calcule ensuite les moments  $V_2$  et  $V_4$  pour chaque phase, ce qui nous permet ensuite d'extraire les paramètres du modèle.

Les résultats sont présentés sur la figure 6.4, qui montre un très bon accord entre ce modèle analytique et les données expérimentales brutes. Les tables 6.1 et 6.2 regroupent respectivement les paramètres du fit et les estimations des paramètres expérimentaux.

| $a^*$ | $a'^*$ | $b^*$ | $b'^*$ |
|-------|--------|-------|--------|
| 1.54  | 0.86   | 0.74  | 0.19   |

TABLE 6.1 – Paramètres du fit.

| $r$  | $\gamma$ | $s = \exp(-2r)$ | $h = \cosh^2(\gamma r)$ | $\eta$ | $\xi$ | $T$ |
|------|----------|-----------------|-------------------------|--------|-------|-----|
| 0.30 | 0.46     | 0.54            | 1.02                    | 0.68   | 0.83  | 0.9 |

TABLE 6.2 – Estimations des paramètres expérimentaux.

Le paramètre le plus important est  $\xi$ , car c'est le seul – avec la transmission  $T$  de la séparatrice – dont dépend la porte expérimentale. Les paramètres  $s$  et  $h$  de l'OPA caractérisent la qualité du chaton initial, et l'efficacité homodyne  $\eta$  n'est pas une caractéristique de l'état obtenu par la porte.

### 6.3.5 Test du modèle de la porte

Afin de vérifier la consistance de ce modèle, on l'utilise maintenant pour prédire le fonctionnement de la porte sur un autre état test, en gardant les paramètres estimés. Cet autre état est obtenu en soustrayant un photon du vide comprimé. Avec l'action de la porte, nous avons donc une double soustraction, produite avec un taux d'environ 6 conditionnements par seconde.

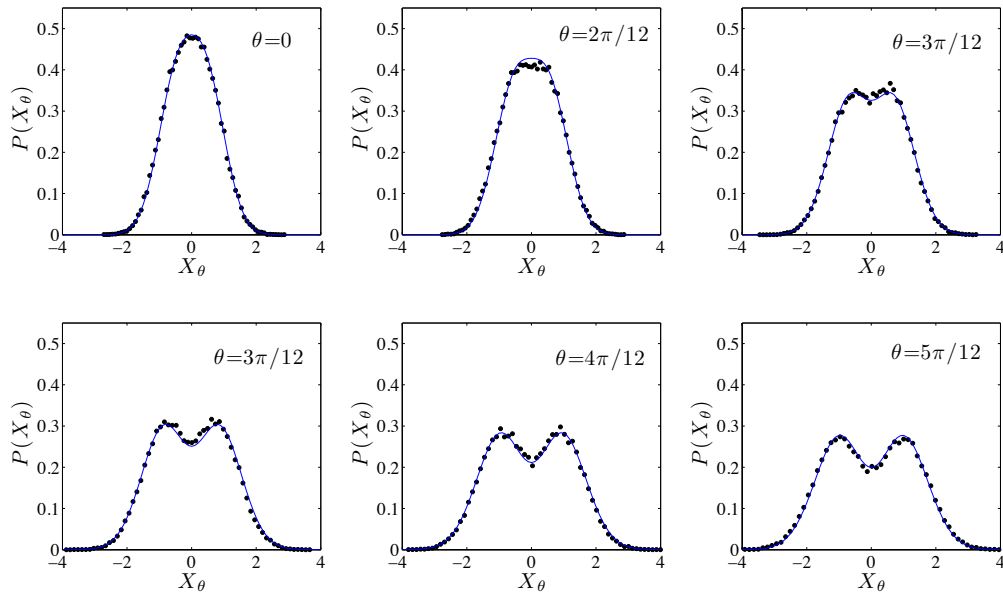


FIGURE 6.4 – Mesures de quadratures pour la porte de phase  $\pi$  appliquée sur le vide comprimé. Les traits pleins correspondent au fit du modèle analytique. Notons que la convention choisie pour définir les quadratures comprimées est différente de [Blandino12a].

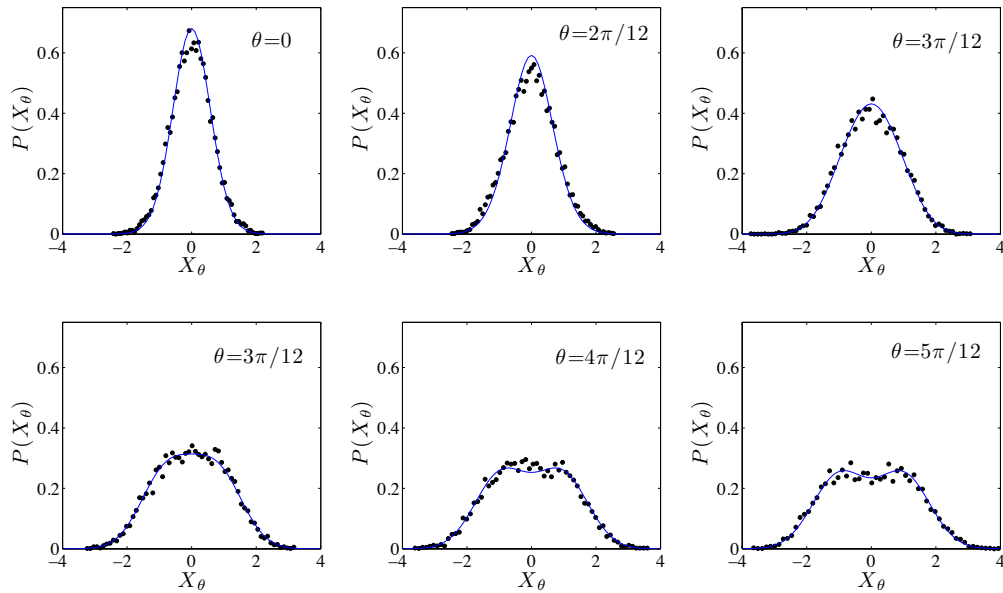


FIGURE 6.5 – Mesures de quadratures pour l'application de la porte de phase  $\pi$  sur le vide comprimé soustrait d'un photon (double soustraction). Les lignes sont les prévisions du modèle avec les paramètres précédemment estimés, sans fit pour cet état.

Nous mesurons 60 000 points, avec deux mesures de 15 000 points et une de 30 000 points. Le dispositif expérimental est contrôlé et au besoin réaligné entre chaque mesure afin de limiter les dérives.

Les résultats expérimentaux sont montrés sur la figure 6.5. Les traits pleins correspondent aux prévisions du modèle utilisant les paramètres précédemment obtenus. Il ne s'agit donc pas d'un fit pour cet état. Nous faisons l'hypothèse que le paramètre  $\xi$  est le même pour les deux APD, ce qui a été confirmé par plusieurs tests. Ceci nous permet d'utiliser le même modèle pour les deux soustractions.

Le bon accord avec les mesures nous permet de valider la consistance de notre modèle, que nous pouvons maintenant utiliser pour caractériser la porte de phase.

### 6.3.6 Incertitude sur l'estimation de $\xi$

Nous verrons que le paramètre  $\xi$  est d'une grande influence sur la qualité de la porte, il est donc nécessaire de quantifier nos incertitudes sur son estimation. Celles ci proviennent des erreurs statistiques lors du calcul des moments  $V_2$  et  $V_4$ , et de l'estimation de la phase lors du tri des données.

**Incertitudes statistiques** Pour chaque intervalle de phase, les  $N=6.6 \times 10^4$  mesures homodynes sont indépendantes et décorréées, et on suppose qu'elles proviennent d'une même variable aléatoire  $X$ , non gaussienne pour l'état conditionné. On note  $V_k = \frac{1}{N} \sum_{i=1}^k x_i^k$  le moment d'ordre  $k$  estimé à partir des mesures. En considérant  $X^k$  comme une variable aléatoire, sa moyenne est égale à  $\mu_k = \langle X^k \rangle$ , et sa variance est égale à  $\sigma^2(X^k) = \langle X^{2k} \rangle - \langle X^k \rangle^2 = \mu_{2k} - \mu_k^2$ .

On suppose que  $N$  est suffisamment grand pour que le théorème central limite s'applique. Par conséquent,  $V_k$  suit une loi normale de moyenne  $\mu_k$  et de variance  $\sigma^2(X^k)/N$ . On peut donc former la variable

$$\mathfrak{U} = \frac{V_k - \mu_k}{\sqrt{\sigma^2(X^k)/N}}, \quad (6.18)$$

qui suit une loi normale centrée réduite. Nous avons rappelé au chapitre sur l'estimation de la discorde (cf. 5.3.3) que l'estimation de la variance d'une variable aléatoire gaussienne de variance  $\sigma^2$ , suit également une loi normale de variance  $2\sigma^4/N$  et de moyenne  $\sigma^2$ . En conséquence, l'estimation de la variance de  $X^k$

$$\Delta^2(X^k) = \frac{1}{N} \sum_{i=1}^N (x_i^k - V_k)^2 = V_{2k} - V_k^2 \quad (6.19)$$

suit également une loi normale de moyenne  $\sigma^2(X^k)$  et de variance  $2[\sigma^2(X^k)]^2/N$ . Ainsi, on peut former une variable aléatoire

$$\chi_N^2 = N \frac{\Delta^2(X^k)}{\sigma^2(X^k)}, \quad (6.20)$$

qui est gaussienne, de moyenne  $N$  et de variance  $2N$ . On peut donc l'interpréter comme une loi du  $\chi^2$  à un nombre infini de degrés de liberté. Par conséquent, la variable

$$\frac{\mathfrak{U}}{\sqrt{\chi_N^2/N}} = \frac{V_k - \mu_k}{\sqrt{\sigma^2(X^k)/N}} \sqrt{\frac{\sigma^2(X^k)}{\Delta^2(X^k)}} = \frac{V_k - \mu_k}{\sqrt{\Delta^2(X^k)/N}} \quad (6.21)$$

suit une loi de Student<sup>1</sup> à un nombre infini de degrés de liberté, qui est égale à une loi normale centrée réduite.

En conclusion, on a donc montré que les estimations des moments  $\mu_2$  et  $\mu_4$  suivent des lois normales respectivement de moyenne  $V_2$  et  $V_4$ , et de variances

$$\Delta^2 V_2 = \frac{V_4 - V_2^2}{N}, \quad (6.22)$$

$$\Delta^2 V_4 = \frac{V_8 - V_4^2}{N}. \quad (6.23)$$

**Incertitudes sur la phase** Compte tenu de la méthode de tri des données, on considère que la phase n'est pas connue à mieux que  $\pi/12$ . Nous faisons l'hypothèse d'une distribution uniforme à l'intérieur de chaque intervalle de phase.

### Propagation des incertitudes

Les incertitudes sont propagées par une méthode Monte Carlo similaire à celle utilisée pour la discorde, afin d'obtenir les incertitudes sur  $\xi$ . On simule un grand nombre de fois les paramètres expérimentaux en prenant en compte leurs incertitudes : pour  $V_2$  (resp.  $V_4$ ), cela revient à ajouter un terme tiré selon une loi normale de moyenne 0 et d'écart-type  $\Delta V_2$  (resp.  $\Delta V_4$ ). Pour la phase, on tire une variable aléatoire uniformément répartie sur  $[0, \pi/12]$ , que l'on ajoute à la borne inférieure de l'intervalle considéré. Pour chaque tirage de ces variables aléatoires, on calcule ensuite la valeur de  $\xi$  correspondante.

On répète cette procédure 50 000 fois, afin d'acquérir un nombre de points suffisant. On obtient alors une distribution des valeurs de  $\xi$ , dont l'écart-type  $\Delta \xi$  correspond à notre incertitude. Le nombre de tirage est ici inférieur à celui de la discorde, pour des raisons de temps de calcul. Cela ne pose pas de problème particulier, car 50 000 tirages assurent déjà une bonne précision. De plus, nous ne cherchons qu'un ordre de grandeur, sans vouloir faire une comparaison précise avec une borne de Cramér-Rao.

Pour les données expérimentales, on obtient

$$\Delta \xi = 0.04. \quad (6.24)$$

## 6.4 Comment caractériser la porte ?

### 6.4.1 Ressemblance des états expérimentaux avec des chats parfaits

Les états produits expérimentalement présentent une certaine ressemblance avec des superpositions d'états cohérents idéales, mais ils n'en restent cependant que des approximations, dont les imperfections sont difficiles à séparer de celles de la porte. La comparaison avec des chats parfaits<sup>2</sup>  $|\alpha\rangle + e^{i\phi} |-\alpha\rangle$  est montrée sur les figures 6.6, 6.7, et 6.8, respectivement pour le vide comprimé, l'état test (vide comprimé soustrait d'un photon), et l'état obtenu en appliquant la porte sur cet état test (avec donc deux soustractions).

1. Soient  $U$  et  $\chi_\nu^2$  des variables indépendantes qui suivent respectivement une loi normale centrée réduite et une loi du  $\chi^2$  à  $\nu$  degrés de liberté. On rappelle que par définition, une variable  $T = \frac{U}{\sqrt{\chi_\nu^2/\nu}}$  suit une loi de Student à  $\nu$  degrés de liberté. Pour  $\nu \rightarrow \infty$ , la loi de Student tend vers une loi normale centrée réduite.

2. En fait, un chat pair de faible amplitude est proche d'un vide comprimé selon la quadrature  $\hat{P}$ , alors qu'avec notre convention pour l'opérateur de squeezing, l'état de l'OPA est comprimé selon  $\hat{X}$ . Il faut donc appliquer une rotation de  $\pi/2$  à l'un des deux états pour pouvoir les comparer entre eux.

Deux problématiques expérimentales sont donc à distinguer : d'une part, la qualité du chaton produit et la ressemblance par rapport à un vrai chaton, et d'autre part la réalisation expérimentale de la porte seule, et sa différence par rapport à la porte idéale, indépendamment de l'état initial utilisé.

### 6.4.2 Tomographie de processus quantique

**Principe** L'action d'un processus quantique peut être entièrement caractérisée par une *tomographie de processus quantique* (QPT) [Mohseni08, Chuang97, Poyatos97], toutefois nous allons voir que cette méthode n'est pas adaptée à la caractérisation de la porte. En appliquant le processus sur un ensemble d'états quantiques servant de sondes, on peut en déduire la transformation d'un état quelconque, sans recourir à une quelconque modélisation. Il n'est donc pas nécessaire d'avoir une connaissance *a priori* du processus, qui peut être une "boîte noire".

Plus précisément, considérons un processus quantique  $\mathcal{E}$  agissant sur une matrice densité quelconque  $\hat{\rho}$ , qui s'écrit sous la forme  $\hat{\rho} = \sum_{n,m} \rho_{nm} |n\rangle\langle m|$ . La linéarité de  $\mathcal{E}$  implique que

$$\mathcal{E}(\hat{\rho}) = \sum_{nm} \rho_{nm} \mathcal{E}(|n\rangle\langle m|). \quad (6.25)$$

En décomposant  $\mathcal{E}(|n\rangle\langle m|) = \sum_{i,j} \mathcal{E}_{ij}^{nm} |i\rangle\langle j|$  sur la même base  $\{|i\rangle\langle j|\}$  que  $\hat{\rho}$ , on a

$$\mathcal{E}(\hat{\rho}) = \sum_{i,j,n,m} \mathcal{E}_{ij}^{nm} \rho_{nm} |i\rangle\langle j|. \quad (6.26)$$

La connaissance des coefficients  $\mathcal{E}_{ij}^{nm}$  caractérise donc complètement le processus  $\mathcal{E}$ , et permet de prédire son action sur un état quelconque. La QPT a pour objet de déterminer expérimentalement ces coefficients : pour un espace de dimension  $d$ , on applique  $\mathcal{E}$  sur  $d^2$  états linéairement indépendants, et on mesure les états produits par tomographie quantique. Plusieurs stratégies existent afin d'adopter la méthode la plus optimale en fonction du processus considéré [Mohseni08]. Grâce à la linéarité de  $\mathcal{E}$ , on peut également obtenir l'évolution des cohérences  $\mathcal{E}(|n\rangle\langle m|)$ , avec  $n \neq m$ , en combinant plusieurs mesures de populations  $\mathcal{E}(|k\rangle\langle k|)$ .

Plusieurs expériences ont montré la faisabilité de la QPT pour des opérations simples à un ou deux qubits [O'Brien04, Mitchell03]. Cette méthode est également généralisable à des processus non déterministes, en tenant compte des probabilités de succès pour chaque état sonde [Bongioanni10].

La préparation expérimentale d'états de Fock est pour l'instant limitée à quelques photons, et ne permet donc pas de caractériser des processus agissant sur des espaces de plus grande dimension. Pour cela, il existe une autre approche, plus complexe, utilisant un ensemble d'états cohérents pour obtenir les coefficients  $\mathcal{E}_{ij}^{nm}$  [Lobino08, Anis12, Rahimi-Keshari11]. En utilisant la décomposition des opérateurs  $|n\rangle\langle m|$  sur les états cohérents avec la fonction  $P$ ,

$$|n\rangle\langle m| = \int d^2\alpha P_{nm}(\alpha) |\alpha\rangle\langle\alpha|, \quad (6.27)$$

on peut obtenir les coefficients  $\mathcal{E}_{ij}^{nm}$  en mesurant  $\mathcal{E}(|\alpha\rangle\langle\alpha|)$  :

$$\mathcal{E}_{ij}^{nm} = \int d^2\alpha P_{nm}(\alpha) \langle i | \mathcal{E}(|\alpha\rangle\langle\alpha|) | j \rangle \quad (6.28)$$

La plupart des états quantiques possèdent une fonction  $P$  hautement singulière qui ne permet pas d'utiliser directement la relation (6.27). Afin de contourner cette difficulté, il est possible de



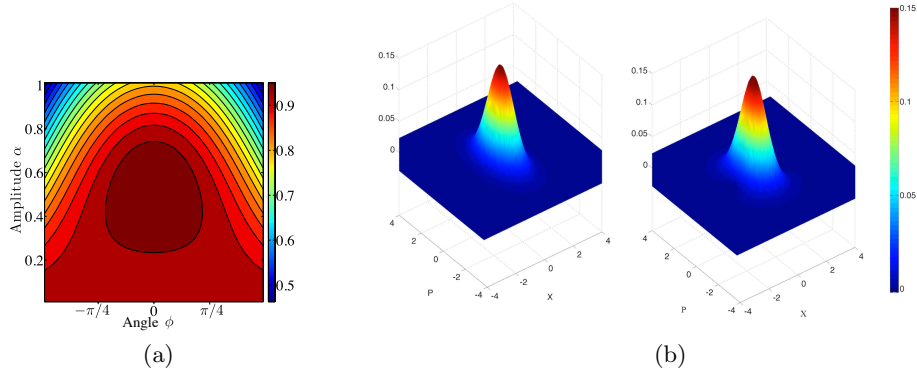


FIGURE 6.6 – (a) Fidélité entre le vide comprimé  $\hat{\rho}_{\text{opa}}$  produit par l'OPA et une superposition  $\mathcal{N}_{\pi,\phi}(|\alpha\rangle + e^{i\phi}|\alpha\rangle)$ ; (b) fonction de Wigner de  $\hat{\rho}_{\text{opa}}$  (gauche), et du chat pair maximisant la fidélité (0.98), d'amplitude  $\alpha=0.55$  (droite).

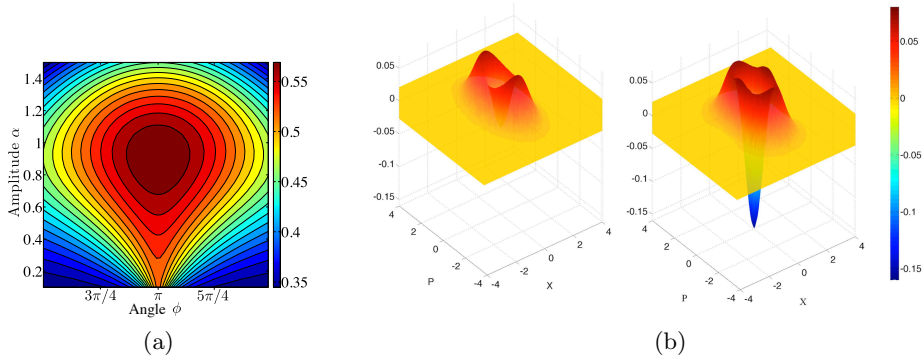


FIGURE 6.7 – (a) Fidélité entre le vide comprimé soustrait d'un photon  $\hat{\rho}_{\text{out}}^1$  et une superposition  $\mathcal{N}_{\pi,\phi}(|\alpha\rangle + e^{i\phi}|\alpha\rangle)$ ; (b) fonction de Wigner de  $\hat{\rho}_{\text{out}}^1$  (gauche), et du chat impair maximisant la fidélité (0.58), d'amplitude  $\alpha=0.92$  (droite).

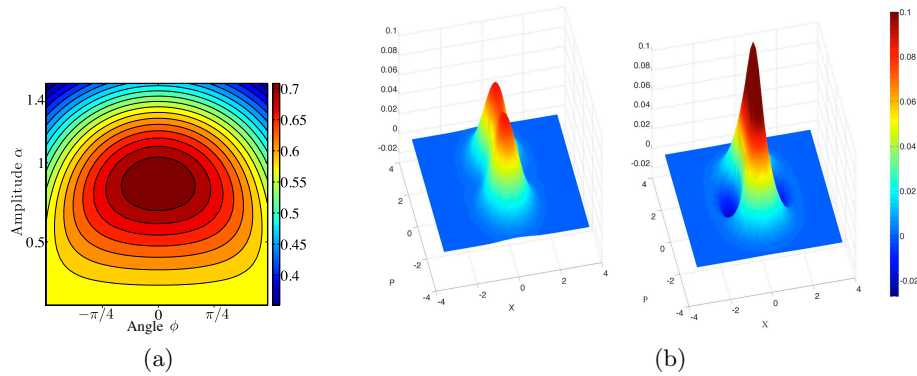


FIGURE 6.8 – (a) Fidélité entre le vide comprimé soustrait de deux photons  $\hat{\rho}_{\text{out}}^2$  et une superposition  $\mathcal{N}_{\pi,\phi}(|\alpha\rangle + e^{i\phi}|\alpha\rangle)$ ; (b) fonction de Wigner de  $\hat{\rho}_{\text{out}}^2$  (gauche), et du chat pair maximisant la fidélité (0.73), d'amplitude  $\alpha=0.87$  (droite).

reconstruire une fonction  $P$  sur laquelle on applique un filtrage ne conservant que les basses fréquences de sa transformée de Fourier [Lobino08], ou bien d'utiliser des relations supplémentaires afin de ne pas avoir explicitement besoin de  $P_{nm}(\alpha)$  [Rahimi-Keshari11].

Cette technique de tomographie peut également s'utiliser pour des processus non déterministes. Même si les états cohérents sont extrêmement faciles à produire, elle présente quand même un inconvénient. Les différentes étapes de reconstruction numérique utilisées dans [Lobino08, Rahimi-Keshari11] ne garantissent pas que le processus obtenu soit physique (c'est-à-dire complètement positif et qui n'augmente pas la trace) [Anis12]. Ce problème peut être résolu en utilisant une reconstruction par maximum de vraisemblance, nécessitant une information *a priori* [Anis12], mais qui ne permet plus de caractériser une boîte noire.

**Limitations** Dans notre cas, plusieurs contraintes limitent sérieusement les possibilités d'une tomographie de processus quantique pour caractériser la porte. D'une part, on ne peut pas utiliser une approche de type qubit, où l'on considère que l'état quantique appartient à un espace de dimension deux. En effet, même si seulement deux états  $|\alpha\rangle$  et  $|- \alpha\rangle$  sont utilisés en théorie pour former un qubit, en pratique les imperfections de la porte font "sortir" l'état de ce sous-espace. De plus, une tomographie limitée à ce sous-espace nécessiterait de pouvoir préparer des superpositions parfaites d'états cohérents, ce qui n'est pas du ressort des possibilités expérimentales actuelles.

Une tomographie avec des états cohérents pourrait être envisageable afin de surmonter ce problème. Elle a d'ailleurs récemment été utilisée pour la tomographie d'une soustraction photon [Kumar12]. En revanche, comme montré dans l'annexe D, la prise en compte du filtrage de la porte, non considérée dans [Kumar12], nécessiterait d'effectuer une tomographie *multimode*, qui ne soit pas seulement limitée au mode spatio-temporel de l'oscillateur local. Cela serait possible en théorie [Rahimi-Keshari11], mais outre les difficultés expérimentales pour la mettre en œuvre, il se pose le problème de savoir quels sont les modes qui devraient être utilisés.

En fin de compte, une tomographie de processus quantique serait donc extrêmement difficile à réaliser expérimentalement, et apporterait également beaucoup plus d'information que nécessaire, puisqu'elle nécessite de sonder une plus grande partie de l'espace des phases que celle qui nous concerne. La méthode que nous présentons offre alors une alternative beaucoup plus simple, en utilisant la modélisation de la porte que nous avons présenté.

### 6.4.3 Utilisation de la modélisation de la porte

Puisque nous ne cherchons pas à caractériser une boîte noire, mais bien l'action de la porte de phase, nous pouvons utiliser une approche plus simple qu'une tomographie de processus quantique, en utilisant notre modélisation de la porte pour simuler son action sur une superposition initiale parfaite d'états cohérents.

De cette manière, on peut s'affranchir des imperfections dues au fait que l'état initial n'est pas parfaitement dans le sous-espace  $\{|\alpha\rangle, |- \alpha\rangle\}$ . Puisque notre caractérisation est orientée vers une réalisation expérimentale de la porte, nous considérons que le qubit initial peut quand même être "entouré" de photons dans des modes non corrélés, qui déclenchent l'APD pour une fraction  $1-\xi$  des conditionnements. Le modèle présenté dans l'annexe D montre comment relier  $\xi$  à un état  $\hat{\rho}_{\text{dh}} \otimes \hat{\rho}_{\text{N}}$ , composé du qubit parfait  $\hat{\rho}_{\text{dh}}$  dans le mode de la détection homodyne, et de photons non corrélés  $\hat{\rho}_{\text{N}}$  dans d'autres modes.

Nous utiliserons deux méthodes pour caractériser la porte : la première est un calcul de fidélité pour chaque superposition initiale possible. Cette méthode nous permet d'étudier la dépendance du fonctionnement de la porte par rapport à l'état initial utilisé. La seconde méthode est basée

sur l'isomorphisme de Jamiolkowski [Jamiolkowski72], qui permet d'associer un état quantique à un processus quantique, et d'utiliser une mesure de distances habituelle, telle que la fidélité, pour comparer la réalisation expérimentale au processus idéal [Gilchrist05].

## Modélisation numérique

Pour les deux méthodes, nous aurons recours à une modélisation numérique plutôt qu'à un calcul analytique. L'extraction des paramètres expérimentaux a été faite en utilisant le modèle analytique simplifié pour l'efficacité du conditionnement, remplaçant l'APD par un opérateur  $\hat{\mathbf{b}}$ , donnant néanmoins de très bons résultats. Cependant, afin d'étudier le comportement de la porte et l'influence des différents paramètres, il est souhaitable de pouvoir changer le type de détecteur en prenant en compte une efficacité  $\lambda$ , ou un compteur parfait  $|1\rangle\langle 1|$ . On veut aussi pouvoir rajouter un déplacement pour simuler l'action d'une porte avec une phase différente de  $\pi$ . Une étude numérique permet donc plus de souplesse, en utilisant la Quantum Optics Toolbox pour Matlab, qui permet d'utiliser directement un formalisme d'opérateurs sur des états quantiques. Elle est de ce fait complémentaire et indépendante des calculs analytiques. De cette manière, on peut très facilement choisir le type d'état initial  $\hat{\rho}$ , qui est selon les cas l'état produit par l'OPA, l'état produit par l'OPA soustrait d'un photon, ou une superposition d'états cohérents parfaite. On peut ensuite choisir le type de détecteur utilisé pour le conditionnement, et calculer la fidélité avec un état cible.

Les différents éléments de la simulation sont résumés ci-dessous.

- On fixe une dimension de l'espace de Hilbert  $N$  suffisamment grande pour négliger les effets dus à la troncature. En raison des nombreuses traces partielles,  $N=12$  assure un temps de calcul raisonnable avec des effets de troncature négligeables.
- Etat produit par l'OPA : l'état produit par l'OPA avant trace partielle sur le mode fictif de l'amplificateur parasite est

$$|\psi_{\text{opa}}\rangle = e^{\gamma r(\hat{a} \otimes \hat{a} - \hat{a}^\dagger \otimes \hat{a}^\dagger)} \left[ e^{\frac{\gamma}{2}(\hat{a}^2 - \hat{a}^{\dagger 2})} \otimes \mathbb{I} \right] |0\rangle \otimes |0\rangle, \quad (6.29)$$

où le deuxième mode correspond au mode fictif  $\hat{c}$  dans la description de la section 3.3.2. L'état produit par l'OPA s'obtient en faisant la trace partielle sur le deuxième mode :

$$\hat{\rho}_{\text{opa}} = \text{Tr}_2\{|\psi_{\text{opa}}\rangle\langle\psi_{\text{opa}}|\} \quad (6.30)$$

- Superposition parfaite d'états cohérents : un état  $|\psi_{\theta,\phi}\rangle$  est obtenu très facilement avec

$$|\psi_{\theta,\phi}\rangle = \mathcal{N}_{\theta,\phi} \left[ \cos(\theta/2) \hat{\mathbf{D}}(\alpha) + e^{i\phi} \sin(\theta/2) \hat{\mathbf{D}}(-\alpha) \right] |0\rangle. \quad (6.31)$$

- Opérateurs utilisés pour le conditionnement : Une lame séparatrice de transmission  $T$  pour le premier mode est modélisée par un opérateur

$$\hat{U}_{\text{BS}}(T) = e^{\text{acos}[\sqrt{T}](\hat{a}^\dagger \otimes \hat{a} - \hat{a} \otimes \hat{a}^\dagger)}. \quad (6.32)$$

Le projecteur  $\hat{\mathbf{P}}$  de la détection est, selon les cas, soit  $\sum_n [1 - (1-\lambda)^n] |n\rangle\langle n|$  pour une APD avec une efficacité  $\lambda < 1$ , soit  $\mathbb{I} - |0\rangle\langle 0|$  pour une APD parfaite, soit  $|1\rangle\langle 1|$  pour un compteur de photon parfait.

Pour le déplacement, on utilise une lame séparatrice de transmission  $T_{\text{dep}}$  et un état cohérent d'amplitude  $\alpha_{\text{dep}} = \zeta / \sqrt{1 - T_{\text{dep}}}$ , avec  $\zeta = -(1 - \sqrt{1 + (\theta\beta)^2}) / (2\theta\beta)$ .

- Conditionnement dans le bon mode : un conditionnement dans le bon mode est finalement modélisé par un opérateur

$$\hat{E} = \left[ \mathbb{I} \otimes \hat{P} \otimes \mathbb{I} \right] \left[ \mathbb{I} \otimes \hat{U}_{BS}(T_{dep}) \right] \left[ \hat{U}_{BS}(T) \otimes \mathbb{I} \right], \quad (6.33)$$

où les modes correspondent respectivement au signal, à la voie de conditionnement, et à l'état de déplacement. En posant  $\hat{\rho}_{tot} = \hat{\rho} \otimes |0\rangle\langle 0| \otimes |\alpha_{dep}\rangle\langle \alpha_{dep}|$ , l'état résultant de ce bon conditionnement est

$$\hat{\rho}'_{out} = \frac{\text{Tr}_{2,3}\{\hat{E}\hat{\rho}_{tot}\hat{E}^\dagger\}}{\text{Tr}\{\hat{E}\hat{\rho}_{tot}\hat{E}^\dagger\}}. \quad (6.34)$$

- Conditionnement dans le mauvais mode : lorsque le conditionnement provient d'un photon décorrélé de  $\hat{\rho}$ , la transformation est simplement due au passage dans la première lame séparatrice de transmission  $T$  :

$$\hat{\rho}^x_{out} = \text{Tr}_2\{\hat{U}_{BS}(T)[\hat{\rho} \otimes |0\rangle\langle 0|]\hat{U}_{BS}^\dagger(T)\} \quad (6.35)$$

- Etat total après le conditionnement : on suppose qu'une fraction  $\xi$  est due à un conditionnement dans le bon mode, et qu'une fraction  $(1-\xi)$  est due à un conditionnement dans un mauvais mode. L'état total après le conditionnement est donc

$$\hat{\rho}_{out} = \xi \hat{\rho}'_{out} + (1-\xi) \hat{\rho}^x_{out} \quad (6.36)$$

Lorsque l'on applique la porte sur l'état de l'OPA soustrait d'un photon, cela revient à appliquer ce conditionnement deux fois sur l'état de l'OPA.

Notons que la quantum optics toolbox est un outils très pratique, mais qui possède quelques inconvénients : lors d'une trace partielle, le résultat est le complexe conjugué du résultat attendu. Il suffit d'en tenir compte pour la phase de l'état cohérent de déplacement.

## 6.5 Fidélité pour un qubit initial parfait

### 6.5.1 Porte expérimentale

Avec cette première méthode de caractérisation, on applique la modélisation de la porte expérimentale sur un qubit initial parfait

$$|\psi_{\theta,\phi}\rangle = \mathcal{N}_{\theta,\phi} \left( \cos(\theta/2)|\alpha\rangle + e^{i\phi} \sin(\theta/2)|-\alpha\rangle \right), \quad (6.37)$$

avec  $\mathcal{N}_{\theta,\phi} = 1/\sqrt{1 + e^{-2|\alpha|^2} \sin \theta \cos \phi}$ . On note  $\hat{\rho}_{\theta,\phi}^{out}$  l'état obtenu. La porte parfaite ajoute une phase  $\varphi$ , et transforme donc  $|\psi_{\theta,\phi}\rangle$  en  $|\psi_{\theta,\phi+\varphi}\rangle$ . On calcule ensuite la fidélité entre les deux états, en faisant varier  $\theta$  et  $\phi$  :

$$\mathcal{F}_{\theta,\phi} = \langle \psi_{\theta,\phi+\varphi} | \hat{\rho}_{\theta,\phi}^{out} | \psi_{\theta,\phi+\varphi} \rangle \quad (6.38)$$

La figure 6.9 montre  $\mathcal{F}_{\theta,\phi}$  pour une porte de phase  $\varphi=\pi$ , en faisant varier  $T$  et  $\xi$  afin de voir l'influence de ces paramètres. La figure 6.10 montre les mêmes calculs, pour une porte de phase  $\varphi=\pi/3$ . Dans tous les cas, le chaton initial a une taille  $\alpha=0.92$ .

La fidélité est très proche de 1 aux pôles, puisqu'on a des états cohérents qui ne sont pas transformés sous l'action de  $\hat{a}$ . La pureté modale  $\xi$  n'a donc pas d'influence pour ces états, qui sont simplement légèrement atténués par la transmission  $T$ . La fidélité décroît ensuite à mesure que  $\theta$  augmente, et est minimale pour  $\theta=\pi/2$ . Pour  $T=0.9$ , on remarque une forte directivité en fonction de  $\phi$ , qui n'est pas la même pour les deux valeurs de  $\varphi$ . Pour une phase  $\varphi=\pi$ , la fidélité

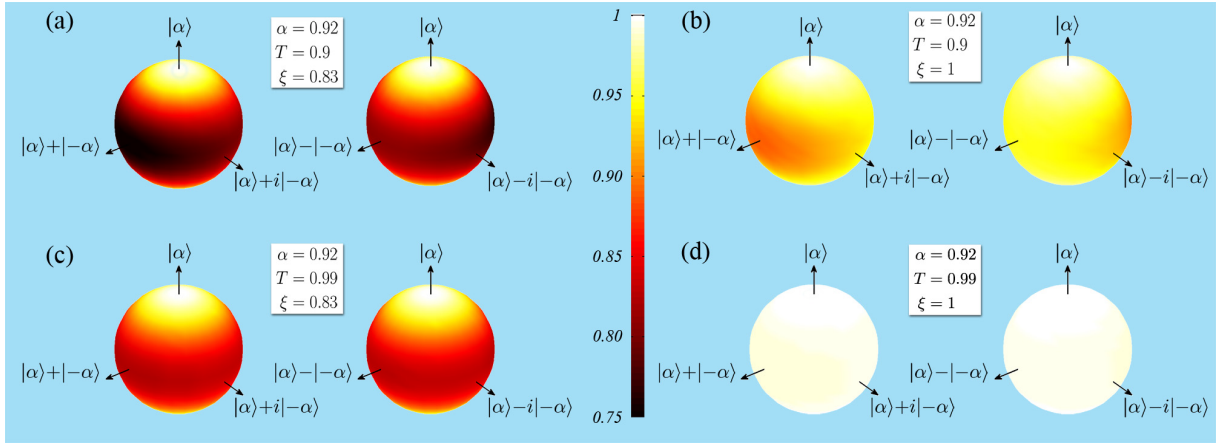


FIGURE 6.9 – Fidélité  $\mathcal{F}_{\theta,\phi}$ , pour une porte de phase  $\varphi=\pi$ . Le détecteur est modélisé par une APD d'efficacité  $\lambda=0.1$ .

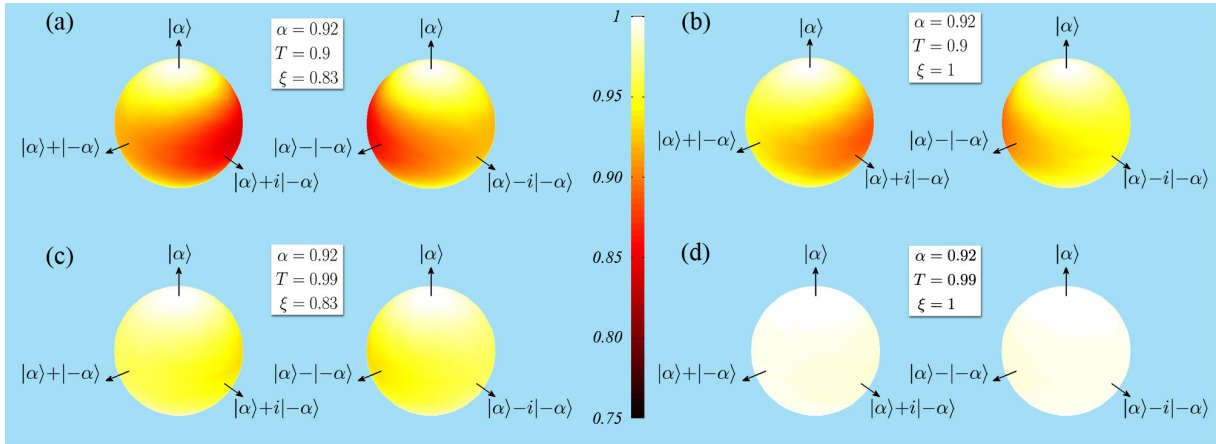


FIGURE 6.10 – Fidélité  $\mathcal{F}_{\theta,\phi}$ , pour une porte de phase  $\varphi=\pi/3$ . Le détecteur est modélisé par une APD d'efficacité  $\lambda=0.1$ .

est minimale lorsque l'état initial est un chat pair ( $\phi=0$ ), avec une valeur d'environ 0.75 pour nos paramètres expérimentaux. Pour  $\varphi=\pi/3$ , c'est pour une phase  $\phi$  comprise entre  $\pi/2$  et  $\pi$ , qui dépend des paramètres de la porte. Cette directivité diminue néanmoins très fortement pour  $T=0.99$ .

On voit enfin que la pureté modale  $\xi$  influe de manière importante sur la fidélité, sans toutefois en changer la directivité. Pour une pureté modale de  $\xi=1$  et une forte transmission  $T=0.99$ , la porte expérimentale donne des résultats très proches de ceux de la porte idéale, avec une fidélité proche de 1 pour toutes les superpositions.

### 6.5.2 Modèle de porte simplifié

Afin d'illustrer l'importance de  $\xi$  et d'expliquer une partie de la dépendance en  $\phi$ , considérons un modèle de porte de phase  $\varphi$  simplifié, agissant toujours sur une superposition initiale parfaite  $|\psi_{\theta,\phi}\rangle$ . Avec une probabilité  $\xi$ , la transformation réalisée est celle de la porte idéale donnant

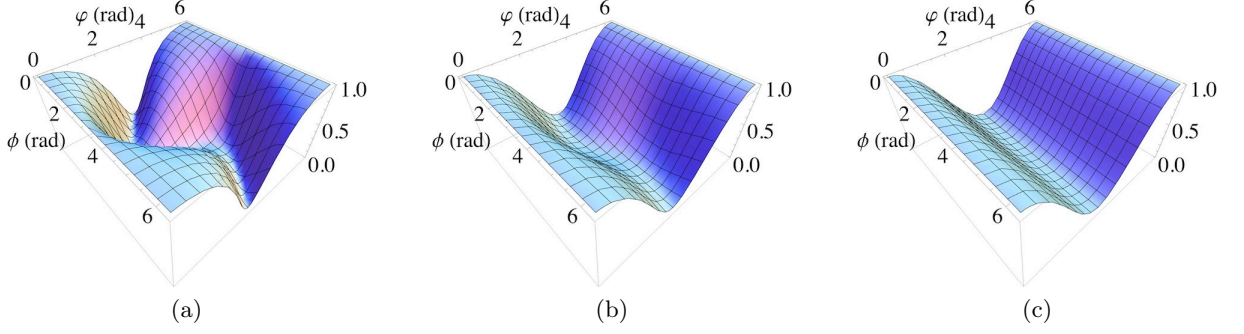


FIGURE 6.11 – Fidélité entre une superposition  $|\psi_{\frac{\pi}{2},\phi}\rangle$  et  $|\psi_{\frac{\pi}{2},\phi+\varphi}\rangle$  (a)  $\alpha=0.5$ ; (b)  $\alpha=0.92$ ; (c)  $\alpha=1.5$ .

$|\psi_{\theta,\phi+\varphi}\rangle$ , et avec une probabilité  $(1-\xi)$  la transformation est l'identité :

$$\hat{\rho}_{\theta,\phi}^{\text{th}} = \xi |\psi_{\theta,\phi+\varphi}\rangle \langle \psi_{\theta,\phi+\varphi}| + (1-\xi) |\psi_{\theta,\phi}\rangle \langle \psi_{\theta,\phi}| \quad (6.39)$$

Ce modèle constitue un cas limite idéal, en ne conservant que la pureté modale comme imperfection. Prenons maintenant  $\theta=\pi/2$ , puisque nous avons vu que c'est là où la porte est la moins efficace. Le calcul de la fidélité avec l'état cible donne donc

$$\langle \psi_{\frac{\pi}{2},\phi+\varphi} | \hat{\rho}_{\frac{\pi}{2},\phi}^{\text{th}} | \psi_{\frac{\pi}{2},\phi+\varphi} \rangle = \xi + (1-\xi) |\langle \psi_{\frac{\pi}{2},\phi} | \psi_{\frac{\pi}{2},\phi+\varphi} \rangle|^2 = \mathcal{F}_{\phi,\varphi}^{\text{th}}. \quad (6.40)$$

Selon ce modèle, la dépendance selon  $\phi$  vient donc du recouvrement  $\langle \psi_{\frac{\pi}{2},\phi} | \psi_{\frac{\pi}{2},\phi+\varphi} \rangle$ , et la fidélité est une fonction affine de  $\xi$ . Une plus grande valeur de  $\xi$  rehausse la fidélité, mais ne modifie pas intrinsèquement sa dépendance par rapport à  $\phi$ , qui est juste amortie d'un facteur  $1-\xi$ . Lorsque ce recouvrement est négligeable, ce qui n'est pas tout à fait le cas pour  $\alpha=0.92$ , la fidélité devient égale à  $\xi$ . Nous expliquons ainsi pourquoi elle est bornée sur les figures 6.9 et 6.10 lorsque  $\xi < 1$ , même en augmentant  $T$ .

Le recouvrement  $\langle \psi_{\frac{\pi}{2},\phi} | \psi_{\frac{\pi}{2},\phi+\varphi} \rangle$  est montré sur la 6.11, pour plusieurs valeurs de  $\alpha$ . Pour  $\alpha=0.5$ , on note une très forte dépendance en  $\phi$ , provenant du recouvrement  $\langle \alpha | -\alpha \rangle$  non négligeable. A mesure que  $\alpha$  augmente, cette dépendance disparaît, et on peut très facilement vérifier que deux superpositions  $|\psi_{\frac{\pi}{2},\phi}\rangle$  et  $|\psi_{\frac{\pi}{2},\phi+\varphi}\rangle$  sont orthogonales pour  $\varphi=\pi$  quelle que soit  $\phi$  lorsque l'on peut négliger le recouvrement  $\langle \alpha | -\alpha \rangle$ .

La figure 6.12 montre la fidélité  $\mathcal{F}_{\phi,\varphi}^{\text{th}}$  pour plusieurs valeurs de  $\varphi$ , pour  $\xi=0.83$ , correspondant à des coupes de la figure 6.11 (b) et (c). C'est pour  $\varphi=\pi$  que la fidélité  $|\langle \psi_{\frac{\pi}{2},\phi} | \psi_{\frac{\pi}{2},\phi+\varphi} \rangle|^2$  est la plus faible. Pour une phase de la porte  $\varphi$  plus faible, cette fidélité est plus importante, et par conséquent  $\mathcal{F}_{\phi,\varphi}^{\text{th}}$  augmente aussi.

### 6.5.3 Fidélités pour des superpositions de même poids

Les figures 6.9 et 6.10 donnent une bonne représentation du comportement général d'une porte de phase. Nous avons vu que c'est pour  $\theta=\pi/2$  que la fidélité est la moins bonne et qu'elle est le plus sensible aux paramètres de la porte. Afin d'étudier en détail leur influence, la figure 6.13 montre maintenant  $\mathcal{F}_{\frac{\pi}{2},\phi}$  en fonction de  $\phi$ , pour plusieurs valeurs de  $\varphi$ ,  $T$ ,  $\xi$ , pour plusieurs types de détecteurs, et pour le modèle simplifié (6.40).

Pour chaque phase  $\varphi$ , on retrouve le fait que la directivité est plus forte pour une plus faible valeur de  $T$ . Cet effet est d'autant plus marqué que le détecteur s'éloigne d'un compteur



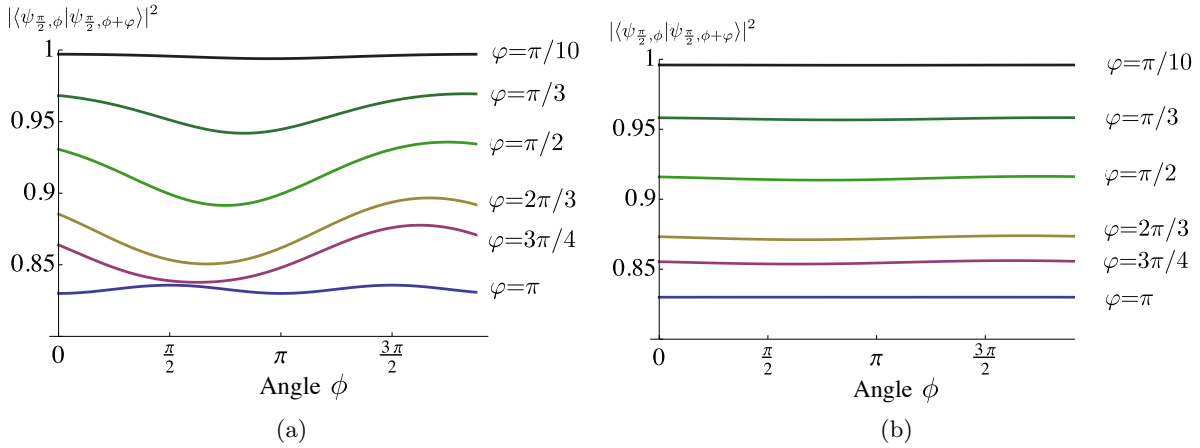


FIGURE 6.12 – Fidélité  $\mathcal{F}_{\phi, \phi}^{\text{th}}$  pour plusieurs valeurs de  $\varphi$ , pour  $\xi=0.83$  et (a)  $\alpha=0.92$  ; (b)  $\alpha=1.5$ .

de photon, en étant maximal pour une APD imparfaite. Un compteur de photon, beaucoup plus contraignant expérimentalement, n'est pourtant pas forcément nécessaire, puisqu'il suffit d'avoir  $T=0.99$  pour être très proche de la modélisation (6.40). Le seul inconvénient est que la probabilité de succès diminue en  $1-T$ .

Comme nous l'avons vu avec le modèle simplifié (6.40), la pureté modale  $\xi$  agit principalement comme un facteur d'échelle sur la fidélité. Elle impose une borne supérieure, même lorsque la soustraction de photon est parfaite.

Ces observations s'appliquent également pour les autres valeurs de  $\varphi$ , pour lesquelles la directivité est davantage liée au recouvrement  $\langle \alpha | -\alpha \rangle \neq 0$  (cf. Fig. 6.12).

La figure 6.14 montre ces mêmes simulations pour  $\alpha=1.5$ . On voit que la directivité par rapport à  $\phi$  a très fortement diminué, ce qui montre qu'elle provient bien principalement du recouvrement  $\langle \alpha | -\alpha \rangle$ . Les autres conclusions sur l'effet des différents paramètres sont similaires : l'utilisation d'un compteur permet d'augmenter la fidélité, mais l'augmentation de  $T$  donne des résultats comparables. On note par contre un plus grand écart entre les fidélités obtenues avec une APD imparfaite et un compteur lorsque  $T=0.9$ .

En conclusion, lorsque  $\alpha$  est trop faible pour pouvoir considérer que  $|\alpha\rangle$  et  $|- \alpha\rangle$  sont orthogonaux, la directivité selon  $\phi$  provient de deux sources : la transmission  $T$ , et le recouvrement  $\langle \alpha | -\alpha \rangle$ . L'utilisation d'un compteur de photon permet de réduire la directivité due à  $T$ , mais pas celle due au recouvrement, qui peut être importante selon  $\varphi$ . Même en utilisant une APD imparfaite, on peut s'approcher très vite du cas idéal (6.40) en augmentant  $T$ . Lorsque  $\alpha$  augmente, la directivité en  $\phi$  est nettement moins importante, et  $T$  agit davantage sur la valeur de la fidélité.

#### 6.5.4 Quelques calculs de fidélité

Nous avons comparé l'action de la porte expérimentale et de la porte parfaite sur une superposition initiale parfaite. Qu'en est-il lorsque l'on utilise nos états expérimentaux ? Même si, comme nous l'avons déjà souligné, cette comparaison est dégradée par les imperfections des états, elle nous permettra néanmoins de montrer quelques propriétés intéressantes.

**Pour le vide comprimé expérimental** La table 6.3 regroupe des valeurs de fidélité entre l'état  $\hat{\rho}_{\text{out}}^1$  obtenu en appliquant la porte de phase  $\pi$  expérimentale sur le vide comprimé, et entre

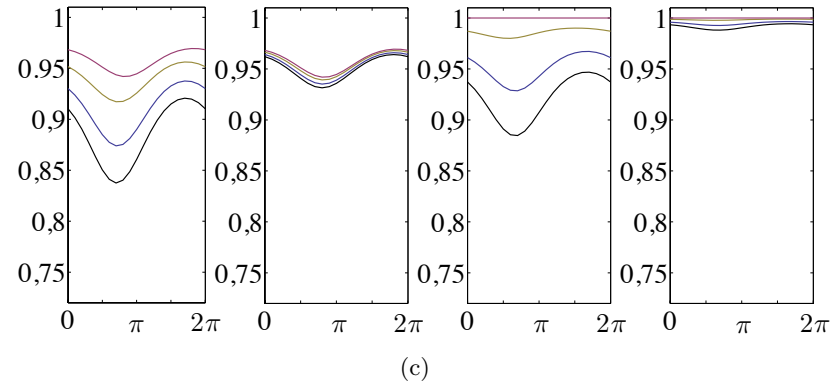
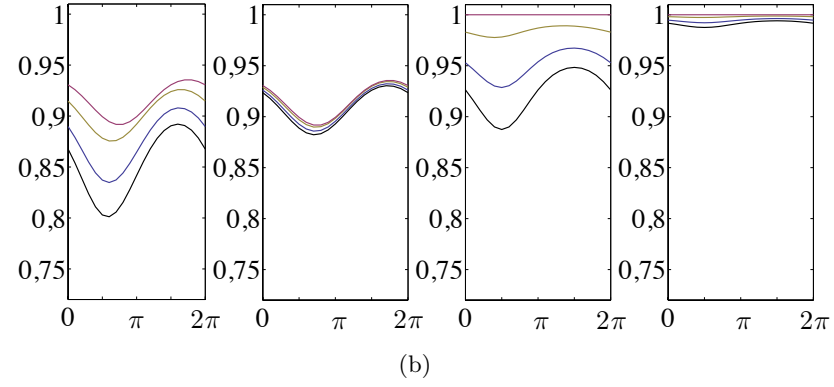
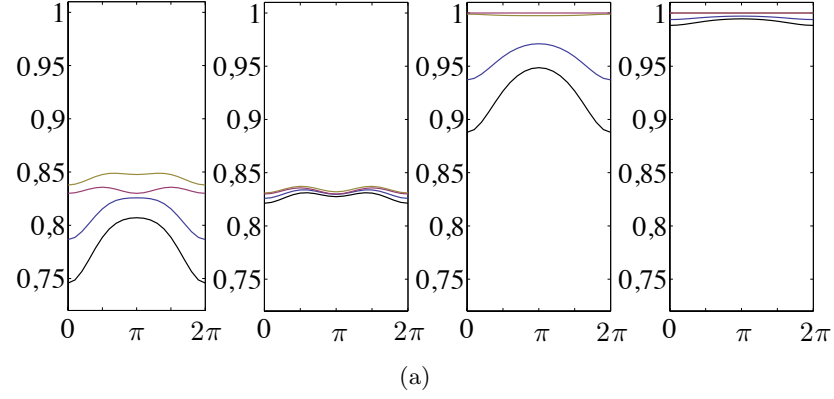


FIGURE 6.13 – Fidélité  $\mathcal{F}_{\frac{\pi}{2}, \phi}$  en fonction de  $\phi$ , pour  $\alpha=0.92$  et de gauche à droite pour chaque sous-figure :  $\xi=0.83$  et  $T=0.9$  ;  $\xi=0.83$  et  $T=0.99$  ;  $\xi=1$  et  $T=0.9$  ;  $\xi=1$  et  $T=0.99$ , pour plusieurs paramètres de la porte : APD avec une efficacité  $\lambda=0.1$  (noir), APD avec une efficacité  $\lambda=1$  (bleu), compteur de photon parfait  $|1\rangle\langle 1|$  (jaune), porte théorique modélisée par (6.40), correspondant à la figure 6.12 (rose), et pour (a)  $\varphi=\pi$  ; (b)  $\varphi=\pi/2$  ; (c)  $\varphi=\pi/3$ .



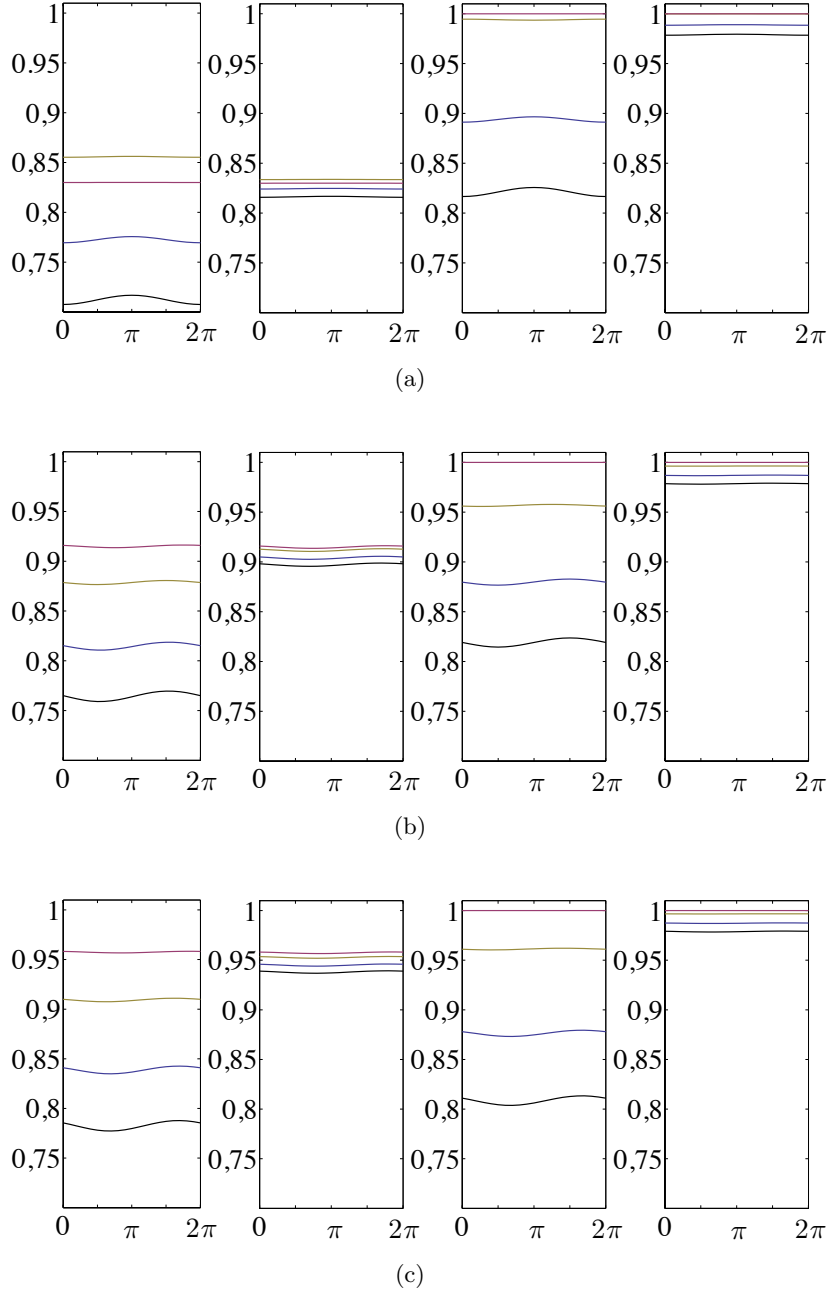


FIGURE 6.14 – Fidélité  $\mathcal{F}_{\frac{\pi}{2}, \phi}$  en fonction de  $\phi$ , pour  $\alpha=1.5$  et de gauche à droite pour chaque sous-figure :  $\xi=0.83$  et  $T=0.9$  ;  $\xi=0.83$  et  $T=0.99$  ;  $\xi=1$  et  $T=0.9$  ;  $\xi=1$  et  $T=0.99$ , pour plusieurs paramètres de la porte : APD avec une efficacité  $\lambda=0.1$  (noir), APD avec une efficacité  $\lambda=1$  (bleu), compteur de photon parfait  $|1\rangle\langle 1|$  (jaune), porte théorique modélisée par (6.40), correspondant à la figure 6.12 (rose), et pour (a)  $\varphi=\pi$  ; (b)  $\varphi=\pi/2$  ; (c)  $\varphi=\pi/3$ .

l'état  $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$  obtenu en appliquant une soustraction parfaite  $\hat{\mathbf{a}}$ , qui est bien sûr normalisé pour le calcul. Comme précédemment, on fait varier les trois paramètres de la porte :  $T$ ,  $\xi$ , et le type de détecteur. A chaque fois, on calcule aussi le chat parfait le plus ressemblant avec  $\hat{\rho}_{\text{out}}^1$ .

On voit que le type de détecteur influe très peu sur la fidélité avec  $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$ . La faible différence observée pour  $T=0.9$  et  $\xi=0.83$  devient insignifiante en augmentant  $T$ , ce qui revient à réduire la probabilité que plus d'un photon soit réfléchi vers l'APD. Ceci nous confirme qu'une APD imparfaite est suffisante compte tenu de la qualité de nos états produits. En revanche, le type de détecteur influe un peu plus sur la fidélité de l'état produit avec un chat parfait, de manière plus importante quand  $T=0.9$  et  $\xi=1$ .

On voit ensuite que la transmission  $T$  a surtout une influence sur la fidélité avec un chat parfait, davantage lorsque l'on utilise une APD avec  $\lambda=0.1$ , mais qu'elle a très peu d'influence sur la fidélité avec  $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$ . Enfin, là encore on retrouve le fait que le paramètre  $\xi$  a le plus d'influence, tant sur la fidélité avec un chat parfait que sur la fidélité avec  $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$ , quel que soit le type de détecteur utilisé.

| Porte expérimentale appliquée sur $\hat{\rho}_{\text{opa}}$ |       |                   | Fidélité avec $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$ | Fidélité maximale avec un chat parfait |                       |
|-------------------------------------------------------------|-------|-------------------|---------------------------------------------------------------------------------|----------------------------------------|-----------------------|
| $T$                                                         | $\xi$ | Méthode           |                                                                                 | Fidélité max                           | $\alpha_{\text{max}}$ |
| 0.9                                                         | 0.83  | APD $\lambda=0.1$ | 0.94                                                                            | 0.58                                   | 0.92                  |
| 0.9                                                         | 0.83  | APD $\lambda=1$   | 0.95                                                                            | 0.60                                   | 0.92                  |
| 0.9                                                         | 0.83  | Compteur          | 0.96                                                                            | 0.63                                   | 0.92                  |
| 0.9                                                         | 1     | APD $\lambda=0.1$ | 0.99                                                                            | 0.69                                   | 0.92                  |
| 0.9                                                         | 1     | APD $\lambda=1$   | 0.99                                                                            | 0.72                                   | 0.92                  |
| 0.9                                                         | 1     | Compteur          | 0.99                                                                            | 0.76                                   | 0.92                  |
| 0.99                                                        | 0.83  | APD $\lambda=0.1$ | 0.96                                                                            | 0.64                                   | 0.96                  |
| 0.99                                                        | 0.83  | APD $\lambda=1$   | 0.96                                                                            | 0.64                                   | 0.96                  |
| 0.99                                                        | 0.83  | Compteur          | 0.96                                                                            | 0.64                                   | 0.96                  |
| 0.99                                                        | 1     | APD $\lambda=0.1$ | 0.99                                                                            | 0.76                                   | 0.96                  |
| 0.99                                                        | 1     | APD $\lambda=1$   | $\sim 1$                                                                        | 0.76                                   | 0.96                  |
| 0.99                                                        | 1     | Compteur          | $\sim 1$                                                                        | 0.77                                   | 0.96                  |
| Porte théorique $\hat{\mathbf{a}}$                          |       |                   | 1                                                                               | 0.77                                   | 0.96                  |

TABLE 6.3 – Calculs de fidélité pour l'action de la porte de phase  $\pi$  expérimentale sur le vide comprimé produit par l'OPA.

### Influence de $\gamma$

Ces calculs de fidélité peuvent paraître un peu contradictoires, étant donnés les résultats de la section 6.5.1 avec des superpositions parfaites en entrée de la porte. Par exemple, pour  $\xi=0.83$  et  $T=0.9$ , la fidélité avec  $\hat{\mathbf{a}}\hat{\rho}_{\text{opa}}\hat{\mathbf{a}}^\dagger$  est de 0.94, alors qu'elle est d'environ 0.74 lorsqu'on applique la même porte expérimentale sur une superposition parfaite  $|+\rangle \propto |\alpha\rangle + |-\alpha\rangle$  (cf. Fig. 6.9).

De manière un peu surprenante, ce comportement peut être imputé au bruit rajouté par  $\gamma$ . Deux tests permettent de s'en assurer : le premier est de rajouter ce bruit à une superposition initiale parfaite  $|+\rangle$  en utilisant l'opérateur  $\hat{S}_2(\gamma r)$  défini par (2.149), donnant un état initial  $\hat{\rho}_\gamma$ . Les fidélités obtenues sont présentées sur la figure 6.15. On remarque par exemple que pour des paramètres  $T=0.9$ ,  $\xi=0.83$ , et une APD avec  $\lambda=0.1$  :

- Sans ajout de bruit, la fidélité entre l'état de sortie de la porte et  $\hat{\mathbf{a}}|+\rangle$  vaut 0.74, comme

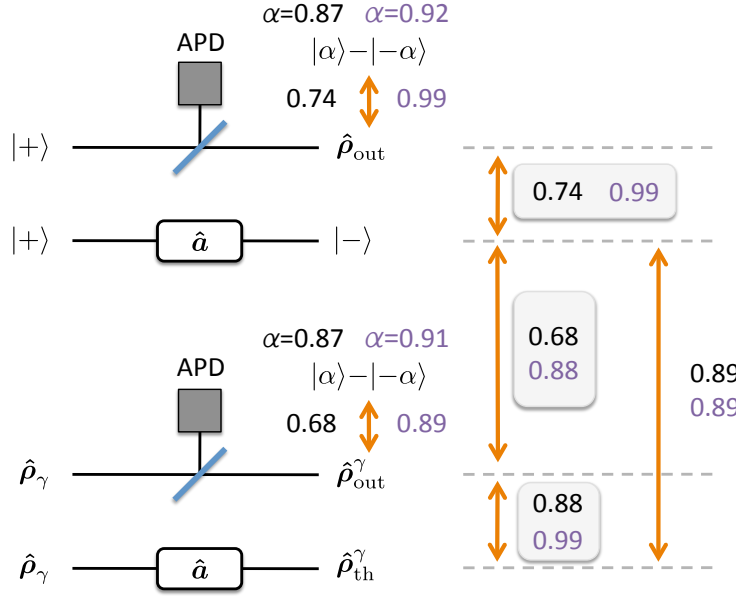


FIGURE 6.15 – Premier exemple illustrant l’effet de  $\gamma$  sur l’action de la porte.. L’état  $\hat{\rho}_\gamma$  est obtenu en utilisant l’opérateur  $\hat{S}_2(\gamma r)$  défini par (2.149), avec  $\gamma r = 0.3 \times 0.46$ . Les flèches oranges correspondent à la fidélité entre les deux états pointés, dont la valeur est donnée pour deux réalisations expérimentales de la porte :  $T=0.9$  et  $\xi=0.83$  (noir), et  $T=0.99$  et  $\xi=1$  (violet). Dans les deux cas, l’APD est d’efficacité  $\lambda=0.1$ . L’amplitude des états  $|+\rangle$  et  $|-\rangle$  est  $\alpha=0.92$ .

le montre la figure 6.9.

- Avec  $\gamma r = 0.46 \times 0.3$ , la fidélité entre l’état de sortie de la porte et  $\hat{a}|+\rangle$  diminue et vaut maintenant 0.68. En revanche, la fidélité avec  $\hat{a}\hat{\rho}_\gamma\hat{a}^\dagger$  vaut 0.88.

Le second exemple est donné par la table 6.4, où l’on calcule la fidélité avec  $\hat{a}\hat{\rho}_{\text{opa}}\hat{a}^\dagger$  pour la porte expérimentale appliquée sur  $\hat{\rho}_{\text{opa}}$ , et le chat parfait maximisant la fidélité avec l’état de sortie. En diminuant  $\gamma$ , on observe que la fidélité avec  $\hat{a}\hat{\rho}_{\text{opa}}\hat{a}^\dagger$  diminue, et atteint une valeur de 0.73 pour  $\gamma=0$ , comparable à celle obtenue pour un état initial  $|+\rangle$ . Dans ce cas, le vide comprimé a une fidélité de 0.999 avec un  $|+\rangle$  d’amplitude  $\alpha=0.55$ . Il est donc normal que les résultats soient comparables à ceux utilisant une superposition initiale parfaite. On remarque également que l’état de sortie est plus proche d’un état  $|-\rangle$  quand  $\gamma$  diminue.

| Porte expérimentale<br>appliquée sur $\hat{\rho}_{\text{opa}}$ |       |                   |          | Fidélité<br>avec<br>$\hat{a}\hat{\rho}_{\text{opa}}\hat{a}^\dagger$ | Fidélité maximale<br>avec un chat parfait |                       |
|----------------------------------------------------------------|-------|-------------------|----------|---------------------------------------------------------------------|-------------------------------------------|-----------------------|
| $T$                                                            | $\xi$ | Méthode           | $\gamma$ |                                                                     | Fidélité max                              | $\alpha_{\text{max}}$ |
| 0.9                                                            | 0.83  | APD $\lambda=0.1$ | 0.46     | 0.94                                                                | 0.58                                      | 0.92                  |
|                                                                |       |                   | 0.30     | 0.91                                                                | 0.66                                      | 0.92                  |
|                                                                |       |                   | 0.20     | 0.87                                                                | 0.70                                      | 0.92                  |
|                                                                |       |                   | 0.1      | 0.82                                                                | 0.72                                      | 0.92                  |
|                                                                |       |                   | 0        | 0.73                                                                | 0.73                                      | 0.92                  |

TABLE 6.4 – Deuxième exemple illustrant l’effet de  $\gamma$  sur l’action de la porte.

En conclusion, ces deux exemples illustrent le fait que  $\gamma$  réduit la fidélité avec l’état cible

de la porte, qui est  $|-\rangle \propto \hat{\mathbf{a}}|+\rangle$ , mais il augmente la fidélité avec l'état produit par l'action de l'opérateur  $\hat{\mathbf{a}}$  sur l'état initial. Dans ce cas, la porte remplit moins bien son rôle de porte de phase, puisque l'état soustrait d'un photon correspond moins à l'état  $|-\rangle$ .

La qualité de la soustraction de photon n'est donc pas toujours liée à la qualité de la réalisation de la porte de phase, sauf quand l'état initial est une superposition d'états cohérents parfaite.

## 6.6 Fidélité avec une porte idéale

### 6.6.1 Principe

#### Associer un état quantique à un processus quantique

La méthode décrite dans la section 6.5 permet de comparer l'action de la porte expérimentale et de la porte théorique, pour un qubit initial donné. Voyons maintenant comment caractériser la porte avec un seul nombre, nous renseignant sur sa qualité globale, indépendamment de l'état initial.

Selon l'isomorphisme de Jamiołkowski [Jamiołkowski72], un processus quantique  $\mathcal{E}$  agissant dans un espace de dimension  $d$ , peut être associé à un état quantique  $\hat{\rho}_{\mathcal{E}}$ , défini par

$$\hat{\rho}_{\mathcal{E}} = [\mathbb{I} \otimes \mathcal{E}] (|\Phi\rangle\langle\Phi|), \quad (6.41)$$

où  $|\Phi\rangle = \frac{1}{\sqrt{d}} \sum_{j=1}^d |j\rangle|j\rangle$  est un état maximalement intriqué, associé à une base  $\{|j\rangle\}$ . Illustrons cette propriété avec l'exemple d'un processus agissant sur un qubit, dans un espace de dimension deux. Un choix possible d'état maximalement intriqué est  $|\Phi^+\rangle = \frac{1}{\sqrt{2}}(|+\rangle|+\rangle + |-\rangle|-\rangle)$ . En notant  $|\mu\rangle\langle\nu| = \hat{\rho}_{\mu\nu}$ , avec  $\mu, \nu = +$  ou  $-$ , l'état  $\hat{\rho}_{\mathcal{E}}$  est alors égal à

$$\hat{\rho}_{\mathcal{E}} = \frac{1}{2} [\hat{\rho}_{++} \otimes \mathcal{E}(\hat{\rho}_{++}) + \hat{\rho}_{--} \otimes \mathcal{E}(\hat{\rho}_{--}) + \hat{\rho}_{+-} \otimes \mathcal{E}(\hat{\rho}_{+-}) + \hat{\rho}_{-+} \otimes \mathcal{E}(\hat{\rho}_{-+})] \quad (6.42)$$

On peut ensuite retrouver l'évolution d'un état quelconque  $|\phi\rangle = x|+\rangle + y|-\rangle$ , en calculant la fidélité entre  $\hat{\rho}_{\mathcal{E}}$  et  $|\phi^*\rangle = x^*|+\rangle + y^*|-\rangle$  :

$$\langle\phi^*|\hat{\rho}_{\mathcal{E}}|\phi^*\rangle \propto |x|^2 \mathcal{E}(\hat{\rho}_{++}) + |y|^2 \mathcal{E}(\hat{\rho}_{--}) + xy^* \mathcal{E}(\hat{\rho}_{+-}) + yx^* \mathcal{E}(\hat{\rho}_{-+}) \quad (6.43)$$

L'isomorphisme de Jamiołkowski est en quelque sorte une application du parallélisme quantique, où  $\hat{\rho}_{\mathcal{E}}$  contient l'évolution de tous les états à la fois. On peut ensuite définir une mesure de distance entre deux processus quantiques  $\mathcal{E}_1$  et  $\mathcal{E}_2$ , en calculant la fidélité entre leurs états correspondant [Gilchrist05] :

$$\Delta(\mathcal{E}_1, \mathcal{E}_2) = \mathcal{F}(\hat{\rho}_{\mathcal{E}_1}, \hat{\rho}_{\mathcal{E}_2}) \quad (6.44)$$

#### Application à la porte de phase

Dans notre cas, on peut définir une action similaire à cet isomorphisme, en se restreignant au sous espace du qubit pour l'état initial. En d'autres termes, même s'il faudrait prendre en compte tout l'espace des phases pour associer un état  $\hat{\rho}_{\mathcal{E}}$  à la porte de phase, la restriction à un état initial contenu dans le sous espace généré par  $|\alpha\rangle$  et  $|-\alpha\rangle$  nous permet de caractériser la porte là où elle est censée fonctionner.

Les états  $|+\rangle = \mathcal{N}_+(|\alpha\rangle + |-\alpha\rangle)$  et  $|-\rangle = \mathcal{N}_-(|\alpha\rangle - |-\alpha\rangle)$  constituent une base orthogonale de ce sous espace, sur laquelle on peut décomposer l'état d'un qubit quelconque

$$x|\alpha\rangle + y|-\alpha\rangle = a|+\rangle + b|-\rangle. \quad (6.45)$$

Nous pouvons donc “sonder” la porte expérimentale en l'appliquant sur un mode de  $|\Phi^+\rangle$ . Le choix de cet état en particulier n'a pas d'influence, car nous montrerons dans la section suivante que les résultats sont indépendants de l'état maximalement intriqué utilisé. Nous noterons  $\hat{\Xi}$  l'état obtenu.

La porte idéale transforme  $|\Phi^+\rangle$  en

$$|\Omega\rangle = \frac{[\mathbb{I} \otimes \hat{\mathbf{a}} + \beta] |\Phi^+\rangle}{\| [\mathbb{I} \otimes (\hat{\mathbf{a}} + \beta)] |\Phi^+\rangle \|}. \quad (6.46)$$

où  $\beta$  fixe la phase  $\varphi$  de la porte, comme montré dans la section 6.2.2. Remarquons que pour  $\varphi = \pi$  (et donc  $\beta = 0$ ), cet état diffère de  $|\Psi^+\rangle = \frac{1}{\sqrt{2}}(|+\rangle|-\rangle + |-\rangle|+\rangle)$ , car  $|\alpha\rangle + |-\alpha\rangle$  et  $|\alpha\rangle - |-\alpha\rangle$  ont des coefficients de normalisation différents. On montre en effet que leur fidélité est donnée par

$$\mathcal{F} = |\langle \Omega | \Psi^+ \rangle|^2 = \frac{1}{2}(1 + \tanh(2\alpha^2)) \quad (6.47)$$

Pour notre valeur  $\alpha = 0.92$ ,  $\mathcal{F} = 0.967$  (figure 6.16).

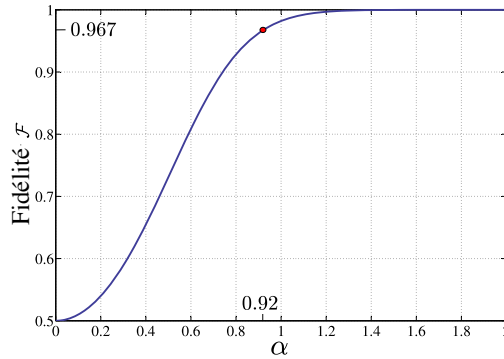


FIGURE 6.16 – Fidélité  $\mathcal{F}$  entre  $|\Omega\rangle$  et  $|\Psi^+\rangle$ .

La fidélité “de processus” entre la porte expérimentale et la porte théorique s’obtient finalement avec

$$F = \langle \Omega | \hat{\Xi} | \Omega \rangle. \quad (6.48)$$

### 6.6.2 Simulations pour la porte de phase

Pour nos valeurs expérimentales de la porte de phase  $\pi$  ( $T = 0.9$ ,  $\xi = 0.83$ , APD avec  $\lambda = 0.1$ ), nous trouvons une fidélité  $F = 0.78 \pm 0.04$ . La figure 6.17 montre les résultats obtenus pour d’autres valeurs de  $\xi$ , en considérant plusieurs configurations expérimentales.

Pour le modèle simplifié (6.40), on voit immédiatement que la fidélité est égale à

$$\langle \Omega | \left( \xi |\Omega\rangle\langle\Omega| + (1-\xi) |\Phi^+\rangle\langle\Phi^+| \right) | \Omega \rangle = \xi + (1-\xi) |\langle \Omega | \Phi^+ \rangle|^2, \quad (6.49)$$

qui est égale à  $\xi$  pour  $\varphi=\pi$  car  $\langle\Omega|\Phi^+\rangle=0$ . La pureté modale  $\xi$  apparaît donc comme une caractéristique fondamentale de la porte, en s'interprétant comme la fidélité entre la porte parfaite, et la "meilleure porte expérimentale possible", compte tenu du filtrage imparfait. La figure 6.17 montre qu'une transmission  $T=0.99$  permet quasiment d'atteindre cette borne supérieure.

Les figures 6.18 et 6.19 montrent le calcul de  $F$  pour deux autres phases  $\varphi$ , respectivement égales à  $\pi/2$  et  $\pi/3$ . Pour une valeur de  $\xi$  inférieure à 1, il existe toujours une borne supérieure pour  $F$ , mais qui est cette fois supérieure à  $\xi$ , puisque  $\langle\Omega|\Phi^+\rangle\neq 0$  si  $\varphi\neq\pi$ .

On s'attend donc à ce que, pour une valeur de  $\xi$  donnée, la borne supérieure de  $F$  augmente quand  $\varphi$  diminue. C'est bien ce qui est observé sur la figure 6.20. On remarque également qu'une valeur de  $T$  trop faible peut réduire la fidélité pour de faibles valeurs de  $\varphi$ . En effet, lorsque  $\varphi$  est faible, l'état en sortie de la porte parfaite est peu différent de celui en entrée, et peut même lui ressembler davantage que celui produit par une soustraction avec une valeur de  $T$  trop faible. Dans ce cas, la contribution du bon conditionnement à  $F$  diminue. En augmentant  $T$  et la qualité de la soustraction, on améliore la qualité de l'état conditionné et on retrouve bien une fidélité  $F$  qui tend vers 1 pour  $\varphi\rightarrow 0$ .

### 6.6.3 Invariance du choix de l'état maximalement intriqué

Montrons maintenant que le calcul de  $F$  ne dépend pas du choix de l'état maximalement intriqué que l'on utilise. Considérons pour cela un état plus général

$$|\Psi\rangle = \frac{1}{\sqrt{2}} \left( |+\rangle|\mu\rangle + e^{i\delta} |-\rangle|\nu\rangle \right), \quad (6.50)$$

où  $|\mu\rangle$  correspond à  $|+\rangle$  ou  $|-\rangle$ , alors que  $|\nu\rangle$  correspond à l'état opposé :  $|\nu\rangle = |-\mu\rangle$ . L'action de la porte  $\mathcal{E}$  décrivant l'action de la porte peut être décomposée en deux parties non linéaires (annexe D), correspondant au bon et au mauvais conditionnement

$$\mathcal{E} = \xi \mathcal{E}^\vee + (1 - \xi) \mathcal{E}^\times \quad (6.51)$$

Considérons maintenant le bon conditionnement  $\mathcal{E}^\vee$ , le raisonnement étant similaire pour  $\mathcal{E}^\times$ . Pour un état initial  $\hat{\rho} = \sum_{x,y=+,-} c_{xy} |x\rangle\langle y|$ , la non linéarité de  $\mathcal{E}^\vee$  ne provient que de la normalisation, qui dépend de  $\hat{\rho}$ . On peut en revanche considérer une action non normalisée  $\tilde{\mathcal{E}}^\vee$ , produisant un état  $\tilde{\mathcal{E}}^\vee(\hat{\rho})$  qui est ensuite normalisé par sa trace :

$$\mathcal{E}^\vee(\hat{\rho}) = \frac{\tilde{\mathcal{E}}^\vee(\hat{\rho})}{\text{Tr}\{\tilde{\mathcal{E}}^\vee(\hat{\rho})\}} \quad (6.52)$$

Appelons  $\zeta_{xy}$  l'action de  $\tilde{\mathcal{E}}^\vee$  sur l'opérateur  $|x\rangle\langle y|$ . Sous l'action de  $\mathbb{I} \otimes \tilde{\mathcal{E}}^\vee$ , l'état  $|\Psi\rangle\langle\Psi|$  est transformé en :

$$\tilde{\chi} = \frac{1}{2} \left( |+\rangle\langle +| \otimes \zeta_{\mu\mu} + |-\rangle\langle -| \otimes \zeta_{\nu\nu} + e^{i\delta} |-\rangle\langle +| \otimes \zeta_{\nu\mu} + e^{-i\delta} |+\rangle\langle -| \otimes \zeta_{\mu\nu} \right) \quad (6.53)$$

En remarquant ensuite que  $(\hat{a}+\beta)|\mu\rangle = (\alpha+\beta) \frac{\mathcal{N}_\mu}{\mathcal{N}_{\mu,\varphi}} |\mu, \varphi\rangle$ , avec  $|+, \varphi\rangle = \mathcal{N}_{+\varphi} (|\alpha\rangle + e^{i\varphi} |-\alpha\rangle)$  et  $|-, \varphi\rangle = \mathcal{N}_{-\varphi} (|\alpha\rangle - e^{i\varphi} |-\alpha\rangle)$ , et en introduisant  $c_{\mu\varphi} = (\frac{\mathcal{N}_\mu}{\mathcal{N}_{\mu,\varphi}})^2$ , l'état produit par la porte idéale s'écrit :

$$|\Omega\rangle = \frac{[\mathbb{I} \otimes (\hat{a}+\beta)]|\Psi\rangle}{\| [\mathbb{I} \otimes (\hat{a}+\beta)]|\Psi\rangle \|} \quad (6.54a)$$

$$= \frac{1}{\sqrt{c_{\mu\varphi} + c_{\nu\varphi}}} \left( \sqrt{c_{\mu\varphi}} |+\rangle|\mu, \varphi\rangle + e^{i\delta} \sqrt{c_{\nu\varphi}} |-\rangle|\nu, \varphi\rangle \right) \quad (6.54b)$$

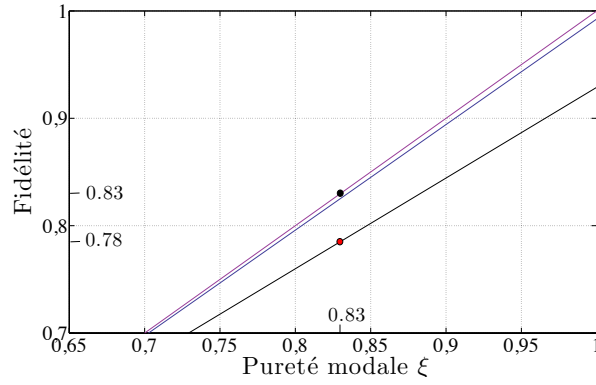


FIGURE 6.17 – Fidélité  $F$  pour une porte de phase  $\varphi=\pi$ , en fonction de  $\xi$ , pour  $T=0.9$  (noir),  $T=0.99$  (bleu) et la porte modélisée par (6.40).  $\alpha=0.92$ . Le point rouge correspond à nos paramètres, et le point noir à la limite atteignable pour notre valeur de  $\xi$ .

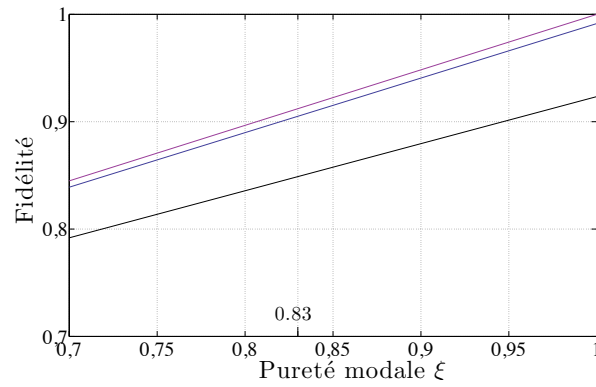


FIGURE 6.18 – Fidélité  $F$  pour une porte de phase  $\varphi=\pi/2$ , en fonction de  $\xi$ , pour  $T=0.9$  (noir),  $T=0.99$  (bleu) et la porte modélisée par (6.40) (violet).  $\alpha=0.92$ .

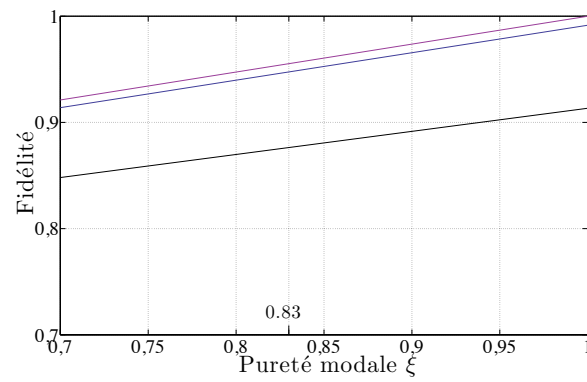


FIGURE 6.19 – Fidélité  $F$  pour une porte de phase  $\varphi=\pi/3$ , en fonction de  $\xi$ , pour  $T=0.9$  (noir),  $T=0.99$  (bleu) et la porte modélisée par (6.40) (violet).  $\alpha=0.92$ .

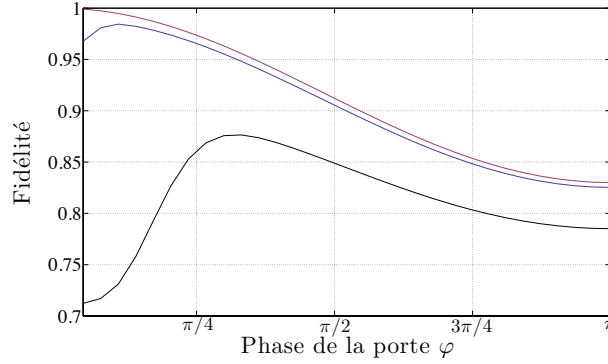


FIGURE 6.20 – Fidélité  $F$  en fonction de  $\varphi$ , pour  $T=0.9$  (noir),  $T=0.99$  (bleu) et la porte modélisée par (6.40) (violet).  $\alpha=0.92$  et  $\xi=0.83$ .

On considère ensuite la fidélité entre l'état normalisé  $\chi=\tilde{\chi}/\text{Tr}\{\tilde{\chi}\}$  et  $|\Omega\rangle$ , qui s'écrit :

$$\langle\Omega|\chi|\Omega\rangle = \tag{6.55a}$$

$$\frac{1}{2\text{Tr}\{\tilde{\chi}\}(c_{\mu\varphi}+c_{\nu\varphi})} (c_{\mu\varphi}\langle\mu,\varphi|\zeta_{\mu\mu}|\mu,\varphi\rangle+c_{\nu\varphi}\langle\nu,\varphi|\zeta_{\nu\nu}|\nu,\varphi\rangle+\sqrt{c_{\mu\varphi}c_{\nu\varphi}}[\langle\nu,\varphi|\zeta_{\nu\mu}|\mu,\varphi\rangle+\langle\mu,\varphi|\zeta_{\mu\nu}|\nu,\varphi\rangle]) \tag{6.55b}$$

Cette expression est symétrique en  $\mu$  et  $\nu$ , et ne dépend pas de  $\delta$ . Le même raisonnement avec  $\tilde{\mathcal{E}}^{\mathbf{x}}$  permet de conclure que l'expression (6.48) est bien indépendante du choix de  $|\Psi\rangle$ .

## 6.7 Conclusion

Le calcul quantique avec des superpositions d'états cohérents est une technique prometteuse, puisqu'elle associe des avantages propres aux variables discrètes et continues. Dans un circuit quantique, les algorithmes sont décomposés en portes quantiques élémentaires agissant sur un ou deux qubits. Une condition nécessaire pour pouvoir implémenter plusieurs portes est qu'elles présentent peu de défauts par rapport à la porte idéale.

Dans ce chapitre, nous avons présenté une méthode simple afin de caractériser une porte de phase. Nous l'avons implémenté expérimentalement pour une phase  $\pi$ , ce qui nous a permis d'extraire les paramètres d'un modèle donnant de très bons résultats. Nous avons ensuite utilisé ce modèle afin de simuler l'action de la porte sur un qubit initial parfait, de manière à séparer les imperfections propres à la porte des imperfections du qubit initial. Cette méthode empirique présente l'avantage de ne pas nécessiter une tomographie de processus quantique, lourde à mettre en œuvre, et qui de plus est difficilement compatible avec le mauvais filtrage des APD. Elle est également relativement indépendante de la qualité des états utilisés pour tester la porte, du moment que l'on dispose d'une modélisation adéquate.

Nous avons caractérisé la porte de deux manières : en calculant la fidélité entre les états produits par la porte modélisée et par la porte théorique, pour une superposition initiale quelconque ; et en considérant l'action de la porte sur une moitié d'état intriqué, afin de lui associer un seul nombre caractérisant sa qualité globale. Nous obtenons une valeur  $F=0.78$ .

Notre étude a permis de souligner l'importance de la pureté modale  $\xi$ , qui apparaît comme le principal facteur limitant la fidélité, et la qualité de la porte. Pour la porte de phase  $\pi$ , nous avons vu que c'est une borne supérieure pour  $F$ .



Son amélioration passe par un système de filtrage plus efficace, ou une diminution des photons “parasites” en améliorant les états produits par l’OPA.

## Troisième partie

# Résultats théoriques : amplificateur sans bruit en cryptographie quantique



# Chapitre 7

## Cryptographie quantique

### Sommaire

---

|            |                                                                  |            |
|------------|------------------------------------------------------------------|------------|
| <b>7.1</b> | <b>Introduction</b>                                              | <b>141</b> |
| 7.1.1      | La cryptographie classique                                       | 142        |
| 7.1.2      | La cryptographie quantique                                       | 143        |
| <b>7.2</b> | <b>Principes généraux de cryptographie quantique</b>             | <b>144</b> |
| 7.2.1      | Hypothèses sur le contrôle de l'environnement                    | 144        |
| 7.2.2      | Types d'attaques                                                 | 145        |
| <b>7.3</b> | <b>Quelques protocoles</b>                                       | <b>145</b> |
| 7.3.1      | Variables discrètes                                              | 145        |
| 7.3.2      | Variables continues & protocoles gaussiens                       | 146        |
| 7.3.3      | Autres protocoles                                                | 146        |
| 7.3.4      | Réseaux et applications commerciales                             | 147        |
| <b>7.4</b> | <b>Le protocole GG02 : de la QKD avec des états "classiques"</b> | <b>147</b> |
| 7.4.1      | Principe                                                         | 147        |
| 7.4.2      | Modélisation à intrication virtuelle                             | 148        |
| 7.4.3      | Preuves de sécurité                                              | 148        |
| 7.4.4      | Expression des taux secrets                                      | 150        |
| <b>7.5</b> | <b>Conclusion</b>                                                | <b>152</b> |

---

### 7.1 Introduction

La cryptographie quantique, ou distribution quantique de clé (QKD), fait certainement partie des applications les plus abouties en information quantique. Le principe est de distribuer une chaîne de bits secrète entre deux acteurs – Alice et Bob – en utilisant un canal quantique contrôlé par un espion – Eve – et des communications classiques [Scarani09, Gisin02]. Depuis le fameux protocole BB84 [Bennett84], de nombreux protocoles ont vu le jour et ont fait l'objet d'études approfondies, à la fois théoriques afin d'estimer les taux secrets accessibles, et expérimentales, conduisant aux premières applications commerciales de l'information quantique.

Tous ces protocoles, basés sur des variables discrètes ou continues, sont fondés sur une idée commune particulièrement simple et élégante. En vertu du théorème de non clonage [Wootters82] et/ou des relations d'incertitude de Heisenberg, Eve ne peut pas intercepter la communication

quantique sans introduire un minimum de bruit sur les mesures d’Alice ou de Bob. De ce fait, l’information qu’elle acquiert peut être estimée et prise en compte pour extraire une clé secrète.

Ce chapitre présente les principes généraux de la distribution quantique de clé, en se concentrant sur le protocole introduit par F. Grosshans et P. Grangier (GG02) [Grosshans02b], pour lequel nous étudierons l’utilisation d’un amplificateur sans bruit dans le chapitre 10. Nous présenterons également quelques autres protocoles de QKD, sans toutefois prétendre à l’exhaustivité. Le lecteur désirant une introduction plus complète pourra consulter les références [Scarani09, Gisin02, Cerf07, Nielsen00].

### 7.1.1 La cryptographie classique

#### Principe

La cryptographie est l’art de coder un message de manière à ce qu’il ne soit déchiffrable que par les personnes auxquelles il est destiné. Elle fait partie de la cryptologie, qui comprend aussi la cryptanalyse, l’art de déchiffrer un message codé. Si son utilisation remonte à plusieurs millénaires, les méthodes ont naturellement beaucoup évoluées avec les technologies disponibles.

Les premières techniques ancestrales de codage reposaient pour la plupart sur le secret de la méthode utilisée. Par exemple, le code de César consistait à décaler les lettres du message d’un pas fixe. Si le même pas est toujours utilisé, il suffit d’avoir une table de correspondance entre les différents caractères pour déchiffrer le message codé. Une telle méthode montre toutefois vite ses limites : lorsque la technique de codage est découverte, il devient facile de déchiffrer tous les anciens messages.

De nos jours, les méthodes de cryptographie ne se basent plus sur le secret de la méthode de cryptage, et sont pour la plupart publiques<sup>1</sup> afin de pouvoir en déceler les limites grâce au travail de la communauté scientifique. Le message est crypté en lui appliquant des modifications qui dépendent d’une clé secrète, produisant un cryptogramme. La sécurité du cryptage repose alors sur l’impossibilité, ou du moins la difficulté, à déchiffrer le cryptogramme sans avoir connaissance de la clé utilisée. On parle de *sécurité inconditionnelle* lorsque le cryptogramme est impossible à déchiffrer sans la clé, quelle que soit la technologie disponible. Dans le cas contraire, le déchiffrement n’est pas impossible, mais seulement “très difficile”, compte tenu de la technologie matérielle et/ou algorithmique disponible. On distingue deux grandes familles de codage, dites symétrique et asymétrique, en fonction du type de clé utilisé.

#### Codage symétrique et asymétrique

Le codage symétrique utilise une clé secrète identique pour l’émetteur et le destinataire. Le code de Vernam (ou code à masque jetable, *One time pad*) en est un exemple très simple, tout en étant inconditionnellement sûr. La clé consiste en une chaîne de bits de la même longueur que le message à transmettre, générés aléatoirement. Les bits du message et de la clé sont ensuite additionnés (modulo 2) à l’aide d’un XOR pour former le cryptogramme. Le déchiffrement se fait de la même manière, en additionnant le cryptogramme et la clé. Il est assez intuitif que ce code soit parfaitement indéchiffrable : puisque chaque bit de la clé est aléatoire, chaque bit du message chiffré l’est également, à condition que la clé ne soit utilisée qu’une seule fois. De manière plus générale, C. Shannon a montré qu’un cryptage ne peut être inconditionnellement sûr que si la clé est au moins aussi longue que le message [Shannon49]. Lorsque cela n’est pas le cas, la sécurité repose sur des limitations technologiques, mais de très bons niveaux de sécurité peuvent

1. Seuls certains protocoles militaires ne sont pas rendus publics.

être atteints avec les méthodes les plus récentes. Par exemple un cryptage symétrique de type AES [NIST01] avec une clé de 128 bits nécessiterait déjà des milliards d'années de calcul avec la puissance des ordinateurs actuels, en utilisant une méthode de type force brute où chaque clé possible est testée.

Un des inconvénients majeur du cryptage symétrique est que la clé doit être au préalable connue de l'émetteur et du destinataire. Des solutions existent : pour les données sensibles exigeant un secret absolu, les clés peuvent par exemple être échangées par valise diplomatique. Mais il est évident que ce dispositif n'est pas le plus efficace pour effectuer ses achats sur internet !

Le codage asymétrique permet de s'affranchir de ce dernier point. Supposons qu'Alice souhaite envoyer un message à Bob. Bob va alors générer deux clés : une première qu'il garde secrètement pour lui, et une seconde qu'il diffuse publiquement. Alice utilise ensuite la clé publique pour crypter son message, qu'elle envoie à Bob. Ce dernier le déchiffre enfin en utilisant sa clé secrète. Grâce à cette méthode, il n'est plus nécessaire de partager une clé secrète au préalable. Le chiffrement asymétrique est couramment utilisé sur internet, avec par exemple le codage RSA [Rivest78].

L'inconvénient du codage RSA est que sa sécurité n'est pas réellement prouvée. Elle repose principalement sur une conjecture mathématique, selon laquelle il ne serait pas possible de factoriser un nombre en un temps polynomial. Mais, contrairement au codage symétrique, une avancée mathématique majeure pourrait mettre à plat sa sécurité. Celle ci serait également fortement compromise si le développement des ordinateurs quantiques atteint un stade leur permettant d'implémenter l'algorithme de Shor [Shor97]. Notons également que le chiffrement asymétrique nécessite des clés plus grandes que le chiffrement symétrique.

### 7.1.2 La cryptographie quantique

#### Principe

La cryptographie quantique, ou distribution quantique de clé (QKD), consiste à utiliser les propriétés de la physique quantique pour échanger une clé secrète entre deux acteurs distants, traditionnellement appelés Alice et Bob, dans un environnement contrôlé par un Espion, Eve. Par hypothèse, Eve n'est limitée que par les lois de la physique, et dispose de tout l'arsenal technologique qu'il puisse être possible de concevoir, en particulier des mémoires quantiques, des amplificateurs limités au bruit de photon, des détecteurs parfaits, etc... Cette hypothèse, sûrement très pessimiste pour Alice et Bob compte tenu de la technologie actuelle, permet de s'assurer que le secret de la clé n'est pas dû à une quelconque limitation technologique qui puisse être surmontée dans le futur.

Tout protocole de cryptographie quantique peut se décomposer en deux étapes (Fig. 7.1). La première est l'étape "quantique", où Alice et Bob s'échangent et mesurent des états quantiques à travers un canal quantique, afin de créer des données corrélées. Ces états ne doivent pas tous être orthogonaux entre eux, afin qu'Eve ne puisse pas les mesurer sans introduire de bruit, ou les discriminer parfaitement. En général, Alice choisit aléatoirement des états parmi un ensemble discret (par exemple, BB84 quatre états de polarisation [Bennett84]) ou continu (par exemple, GG02 utilise une modulation gaussienne d'états cohérents [Grosshans02b]). Cette première étape quantique peut être réalisée en utilisant différentes propriétés du champ électromagnétique, comme son amplitude, sa polarisation, ou encore le nombre de photons, donnant lieu aux différents protocoles.

La seconde étape est une étape purement classique. Alice et Bob échangent des informations à propos de leurs données à travers un canal de communication classique, afin de les transformer en

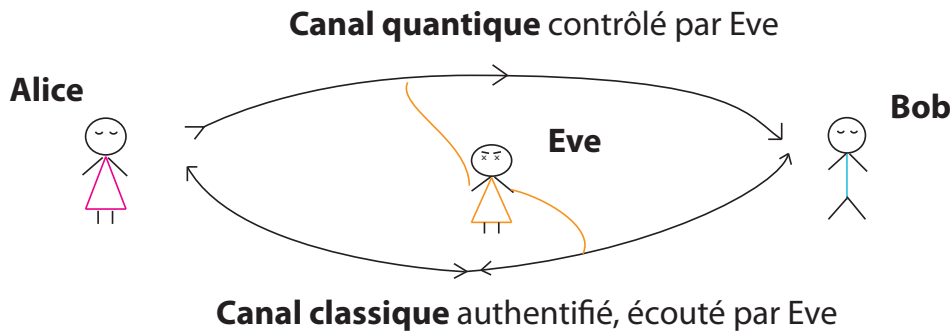


FIGURE 7.1 – Schéma général d'un protocole de QKD.

une chaîne de bits parfaitement secrète et inconnue d'Eve. Le canal classique doit être authentifié, ce qui signifie qu'Eve peut seulement écouter la conversation, mais ne peut pas se faire passer pour Alice ou Bob.

Terminons cette introduction par une remarque importante. Un protocole de QKD n'est utilisé que pour transmettre une clé, servant ensuite à coder le message, mais la clé en elle-même ne contient pas d'information. Un protocole inconditionnellement sûr signifie que l'espion n'a aucune information sur cette clé, mais cela ne garantit pas une sécurité inconditionnelle du cryptage, qui ne sera atteinte que si la clé est aussi longue que le message, et n'est utilisée qu'une seule fois.

## 7.2 Principes généraux de cryptographie quantique

Afin de pouvoir estimer les taux de clés qu'il est possible de transmettre, il est nécessaire de formaliser les actions potentielles d'Eve. Ces hypothèses sont communes à tous les protocoles de QKD.

### 7.2.1 Hypothèses sur le contrôle de l'environnement

Les premières hypothèses concernent le contrôle de l'environnement utilisé par Alice et Bob :

- Contrôle complet du canal quantique : Eve peut par exemple remplacer le canal quantique imparfait (avec des pertes, du bruit, etc...) par un canal quantique parfait. Elle peut ensuite effectuer des opérations (interaction avec les états envoyés par Alice, mesures,...) qui introduiront les mêmes imperfections que le canal quantique imparfait du point de vue d'Alice et Bob, mais qui lui fourniront des informations sur la communication quantique. Pour cette raison, les imperfections du canal sont par hypothèse toujours à l'avantage d'Eve.
- Pas de limite sur les capacités de calcul : Eve peut extraire le maximum d'information de ses mesures sans être limitée par l'efficacité d'algorithmes ou par la puissance de calcul pour y parvenir.
- Espionnage du canal classique sans modifier les messages : Eve écoute la communication classique, mais ne peut pas la modifier. Cette hypothèse est indispensable et ne pose aucun problème, il suffit qu'Alice et Bob disposent d'une petite clé secrète au préalable pour l'authentification [Renner05b, Weedbrook12].
- Pas d'accès aux installations d'Alice et de Bob : naturellement, si Eve peut accéder à l'ordinateur d'Alice ou de Bob, toute communication cryptée devient inutile car il n'y a plus

rien à cacher. Moins naïvement, cela signifie aussi que les imperfections des détecteurs d'Alice et de Bob ne sont pas attribuées à Eve. Cette hypothèse, relativement raisonnable permet même d'augmenter le taux secret en ajoutant du bruit chez Bob [García-Patrón09]. Notons cependant qu'en pratique, il est possible qu'Eve utilise des imperfections technologiques pour acquérir de l'information. Elle peut par exemple endommager le détecteur de Bob [Gerhardt11], ou surveiller le réseau électrique pour y déceler l'action d'un composant électronique révélant des informations sur la communication quantique. Ces "side channels" peuvent être pris en compte dans les preuves de sécurité, mais leur diversité rend la tâche particulièrement complexe.

### 7.2.2 Types d'attaques

Les autres hypothèses concernent le type d'attaques qu'Eve peut implémenter. On en distingue trois :

- Attaques individuelles : ce sont les attaques les plus simples. Eve interagit avec chaque état envoyé par Alice, et dispose d'une mémoire quantique pour stocker ses états et les mesurer individuellement. Cette mémoire permet d'attendre que Bob révèle le type de mesure qu'il a effectué si le protocole de QKD le nécessite.
- Attaques collectives : Eve interagit toujours individuellement avec chaque état envoyé par Alice, mais peut attendre que le post-processing classique soit terminé pour faire des mesures collectives sur ses états stockés dans une mémoire quantique.
- Attaques cohérentes : ce sont les attaques les plus puissantes, où Eve n'est limitée que par les lois de la physique quantique. Elle peut préparer un ensemble de ressources intriquées stockées dans une mémoire quantique, qu'elle fait interagir avec chaque état envoyé par Alice, et ensuite attendre que le post-processing soit terminé et faire des mesures collectives. La sécurité contre ce type d'attaques garantit une sécurité inconditionnelle du protocole de QKD. Dans ce cas, on dit aussi qu'un protocole est inconditionnellement sûr.

Nous détaillerons le calcul du taux secret pour le protocole GG02 dans la section suivante. Pour les autres protocoles, le lecteur pourra consulter les références [Scarani09, Weedbrook12, García-Patrón07] pour une présentation plus détaillée des protocoles gaussiens.

## 7.3 Quelques protocoles

### 7.3.1 Variables discrètes

Les protocoles avec des variables discrètes offrent en général de meilleures performances que les protocoles à variables continues que nous présentons ci-dessous, car ils sont moins sensibles au bruit. Lorsqu'un photon est perdu à cause des pertes, il ne va simplement pas contribuer à établir la clé, alors que pour les variables continues, les pertes se traduisent par une contribution du vide lors des mesures homodynes. De ce fait, un échange de clé est possible pour des distances de l'ordre de la centaine de kilomètres. Les protocoles les plus connus sont par exemple BB84 [Bennett84], ou E91 basé sur une distribution d'intrication [Ekert91]. Pour BB84, Alice envoie aléatoirement à Bob un photon dans un des quatre états de polarisation  $|H\rangle$ ,  $|V\rangle$ ,  $|+\rangle = \frac{1}{\sqrt{2}}(|H\rangle + |V\rangle)$ , ou  $|-\rangle = \frac{1}{\sqrt{2}}(|H\rangle - |V\rangle)$ . Ce dernier choisit aléatoirement une base de mesure  $\{|H\rangle, |V\rangle\}$  ou  $\{|-\rangle, |+\rangle\}$ , qu'il révèle ensuite à Alice, en gardant le résultat de sa mesure secret. Les communications classiques permettent ensuite d'établir la clé secrète. Citons également le protocole SARG04 [Scarani04], d'un principe similaire, mais qui est plus résistant à certains types d'attaques. Ces protocoles sont prouvés inconditionnellement sûrs [Scarani09].



### 7.3.2 Variables continues & protocoles gaussiens

Comme pour la plupart de leurs utilisations en information quantique, les variables discrètes posent de nombreux problèmes expérimentaux, pour produire et détecter les photons uniques. L'utilisation de variables continues permet de simplifier grandement le dispositif expérimental, en particulier la détection, où les compteurs de photons sont remplacés par des détecteurs homodynes ou hétérodynes. En contrepartie, le vide n'est plus filtré et les imperfections réduisent davantage les corrélations entre Alice et Bob. L'étape de post-processing classique est plus lourde à mettre en oeuvre, tant du point de vue logiciel que matériel, et constitue un des principaux facteurs limitant les performances. Des distances de 80 km ont été atteintes pour le protocole GG02 que nous présenterons plus loin en détail [Jouguet12c].

Les premiers protocoles à variables continues utilisaient une modulation discrète d'états gaussiens généralisant BB84 [Ralph99, Hillery00], ou des états EPR [Reid00]. Puis plusieurs protocoles utilisant une modulation continue et gaussienne ont été proposés, d'abord avec des états comprimés [Cerf01], puis avec des états cohérents (protocole GG02) [Grosshans02b, Grosshans03b], conjointement avec une détection homodyne chez Bob mesurant aléatoirement la quadrature  $\hat{X}$  ou  $\hat{P}$ . Pour ces deux types d'états utilisés par Alice, Bob peut également utiliser une détection hétérodyne. D'un point de vue pratique, cela évite l'utilisation d'un générateur de nombres aléatoires pour choisir la quadrature qu'il mesure, mais nécessite en contrepartie deux détecteurs homodynes. Par rapport à leurs homologues avec détection homodyne, les performances sont relativement comparables avec des états cohérents [Weedbrook04, Weedbrook06, Lance05]. Pour des états comprimés en revanche, les performances peuvent être supérieures, notamment en terme de tolérance au bruit [García-Patrón09]. La production d'états comprimés reste toutefois plus complexe que la production d'états cohérents.

### 7.3.3 Autres protocoles

**“Distributed phase reference” QKD** Pour ces protocoles, la clé est établie à partir d'une variable discrète, et encodée dans une séquence d'impulsions atténuées  $|\mu\rangle$ . Pour le protocole “Differential Phase Shift” [Inoue02, Wang12], Alice encode l'information dans la différence de phase  $\phi_k$  (0 ou  $\pi$ ) de deux impulsions successives, en envoyant un état

$$|\Psi\rangle = \dots |e^{i\phi_{k-1}}\mu\rangle |e^{i\phi_k}\mu\rangle |e^{i\phi_{k+1}}\mu\rangle \dots \quad (7.1)$$

Pour l'impulsion  $k$ , le bit vaut 0 si  $e^{i\phi_k} = e^{i\phi_{k+1}}$ , et 1 autrement. Pour l'autre protocole appelé “Coherent One Way” [Stucki05, Stucki09], chaque bit est encodé dans une séquence de deux impulsions successives contenant ou non des photons :

$$|0_k\rangle = |\mu\rangle_{2k-1} |0\rangle_{2k} \quad \text{et} \quad |1_k\rangle = |0\rangle_{2k-1} |\mu\rangle_{2k} \quad (7.2)$$

Un des inconvénients de ces protocoles est que les preuves de sécurité sont particulièrement difficiles à établir [Scarani09], avec toutefois des avancées récentes [Moroder12].

**“Device-Independent” QKD** Les protocoles présentés jusqu'à maintenant nécessitent une prise en compte exhaustive de toutes les imperfections expérimentales afin de réellement garantir leur sécurité. Cette tâche peut s'avérer difficile, tant pour le calcul du taux secret, que pour garantir que toutes les imperfections ont bien été prises en compte. Une solution consiste à utiliser des protocoles qui ne dépendent pas explicitement du fonctionnement des appareils utilisés, en se basant par exemple sur la violation d'une inégalité de Bell (“Device Independent”

QKD)[Acín06, Acín07, Gisin10]. Une telle violation est toutefois assez exigeante expérimentalement : elle requiert notamment des détecteurs de très forte efficacité et des faibles pertes afin de ne pas introduire de *detection loophole* [Scarani09].

**Protocoles à variables continues avec une modulation non gaussienne** Afin de simplifier le post-processing classique tout en bénéficiant des avantages des variables continues, certains protocoles proposent d'utiliser une modulation non gaussienne d'états cohérents [Leverrier09, Leverrier11, Sych10], qui peuvent offrir de meilleures performances qu'avec une modulation gaussienne, mais pour lesquels les preuves de sécurité sont plus difficiles à établir.

### 7.3.4 Réseaux et applications commerciales

La QKD commence à atteindre un stade commercial depuis quelques années. Plusieurs entreprises<sup>2</sup> proposent des modules en général basés sur BB84 ou GG02, pouvant aussi être couplés à des systèmes de cryptographie classique de type AES. En 2008, le projet international SECOQC [Peev09] a permis de démontrer la faisabilité d'un réseau de QKD à plusieurs noeuds, pour la plupart distants de plus de 20 km. Plusieurs protocoles furent utilisés : BB84 et sa version avec états intriqués BBM92 [Bennett92], SARG04, le protocole "Coherent One Way", et le protocole à états cohérents GG02. Tous les noeuds étaient reliés par des fibres optiques, sauf un qui était en espace libre. Un autre projet de démonstration est en cours à Tokyo [Sasaki11], utilisant BB84 et BBM92, SARG04, et le protocole "Differential Phase Shift".

## 7.4 Le protocole GG02 : de la QKD avec des états "classiques"

Nous détaillons maintenant plus en détail le protocole GG02 qui possède de nombreux avantages techniques en faveur d'un développement commercial, comme une mise en œuvre avec des composants telecom standards, ou la compatibilité avec le multiplexage en longueur d'onde [Qi10].

### 7.4.1 Principe

Dans la description "Prepare & Measure" (P&M), Alice utilise les quadratures d'états cohérents pour encoder l'information. Comme tout protocole de QKD, GG02 comporte une partie quantique et une partie classique. Les étapes de la partie quantique sont les suivantes :

- Préparation des états quantiques : Alice choisit deux nombres  $x_A$  et  $p_A$  selon une modulation gaussienne de variance<sup>3</sup>  $V_A$  et de moyenne nulle, et prépare ensuite un état cohérent centré en  $(x_A, p_A)$ . Elle répète cette procédure autant de fois que nécessaire.
- Transmission des états quantiques : Alice envoie ses états cohérents à Bob à travers le canal quantique. Pour des raisons qui seront plus claires par la suite, on suppose que le canal est gaussien. Il est décrit par une transmission  $T$  et un excès de bruit ramené à l'entrée  $\epsilon$ .
- Mesure des états quantiques : pour chaque état reçu, Bob choisit aléatoirement de mesurer la quadrature  $\hat{X}$  ou  $\hat{P}$  avec sa détection homodyne. Il obtient une valeur  $x_B$  ou  $p_B$  corrélée avec la quadrature d'Alice correspondante. Le degré de corrélation dépend de la variance de modulation  $V_A$  et des caractéristiques du canal quantique.

2. Par exemple ID Quantique (idquantique.com), MagiQ (magiqtech.com/MagiQ/), Quintessence Labs (qlabsusa.com) ou SeQureNet (sequirenet.fr).

3. Dans ce chapitre, nous utilisons la convention  $N_0=1$ .

Alice et Bob utilisent ensuite des communications et des algorithmes classiques pour extraire une clé, si les conditions le permettent. Les étapes de cette partie classique sont résumées ci dessous :

- “Sifting” : pour chaque état envoyé, Alice possède deux valeurs  $x_A$  et  $p_A$  mais Bob n’en mesure qu’une. Bob communique donc à Alice les quadratures qu’il a mesuré (sans révéler les résultats de ses mesures), qui ne garde que les quadratures correspondantes.
- Estimation des paramètres : une partie aléatoire des données est utilisée pour estimer les paramètres du canal quantique et borner l’information qu’Eve a pu obtenir.
- Réconciliation et correction d’erreurs : Alice et Bob communiquent des syndromes d’erreur pour obtenir la même séquence de bits, qui n’est à ce stade pas secrète. Cette étape est un des facteurs limitant de la QKD avec des variables continues. Lorsque les données d’Alice servent de référence et que Bob corrige les siennes, on parle de *réconciliation directe*. Dans le cas contraire, lorsque ce sont les données de Bob qui servent de référence et qu’Alice corrige les siennes, on parle de *réconciliation inverse*.
- Amplification de confidentialité : la séquence de bits est transformée en une clé secrète plus petite, mais inconnue d’Eve.

Le protocole GG02 permet une mise en oeuvre expérimentale relativement simple. Il fonctionne maintenant de manière entièrement fibrée [Lodewyck07, Fossier09b], avec une bonne fiabilité sur une longue durée [Jouguet12b]. Récemment, P. Jouguet *et al.* ont réussi à transmettre une clé secrète pour une distance de 80 km avec un débit d’environ 1 kbit/s [Jouguet12c], en se basant sur une amélioration des algorithmes de réconciliation [Jouguet11].

### 7.4.2 Modélisation à intrication virtuelle

Ce protocole peut être reformulé en une description à intrication virtuelle, dite “Entanglement-Based” (EB), complètement équivalente à la version P&M, mais pour laquelle les preuves de sécurité sont plus faciles à établir [Grosshans03a, Weedbrook12].

Au lieu de préparer et d’envoyer des états cohérents, Alice effectue une mesure hétérodyne sur un mode d’un état EPR  $|\lambda\rangle = \sqrt{1-\lambda^2} \sum_n \lambda^n |n\rangle |n\rangle$ , alors que l’autre mode est envoyé à Bob (Fig. 7.2). Comme nous le montrons en annexe I, cette mesure projette le mode de Bob sur un état cohérent, avec une amplitude proportionnelle au résultat de la mesure d’Alice, et avec une variance de modulation  $V_A$  égale à

$$V_A = V - 1, \quad (7.3)$$

où  $V = \frac{1+\lambda^2}{1-\lambda^2}$  est la variance de l’état thermique obtenu en traçant un des modes de  $|\lambda\rangle$ . Même si l’implémentation expérimentale n’utilise pas d’intrication, les taux secrets sont les mêmes dans les descriptions P&M et EB. L’avantage de la description EB est que l’on peut simplifier théoriquement la modélisation du protocole en “retardant” la mesure hétérodyne d’Alice, puisque cette dernière commute avec l’action du canal quantique et la mesure de Bob. Nous montrerons alors que l’on peut obtenir une borne sur le taux secret qui ne dépend que de la matrice de covariance entre Alice et Bob (A et B sur la figure 7.2) avant leurs détecteurs.

### 7.4.3 Preuves de sécurité

En estimant l’information qu’Eve peut acquérir au cours de la communication, il est possible de calculer le taux de clé secrète théoriquement accessible pour un protocole donné. Nous indi-

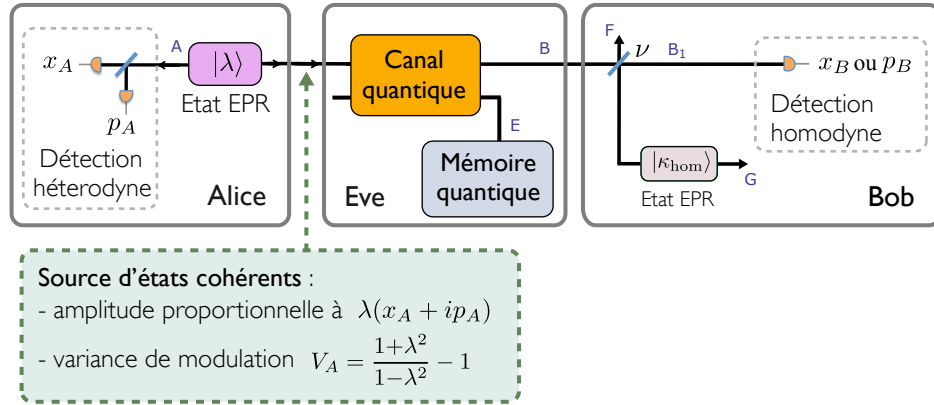


FIGURE 7.2 – Modélisation à intrication virtuelle pour le protocole GG02. Pour Bob, la lame séparatrice  $\nu$  et l'état EPR  $|\kappa_{\text{hom}}\rangle$  modélisent respectivement l'efficacité et le bruit électronique de la détection homodyne.

querons seulement les principales étapes permettant de calculer le taux secret pour le protocole GG02, sans détailler toutes les étapes de calcul. Pour obtenir plus de détails sur ce protocole, ou pour un protocole quelconque, le lecteur pourra notamment consulter les références [Scarani09] et [Weedbrook12], la thèse de R. García-Patrón [García-Patrón07] ou encore la thèse de R. Renner [Renner05a]. Le calcul du taux secret se fait en deux étapes. La première est de calculer l'information mutuelle entre Alice et Bob. La seconde est d'estimer l'information qu'Eve a pu acquérir, en fonction du type d'attaque. Le taux secret est ensuite simplement la différence de ces deux quantités.

### Optimalité des attaques gaussiennes

Pour des protocoles gaussiens, la sécurité contre des attaques collectives suffit à garantir la sécurité contre des attaques arbitraires à la limite asymptotique [Renner09], c'est à dire lorsque le nombre de mesures tend vers l'infini. Ce résultat simplifie considérablement l'analyse, car les attaques gaussiennes sont optimales parmi les attaques collectives pour les protocoles gaussiens [García-Patrón06, Navascues06, Pirandola08, Leverrier10a].

De plus, si jamais l'état reçu par Bob n'est pas parfaitement gaussien, le taux secret est minimisé en considérant l'état gaussien de même matrice de covariance. Ainsi, on peut considérer que le canal et l'état reçu par Bob sont gaussiens pour estimer une borne inférieure du taux secret. Pour un canal linéaire, symétrique et sans mémoire, le cas le plus courant expérimentalement, deux paramètres sont nécessaires pour le caractériser : sa transmission  $T$  et son excès de bruit  $\epsilon$ , comme nous l'avons vu dans la section 4.4.

Afin de calculer l'information de l'espion, les paramètres du canal et la matrice de covariance doivent être estimés à partir d'un échantillon des données partagées par Bob et Alice. A la limite asymptotique, ces paramètres peuvent être parfaitement estimés. Pour un nombre de mesures fini en revanche, les incertitudes statistiques doivent être prises en compte dans les preuves de sécurité et compliquent considérablement la tâche. Il a été démontré que le protocole GG02 est sûr contre les attaques collectives pour un nombre fini de mesures en supposant que l'état est gaussien [Leverrier10b], et que la sécurité contre les attaques collectives implique la sécurité contre les attaques cohérentes sans faire d'hypothèse sur l'état [Leverrier13].

## Efficacité de réconciliation

L'information mutuelle donnée par la théorie de Shannon est une limite théorique qu'il est difficile d'atteindre, car elle nécessite des algorithmes extrêmement performants et d'importantes ressources matérielles pour les implémenter. En pratique, Alice et Bob ne disposent alors pas d'une information mutuelle  $I(A:B)$ , mais de  $\beta I(A:B)$ , où  $0 < \beta < 1$  est l'*efficacité de réconciliation*.

Obtenir une bonne efficacité de réconciliation est un des challenges de la QKD avec des variables continues. De nombreux progrès ont été faits ces dernières années [Bloch06, Leverrier08, Jouguet12a] et les procédures actuelles permettent d'obtenir  $\beta \approx 95\%$  pour une grande plage de rapport signal-à-bruit [Jouguet11].

### 7.4.4 Expression des taux secrets

Nous nous plaçons dans le cas général d'une détection homodyne imparfaite, avec une efficacité  $\nu$  et un bruit électronique  $\kappa$  (sur les mesures). Le canal quantique gaussien est caractérisé par sa transmission  $T$ , et un excès de bruit ramené à l'entrée  $\epsilon$ . Rappelons les notations utilisées dans la section 4.4 :  $\chi_{\text{line}} = \frac{1-T}{T} + \epsilon$  est le bruit ajouté par le canal ramené à l'entrée incluant l'effet des pertes,  $\chi_{\text{hom}} = \frac{1-\nu}{\nu} + \frac{\kappa}{\nu}$  est le bruit ajouté par la détection homodyne ramené à l'entrée de celle-ci, et  $\chi_{\text{tot}} = \chi_{\text{line}} + \frac{\chi_{\text{hom}}}{T}$  est le bruit ajouté total ramené à l'entrée du canal.

Pour un tel canal, la matrice de covariance d'Alice et Bob avant leurs détections est donnée par [Weedbrook12]

$$\Gamma_{AB} = \begin{pmatrix} V(\lambda)\mathbb{I}_2 & \sqrt{T(V(\lambda)^2-1)}\mathbb{Z} \\ \sqrt{T(V(\lambda)^2-1)}\mathbb{Z} & T(V(\lambda) + \chi_{\text{tot}})\mathbb{I}_2 \end{pmatrix}, \quad (7.4)$$

avec  $\mathbb{I}_2 = \text{diag}(1, 1)$ ,  $\mathbb{Z} = \text{diag}(1, -1)$ , et  $V(\lambda) = \frac{1+\lambda^2}{1-\lambda^2}$  est la variance de l'état thermique  $\text{Tr}_A |\lambda\rangle\langle\lambda|$ .

### Information mutuelle entre Alice et Bob

En utilisant (4.32) avec  $\langle X_A^2 \rangle = V - 1$  et  $\langle X_N^2 \rangle = 1 + \chi_{\text{tot}}$ , l'information mutuelle  $I_{AB} := I(A:B)$  entre Alice et Bob est donnée par :

$$I_{AB} = \frac{1}{2} \log_2 \left[ \frac{V + \chi_{\text{tot}}}{1 + \chi_{\text{tot}}} \right] \quad (7.5a)$$

$$= \frac{1}{2} \log_2 \left[ 1 + \frac{2T\nu\lambda^2}{(1-\lambda^2)(1+T\nu\epsilon+\kappa)} \right] \quad (7.5b)$$

### Attaques individuelles

Lorsqu'Eve est limitée aux attaques individuelles, l'information qu'elle peut acquérir est donnée par l'information mutuelle de Shannon. Si la réconciliation est directe, elle acquiert  $I_{BA} := I(E:A)$  bits d'information sur les données d'Alice. Si la réconciliation est inverse, elle acquiert  $I_{BE} := I(E:B)$  bits d'information sur les mesures de Bob. Nous ne considérerons que ce dernier cas dans la suite de ce manuscrit. Dans ce cas  $I_{BE}$  est donnée par [Lodewyck07,

García-Patrón07] :

$$I_{BE} = \frac{1}{2} \log_2 \frac{V_B}{V_{B|E}} \quad (7.6a)$$

$$= \frac{1}{2} \log_2 \left[ \frac{T^2 (V + \chi_{\text{tot}}) \left( \frac{1}{V} + \chi_{\text{line}} \right)}{1 + T \chi_{\text{hom}} \left( \frac{1}{V} + \chi_{\text{line}} \right)} \right] \quad (7.6b)$$

$$= \frac{1}{2} \log_2 \left[ \frac{(\lambda^2 + T [\lambda^2(\epsilon - 2) + \epsilon] + 1) ([1 + \kappa] [\lambda^2 - 1] + \nu T [\lambda^2(\epsilon - 2) - \epsilon])}{(\lambda^2 - 1) ([\kappa + 1] [\lambda^2 + 1] + T [\kappa - \nu + 1] [\lambda^2(\epsilon - 2) + \epsilon])} \right] \quad (7.6c)$$

Le taux secret contre les attaques individuelles est finalement égal à :

$$\Delta I_{\text{Ind}} = \beta I_{AB} - I_{BE} \quad (7.7)$$

### Attaques collectives

Lorsqu'Eve peut faire des attaques collectives, l'information qu'elle peut acquérir n'est plus donnée par la théorie de Shannon, mais par la borne de Holevo. Plus précisément, si Bob fait une mesure donnant une valeur  $x_B$  avec une densité de probabilité  $p(x_B)$ , l'information accessible à Eve est bornée supérieurement par [Devetak05, Scarani09]

$$S(E:B) \leq \chi_{BE} = S(\hat{\rho}_E) - \int dx_B p(x_B) S(\hat{\rho}_E^{x_B}), \quad (7.8)$$

où  $\hat{\rho}_E^{x_B}$  est l'état de Eve conditionné par la mesure  $x_B$  de Bob. Cette borne peut se simplifier afin d'obtenir une expression analytique. Par hypothèse, Eve purifie l'état d'Alice et Bob avant leurs détections : l'état  $\hat{\rho}_{ABE}$  est pur. Comme montré dans la section 4.3.2, on a alors  $S(\hat{\rho}_E) = S(\hat{\rho}_{AB})$ .

De plus, après la mesure de Bob, l'état  $\hat{\rho}_{AEFG}$  est également pur, et donc  $S(\hat{\rho}_E^{x_B}) = S(\hat{\rho}_{AFG}^{x_B})$  (voir la figure 7.2). On montre ensuite que la matrice de covariance d'un état suite à une mesure homodyne ne dépend pas du résultat de la mesure [Eisert02]. Puisque l'entropie d'un état gaussien, donnée par (4.14), ne dépend que de sa matrice de covariance, ceci implique que  $S(\hat{\rho}_{AFG}^{x_B})$  est indépendant de  $x_B$ , et peut être sorti de l'intégrale dans (7.9), qui devient alors

$$\chi_{BE} = S(\hat{\rho}_{AB}) - S(\hat{\rho}_{AFG}^{x_i}). \quad (7.9)$$

Cette expression ne dépend donc que de la matrice de covariance d'Alice et Bob et des paramètres de la détection homodyne de Bob. Elle est donnée par [Lodewyck07] :

$$\chi_{BE} = G \left[ \frac{\mu_1 - 1}{2} \right] + G \left[ \frac{\mu_2 - 1}{2} \right] - G \left[ \frac{\mu_3 - 1}{2} \right] - G \left[ \frac{\mu_4 - 1}{2} \right] \quad (7.10)$$

avec

$$G[x] = (x + 1) \log_2[x + 1] - x \log_2 x \text{ si } x \neq 0, \text{ et } G[0] = 0 \quad (7.11a)$$

$$\mu_{1,2}^2 = \frac{1}{2} \left( A \pm \sqrt{A^2 - 4E} \right) \quad \mu_{3,4}^2 = \frac{1}{2} \left( C \pm \sqrt{C^2 - 4D} \right) \quad (7.11b)$$

$$A = V^2(1 - 2T) + 2T + T^2(V + \chi_{\text{line}})^2 \quad E = T^2(V \chi_{\text{line}} + 1)^2 \quad (7.11c)$$

$$C = \frac{V\sqrt{E} + T(V + \chi_{\text{line}}) + A\chi_{\text{hom}}}{T(V + \chi_{\text{tot}})} \quad D = \sqrt{E} \frac{V + \sqrt{E}\chi_{\text{hom}}}{T(V + \chi_{\text{tot}})} \quad (7.11d)$$

Le taux secret contre les attaques collectives est finalement donné par :

$$\Delta I_H = \beta I_{AB} - \chi_{BE} \quad (7.12)$$

## 7.5 Conclusion

Dans ce chapitre, nous avons présenté les principes généraux de la distribution quantique de clé, et les principaux protocoles utilisés. Nous avons également présenté les principaux outils permettant de calculer le taux de clé secret qu'il est possible d'atteindre, selon différentes hypothèses sur les capacités d'un espion tentant d'acquérir de l'information sur la clé.

Le protocole GG02 utilisant des états cohérents présente l'avantage d'être l'un des plus simples à réaliser expérimentalement, tout en offrant des performances intéressantes. Cependant, comparées à d'autres protocoles, les distances atteignables restent encore un peu plus faibles, principalement en raison d'un traitement classique qui nécessite des algorithmes très efficaces.

L'amélioration de cette distance maximale de transmission fait l'objet des deux chapitres suivants, où l'on montrera l'utilité d'un amplificateur quantique spécial, capable d'amplifier un signal sans en amplifier le bruit quantique.

## Chapitre 8

# L'amplificateur sans bruit non déterministe

### Sommaire

---

|            |                                                                                      |            |
|------------|--------------------------------------------------------------------------------------|------------|
| <b>8.1</b> | <b>Introduction</b>                                                                  | <b>153</b> |
| <b>8.2</b> | <b>Amplificateurs déterministes</b>                                                  | <b>154</b> |
| 8.2.1      | Bruit minimal ajouté par un amplificateur déterministe                               | 154        |
| 8.2.2      | Amplificateur indépendant de la phase                                                | 155        |
| 8.2.3      | Amplificateur dépendant de la phase                                                  | 156        |
| <b>8.3</b> | <b>Principe et propriétés de base de l'amplificateur sans bruit non déterministe</b> | <b>157</b> |
| 8.3.1      | Principe                                                                             | 157        |
| 8.3.2      | Transformation de quelques états gaussiens                                           | 159        |
| <b>8.4</b> | <b>Implémentations théoriques et réalisations expérimentales</b>                     | <b>161</b> |
| 8.4.1      | Les ciseaux quantiques                                                               | 161        |
| 8.4.2      | Autres réalisations                                                                  | 164        |
| <b>8.5</b> | <b>Applications</b>                                                                  | <b>165</b> |
| 8.5.1      | Préparation d'états quantiques                                                       | 165        |
| 8.5.2      | Communications quantiques                                                            | 166        |
| <b>8.6</b> | <b>Conclusion</b>                                                                    | <b>166</b> |

---

### 8.1 Introduction

Les communications quantiques sont en plein essor et commencent même à être technologiquement réalisables pour des systèmes simples. De manière générale, l'information est encodée sur un état quantique, qui est ensuite envoyé dans un canal quantique à un destinataire. Le canal de transmission est "quantique", dans le sens où il préserve les superpositions quantiques de l'état envoyé. En optique quantique, il peut s'agir d'une fibre optique, ou voir simplement de l'espace libre. En pratique cependant, des pertes et/ou du bruit vont contribuer à dégrader l'état et donc l'information encodée. Un *amplificateur quantique*, agissant directement sur les états quantiques, peut alors être nécessaire de manière à amplifier les composantes de l'état quantique codant l'information.



Lorsque les états quantiques sont de faibles amplitudes, la qualité des appareils de mesure est extrêmement importante. Si le signal mesuré par une détection homodyne est entaché d'un bruit électronique du même ordre de grandeur, l'information sera très difficile à extraire, même en utilisant un amplificateur classique *après* la mesure, puisque le signal sera “noyé” dans le bruit.

Un amplificateur quantique peut donc être particulièrement utile pour pallier les défauts des appareils de mesure ou du canal de transmission, en agissant directement sur les états quantiques *avant* que les mesures ne soient effectuées. Encore faut-il que cet amplificateur préserve les cohérences quantiques, et qu'il n'introduise pas trop de bruit sur les états en les amplifiant. Malheureusement, une propriété fondamentale imposée par les lois de la physique quantique limite fortement ce dernier point : un amplificateur quantique linéaire et déterministe, fonctionnant indépendamment de la phase de l'état d'entrée, *doit* ajouter une quantité minimale de bruit, quelle que soit la technologie utilisée [Caves82, Clerk10]. L'état amplifié est alors bruité selon deux contributions : le bruit initial amplifié, et le bruit ajouté par l'amplificateur.

En renonçant à un fonctionnement déterministe, T. Ralph et A. Lund ont proposé un nouveau type d'amplificateur quantique possédant des propriétés bien particulières [Ralph08], connu sous le nom d'amplificateur sans bruit non déterministe (ou *heralded noiseless linear amplifier*, que nous abrègerons par *NLA*) : *sans bruit* car non seulement il n'ajoute pas de bruit bien qu'étant indépendant de la phase, mais il est de plus capable d'amplifier un état sans amplifier le bruit quantique associé, en transformant un état cohérent  $|\alpha\rangle$  en un état amplifié  $|g\alpha\rangle$  ; *non déterministe* car cette amplification ne peut être réalisée qu'avec une probabilité de succès inférieure à 1.

Ce chapitre est consacré à la présentation détaillée de cet amplificateur non déterministe, de ses propriétés et des implémentations expérimentales réalisées par plusieurs groupes. Nous commencerons par montrer les limites d'une amplification déterministe, qu'elle soit dépendante ou indépendante de la phase. Puis nous introduirons ensuite le NLA d'un point de vue théorique, en étudiant ses propriétés de base et la transformation de quelques états courants. Nous passerons ensuite en revue plusieurs méthodes d'implémentations approchées du NLA ainsi que les réalisations expérimentales associées. Enfin, nous terminerons ce chapitre en présentant quelques applications de cet amplificateur dans des domaines variés.

Son utilisation en cryptographie quantique fera l'objet d'une étude détaillée dans le chapitre 10.

## 8.2 Amplificateurs déterministes

### 8.2.1 Bruit minimal ajouté par un amplificateur déterministe

Considérons un amplificateur quantique, linéaire, et unitaire (donc déterministe). Dans le cas général, son fonctionnement peut dépendre de la phase, si bien que les quadratures  $\hat{X}$  et  $\hat{P}$  peuvent être amplifiées avec des gains différents  $g_X$  et  $g_P$ . Si on suppose que l'amplificateur ne couple pas les quadratures entre elles, leurs valeurs moyennes sont transformées en :

$$\langle \hat{X} \rangle \rightarrow g_X \langle \hat{X} \rangle \quad (8.1a)$$

$$\langle \hat{P} \rangle \rightarrow g_P \langle \hat{P} \rangle \quad (8.1b)$$

On pourrait alors être tenté d'écrire la transformation des opérateurs selon  $\hat{X}_{\text{out}} = g_X \hat{X}_{\text{in}}$  et  $\hat{P}_{\text{out}} = g_P \hat{P}_{\text{in}}$ . Un rapide calcul montre cependant qu'une telle transformation ne préserverait pas

le commutateur de  $\hat{\mathbf{X}}_{\text{out}}$  et  $\hat{\mathbf{P}}_{\text{out}}$ , alors que cela est requis par l'unitarité. Il est donc nécessaire d'introduire des opérateurs  $\hat{\mathbf{B}}_{\text{X}}$  et  $\hat{\mathbf{B}}_{\text{P}}$  tels que :

$$\hat{\mathbf{X}}_{\text{out}} = g_{\text{X}}\hat{\mathbf{X}}_{\text{in}} + \hat{\mathbf{B}}_{\text{X}} \quad (8.2a)$$

$$\hat{\mathbf{P}}_{\text{out}} = g_{\text{P}}\hat{\mathbf{P}}_{\text{in}} + \hat{\mathbf{B}}_{\text{P}} \quad (8.2b)$$

Les opérateurs  $\hat{\mathbf{B}}_{\text{X}}$  et  $\hat{\mathbf{B}}_{\text{P}}$  doivent être de moyenne nulle afin que l'on puisse avoir la transformation (8.1). La relation de commutation  $[\hat{\mathbf{X}}_{\text{out}}, \hat{\mathbf{P}}_{\text{out}}] = 2iN_0$  peut maintenant être vérifiée si

$$[\hat{\mathbf{B}}_{\text{X}}, \hat{\mathbf{B}}_{\text{P}}] = (1 - g_{\text{X}}g_{\text{P}})2iN_0. \quad (8.3)$$

Ce commutateur entraîne alors une relation d'incertitude de Heisenberg

$$\Delta B_{\text{X}}\Delta B_{\text{P}} \geq |g_{\text{X}}g_{\text{P}}-1|N_0. \quad (8.4)$$

Si  $g_{\text{P}} \neq 1/g_{\text{X}}$ , les opérateurs  $\hat{\mathbf{B}}_{\text{X}}$  et  $\hat{\mathbf{B}}_{\text{P}}$  ne peuvent pas avoir simultanément une variance nulle : ils vont donc nécessairement ajouter du *bruit* sur le signal quantique.

### 8.2.2 Amplificateur indépendant de la phase

Un tel amplificateur agit de la même façon sur les quadratures  $\hat{\mathbf{X}}$  et  $\hat{\mathbf{P}}$  :  $g_{\text{P}}=g_{\text{X}}$ . Selon la relation (8.4), on a alors  $\Delta B_{\text{X}}\Delta B_{\text{P}} \geq |g^2-1|N_0$ . On peut raisonnablement supposer que l'amplificateur introduit un bruit identique sur les deux quadratures, avec  $\Delta B_{\text{X}}=\Delta B_{\text{P}}$ . On a donc :

$$\Delta^2 B_{\text{X}} \geq |g^2-1|N_0 \quad \text{et} \quad \Delta^2 B_{\text{P}} \geq |g^2-1|N_0 \quad (8.5)$$

Le bruit ajouté par l'amplificateur aura une variance au moins égale à  $|g^2-1|N_0$ . L'inégalité est saturée pour un amplificateur *idéal*. Un tel amplificateur transforme donc une quadrature  $\hat{\mathbf{Q}}_{\text{in}}$  en

$$\hat{\mathbf{Q}}_{\text{out}} = g^2\hat{\mathbf{Q}}_{\text{in}} + |g^2-1|\hat{\mathbf{Q}}_0, \quad (8.6)$$

où  $\hat{\mathbf{Q}}_0$  est la quadrature d'un mode vide. En termes d'amplitude et de variance, cette transformation est équivalente à

$$\begin{aligned} \langle Q_{\text{out}} \rangle &= g\langle Q_{\text{in}} \rangle \\ \Delta^2 Q_{\text{out}} &= g^2\Delta^2 Q_{\text{in}} + |g^2-1|N_0 \end{aligned}$$

(8.7)

Pour un état cohérent  $|\alpha\rangle$ , l'état amplifié aura donc une variance  $g^2N_0$  due à l'amplification du bruit initial, plus la contribution minimale  $|g^2-1|N_0$  de l'amplificateur (Fig. 8.1). Le bruit ajouté ramené à l'entrée à une variance de  $|1-\frac{1}{g^2}|N_0$ . Un amplificateur de grand gain va donc ajouter une unité de bruit de photon à l'état initial. Pour des états cohérents de grandes amplitudes devant la variance du vide, cette contribution sera toutefois négligeable : on retrouve la possibilité d'amplifier un champ lumineux classique sans ajouter de bruit. En revanche, pour nos états quantiques de faibles amplitudes, cet ajout de bruit pourra s'avérer problématique et limiter l'utilité d'un tel amplificateur.

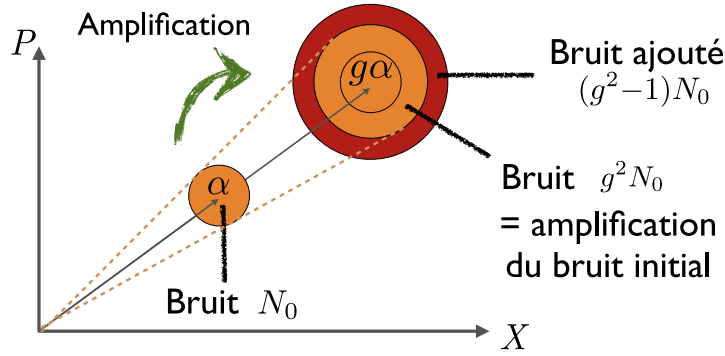


FIGURE 8.1 – Principe d'un amplificateur déterministe indépendant de la phase.

Nous avons gardé les valeurs absolues dans (8.5) afin de pouvoir traiter le cas d'un amplificateur de gain inférieur à 1, correspondant à un canal à pertes, déjà traité dans la section 4.4. Le bruit ajouté sur l'état amplifié est dans ce cas  $(1-g^2)N_0$ . Pour un état cohérent, le bruit total  $(1-g^2)N_0 + g^2N_0 = N_0$  reste égal au bruit de photon : on retrouve le fait qu'un état cohérent atténué reste un état cohérent.

De nombreux systèmes physiques permettent de réaliser un amplificateur déterministe indépendant de la phase (voir la référence [Caves12] pour quelques exemples). Il est possible de se passer d'une interaction non linéaire en utilisant un état comprimé préparé *off-line*, une détection homodyne et un déplacement [Filip05], ou simplement une détection hétérodyne et un déplacement [Josse06]. Toutes les diverses implémentations, saturant ou non (8.5), sont équivalentes à un squeezer bimode dont un des modes d'entrée est un état  $\hat{\sigma}$  qui dépend du bruit ajouté par l'amplificateur [Caves12, Jiang12].

### 8.2.3 Amplificateur dépendant de la phase

Tous les amplificateurs quantiques ne rajoutent pas forcément du bruit. Si l'on renonce à un fonctionnement indépendant de la phase, en amplifiant une quadrature avec un gain  $g$  et en atténuant l'autre avec un gain  $1/g$ , comme montré sur la figure 8.2, les opérateurs  $\hat{B}$  ne sont plus nécessaires car le commutateur de  $\hat{X}_{\text{out}}$  et  $\hat{P}_{\text{out}}$  est bien conservé<sup>1</sup>. Un amplificateur dépendant de la phase est par exemple réalisé en utilisant un squeezer monomode, avec un gain  $g=e^r$  pour une quadrature et  $1/g=e^{-r}$  pour l'autre. Seul le bruit initial est amplifié pour une quadrature et atténué pour l'autre : l'état amplifié sera donc *comprimé*. Puisque le bruit et l'amplitude du signal sont amplifiés de la même façon, le rapport signal-à-bruit est conservé, pour les deux quadratures.

On comprend intuitivement que ce type d'amplificateur, bien que plus performant que l'amplificateur indépendant de la phase, ne trouve qu'un intérêt limité lorsque l'on effectue une mesure avec un détecteur parfait, supposé pouvoir mesurer des amplitudes avec une précision arbitraire. En revanche, il peut être utile lorsque les appareils sont imparfaits, notamment en cryptographie quantique afin de compenser les imperfections de la détection homodyne [Fossier09a].

1. Un tel amplificateur est pour cette raison quelquefois appelé "amplificateur sans bruit" dans la littérature, mais nous garderons cette dénomination pour le véritable amplificateur sans bruit présenté dans la prochaine section.

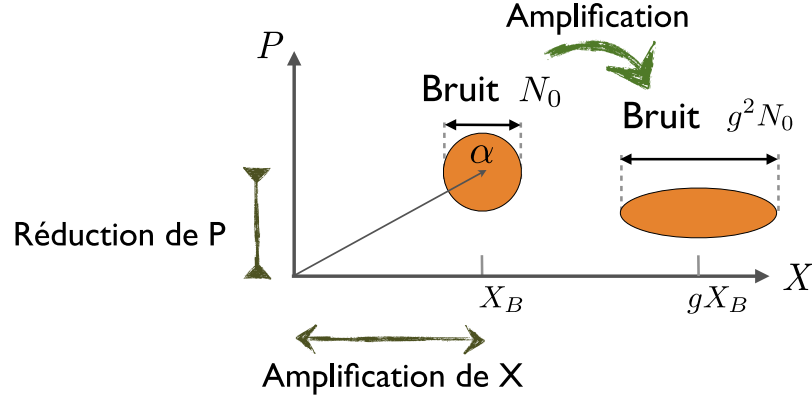


FIGURE 8.2 – Principe d'un amplificateur déterministe dépendant de la phase.

### 8.3 Principe et propriétés de base de l'amplificateur sans bruit non déterministe

Nous avons vu qu'une amplification unitaire ajoute toujours un minimum de bruit si elle est indépendante de la phase. Au mieux, elle n'amplifie que le bruit de l'état initial en étant dépendante de la phase. Il est même assez difficile d'imaginer comment il pourrait en être autrement : il faudrait que l'amplificateur puisse "distinguer" le bruit du signal pour faire une amplification sélective. C'est pourtant une opération surprenante qu'il est possible de faire en renonçant à l'unitarité, avec *l'amplificateur sans bruit non déterministe*.

#### 8.3.1 Principe

Cet amplificateur, introduit par T. Ralph et A. Lund [Ralph08], est indépendant de la phase et réalise la transformation :

$$|\alpha\rangle \rightarrow |g\alpha\rangle \quad (8.8)$$

Un état cohérent est amplifié en un autre état cohérent, en restant au bruit de photon (Fig. 8.3). Cette transformation est différente d'un déplacement car elle fonctionne sans que l'on ait besoin de connaître  $\alpha$ .

#### Un amplificateur non unitaire

En vertu de (8.3), nous savons déjà que cette opération ne peut pas être unitaire. On peut également démontrer cette impossibilité en invoquant le violation de plusieurs propriétés fondamentales. Il serait par exemple possible de violer le théorème de non clonage [Scarani05, Wootters82], en transformant un état  $|\alpha\rangle$  en  $|\sqrt{2}\alpha\rangle$  puis en le séparant sur une lame séparatrice. Un autre argument a été proposé dans la référence [Ralph08] : supposons qu'il existe un opérateur unitaire  $\hat{T}$  permettant de réaliser la transformation  $\hat{T}|\alpha\rangle = e^{i\theta}|g\alpha\rangle$ , où  $\theta$  est une phase quelconque et  $g > 1$ . On a donc :

$$\hat{T}\hat{a}|\alpha\rangle = \hat{T}\hat{a}\hat{T}^\dagger\hat{T}|\alpha\rangle = \hat{T}\hat{a}\hat{T}^\dagger e^{i\theta}|g\alpha\rangle = \alpha e^{i\theta}|g\alpha\rangle \quad (8.9)$$

$|g\alpha\rangle$  est donc vecteur propre de l'opérateur  $\hat{b} := \hat{T}\hat{a}\hat{T}^\dagger$  avec la valeur propre  $\alpha$ . On doit alors avoir  $\hat{b} = (1/g)\hat{a}$ . Puisque  $[\hat{a}, \hat{a}^\dagger] = 1$ , ceci implique que  $[\hat{b}, \hat{b}^\dagger] = 1/g^2$ . Mais on a aussi  $[\hat{b}, \hat{b}^\dagger] = \hat{T}[\hat{a}, \hat{a}^\dagger]\hat{T}^\dagger = 1$ .

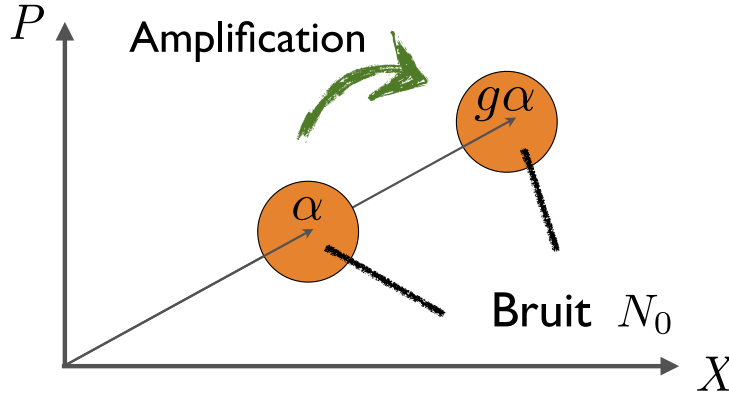


FIGURE 8.3 – Principe de l'amplificateur sans bruit non déterministe.

Il y a donc une contradiction, et on doit renoncer à l'unitarité de  $\hat{T}$ , qui ne peut donc être qu'une opération non déterministe.

D'autres opérations, comme le clonage, impossibles à réaliser de manière déterministe, peuvent également l'être approximativement de manière probabiliste [Fiurášek04]. Le lien entre l'amplification sans bruit et le clonage est même direct : le clonage probabiliste  $|\alpha\rangle \rightarrow |\alpha\rangle^{\otimes M}$ , suivi d'une transformation unitaire de  $|\alpha\rangle^{\otimes M}$  en  $|\sqrt{M}\alpha\rangle \otimes [|0\rangle^{\otimes M-1}]$  [Fiurášek01] correspondrait à l'amplification sans bruit (8.8).

### Opérateur décrivant l'amplification et probabilité de succès

Quel opérateur peut-on associer à la transformation (8.8), pour les cas où l'amplification est réussie ? Puisque l'état final doit être normalisé, on pourrait imaginer qu'il existe une certaine liberté sur l'opérateur lié au NLA. Par exemple, une transformation  $\hat{E}(|\alpha\rangle\langle\alpha|) = |g\alpha\rangle\langle g\alpha|$  suffirait à remplir le cahier des charges défini à partir de l'amplification d'un état cohérent. En fait, les implémentations que nous détaillerons par la suite réalisent – approximativement – une transformation  $\hat{T}$  un peu différente, transformant un état de Fock  $|n\rangle$  en  $\hat{T}|n\rangle = g^n|n\rangle$ , et qui peut être décrite par :

$$\hat{T} = g^{\hat{n}} \quad (8.10)$$

Comme attendu, cet opérateur n'est pas unitaire, puisque  $\hat{T}^\dagger = \hat{T}$  et  $\hat{T}^{-1} = (1/g)^{\hat{n}}$ . Une certaine attention est nécessaire lors de l'utilisation de  $\hat{T}$ , car c'est un opérateur non borné qui peut, pour certains états, conduire à des états amplifiés divergents. Nous aurons l'occasion de revenir sur ce point dans la suite de ce chapitre. Il ne conserve pas non plus la trace, qui peut augmenter : la probabilité de succès ne peut pas être obtenue en prenant simplement la norme de l'état amplifié.

Alors, avec quelle probabilité de succès peut-on réaliser une amplification sans bruit ? Tout dépend du degré de perfection que l'on cherche à obtenir...car strictement parlant, il n'est possible d'implémenter parfaitement  $\hat{T}$  qu'avec une probabilité de succès nulle [Menzies09] ! Cette conclusion se retrouve d'ailleurs dans la référence [Fiurášek04], puisqu'un clonage parfait n'est lui aussi possible qu'avec une probabilité de succès nulle. Toutefois, cela ne rend pas le NLA inutile pour autant : la probabilité de succès n'est nulle que si l'on veut pouvoir amplifier *n'importe quel état*, avec une amplitude arbitrairement grande. Si l'on se restreint à des états correctement

décrits par un espace de Hilbert de dimension finie limité à  $N$  photons, les conditions sont moins drastiques.

Dans ces conditions, il est possible de modéliser le NLA par un POVM agissant dans un espace de Hilbert tronqué à  $N$  photons, d'éléments  $\{\hat{M}_{\text{suc}}^{N\dagger}\hat{M}_{\text{suc}}^N, \hat{M}_{\text{fail}}^{N\dagger}\hat{M}_{\text{fail}}^N\}$ . L'opérateur  $\hat{M}_{\text{suc}}^N$  est associé à une amplification réussie, et est défini par :

$$\hat{M}_{\text{suc}}^N = \frac{1}{g^N} \sum_{n=0}^N g^n |n\rangle\langle n| \quad (8.11)$$

Une amplification “ratée” est alors modélisée par l'opérateur  $\hat{M}_{\text{fail}}^N$ , défini de telle sorte que  $\hat{M}_{\text{suc}}^{N\dagger}\hat{M}_{\text{suc}}^N + \hat{M}_{\text{fail}}^{N\dagger}\hat{M}_{\text{fail}}^N = \mathbb{I}$ .

Si l'espace n'est pas limité à  $N$  photons, on peut définir un opérateur similaire sans troncature qui déforme moins l'état que (8.11) [Leverrier12] :

$$\hat{E}_{\text{suc}}^N = \frac{1}{g^N} \sum_{n=0}^N g^n |n\rangle\langle n| + \sum_{k=N+1}^{\infty} |k\rangle\langle k| \quad (8.12)$$

L'opérateur  $\hat{E}_{\text{fail}}^N$  associé à une mauvaise amplification est ensuite défini de manière à vérifier  $\hat{E}_{\text{suc}}^{N\dagger}\hat{E}_{\text{suc}}^N + \hat{E}_{\text{fail}}^{N\dagger}\hat{E}_{\text{fail}}^N = \mathbb{I}$ .

Lorsque  $N$  est fini mais suffisamment grand pour que la troncature introduite sur un état décrit par des variables continues soit négligeable, on peut raisonnablement supposer que l'amplification correspondant à l'un de ces deux POVM produit la transformation (8.8). Nous considérerons donc, comme très souvent dans la littérature, que le NLA est parfait et décrit par  $\hat{T}$ , avec toutefois une probabilité de succès non nulle. Nous montrerons au chapitre suivant que dans ce cas, plusieurs arguments permettent de borner supérieurement la probabilité de succès par  $1/g^2$ .

Enfin, puisque le NLA est une opération principalement destinée aux états décrits par des variables continues, on peut se demander comment en réaliser une implémentation avec une description elle même continue. Nous montrons dans l'annexe E que l'on peut définir une version approchée de  $\hat{M}_{\text{suc}}^N$  en utilisant une décomposition sur les opérateurs de déplacement :

$$\hat{M}_{\text{suc}}^{N, \alpha_m} = \frac{1}{g^N} \sum_{n=0}^N \left[ \sum_{k=0}^N g^k \int_0^{\alpha_m^2} dz e^{-z} L_k(z) L_n(z) \right] |n\rangle\langle n| \quad (8.13a)$$

$$= \frac{1}{\pi g^N} \sum_{k=0}^N g^k \int_{|\alpha| \leq \alpha_m} d^2\alpha e^{-\frac{1}{2}|\alpha|^2} L_k(|\alpha|^2) \hat{D}^\dagger(\alpha) \quad (8.13b)$$

Cette méthode semble toutefois loin d'être envisageable expérimentalement.

### 8.3.2 Transformation de quelques états gaussiens

Vérifions pour commencer que  $\hat{T}$ , défini par (8.10), effectue bien l'amplification recherchée sur un état cohérent :

$$\hat{T}|\alpha\rangle = e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} g^n |n\rangle = e^{\frac{|\alpha|^2}{2}(g^2-1)} |g\alpha\rangle \quad (8.14)$$

Il s'agit bien d'un état cohérent d'amplitude plus grande, avec un facteur  $e^{\frac{|\alpha|^2}{2}(g^2-1)}$  supérieur à 1 (puisque  $g \geq 1$ ). Notons que  $\hat{T}$  transforme toujours un état cohérent en un autre état cohérent physique, sans restriction sur  $g$ .

### Etat gaussien quelconque

Le NLA est une transformation gaussienne : un état gaussien est toujours transformé en un état gaussien, si la transformation reste physique. En utilisant la fonction  $P$ , l'amplification d'un état gaussien  $\hat{\rho}$  quelconque, s'écrit :

$$\hat{T}\hat{\rho}\hat{T} = \int d^2\alpha e^{|\alpha|^2(g^2-1)} P(\alpha) |g\alpha\rangle\langle g\alpha| \quad (8.15)$$

Le changement de variable  $u=g\alpha=u_x+iu_y$  donne ensuite  $d^2\alpha=d^2u/g^2$ , et

$$\boxed{\hat{T}\hat{\rho}\hat{T} = \frac{1}{g^2} \int d^2u P(u/g) e^{\frac{g^2-1}{g^2}|u|^2} |u\rangle\langle u|} \quad (8.16)$$

La transformation (8.16) nous permet de dégager plusieurs propriétés générales du NLA : l'amplification est non seulement indépendante de la phase, mais également du signe de  $u$ . Ainsi, un état de moyenne nulle reste de moyenne nulle, et le NLA ne peut pas comprimer un état initial qui ne le serait pas initialement, ou transformer un état comprimé en un état non comprimé. Nous voyons donc qu'un état thermique ou un état comprimé seront transformés en des états de même nature.

Pour un état gaussien quelconque en revanche, la transformation est beaucoup moins intuitive. L'état initial et l'état amplifié doivent chacun pouvoir s'écrire comme des états thermiques comprimés et déplacés puisqu'ils sont gaussiens [Ferraro05], mais chacun des nouveaux paramètres de l'état amplifié (compression, déplacement,...) dépend très souvent de tous les paramètres de l'état initial d'une façon non triviale.

### Etat EPR

Un état EPR de paramètre  $\lambda$  est transformé en un autre état EPR de paramètre  $g\lambda$  lorsque le NLA est utilisé sur un des modes :

$$(\mathbb{I} \otimes \hat{T})|\lambda\rangle = (\hat{T} \otimes \mathbb{I})|\lambda\rangle = \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} \lambda^n g^n |n, n\rangle = \sqrt{\frac{1-\lambda^2}{1-g^2\lambda^2}} |g\lambda\rangle \quad (8.17)$$

Afin que  $|g\lambda\rangle$  ne diverge pas, il est nécessaire que  $g\lambda < 1$ . Dans le cas contraire, le fait que  $\hat{T}$  soit non borné se manifeste et rend la transformation impossible.

### Etat thermique

Un état thermique  $\hat{\rho}_{\text{th}}(\lambda)$  est transformé en un autre état thermique de paramètre  $g\lambda$  :

$$\hat{T}\hat{\rho}_{\text{th}}(\lambda)\hat{T} = (1-\lambda^2) \sum_{n=0}^{\infty} g^{2n} \lambda^{2n} |n\rangle\langle n| = \frac{1-\lambda^2}{1-g^2\lambda^2} \hat{\rho}_{\text{th}}(g\lambda) \quad (8.18)$$

Là encore, il faut que  $g\lambda < 1$  pour que l'état ne diverge pas. Nous voyons ici que la variance des quadratures, donc le bruit, est augmentée pour un état thermique. Une amplification "sans bruit" ne sera possible que si l'état initial est limité au bruit de photon.

### Etat comprimé

Un état comprimé  $|r\rangle$  est transformé en un autre état comprimé de paramètre  $r'$ , avec  $\tanh r' = g^2 \tanh r$  [Gagatsos12] :

$$\hat{T}|r\rangle = \frac{1}{\sqrt{\cosh r}} \sum_{n=0}^{\infty} \frac{1}{n!} \sqrt{(2n)!} \left(-\frac{1}{2} \tanh r\right)^n g^{2n} |2n\rangle = \sqrt{\frac{\cosh r'}{\cosh r}} |r'\rangle \quad (8.19)$$

Pour que l'état amplifié reste physique, il faut que  $g^2 \tanh r < 1$ .

## 8.4 Implémentations théoriques et réalisations expérimentales

Le NLA a fait l'objet de nombreuses recherches dans la communauté scientifique ces dernières années. Plusieurs propositions théoriques permettent d'en implémenter une version approchée, en utilisant des techniques différentes [Ralph08, Fiurášek09, Marek10a, Menzies09]. Elles ont pour la plupart été réalisées expérimentalement par plusieurs groupes, dont le notre à Palaiseau au cours de ma première année de thèse [Ferreyrol10, Ferreyrol11b, Ferreyrol11a, Zavatta11, Xiang10, Usuga10].

Dans cette section nous détaillons ces différentes réalisations dans les grandes lignes. Le lecteur pourra également consulter la référence [Barbieri11] pour une comparaison des différentes expériences.

### 8.4.1 Les ciseaux quantiques

La première implémentation proposée par T. Ralph et A. Lund dans l'article qui a introduit le NLA [Ralph08] est basée sur les "ciseaux quantiques" [Pegg98, Babichev03]. L'idée est de décomposer un état cohérent  $|\alpha\rangle$  quelconque en  $N$  états cohérents  $|\alpha/\sqrt{N}\rangle$  de plus faible amplitude. Une telle transformation (N-splitter) peut être réalisée unitairement [Fiurášek01]. Chaque état  $|\alpha/\sqrt{N}\rangle$  subit ensuite la transformation décrite par la figure 8.4, constituant *un étage* de l'amplificateur, qui le transforme approximativement en un état amplifié  $|g\alpha/\sqrt{N}\rangle$ . Puis les états de sortie sont ensuite recombinaison de manière probabiliste dans une seule voie de sortie (Fig. 8.5). Le N-splitter est certes une opération unitaire, donc réversible, et redonnerait le vide sur  $(N-1)$  ports de sortie avec  $N$  états cohérents en entrée. Ici cependant, les états amplifiés ne sont pas rigoureusement cohérents et il faut s'assurer par une mesure que l'on a bien le vide sur  $(N-1)$  ports de sortie pour reconstituer effectivement un état cohérent, d'où l'aspect probabiliste de la recombinaison.

### Un étage

Commençons par étudier le montage décrit par la figure 8.4. Chaque étage est composé de deux lames séparatrices, l'une étant asymétrique (A-BS), de réflectivité  $r$  en amplitude, et l'autre étant symétrique (S-BS), de transmission  $1/\sqrt{2}$  en amplitude. Un état de Fock  $|1\rangle$  est nécessaire comme ressource dans un des modes d'entrée de A-BS. L'autre mode est vide. La partie réfléchie du photon unique va ensuite interférer avec l'état cohérent à amplifier au niveau de S-BS. Le conditionnement est ensuite réussi lorsque le compteur de photon D2 ou D1 détecte un photon, alors que l'autre n'en détecte pas. Dans ce cas, l'état de sortie est proportionnel à

$$e^{-\frac{|\alpha|^2}{2}} \sqrt{\frac{r^2}{2}} \left( |0\rangle \pm \sqrt{\frac{1-r^2}{r^2}} \alpha |1\rangle \right). \quad (8.20)$$



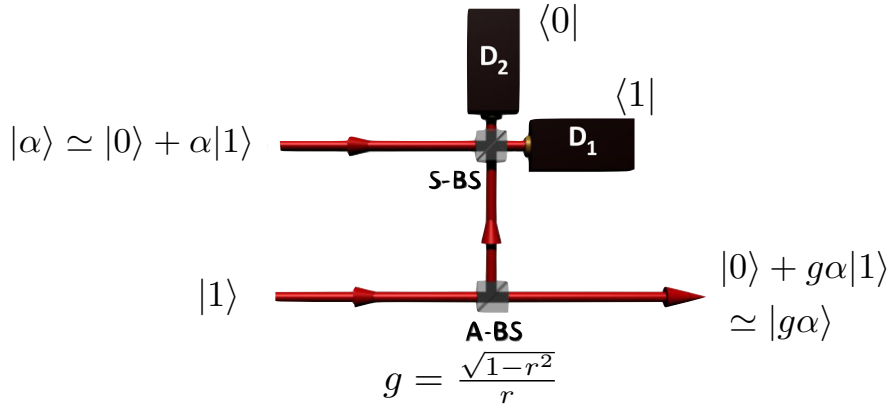


FIGURE 8.4 – Implémentation d’une version approchée de l’amplificateur sans bruit basée sur les “ciseaux quantiques” [Pegg98, Ralph08].

Le signe  $\pm$  dépend du compteur ayant détecté le photon, et on peut sans perte de généralité ne garder que le signe  $+$  en utilisant une transformation de phase si nécessaire. En définissant le gain  $g = \sqrt{\frac{1-r^2}{r^2}}$ , l’état de sortie est donc proportionnel à

$$|0\rangle + g\alpha|1\rangle, \quad (8.21)$$

qui est le développement d’un état cohérent amplifié  $|g\alpha\rangle$  tronqué à 1 photon. Si  $g\alpha \ll 1$ , on peut négliger les termes d’ordres supérieurs et on a bien obtenu l’amplification recherchée  $|\alpha\rangle \rightarrow |g\alpha\rangle$ . L’appellation “ciseaux quantiques”, originellement introduite dans la référence [Pegg98], est due au fait que pour  $r=1/\sqrt{2}$ , et donc  $g=1$ , l’état de sortie est une réplique de l’état d’entrée tronquée à au plus 1 photon, comme si le reste de l’état avait été “coupé”. Notons que le fonctionnement de cet étage est très similaire à une téléportation quantique.

### N étages

L’amplification d’un état  $|\alpha\rangle$  d’amplitude quelconque est possible en le divisant en  $N$  états d’amplitudes suffisamment faibles pour que  $g\alpha/\sqrt{N} \ll 1$ . Les états sont chacun amplifiés à l’aide d’un étage, et sont recombinaés, lorsque toutes les amplifications ont simultanément réussies, en un état

$$e^{-\frac{|\alpha|^2}{2}} r^N \left(1 + \frac{g\alpha}{N} \hat{a}^\dagger\right)^N |0\rangle, \quad (8.22)$$

où  $\hat{a}$  est l’opérateur de destruction du mode de sortie de l’amplificateur. En base de Fock, cette transformation correspond à [Ralph08]

$$|n\rangle \rightarrow r^N \frac{N!}{(N-n)! N^n} g^n |n\rangle. \quad (8.23)$$

Lorsque le nombre d’étages  $N$  est grand,  $\frac{N!}{(N-n)! N^n} \simeq 1$ , et (8.23) tend donc vers  $|n\rangle \rightarrow r^N g^n |n\rangle$ . Le dispositif à  $N$  étages permet d’obtenir une bonne approximation de l’opérateur  $\hat{T}$  lorsque  $N$  augmente, mais au prix d’une probabilité de succès de plus en plus faible à cause du terme  $r^N$ .

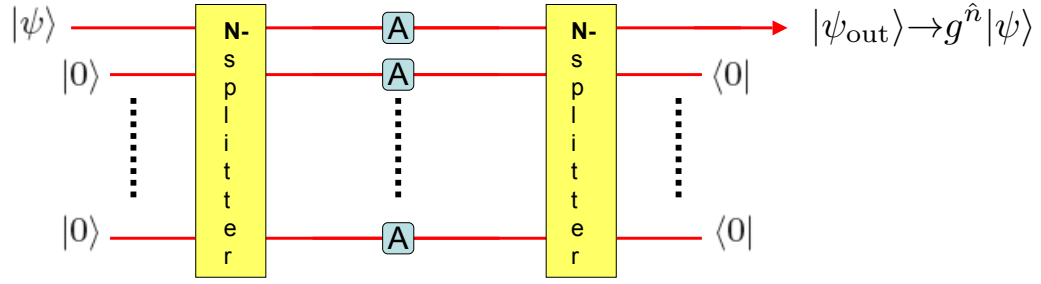


FIGURE 8.5 – Implémentation de l’amplificateur sans bruit basée sur un grand nombre d’étages. Chaque élément A correspond à un étage de la figure 8.4. L’illustration est tirée de la référence [Ralph08].

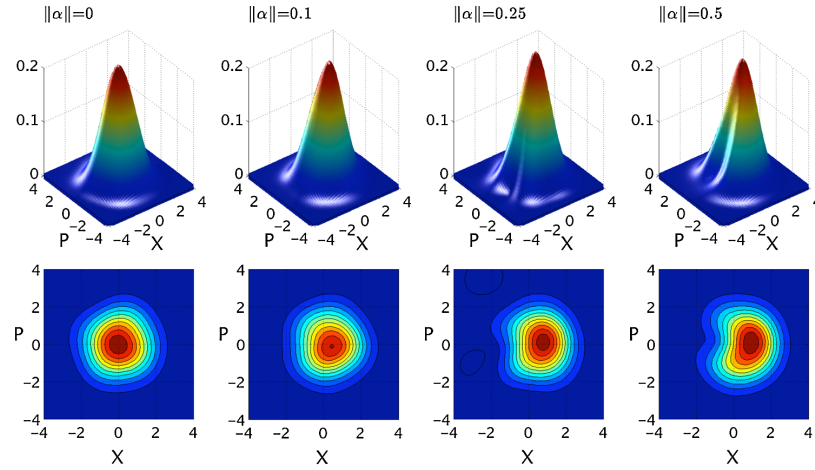


FIGURE 8.6 – Fonction de Wigner d’un état amplifié, obtenue avec notre montage expérimental [Ferreyrol10, Ferreyrol11b].

## Expérience de Palaiseau

L’expérience réalisée dans notre groupe [Ferreyrol10, Ferreyrol11b, Ferreyrol11a] est une réalisation directe d’un étage de l’amplificateur basé sur les ciseaux quantiques. Nous avons amplifié des états cohérents de faible amplitude  $\alpha \leq 0.5$  pour un gain  $g=2$ , puis reconstruits les fonctions de Wigner des états amplifiés par tomographie quantique (Fig. 8.6). L’état amplifié est très proche d’un état gaussien jusqu’à une amplitude d’environ  $\alpha=0.1$ . Des distorsions apparaissent ensuite rapidement pour des valeurs plus grandes dues au caractère non gaussien de l’état produit. Le caractère “sans bruit” de l’amplificateur a également pu être vérifié en comparant le bruit équivalent ramené à l’entrée de l’état amplifié à la quantité minimale (8.5) imposée par un amplificateur déterministe.

Enfin, une modélisation analytique complète développée par F. Ferreyrol a permis de retrouver nos résultats expérimentaux avec un très bon accord [Ferreyrol11b, Ferreyrol11a], et sera certainement utile afin d’étudier l’intégration de ce montage dans un dispositif plus complexe.

## Expérience de Brisbane

L'expérience réalisée à Brisbane dans le groupe de G. Pryde [Xiang10] est également basée sur les ciseaux quantiques. Elle réalise l'amplification d'un mélange statistique d'états cohérents d'amplitude fixée, mais de phases aléatoires, qui peut s'approximer par un mélange statistique de vide et d'un photon unique. L'état amplifié est ensuite caractérisé par interférométrie, en utilisant une APD au lieu d'une détection homodyne. Ce montage permet entre autre une meilleure caractérisation de la linéarité du gain, mais au prix d'une caractérisation limitée de l'état produit.

### 8.4.2 Autres réalisations

#### Expansion polynomiale

Il s'agit d'une méthode introduite par J. Fiurášek [Fiurášek09], différente de celle basée sur les ciseaux quantiques. L'idée est d'approximer  $\hat{T}$  par une série tronquée à un certain ordre  $N$  :

$$\hat{T}=e^{\hat{n} \ln g} \simeq \sum_{k=0}^N \frac{(\ln g)^k}{k!} \hat{n}^k \quad (8.24)$$

Cette décomposition polynomiale en  $\hat{n}=\hat{a}^\dagger \hat{a}$  peut être implémentée à l'aide de soustractions et d'additions de photons. Des considérations sur le gain effectivement obtenu amènent à utiliser une décomposition légèrement différente, qui donne pour  $N=1$ ,  $\hat{T} \simeq (g-1)\hat{n}+1$ . Pour un gain  $g=2$ , cette superposition se réduit alors simplement à  $\hat{a}\hat{a}^\dagger$ . Une mise en oeuvre expérimentale a été réalisée dans le groupe de M. Bellini [Zavatta11], qui avait déjà réalisé des combinaisons de soustractions et d'additions de photons [Parigi07, Zavatta09].

Comparée aux ciseaux quantiques, cette méthode possède l'avantage de ne pas tronquer l'état amplifié, estompant un peu les effets de la non gaussianité introduite par une petite valeur de  $N$ . Il en résulte une plus grande zone de fonctionnement pour un gain donné avant l'apparition des distorsions.

Une technique relativement similaire a été proposée par M. Partanen *et al.* [Partanen12], exploitant le fait qu'un état cohérent est état propre de  $\hat{a}$ , de manière à remplacer  $\hat{a}\hat{a}^\dagger$  par  $\hat{a}\hat{a}^\dagger \hat{a}$ . En réalisant la première soustraction avec une mesure non destructive [Grangier98, Munro05]  $|1\rangle\langle 1|$  plutôt qu'avec une APD, le photon peut ensuite être réutilisé pour implémenter  $\hat{a}^\dagger$ . De ce fait, il n'y a pas besoin d'un photon unique comme ressource comme c'est le cas pour les autres implémentations, mais cette technique semble quand même limitée : il faut réaliser une mesure non destructive, faire un conditionnement sur l'absence de photon, et de plus la soustraction supplémentaire pourrait biaiser une superposition d'états cohérents d'amplitudes différentes, et ainsi produire une transformation différente de  $\hat{T}$ .

#### Concentration de phase

Les implémentations présentées jusqu'ici sont toutes relativement couteuses en ressources expérimentales. Le protocole proposé par P. Marek et R. Filip [Marek10a], et démontré expérimentalement dans le groupe d'U. Andersen [Usuga10], permet d'utiliser un montage plus simple, en remplaçant l'ajout de photon de [Zavatta11, Fiurášek09] par un ajout de bruit thermique. Le fonctionnement est le suivant [Marek10a] : supposons que l'état initial soit un état cohérent  $\alpha$ . L'ajout de bruit thermique produit alors un état thermique déplacé de  $\alpha$ . Or puisqu'un état thermique correspond à un mélange statistique d'états cohérents de moyenne nulle (cf. (2.168)), cet état thermique déplacé correspond à un mélange statistique d'états cohérents, distribués

autour de  $|\alpha\rangle$  avec une certaine variance. La soustraction de photon, de par son coefficient  $\beta$  lorsqu'elle est appliquée sur un état cohérent  $|\beta\rangle$ , va alors agir comme un filtre en “favorisant” les états de plus grandes amplitudes dans le mélange statistique.

Même si l'état résultant est relativement déformé et ne correspond pas vraiment à un état cohérent amplifié, il possède une phase moyenne identique à celle de  $\alpha$ , mais dont l'incertitude est réduite, d'où le terme “concentration de phase”. Ce type d'amplification permet donc d'améliorer l'estimation de la phase d'un état, et pourrait trouver des applications en métrologie.

Ce dispositif a également été utilisé pour faire une approximation expérimentale de clonage quantique sans nécessiter de référence de phase, en séparant l'état amplifié en deux états avec une lame séparatrice [Müller12].

L'ajout de bruit peut aussi être remplacé par un amplificateur déterministe indépendant de la phase : nous avons vu qu'un tel amplificateur ajoute lui aussi du bruit, mais il amplifie en plus le signal. Ainsi, J. Jeffers a montré que cette modification offre des performances supérieures au protocole de P. Marek et R. Filip en termes de concentration de phase [Jeffers11].

## Amplification basée sur la discrimination d'états quantiques

Si une amplification sans bruit *parfaite* fonctionnant pour un état d'entrée quelconque est impossible à réaliser, cela n'est plus le cas lorsque l'état provient d'un ensemble fini et connu, comme récemment montré dans la référence [Dunjko12]. La connaissance *a priori* permet d'adapter le processus d'amplification afin de pouvoir obtenir une amplification parfaite avec une probabilité non nulle. Le principe est en fait relié à la discrimination d'états quantiques [Barnett09], qui peut être parfaite et avec une probabilité non nulle si l'état à estimer appartient à un ensemble fini et connu d'états linéairement indépendants [Chefles98]. Cette possibilité existe aussi pour le clonage quantique [Duan98].

Ces conditions de fonctionnement assez restrictives sont néanmoins respectées pour certains protocoles, tel que le protocole de cryptographie quantique utilisant quatre états cohérents modulés en phase [Leverrier11], pour lequel cet amplificateur pourrait trouver une utilité.

## 8.5 Applications

L'amplificateur sans bruit est maintenant passé du stade de curiosité théorique à celui de véritable outils pouvant trouver de nombreuses applications en information quantique. Ses propriétés le rendent particulièrement utile pour deux applications : la préparation d'états quantiques, et les communications quantiques.

### 8.5.1 Préparation d'états quantiques

Le NLA pourrait être utile pour amplifier des états difficiles à produire expérimentalement. C'est le cas des états chats de Schrödinger du type  $|\alpha\rangle + |-\alpha\rangle$ , dont la réalisation expérimentale est pour l'instant limitée à de faibles amplitudes [Ourjoumtsev06b, Ourjoumtsev07c]. Le NLA pourrait donc permettre de produire un chat de plus grande amplitude  $|g\alpha\rangle + | -g\alpha\rangle$ . Il pourrait bien sûr être utile pour amplifier des superpositions de plus de deux états cohérents, à condition toutefois que les états aient tous une amplitude de même module, pour ne pas que la superposition devienne biaisée par le facteur  $\exp[(g^2-1)|\alpha|^2/2]$  (cf. équation (8.14)).

Comme nous l'avons vu avec l'équation (8.17), il est possible d'augmenter l'intrication d'un état EPR [Ralph08, Yang12]. Cette application est également intéressante d'un point de vue expérimental car de fortes compressions nécessitent de fortes non linéarités qui sont difficiles

à obtenir. Toujours dans l'augmentation de la compression d'un état, le NLA peut également servir de squeezer monomode indépendant de la phase de l'état à amplifier [Gagatsos12]. Un état avec une compression selon une direction quelconque s'en trouvera davantage comprimé, alors qu'un squeezer traditionnel ne peut comprimer que selon une direction donnée.

Enfin, l'approximation du NLA à un étage pourrait être utilisée pour un test des inégalités de Bell sans échappatoire [Brask12], ainsi que pour transférer l'état d'un qubit encodé sur états chats pairs et impairs à un qubit encodé en polarisation [Ferreyrol11a].

### 8.5.2 Communications quantiques

L'autre domaine d'application du NLA concerne les communications quantiques, où il permet notamment de diminuer les pertes effectives subies par un état EPR amplifié [Ralph08] : un état EPR, dont un des modes est envoyé dans un canal de transmission  $T$ , puis amplifié, est équivalent à un état EPR initial davantage comprimé, envoyé dans un canal de transmission  $T'$  supérieure à  $T$ , sans utiliser de NLA. En utilisant cette propriété, T. Ralph a proposé un protocole de correction d'erreur pour des états gaussiens [Ralph11]. Nous développerons cette équivalence en détail dans le chapitre 9, en prenant en compte l'effet du bruit introduit par le canal, et en montrant une équivalence avec un système effectif simplifiant les calculs.

Combiné avec un *atténuateur* sans bruit, réalisant une amplification sans bruit avec un gain inférieur à 1, M. Mičuda *et al.* ont montré qu'il est possible de transformer un canal introduisant des pertes en un canal d'une transmission aussi proche de 1 que l'on veut [Mičuda12], au prix bien sûr d'une probabilité de succès de plus en plus faible. Nous reviendrons également sur ce résultat au cours du chapitre 9.

Enfin, la cryptographie quantique est certainement le premier domaine où l'amplificateur sans bruit est utilisé en tant qu'outil permettant une amélioration des performances [Blandino12c, Fiurášek12, Walk13]. Avec des protocoles à variables continues, nous verrons dans le chapitre 10 qu'il permet d'augmenter la distance maximale de transmission d'une clé secrète, ainsi que la résistance au bruit.

La version à un étage du NLA trouve également une utilité en cryptographie quantique avec des variables discrètes, où là encore il permet d'effectuer une compensation des pertes subies par l'état distribué dans le canal quantique [Gisin10, Kocsis13, Osorio12].

## 8.6 Conclusion

Dans ce chapitre, nous avons présenté les principales propriétés d'un amplificateur sans bruit non déterministe. Nous avons vu que cet amplificateur permet de réaliser des transformations impossibles de manière déterministe, qui ouvrent la voie à de nombreuses applications en information quantique.

## Chapitre 9

# Propriétés de l'amplificateur sans bruit non déterministe

### Sommaire

---

|            |                                                                             |            |
|------------|-----------------------------------------------------------------------------|------------|
| <b>9.1</b> | <b>Introduction</b>                                                         | <b>167</b> |
| <b>9.2</b> | <b>Bornes supérieures pour la probabilité de succès</b>                     | <b>168</b> |
| 9.2.1      | NLA parfait et probabilité de succès non nulle                              | 168        |
| 9.2.2      | Bornes obtenues par non diminution de la fidélité                           | 169        |
| 9.2.3      | Borne de Vidal en dimension finie                                           | 171        |
| <b>9.3</b> | <b>Amplification après un canal quantique : système effectif équivalent</b> | <b>173</b> |
| 9.3.1      | Amplificateur sans bruit après un canal quantique                           | 174        |
| 9.3.2      | Amplificateur sans bruit en amont d'un canal quantique effectif             | 177        |
| 9.3.3      | Paramètres effectifs                                                        | 178        |
| 9.3.4      | Trois types de canaux effectifs                                             | 179        |
| <b>9.4</b> | <b>Application aux protocoles de cryptographie quantique</b>                | <b>180</b> |
| 9.4.1      | Simplification du système effectif pour $\eta \leq 1$                       | 181        |
| 9.4.2      | Allure des paramètres effectifs                                             | 181        |
| 9.4.3      | Sens physique des paramètres effectifs                                      | 185        |
| 9.4.4      | Comportements limites des paramètres effectifs                              | 187        |
| 9.4.5      | Vérification numérique                                                      | 188        |
| 9.4.6      | Prise en compte d'une troncature                                            | 189        |
| <b>9.5</b> | <b>Application aux communications quantiques</b>                            | <b>190</b> |
| 9.5.1      | Suppression des pertes et atténuateur sans bruit                            | 190        |
| 9.5.2      | "Concentration de phase"                                                    | 192        |
| <b>9.6</b> | <b>Conclusion</b>                                                           | <b>193</b> |

---

### 9.1 Introduction

L'amplificateur sans bruit est sans conteste un outil prometteur afin d'améliorer les performances des communications quantiques. De nouvelles possibilités d'applications sont régulièrement découvertes, en considérant aussi bien une implémentation réaliste avec quelques étages ainsi que les imperfections expérimentales associées, ou un amplificateur parfait décrit par  $g^{\hat{n}}$ .

Dans le deuxième cas, au moins deux problématiques se posent : quelle probabilité de succès attribuer à l'amplification, et comment calculer simplement l'état amplifié.

Ce chapitre est consacré à l'étude de ces deux questions, en vue d'une application en cryptographie quantique. Nous verrons que même si un NLA parfait doit en théorie avoir une probabilité de succès nulle, plusieurs arguments permettent quand même de considérer une implémentation "optimiste", décrite par un NLA parfait, et avec une probabilité de succès égale à  $1/g^2$  lorsque l'état à amplifier est un état thermique. Cette probabilité nous sera utile au chapitre 10 afin de borner les bénéfices apportés par le NLA en termes de taux secrets.

Le NLA possède sans doute un fort potentiel pour compenser les pertes subies dans un canal quantique. Nous avons présenté plusieurs applications basées sur cette propriété au chapitre précédent. Cependant, bien que certains états gaussiens se transforment simplement, l'amplification d'un état quelconque en sortie d'un canal ajoutant du bruit n'est pas triviale à calculer. Dans cette optique, nous montrerons un résultat qui simplifie considérablement cette tâche : la combinaison d'un canal quantique suivi d'un NLA est équivalente à un NLA de gain différent, placé en amont d'un canal ayant une transmission plus grande et ajoutant davantage de bruit. En n'étant pas limités aux états gaussiens, cette équivalence est donc potentiellement utilisable avec n'importe quel protocole d'information quantique.

Nous utiliserons ensuite ces résultats dans le chapitre 10, afin de montrer l'intérêt d'un NLA en cryptographie quantique avec des variables continues.

## 9.2 Bornes supérieures pour la probabilité de succès

### 9.2.1 NLA parfait et probabilité de succès non nulle

La question de la probabilité de succès d'un NLA parfait est un sujet délicat. Strictement parlant, nous avons déjà vu au chapitre précédent que la transformation parfaite  $g^{\hat{n}}$  ne peut avoir qu'une probabilité de succès nulle. Un tel amplificateur est non physique car il doit fonctionner pour un état d'entrée quelconque, quelle que soit son amplitude, ce qui conduit à des énergies infinies et des divergences. Si les protocoles de QKD à variables continues tels que GG02 utilisent une modulation qui, *en théorie*, possède un support non borné, toute réalisation expérimentale introduit quand même nécessairement une coupure dans la modulation. Pour un NLA tronqué à un ordre suffisamment élevé, correspondant par exemple à (8.11), nous pouvons donc supposer que l'état amplifié reste raisonnablement gaussien. Pour une étude plus orientée vers une réalisation expérimentale, avec une troncature à quelques photons, on pourra si besoin adapter les résultats que nous présenterons dans la section 9.4, et utiliser le théorème d'optimalité gaussienne [García-Patrón06, Navascués06, Pirandola08, Leverrier10a]. Sous ces hypothèses, un NLA "presque parfait" pour une certaine zone de fonctionnement peut donc avoir une probabilité de succès non nulle.

Nous pouvons alors chercher une borne supérieure pour la probabilité de succès, basée uniquement sur des contraintes théoriques, et qui nous permettra de nous placer dans un régime de fonctionnement très optimiste. En général, on peut distinguer deux régimes d'utilisation du NLA pour un protocole quantique : le premier est un régime où il permet au protocole de fonctionner, alors que cela n'est pas le cas sans lui. Cela sera par exemple le cas en cryptographie quantique, où nous montrerons qu'un taux secret négatif (*i.e.* inexploitable, puisqu'il n'y a aucune information secrète) peut être rendu positif en utilisant un NLA. Dans ce cas, la probabilité de succès n'agit que comme un facteur multiplicatif qui influe sur le taux secret, mais pas sur la distance de transmission ou la résistance au bruit.

Le deuxième régime d'utilisation est lorsque le protocole fonctionne sans NLA. Dans ce cas,



le NLA peut apporter une amélioration, mais qui est naturellement pondérée par la probabilité de succès, laquelle dépend de son implémentation expérimentale. Une borne supérieure permet alors d'étudier le fonctionnement le plus optimiste qui soit : si même dans ces conditions le NLA n'apporte pas d'amélioration, on pourra en conclure que toute implémentation plus réaliste ne sera pas plus utile. Ainsi, nous verrons au chapitre 10 qu'un NLA donne toujours un taux secret inférieur au taux obtenu sans NLA, lorsque qu'Eve est restreinte à des attaques individuelles.

Nous présentons maintenant deux bornes obtenues par des considérations différentes : une borne supérieure obtenue par non diminution de la fidélité entre deux états quantiques, et une borne théoriquement atteignable, basée sur la méthode de G. Vidal [Vidal99]. La condition de non diminution de la fidélité va conduire à une borne supérieure constante, égale à  $1/g^2$ . Nous l'obtiendrons analytiquement en considérant la fidélité entre un état thermique et le vide, et numériquement en considérant la fidélité entre deux états cohérents dont l'un tend vers l'autre. La borne de Vidal est une borne supérieure atteignable valable pour des états bimodes pris dans un espace de Hilbert de dimension finie. Nous la comparerons ensuite à la borne obtenue par non diminution de la fidélité pour des états tronqués, et nous verrons que cette dernière est, comme il se doit, toujours supérieure à la borne de Vidal.

### 9.2.2 Bornes obtenues par non diminution de la fidélité

#### Principe

Deux états quelconques  $\hat{\rho}$  et  $\hat{\sigma}$  ne peuvent pas devenir plus discernables sous l'action d'une opération quantique  $\mathcal{T}$  qui préserve la trace [Nielsen00]<sup>1</sup>. Ceci se traduit par une non diminution de la fidélité (définie par (2.86)) :

$$\mathcal{F}(\hat{\rho}, \hat{\sigma}) \leq \mathcal{F}(\mathcal{T}[\hat{\rho}], \mathcal{T}[\hat{\sigma}]) \quad (9.1)$$

L'opérateur  $\hat{T}$  n'est bien sûr pas une opération qui préserve la trace, en revanche on peut en former une en ne faisant pas de post-sélection et en tenant compte des amplifications non réussies. Considérons que le NLA produit un état amplifié  $\hat{T}\hat{\rho}\hat{T} / \text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}$  avec une probabilité de succès  $P(\hat{\rho})$ , où  $\hat{\rho}$  est un état quelconque. On peut sans perte de généralité supposer que l'état produit est le vide lorsque l'amplification ne marche pas (il suffit de bloquer le faisceau lorsque le conditionnement n'est pas bon). Sans post-sélection, le NLA peut donc être représenté par une opération  $\mathcal{T}$  qui préserve la trace :

$$\mathcal{T}[\hat{\rho}] = P(\hat{\rho}) \frac{\hat{T}\hat{\rho}\hat{T}}{\text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}} + [1 - P(\hat{\rho})] |0\rangle\langle 0| \quad (9.2)$$

On montre dans l'annexe F comment relier la définition (9.2) à une opération quantique "physique", afin de s'affranchir en particulier de la non linéarité due au dénominateur.

Il nous faut maintenant un deuxième état quantique  $\hat{\sigma}$  afin de pouvoir utiliser l'inégalité (9.1). Le vide est particulièrement intéressant car il n'est pas transformé par le NLA :  $\mathcal{T}[|0\rangle\langle 0|] = |0\rangle\langle 0|$ . On obtient ainsi :

$$\mathcal{F}(\hat{\rho}, |0\rangle\langle 0|) \leq \mathcal{F}(\mathcal{T}[\hat{\rho}], |0\rangle\langle 0|) \quad (9.3)$$

1. Voir notamment le théorème 9.6 page 414.



Puisque  $\mathcal{F}(\hat{\mathbf{A}}, |0\rangle\langle 0|) = \langle 0|\hat{\mathbf{A}}|0\rangle$ , l'inégalité (9.3) peut s'écrire

$$\langle 0|\hat{\rho}|0\rangle \leq P(\hat{\rho}) \langle 0|\frac{\hat{\mathbf{T}}\hat{\rho}\hat{\mathbf{T}}}{\text{Tr}\{\hat{\mathbf{T}}\hat{\rho}\hat{\mathbf{T}}\}}|0\rangle + 1 - P(\hat{\rho}), \quad (9.4)$$

ce qui implique :

$$P^{\max}(\hat{\rho}) = \frac{1 - \langle 0|\hat{\rho}|0\rangle}{1 - \langle 0|\frac{\hat{\mathbf{T}}\hat{\rho}\hat{\mathbf{T}}}{\text{Tr}\{\hat{\mathbf{T}}\hat{\rho}\hat{\mathbf{T}}\}}|0\rangle} \quad (9.5)$$

Cette borne est une borne théorique *supérieure* : rien ne nous garanti qu'il soit effectivement possible de l'atteindre. En fonction du type d'état initial (état thermique, état EPR, ou état cohérent), nous pourrions trouver différentes valeurs ainsi que différents comportements : il n'y a pas d'incompatibilité entre les résultats, les bornes plus grandes que la plus petite borne obtenue sont simplement plus optimistes.

Notre borne de probabilité de succès est indépendante de la réalisation explicite de  $\mathcal{T}$ . D'une certaine manière, nous maximisons la probabilité de succès sur les différentes réalisations, sans garantir toutefois qu'elles soient toutes effectivement réalisables.

### Borne pour des états thermiques

Dans le protocole GG02, Bob n'a aucune information sur l'état qu'a envoyé Alice. De son point de vue, il reçoit donc un état thermique

$$\hat{\rho}_B(\lambda^*) = [1 - (\lambda^*)^2] \sum_{n=0}^{\infty} (\lambda^*)^{2n} |n\rangle\langle n|, \quad (9.6)$$

où  $\lambda^*$  dépend des paramètres du canal, et de la variance de modulation d'Alice (voir l'annexe H). C'est donc l'état le plus légitime que l'on puisse utiliser pour une utilisation du NLA avec le protocole GG02. En utilisant le fait que  $\langle 0|\hat{\rho}(\lambda^*)|0\rangle = 1 - (\lambda^*)^2$ , la borne supérieure (9.5) est :

$$P_{\text{th}}^{\max} = \frac{1}{g^2} \quad (9.7)$$

Cette borne est indépendante de  $\lambda^*$  et des paramètres du canal quantique. Elle est particulièrement intéressante car nous verrons que c'est la limite à partir de laquelle le NLA pourrait *théoriquement* donner un meilleur taux secret en considérant les attaques individuelles. Puisqu'elle est constante, elle rend également les calculs plus simples. C'est elle qui sera utilisée par la suite, en précisant encore une fois qu'elle est extrêmement optimiste. Remarquons que l'on obtient la même borne en utilisant un état EPR  $|\lambda^*\rangle = \sqrt{1 - (\lambda^*)^2} \sum_{n=0}^{\infty} (\lambda^*)^n |n\rangle|n\rangle$  purifiant  $\hat{\rho}_B(\lambda^*)$ , car  $\langle 0|\hat{\rho}(\lambda^*)|0\rangle = \langle 0|\lambda^*\rangle\langle\lambda^*|0\rangle$ .

### Borne pour des états cohérents

Plutôt que de considérer un état thermique, on peut également calculer la borne (9.5) pour des états cohérents. A défaut de pouvoir l'utiliser directement en cryptographie quantique, on peut toutefois vérifier qu'elle n'est pas incompatible avec (9.7). En utilisant  $|\langle 0|\alpha\rangle|^2 = e^{-|\alpha|^2}$  et  $|\langle 0|g\alpha\rangle|^2 = e^{-|g\alpha|^2}$ , on obtient :

$$P_{\alpha}^{\max} = \frac{1 - e^{-|\alpha|^2}}{1 - e^{-g^2|\alpha|^2}} \quad (9.8)$$

Cette borne tend vers  $1/g^2$  lorsque  $\alpha$  tend vers 0, et vers 1 lorsque  $\alpha$  tend vers l'infini. Elle vaut également 1 lorsque  $g=1$ , et tend vers  $1 - e^{-|\alpha|^2}$  lorsque  $g \rightarrow \infty$ . Elle n'est donc pas incompatible avec  $P_{\text{th}}^{\text{max}}$ , mais simplement plus optimiste pour les grandes valeurs de  $\alpha$  et de  $g$ . En fait, au lieu de prendre  $\hat{\sigma}=|0\rangle\langle 0|$ , on peut calculer numériquement avec Matlab la borne obtenue pour  $\hat{\sigma}=|a\alpha\rangle\langle a\alpha|$ , où  $a$  est un paramètre réel variant entre 0 et 1 (Fig. 9.1). Lorsque  $a=0$  on retrouve la borne (9.8), en revanche lorsque  $a \rightarrow 1$  on obtient à nouveau  $1/g^2$  quelle que soit la valeur de  $\alpha$ , ce qui confirme les résultats obtenus pour les états états thermiques.

### Simulation numérique

- Les états sont générés à l'aide de la Quantum Optics Toolbox, en prenant un espace de Hilbert suffisamment grand ( $N=30$ ).
- On suppose que la probabilité de succès  $P$  est la même pour les deux états d'entrée  $\hat{\rho}=|\alpha\rangle\langle\alpha|$  et  $\hat{\sigma}=|a\alpha\rangle\langle a\alpha|$ . Cette hypothèse est peut être discutable pour une valeur de  $a$  quelconque, en revanche elle est raisonnable pour les deux cas qui nous intéressent,  $a=0$  et  $a \rightarrow 1$ . Les états produits par le NLA sont :

$$\hat{\rho}(P)=P|g\alpha\rangle\langle g\alpha| + (1-P)|0\rangle\langle 0| \text{ et } \hat{\sigma}(P)=P|ag\alpha\rangle\langle ag\alpha| + (1-P)|0\rangle\langle 0| \quad (9.9)$$

- On calcule ensuite la fidélité entre  $\hat{\rho}$  et  $\hat{\sigma}$ , et entre  $\hat{\rho}(P)$  et  $\hat{\sigma}(P)$  :

$$\mathcal{F}_{\alpha}^a = |\langle a\alpha|\alpha\rangle|^2 \text{ et } \mathcal{F}_{\alpha,\text{NLA}}^a(P) = \mathcal{F}(\hat{\rho}(P), \hat{\sigma}(P)) \quad (9.10)$$

- La borne de probabilité  $P_{\alpha,\text{opt}}^{\text{max}}$  est atteinte lorsque  $\mathcal{F}_{\alpha}^a = \mathcal{F}_{\alpha,\text{NLA}}^a(P_{\alpha,\text{opt}}^{\text{max}})$ . Pour trouver cette valeur, on recherche le minimum de la fonction  $L(P) = |\mathcal{F}_{\alpha,\text{NLA}}^a(P) - \mathcal{F}_{\alpha}^a|$ , sous la contrainte  $\mathcal{F}_{\alpha}^a \leq \mathcal{F}_{\alpha,\text{NLA}}^a(P)$ .

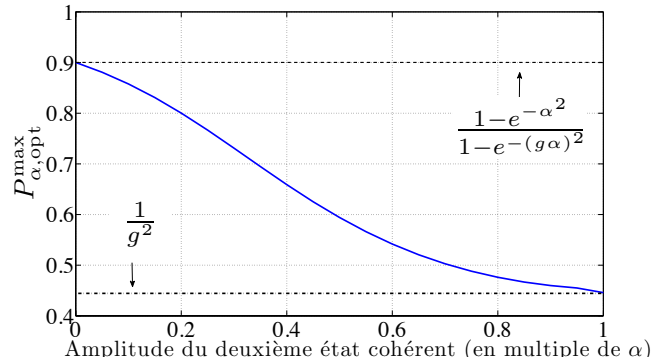


FIGURE 9.1 –  $P^{\text{max}}$  calculée avec la non augmentation de fidélité pour deux états cohérents  $|\alpha\rangle$  et  $|a\alpha\rangle$ , avec  $\alpha=1.5$ ,  $g=1.5$  et  $a$  variant entre 0 et 1.

### 9.2.3 Borne de Vidal en dimension finie

#### Borne de Vidal

G. Vidal a proposé un protocole [Vidal99] permettant de calculer la probabilité maximale avec laquelle on peut transformer un état pur bipartite  $|\phi\rangle$  de dimension finie en un autre état pur bipartite  $|\psi\rangle$  en utilisant seulement des communications classiques et des opérations locales (LOCC, *Local Operations Classical Communications*). Cette borne a la particularité d'être atteignable par un protocole explicitement donné.

Considérons un état EPR parfait  $|\lambda\rangle = \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} \lambda^n |n\rangle |n\rangle$ . Cet état appartient à un espace de Hilbert de dimension infinie, mais nous pouvons définir l'état tronqué dans un espace de dimension  $N$  contenant au plus  $N-1$  photons :

$$|\lambda_N\rangle = \sqrt{\frac{1-\lambda^2}{1-\lambda^{2N}}} \sum_{n=0}^{N-1} \lambda^n |n\rangle |n\rangle := \sum_{n=0}^{N-1} \sqrt{\alpha_n} |n\rangle |n\rangle \quad (9.11)$$

On cherche ensuite la probabilité maximale avec laquelle on peut transformer l'état  $|\lambda_N\rangle$  en un état amplifié

$$|g\lambda_N\rangle = \sqrt{\frac{1-(g\lambda)^2}{1-(g\lambda)^{2N}}} \sum_{n=0}^{N-1} (g\lambda)^n |n\rangle |n\rangle := \sum_{n=0}^{N-1} \sqrt{\beta_n} |n\rangle |n\rangle, \quad (9.12)$$

en utilisant seulement des LOCC. G. Vidal a montré que la probabilité maximale  $P(\phi \rightarrow \psi)$  avec laquelle on peut transformer un état  $|\phi\rangle = \sum_{n=0}^{N-1} \sqrt{\alpha_n} |n\rangle |n\rangle$  en  $|\psi\rangle = \sum_{n=0}^{N-1} \sqrt{\beta_n} |n\rangle |n\rangle$  est égale à :

$$P(\phi \rightarrow \psi) = \min_{l \in [0, N-1]} \frac{\sum_{i=l}^{N-1} \alpha_i}{\sum_{i=l}^{N-1} \beta_i} \quad (9.13)$$

Pour la transformation  $|\lambda_N\rangle \rightarrow |g\lambda_N\rangle$ , cette probabilité vaut :

$$P(\lambda \rightarrow g\lambda) = \frac{1-\lambda^2}{1-(g\lambda)^2} \frac{1-(g\lambda)^{2N}}{1-\lambda^{2N}} \min_{l \in [0, N-1]} \frac{\sum_{i=l}^{N-1} \lambda^{2i}}{\sum_{i=l}^{N-1} (g\lambda)^{2i}} \quad (9.14)$$

On montre facilement que le minimum de  $\frac{\sum_{i=l}^{N-1} \lambda^{2i}}{\sum_{i=l}^{N-1} (g\lambda)^{2i}}$  est atteint pour  $l=N-1$ . En effet :

$$\frac{\sum_{i=l}^{N-1} \lambda^{2i}}{\sum_{i=l}^{N-1} (g\lambda)^{2i}} = \frac{1}{g^{2l}} \frac{\lambda^{2l} \sum_{i=0}^{N-1-l} \lambda^{2i}}{\lambda^{2l} \sum_{i=0}^{N-1-l} (g\lambda)^{2i}} \quad (9.15)$$

$$= \frac{1}{g^{2(N-1)}} \frac{g^{2(N-1-l)} \sum_{i=0}^{N-1-l} \lambda^{2i}}{\sum_{i=0}^{N-1-l} (g\lambda)^{2i}} \quad (9.16)$$

Or  $g \geq 1$ , donc  $g^{2(N-1-l)} \geq g^{2i}$  pour  $i \in [0, N-1-l]$ . Donc  $g^{2(N-1-l)} \sum_{i=0}^{N-1-l} \lambda^{2i} \geq \sum_{i=0}^{N-1-l} (g\lambda)^{2i}$ . Cette inégalité est saturée pour  $l=N-1$ , ce qui correspond donc au minimum recherché.

On a donc finalement,

$$\frac{\sum_{i=l}^{N-1} \lambda^{2i}}{\sum_{i=l}^{N-1} (g\lambda)^{2i}} \geq \frac{1}{g^{2(N-1)}}, \quad (9.17)$$

ce qui donne la valeur maximale de  $P(\lambda \rightarrow g\lambda)$  :

$$\boxed{P(\lambda \rightarrow g\lambda) = \frac{1}{g^{2(N-1)}} \frac{1-\lambda^2}{1-(g\lambda)^2} \frac{1-(g\lambda)^{2N}}{1-\lambda^{2N}}} \quad (9.18)$$

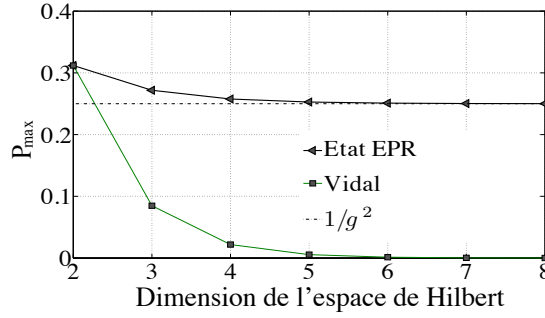


FIGURE 9.2 – Comparaison entre la borne de Vidal et la borne obtenue pour un état EPR tronqué.  $g=2$  et  $\lambda=0.3$ .

### Comparaison avec la borne obtenue pour les états thermiques

On peut comparer la borne de Vidal (9.18) et la borne  $P_N^{\max}(\lambda)$  obtenue par un raisonnement similaire à la section 9.2.2 en utilisant la formule (9.5) pour un EPR tronqué<sup>2</sup> (9.11) :

$$P_N^{\max}(\lambda) = \frac{1}{g^2} \frac{1 - \lambda^{2(N-1)}}{1 - (g\lambda)^{2(N-1)}} \frac{1 - (g\lambda)^{2N}}{1 - \lambda^{2N}} \quad (9.19)$$

Ces deux bornes sont tracées sur la figure 9.2 en fonction de la dimension de l'espace de Hilbert. La borne de Vidal décroît très rapidement et tend vers 0 quand  $N$  augmente, alors  $P_N^{\max}$  tend vers  $1/g^2$ . Comme on pouvait s'y attendre, alors que  $P_N^{\max}$  est toujours supérieure à la borne de Vidal quelle que soit la dimension  $N$ .

En conclusion, plusieurs méthodes nous ont montré que la probabilité de succès du NLA peut être bornée supérieurement par  $1/g^2$ . Bien que cette borne soit extrêmement optimiste, elle présente le très gros avantage d'être constante et indépendante de l'état à amplifier. De ce fait elle n'aura qu'un rôle secondaire dans nos calculs, en étant simplement un facteur multiplicatif. Parmi toutes les bornes qu'il est possible d'obtenir, il n'y a pas d'incompatibilité entre elles. A part la borne de Vidal, toutes sont des bornes supérieures qui ne sont pas forcément atteignables. La plus faible est simplement la plus restrictive, les autres étant trop optimistes. D'ailleurs, il est tout à fait envisageable de pouvoir trouver une borne inférieure à  $1/g^2$ .

## 9.3 Amplification après un canal quantique : système effectif équivalent

Intéressons nous maintenant à l'amplification d'un état après un canal quantique. Nous allons montrer qu'un canal quantique gaussien suivi d'un NLA est équivalent à un NLA de gain différent placé en amont d'un canal quantique effectif ayant lui aussi des paramètres différents (Fig. 9.3). Pour cela, nous calculerons les deux états obtenus avec les deux configurations, pour un état initial quelconque, et nous obtiendrons les valeurs des paramètres effectifs qui permettent de les égaliser. Cette équivalence simplifiera grandement les calculs puisque l'amplification d'un état pur est relativement facile à calculer. Elle l'est d'autant plus lorsque l'état est gaussien,

2. Compte tenu de la remarque à la fin de la section 9.2.2, on obtient le même résultat qu'en considérant l'état thermique tronqué  $\hat{\rho}_N = \frac{1-\lambda^2}{1-\lambda^{2N}} \sum_{n=0}^{N-1} \lambda^n |n\rangle\langle n|$

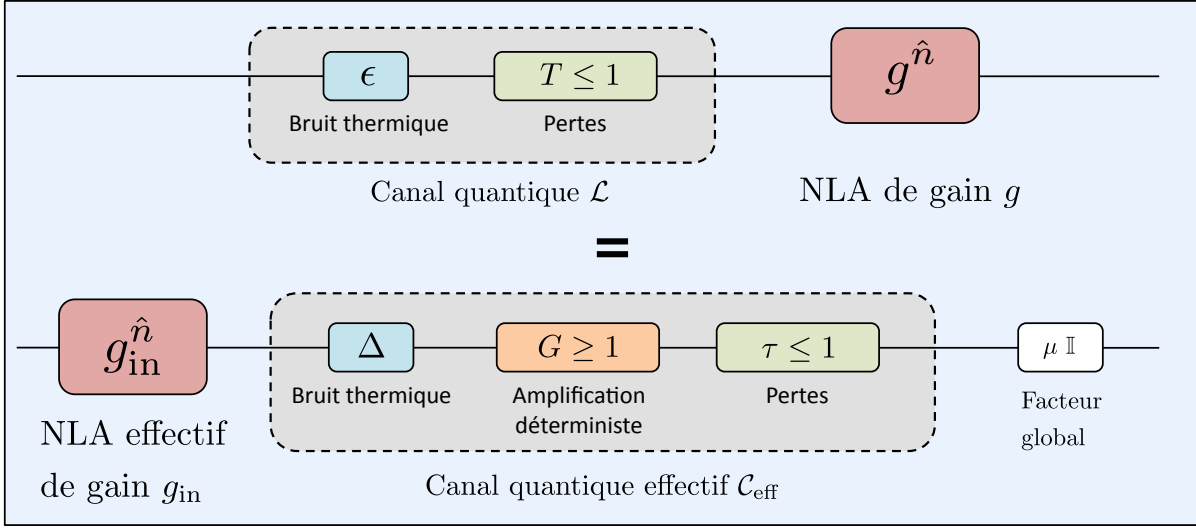


FIGURE 9.3 – *Equivalence qui sera démontrée dans ce chapitre : un canal quantique  $\mathcal{L}$  de transmission  $T$ , et d'excès de bruit  $\epsilon$ , suivi d'un NLA de gain  $g$ , est équivalent à un NLA de gain  $g_{\text{in}}$  placé en amont d'un canal quantique  $\mathcal{C}_{\text{eff}}$  constitué d'un excès de bruit  $\Delta$ , d'un amplificateur déterministe indépendant de la phase de gain  $G$  en intensité, et de pertes  $\tau$ .  $\mu$  est un facteur global constant dépendant uniquement de  $g$ ,  $T$  et de  $\epsilon$ , disparaissant lors de la normalisation.*

comme c'est le cas pour les protocoles de QKD considérés dans ce manuscrit. Nous analyserons ensuite en détail sous quelles conditions ces nouveaux paramètres effectifs prennent des valeurs physiques.

Remarquons que N. Walk *et al.* ont récemment développé un formalisme permettant de calculer l'amplification d'un état gaussien quelconque [Walk12]. Outre le fait que ces résultats n'étaient pas encore établis lors de notre étude, ils sont limités aux états gaussiens, et font appel à un arsenal mathématique relativement complexe.

### 9.3.1 Amplificateur sans bruit après un canal quantique

Commençons par étudier l'action d'un NLA placé après un canal  $\mathcal{L}$  introduisant du bruit thermique et des pertes, tel que schématisé sur la figure 9.3. Nous supposons que ce canal est gaussien, linéaire, et symétrique, de transmission  $T$ , et d'excès de bruit équivalent à l'entrée  $\epsilon$  (que nous appellerons par abus de langage le “bruit ajouté”). Un état initial de variance  $V$  est donc transformé en un état de variance<sup>3</sup>  $T(V+\epsilon)+1-T$  (cf. (8.7)).

Nous associons à ce canal un opérateur  $\mathfrak{L}$ . Puisque l'amplification d'un état cohérent est particulièrement simple, la fonction  $P$  (cf. section 2.3.2) est l'outil naturel à utiliser afin de calculer l'action du NLA en sortie du canal.

#### Action du canal $\mathfrak{L}$

Considérons un état  $\hat{\rho}_{\text{in}}$  quelconque, pas forcément gaussien, envoyé dans le canal. Cet état se décompose selon

$$\hat{\rho}_{\text{in}} = \int d^2\gamma P_{\text{in}}(\gamma) |\gamma\rangle\langle\gamma|. \quad (9.20)$$

3. Sauf mention contraire, nous utiliserons  $N_0=1$ .

En utilisant la linéarité de  $\mathfrak{L}$ , l'état en sortie du canal s'écrit

$$\hat{\rho}_{\text{out}} = \mathfrak{L}[\hat{\rho}_{\text{in}}] = \int d^2\gamma P_{\text{in}}(\gamma) \mathfrak{L}[|\gamma\rangle\langle\gamma|]. \quad (9.21)$$

On applique ensuite le NLA pour produire l'état amplifié  $\hat{\rho}_{\text{out}}^{\text{NLA}}$  (non normalisé) :

$$\hat{\rho}_{\text{out}}^{\text{NLA}} = \hat{T} \hat{\rho}_{\text{out}} \hat{T} \quad (9.22a)$$

$$= \int d^2\gamma P_{\text{in}}(\gamma) \hat{T} \mathfrak{L}[|\gamma\rangle\langle\gamma|] \hat{T} \quad (9.22b)$$

Ainsi, grâce à la linéarité du canal et à celle du NLA, nous n'avons besoin que de connaître la transformation que subit chaque état cohérent  $|\gamma\rangle\langle\gamma|$ . L'évolution due au canal est particulièrement intuitive : en premier lieu, les pertes ont pour effet de transformer l'amplitude  $\gamma$  en  $\sqrt{T}\gamma$ . Ensuite, la variance des quadratures est transformée en  $T(1+\epsilon)+1-T=1+T\epsilon$ . Puisque le canal est gaussien, l'état  $\mathfrak{L}[|\gamma\rangle\langle\gamma|]$  est donc un état thermique  $\hat{\rho}_{\text{th}}(\lambda_{\text{ch}})$  déplacé de  $\sqrt{T}\gamma$  :

$$\mathfrak{L}[|\gamma\rangle\langle\gamma|] = \hat{D}(\sqrt{T}\gamma) \hat{\rho}_{\text{th}}(\lambda_{\text{ch}}) \hat{D}^\dagger(\sqrt{T}\gamma) \quad (9.23)$$

Le paramètre  $\lambda_{\text{ch}}$  est tel que la variance de  $\hat{\rho}_{\text{th}}(\lambda_{\text{ch}})$  soit égale à  $1+T\epsilon$ , ce qui donne

$$\lambda_{\text{ch}}^2 = \frac{T\epsilon}{2+T\epsilon}. \quad (9.24)$$

### Amplification d'un état thermique déplacé

Afin de calculer l'action du NLA, il est utile de recourir à nouveau à une décomposition utilisant la fonction  $P$  pour l'état thermique déplacé  $\mathfrak{L}[|\gamma\rangle\langle\gamma|]$  :

$$\mathfrak{L}[|\gamma\rangle\langle\gamma|] = \int d^2\alpha P_\gamma(\alpha) |\alpha\rangle\langle\alpha| \quad (9.25)$$

Nous avons déjà calculé la fonction  $P$  d'un état thermique déplacé au chapitre 2. Elle est donnée par l'équation (2.170), que l'on rappelle ici :

$$P_\gamma(\alpha) = \frac{1}{\pi\bar{n}} e^{-\frac{1}{\bar{n}}|\alpha - \sqrt{T}\gamma|^2} \quad (9.26)$$

où  $\bar{n}$  est le nombre moyen de photons de l'état thermique, vérifiant  $2\bar{n}+1 = \frac{1+\lambda_{\text{ch}}^2}{1-\lambda_{\text{ch}}^2}$ , ce qui implique que  $\frac{1}{\bar{n}} = \frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}$ . En posant  $\alpha = \alpha_x + i\alpha_y$ , l'expression (9.26) se décompose donc en un produit de deux fonctions,

$$P_\gamma(\alpha) = P_{\gamma_x}(\alpha_x) P_{\gamma_y}(\alpha_y), \quad (9.27)$$

avec

$$P_{\gamma_x}(\alpha_x) = \frac{1}{\sqrt{\pi}} \sqrt{\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}} e^{-\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}(\alpha_x - \sqrt{T}\gamma_x)^2}, \quad (9.28a)$$

$$P_{\gamma_y}(\alpha_y) = \frac{1}{\sqrt{\pi}} \sqrt{\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}} e^{-\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}(\alpha_y - \sqrt{T}\gamma_y)^2}. \quad (9.28b)$$

En utilisant à nouveau la linéarité du NLA, l'amplification d'un état thermique déplacé est donnée par :

$$\hat{T}\mathfrak{L}[|\gamma\rangle\langle\gamma|]\hat{T} = \int d^2\alpha P_\gamma(\alpha)\hat{T}|\alpha\rangle\langle\alpha|\hat{T} \quad (9.29a)$$

$$= \int d^2\alpha P_\gamma(\alpha)e^{(g^2-1)|\alpha|^2}|g\alpha\rangle\langle g\alpha| \quad (9.29b)$$

Le changement de variable  $u=g\alpha=u_x+iu_y$  donne ensuite  $d^2\alpha=d^2u/g^2$ , et

$$\hat{T}\mathfrak{L}[|\gamma\rangle\langle\gamma|]\hat{T} = \int \frac{1}{g^2}d^2u P_\gamma(u/g)e^{\frac{g^2-1}{g^2}|u|^2}|u\rangle\langle u| \quad (9.30)$$

Comme précédemment, on voit sans difficulté que l'on peut séparer les composantes selon  $u_x$  et  $u_y$ . Intéressons nous maintenant à la variable  $u_x$ , le raisonnement étant identique pour  $u_y$ . Afin d'identifier l'état amplifié, on remarque que

$$\frac{1}{g}P_{\gamma_x}(u_x/g)e^{\frac{g^2-1}{g^2}u_x^2} = \frac{1}{g}\frac{1}{\sqrt{\pi}}\sqrt{\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}}e^{-\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}\left(\frac{u_x}{g}-\sqrt{T}\gamma_x\right)^2+\frac{g^2-1}{g^2}u_x^2} \quad (9.31a)$$

$$= \sqrt{\frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}}\frac{1}{\sqrt{\pi}}\sqrt{\frac{1-g^2\lambda_{\text{ch}}^2}{g^2\lambda_{\text{ch}}^2}}e^{-\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}\left(\frac{u_x}{g}-\sqrt{T}\gamma_x\right)^2+\frac{g^2-1}{g^2}u_x^2}. \quad (9.31b)$$

On peut ensuite réécrire le terme de l'exponentielle comme

$$-\frac{1-\lambda_{\text{ch}}^2}{\lambda_{\text{ch}}^2}\left(\frac{u_x}{g}-\sqrt{T}\gamma_x\right)^2 + \frac{g^2-1}{g^2}u_x^2 = \underbrace{-\frac{1-g^2\lambda_{\text{ch}}^2}{g^2\lambda_{\text{ch}}^2}}_{\text{Etat thermique de paramètre } g\lambda_{\text{ch}}} \underbrace{\left(u_x-\sqrt{T}\gamma_x g\frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}\right)^2}_{\text{Gain effectif}} + \underbrace{T\gamma_x^2\frac{(g^2-1)(1-\lambda_{\text{ch}}^2)}{1-g^2\lambda_{\text{ch}}^2}}_{\text{Terme de normalisation indépendant de } u_x}. \quad (9.32)$$

Mis à part le terme de normalisation, nous reconnaissons donc la signature d'un état thermique de paramètre  $g\lambda_{\text{ch}}$  et de variance

$$\frac{1+g^2\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2} = \frac{2+T\epsilon(1+g^2)}{2+T\epsilon(1-g^2)}, \quad (9.33)$$

déplacé de  $g\frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}\sqrt{T}\gamma_x$ . Le NLA amplifie donc l'amplitude moyenne de l'état avec un gain

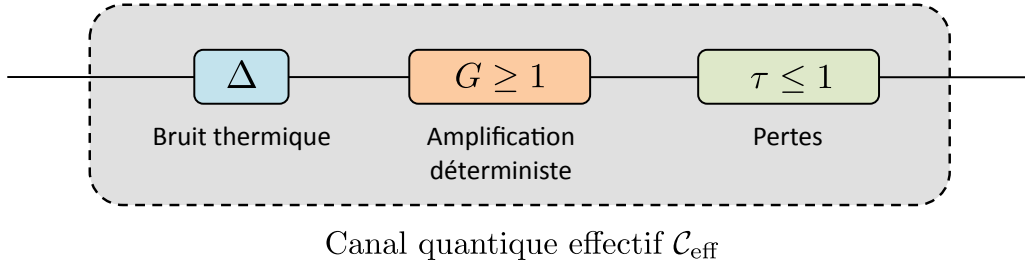
$$\tilde{g} = g\frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2} \quad (9.34)$$

supérieur à  $g$ , puisque  $g\lambda_{\text{ch}}$  doit rester inférieur à 1 pour que l'état amplifié soit physique.

En conclusion, l'amplification (non normalisée) d'un état thermique déplacé est donc donnée par :

$$\boxed{\hat{T}\mathfrak{L}[|\gamma\rangle\langle\gamma|]\hat{T} = \hat{D}\left(\tilde{g}\sqrt{T}\gamma\right)\hat{\rho}_{\text{th}}(g\lambda_{\text{ch}})\hat{D}^\dagger\left(\tilde{g}\sqrt{T}\gamma\right) \times \left(\frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}\right)e^{T|\gamma|^2\frac{(g^2-1)(1-\lambda_{\text{ch}}^2)}{1-g^2\lambda_{\text{ch}}^2}}} \quad (9.35)$$

Remarquons que comme mentionné précédemment, les nouveaux paramètres de l'état amplifié (amplitude moyenne et variance) dépendent de façon peu intuitive des anciens.

FIGURE 9.4 – Le canal quantique effectif  $\mathcal{C}_{\text{eff}}$ .

### Etat amplifié

En injectant (9.35) dans (9.22b), on obtient finalement l'état (non normalisé) amplifié après le canal quantique, pour l'état d'entrée  $\hat{\rho}_{\text{in}}$  :

$$\hat{\rho}_{\text{out}}^{\text{NLA}} = \left( \frac{1 - \lambda_{\text{ch}}^2}{1 - g^2 \lambda_{\text{ch}}^2} \right) \int d^2\gamma P_{\text{in}}(\gamma) \hat{D}(\tilde{g}\sqrt{T}\gamma) \hat{\rho}_{\text{th}}(g\lambda_{\text{ch}}) \hat{D}^\dagger(\tilde{g}\sqrt{T}\gamma) e^{T|\gamma|^2 \frac{(g^2-1)(1-\lambda_{\text{ch}}^2)}{1-g^2\lambda_{\text{ch}}^2}} \quad (9.36)$$

### 9.3.2 Amplificateur sans bruit en amont d'un canal quantique effectif

#### Canal quantique effectif

Dans l'expression (9.36), les états cohérents  $\gamma$  de grandes amplitudes sont favorisés à cause du facteur exponentiel dépendant de leur valeur absolue. Ce comportement est en tout point similaire à l'amplification d'un état gaussien quelconque (8.15), ce qui laisse supposer que l'on pourrait obtenir le terme exponentiel de (9.36) en appliquant un NLA effectif de gain  $g_{\text{in}}$  directement sur l'état initial  $\hat{\rho}_{\text{in}}$ , avec

$$g_{\text{in}}^2 - 1 = T \frac{(g^2 - 1)(1 - \lambda_{\text{ch}}^2)}{1 - g^2 \lambda_{\text{ch}}^2}. \quad (9.37)$$

Ce NLA modifie également l'amplitude des états cohérents dans la décomposition (9.20) : afin de retrouver (9.36), on s'attend donc à ce que le canal quantique qui le succède soit différent du canal  $\mathcal{L}$  de transmission  $T$  et de bruit ajouté  $\epsilon$ .

Raisonnons maintenant plus quantitativement : on cherche donc à obtenir le même état (9.36) en amplifiant d'abord l'état initial avec un NLA effectif de gain  $g_{\text{in}}$ , puis en envoyant l'état amplifié dans un canal effectif  $\mathcal{C}_{\text{eff}}$  décrit par un opérateur  $\mathfrak{L}_g$ . Pour se placer dans le cas le plus général, on suppose que  $\mathcal{C}_{\text{eff}}$  est constitué des trois éléments décrits sur la figure 9.3, que l'on rappelle sur la figure 9.4 :

- Ajout d'un bruit gaussien de moyenne nulle et de variance  $\Delta$  en entrée du canal.
- Amplificateur déterministe indépendant de la phase, de gain  $G \geq 1$  en intensité, sans excès de bruit.
- Canal de transmission  $\tau \leq 1$  en intensité, sans excès de bruit.

Intuitivement, les valeurs moyennes des quadratures de l'état avant le canal  $\mathcal{C}_{\text{eff}}$  sont amplifiées d'un facteur  $\sqrt{\tau G}$ , qui peut être inférieur, égal, ou supérieur à 1. Une variance égale à  $V$  en entrée de ce canal est quand à elle transformée selon :



$$V \xRightarrow[\text{Ajout de bruit thermique}]{\Rightarrow} V+\Delta \xRightarrow[\text{Amplification déterministe}]{\Rightarrow} G(V+\Delta) + G-1 \xRightarrow[\text{Pertes}]{\Rightarrow} V_{\text{out}} = \tau \left[ G(V+\Delta) + G-1 \right] + 1 - \tau \quad (9.38)$$

### Calcul de l'état amplifié

On cherche donc les paramètres  $g_{\text{in}}$ ,  $\Delta$ ,  $G$ , et  $\tau$  donnant l'équivalence

$$\boxed{\hat{T} \mathfrak{L}[\hat{\rho}_{\text{in}}] \hat{T} = \mu \mathfrak{L}_g[\hat{T}_{\text{in}} \hat{\rho}_{\text{in}} \hat{T}_{\text{in}}]}, \quad (9.39)$$

où  $\mu$  est un facteur constant, indépendant de  $\hat{\rho}_{\text{in}}$ , que l'on s'autorise à avoir car il est sans conséquences physiques, et où  $\hat{T}_{\text{in}} = g_{\text{in}}^{\hat{n}}$  est l'opérateur du NLA effectif. Commençons par écrire l'amplification de  $\hat{\rho}_{\text{in}}$  :

$$\hat{T}_{\text{in}} \hat{\rho}_{\text{in}} \hat{T}_{\text{in}} = \int d^2\gamma P_{\text{in}}(\gamma) \hat{T}_{\text{in}} |\gamma\rangle \langle \gamma| \hat{T}_{\text{in}} \quad (9.40a)$$

$$= \int d^2\gamma P_{\text{in}}(\gamma) |g_{\text{in}}\gamma\rangle \langle g_{\text{in}}\gamma| e^{(g_{\text{in}}^2-1)|\gamma|^2} \quad (9.40b)$$

Puisque  $\mathcal{C}_{\text{eff}}$  est un canal gaussien et symétrique, un état cohérent  $|g_{\text{in}}\gamma\rangle$  est simplement transformé en un état de moyenne  $g_{\text{in}}\sqrt{\tau G}\gamma$ , et de variance  $V_{\text{out}} = 1 + \tau G\Delta + 2\tau(G-1)$  : il peut donc également s'écrire comme un état thermique déplacé  $\hat{D}(g_{\text{in}}\sqrt{\tau G}\gamma) \hat{\rho}_{\text{th}}(\lambda_{\text{ch}}^g) \hat{D}^\dagger(g_{\text{in}}\sqrt{\tau G}\gamma)$ , où  $\lambda_{\text{ch}}^g$  est tel que

$$\frac{1 + (\lambda_{\text{ch}}^g)^2}{1 - (\lambda_{\text{ch}}^g)^2} = 1 + \tau G\Delta + 2\tau(G-1) \Rightarrow \lambda_{\text{ch}}^g = \sqrt{\frac{\tau(\Delta G + 2G - 2)}{\Delta\tau G + 2\tau G + 2(1-\tau)}}. \quad (9.41)$$

Après le canal effectif  $\mathcal{C}_{\text{eff}}$ , (9.40b) devient donc :

$$\boxed{\mathfrak{L}_g[\hat{T}_{\text{in}} \hat{\rho}_{\text{in}} \hat{T}_{\text{in}}] = \int d^2\gamma P_{\text{in}}(\gamma) \hat{D}(g_{\text{in}}\sqrt{\tau G}\gamma) \hat{\rho}_{\text{th}}(\lambda_{\text{ch}}^g) \hat{D}^\dagger(g_{\text{in}}\sqrt{\tau G}\gamma) e^{(g_{\text{in}}^2-1)|\gamma|^2}} \quad (9.42)$$

### 9.3.3 Paramètres effectifs

#### Conditions à respecter

Il ne reste plus qu'à trouver les expressions des paramètres permettant de satisfaire l'égalité (9.39). En comparant (9.36) et (9.42), on obtient immédiatement trois équations :

$$g_{\text{in}}\sqrt{\tau G} = \tilde{g}\sqrt{T} \quad (9.43a)$$

$$\lambda_{\text{ch}}^g = g\lambda_{\text{ch}} \quad (9.43b)$$

$$g_{\text{in}}^2 - 1 = T \frac{(g^2 - 1)(1 - \lambda_{\text{ch}}^2)}{1 - g^2 \lambda_{\text{ch}}^2} \quad (9.43c)$$

La résolution de ce système se fait sans difficulté, et nous obtenons ainsi les paramètres recherchés :

$$\begin{aligned} g_{\text{in}} &= \sqrt{\frac{2 + (g^2 - 1)(2 - \epsilon)T}{2 - (g^2 - 1)\epsilon T}} \\ \tau G &= \frac{g^2 T}{1 + (g^2 - 1)T[\frac{1}{4}(g^2 - 1)(\epsilon - 2)\epsilon T - \epsilon + 1]} := \eta \\ \Delta &= \frac{2}{G} + \frac{2 - \epsilon}{2} [(g^2 - 1)T\epsilon - 2] \end{aligned} \quad (9.44)$$

On obtient également

$$\mu = \frac{1 - \lambda_{\text{ch}}^2}{1 - g^2 \lambda_{\text{ch}}^2}, \quad (9.45)$$

qui est bien indépendant de  $\hat{\rho}_{\text{in}}$  et qui disparaît lors de la normalisation finale.

### Dégénérescence du canal effectif

Les conditions (9.43a) et (9.43b) déterminent le canal effectif  $\mathcal{C}_{\text{eff}}$ , en fixant le gain en amplitude et le bruit ajouté *en sortie* du canal. Il y a donc une certaine dégénérescence : plusieurs combinaisons  $(\Delta, G, \eta)$  peuvent permettre de satisfaire ces conditions. Ainsi, c'est pourquoi (9.44) détermine seulement l'expression de  $\tau G := \eta$ , et d'un bruit effectif  $\Delta$  qui dépend de  $G$ . Cette dépendance est en fait tout à fait normale : puisque l'amplificateur déterministe introduit du bruit, l'excès de bruit  $\Delta$  est d'autant plus faible que  $G$  est grand, pour une valeur de  $\eta$  et un bruit total en sortie fixés.

En remarquant que la variance (9.38) peut également se mettre sous la forme

$$V_{\text{out}} = \tau \left[ G(V + \Delta) + (G - 1) \right] + 1 - \tau \quad (9.46a)$$

$$= \tau G \left( V + \Delta + \frac{G - 1}{G} + \frac{1 - \tau}{\tau G} \right), \quad (9.46b)$$

on peut donc définir un bruit ajouté *total* ramené à l'entrée

$$\chi_{\text{tot}} = \Delta + \chi_{\text{ch}}, \quad (9.47)$$

composé du bruit  $\Delta$ , et du bruit dû à l'amplification déterministe et aux pertes

$$\chi_{\text{ch}} = \frac{G - 1}{G} + \frac{1 - \tau}{\tau G} = \frac{\tau(G - 2) + 1}{\tau G}. \quad (9.48)$$

Enfin, insistons sur le fait que le gain  $g_{\text{in}}$  défini par (9.44) ne dépend pas du choix de  $\mathcal{C}_{\text{eff}}$ .

### 9.3.4 Trois types de canaux effectifs

En fonction de la valeur de la transmission globale  $\eta$  (fixée par  $T$ ,  $\epsilon$  et  $g$ ), on peut utiliser la dégénérescence afin de simplifier le canal effectif  $\mathcal{C}_{\text{eff}}$ , en posant  $\tau = 1$  ou  $G = 1$  :

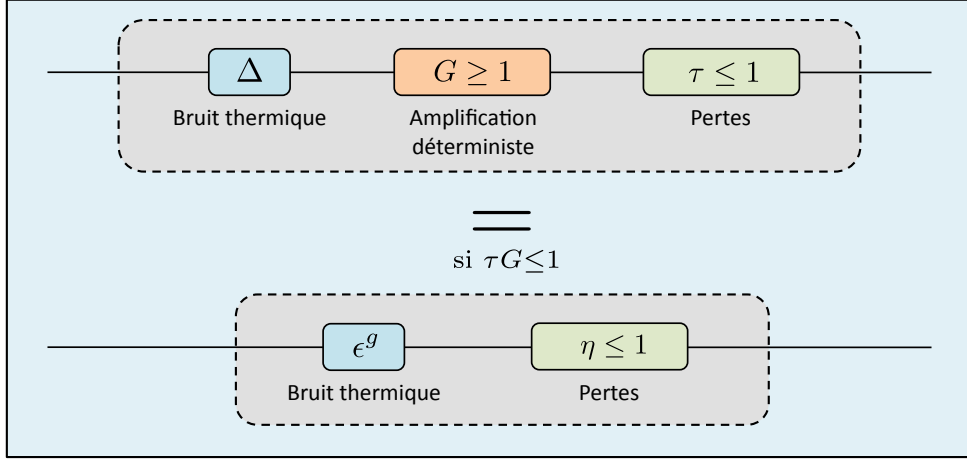


FIGURE 9.5 – Réalisations équivalentes du canal effectif  $\mathcal{C}_{\text{eff}}$  lorsque  $\tau G \leq 1$ .

- si  $\eta \leq 1$  : on peut poser  $G=1$ , et  $\tau=\eta$ . Dans ce cas,

$$\chi_{\text{ch}} = \frac{1-\eta}{\eta}, \quad (9.49)$$

et le canal  $\mathcal{C}_{\text{eff}}$  est donc équivalent à un canal de transmission  $\eta$ , avec un excès de bruit  $\epsilon^g := \Delta_{G=1}$  (Fig. 9.5).

- si  $\eta = 1$  : on peut poser  $G=\tau=1$ . Dans ce cas  $\chi_{\text{ch}}=0$ , et le canal  $\mathcal{C}_{\text{eff}}$  est donc simplement équivalent à un ajout de bruit  $\Delta_{G=1}$ .
- si  $\eta \geq 1$  : on peut poser  $\tau=1$ , et  $G=\eta$ . Dans ce cas,

$$\chi_{\text{ch}} = \frac{\eta-1}{\eta}, \quad (9.50)$$

et le canal  $\mathcal{C}_{\text{eff}}$  est équivalent à un amplificateur déterministe de gain  $\eta$ , avec un excès de bruit  $\Delta$ .

Soulignons le fait que  $\Delta$  n'a pas la même valeur dans ces trois cas, puisqu'il dépend en particulier de  $G$ . Remarquons également que l'on pourrait étendre notre étude au cas où le canal  $\mathcal{L}$  précédant le NLA est constitué d'un amplificateur déterministe à la place de pertes, en suivant exactement la même démarche que celle que nous avons présenté ci-dessus.

Enfin, l'équivalence que nous avons montré est tout à fait générale, et n'a pas fait intervenir de normalisation à aucune étape. Puisque toutes les transformations sont linéaires, nos résultats s'appliquent également pour un état multimode dont un des modes est envoyé dans le canal  $\mathcal{L}$  suivi du NLA.

## 9.4 Application aux protocoles de cryptographie quantique

Pour notre application en cryptographie quantique, nous pouvons nous limiter au cas où  $\eta \leq 1$ . Nous verrons en effet au chapitre suivant qu'un NLA n'apporte une amélioration que pour une certaine plage de gains, pour laquelle  $\eta$  est inférieur à 1. Comme nous l'avons indiqué précédemment, le canal effectif  $\mathcal{C}_{\text{eff}}$  peut alors se mettre sous la forme d'un canal composé d'un ajout de bruit  $\epsilon^g$ , et de pertes  $\eta$ , en posant  $G=1$  (Fig. 9.5).

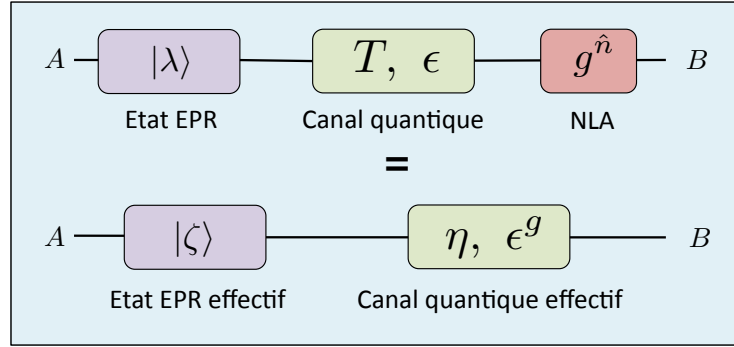


FIGURE 9.6 – Application à la cryptographie quantique à variables continues : système effectif pour la représentation à intrication virtuelle.

#### 9.4.1 Simplification du système effectif pour $\eta \leq 1$

Pour les protocoles de cryptographie quantique gaussiens modélisés par une intrication virtuelle, l'état initial envoyé dans le canal est toujours un état EPR  $|\lambda\rangle$ . Le NLA effectif en amont du canal transforme alors simplement cet état en un autre état EPR  $|\zeta\rangle = |g_{\text{in}}\lambda\rangle$  en vertu de la transformation (8.17). Nous retrouvons finalement les conclusions de [Blandino12c] : l'amplification sans bruit chez Bob d'un mode d'un état EPR  $\lambda$  envoyé dans un canal de transmission  $T$  et de bruit ajouté  $\epsilon$ , produit un état identique à un état EPR effectif

$$|\zeta\rangle = |g_{\text{in}}\lambda\rangle, \quad (9.51)$$

envoyé à dans un canal quantique effectif de transmission  $\eta$  et de bruit ajouté  $\epsilon^g$ , définis par

$$\eta = \frac{g^2 T}{1 + (g^2 - 1) T [\frac{1}{4} (g^2 - 1) (\epsilon - 2) \epsilon T - \epsilon + 1]} \quad (9.52a)$$

$$\epsilon^g = \epsilon + \frac{1}{2} (g^2 - 1) (2 - \epsilon) \epsilon T \quad (9.52b)$$

comme montré sur la figure 9.6.

Les résultats présentés dans ce chapitre présentent l'avantage d'être plus généraux que nos premiers résultats [Blandino12c], que nous avons obtenu avec une méthode différente décrite dans l'annexe H. Ici nous ne faisons pas d'hypothèse sur  $\hat{\rho}_{\text{in}}$  afin de montrer l'équivalence de la figure 9.3. En revanche, nous supposons dans [Blandino12c] que l'état initial est un état EPR  $|\lambda\rangle$ , et nous cherchons un système effectif constitué d'un état EPR effectif  $|\zeta\rangle$  envoyé dans un canal effectif  $(\eta, \epsilon^g)$ . La présence du NLA effectif en amont est donc masquée, puisque directement intégrée dans l'expression de  $\zeta$ , et on ne pouvait pas directement conclure que le canal effectif était bien indépendant de  $\lambda$ , bien que l'on pouvait déjà s'en douter puisque  $\eta$  et  $\epsilon^g$  n'en dépendent pas.

#### 9.4.2 Allure des paramètres effectifs

Les paramètres effectifs sont montrés sur les figures 9.7, 9.8, et 9.9, correspondant respectivement à  $\eta$ ,  $\epsilon^g$  et  $g_{\text{in}}$ .

De manière qualitative, on peut conclure que tous les paramètres effectifs augmentent avec  $T$ ,  $\epsilon$ , et  $g$ .

### Transmission effective $\eta$

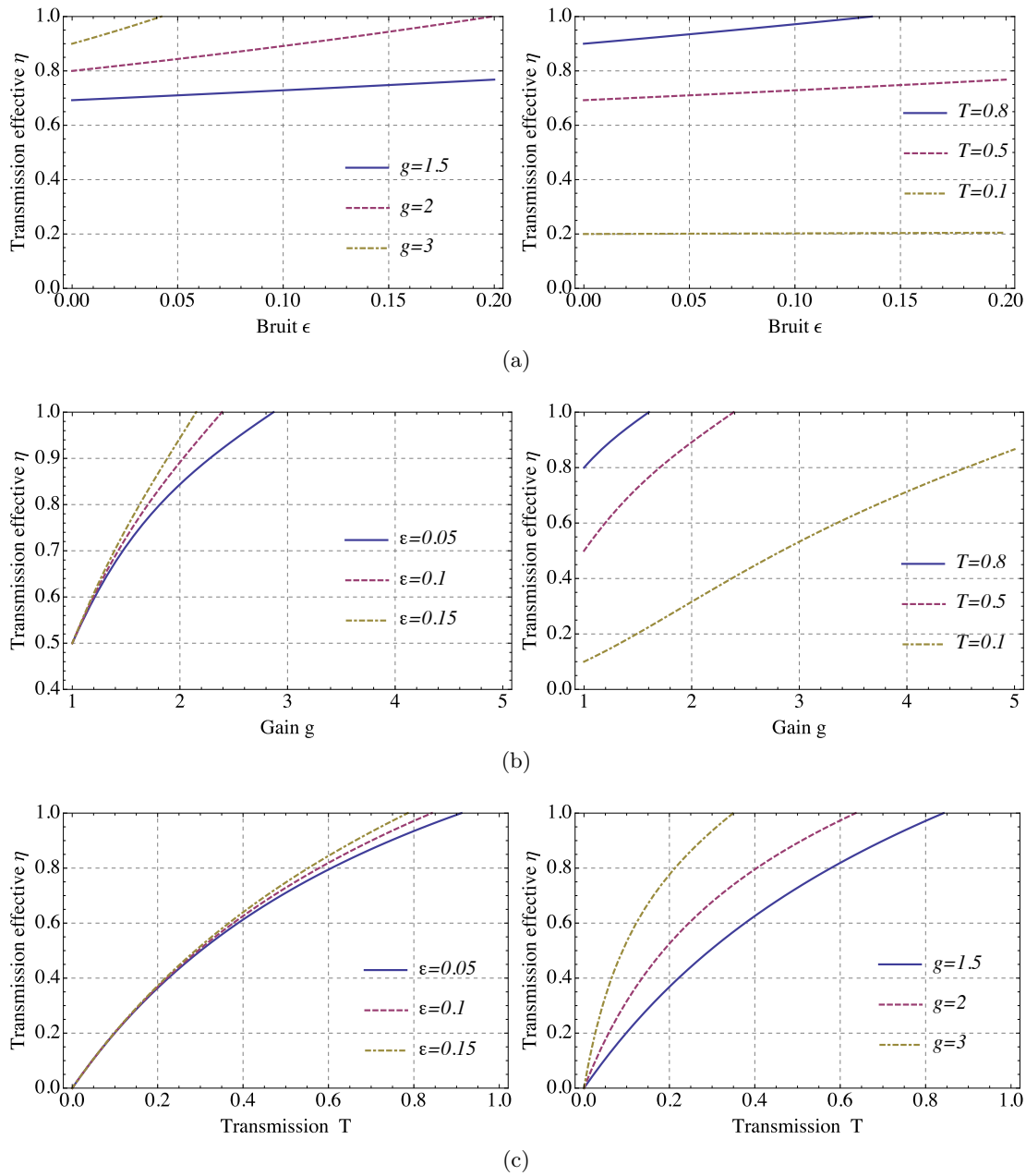


FIGURE 9.7 – Transmission effective  $\eta$  : (a) en fonction du bruit  $\epsilon$ , pour  $T=0.5$  (gauche) et  $g=1.5$  (droite); (b) en fonction du gain  $g$ , pour  $T=0.5$  (gauche) et  $\epsilon=0.1$  (droite); (c) en fonction de la transmission  $T$ , pour  $g=1.5$  (gauche) et  $\epsilon=0.1$  (droite).

La transmission effective  $\eta$  est montrée sur la figure 9.7. On voit qu'elle est relativement peu sensible au bruit  $\epsilon$  pour de faibles gains. Elle est toujours supérieure à  $T$ , et ce d'autant plus que  $g$  est important. D'ailleurs dès que  $\epsilon > 0$ ,  $\eta$  devient égale à 1 pour un gain fini  $g^{\max}$  qui dépend de  $T$  et de  $\epsilon^g$ . Ce gain est d'autant plus petit que  $\epsilon$  est grand : pour un gain  $g$  donné, on peut ainsi diminuer les pertes effectives en ajoutant du bruit à l'entrée du canal. Nous reviendrons sur ce point dans la section 9.5.1.

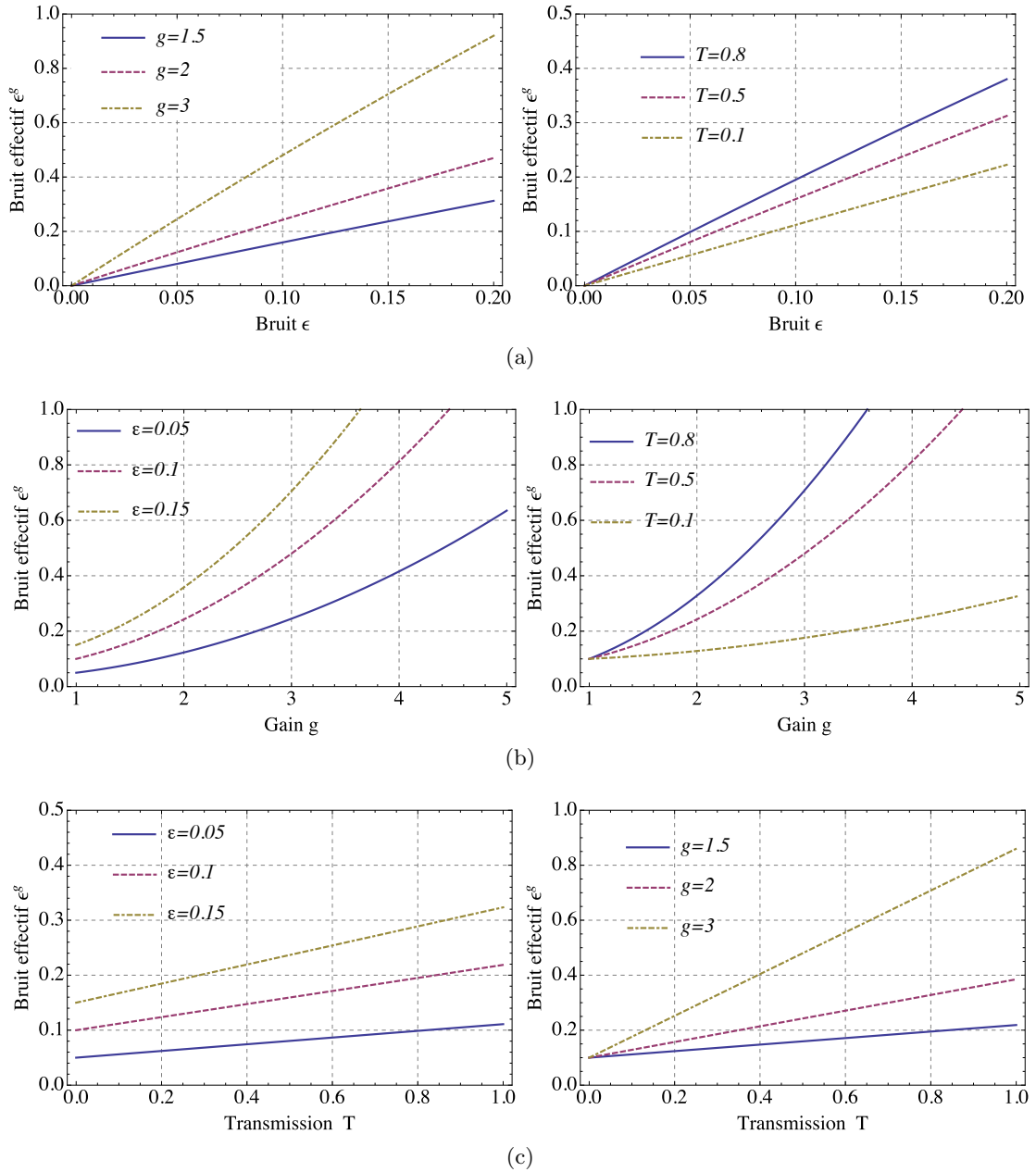
**Bruit effectif  $\epsilon^g$** 

FIGURE 9.8 – Bruit effectif  $\epsilon^g$  : (a) en fonction du bruit  $\epsilon$ , pour  $T=0.5$  (gauche) et  $g=1.5$  (droite); (b) en fonction du gain  $g$ , pour  $T=0.5$  (gauche) et  $\epsilon=0.1$  (droite); (c) en fonction de la transmission  $T$ , pour  $g=1.5$  (gauche) et  $\epsilon=0.1$  (droite).

Le bruit effectif  $\epsilon^g$  est montré sur la figure 9.8. Comme on peut le constater depuis l'expression (9.52), la dépendance est linéaire en  $T$ , et l'est quasiment pour de faibles valeurs de  $\epsilon$ . Il n'y a pas de comportement problématique, retenons simplement que c'est le gain  $g$  qui influe le plus sur l'augmentation de  $\epsilon^g$ , puisque la dépendance est quadratique.

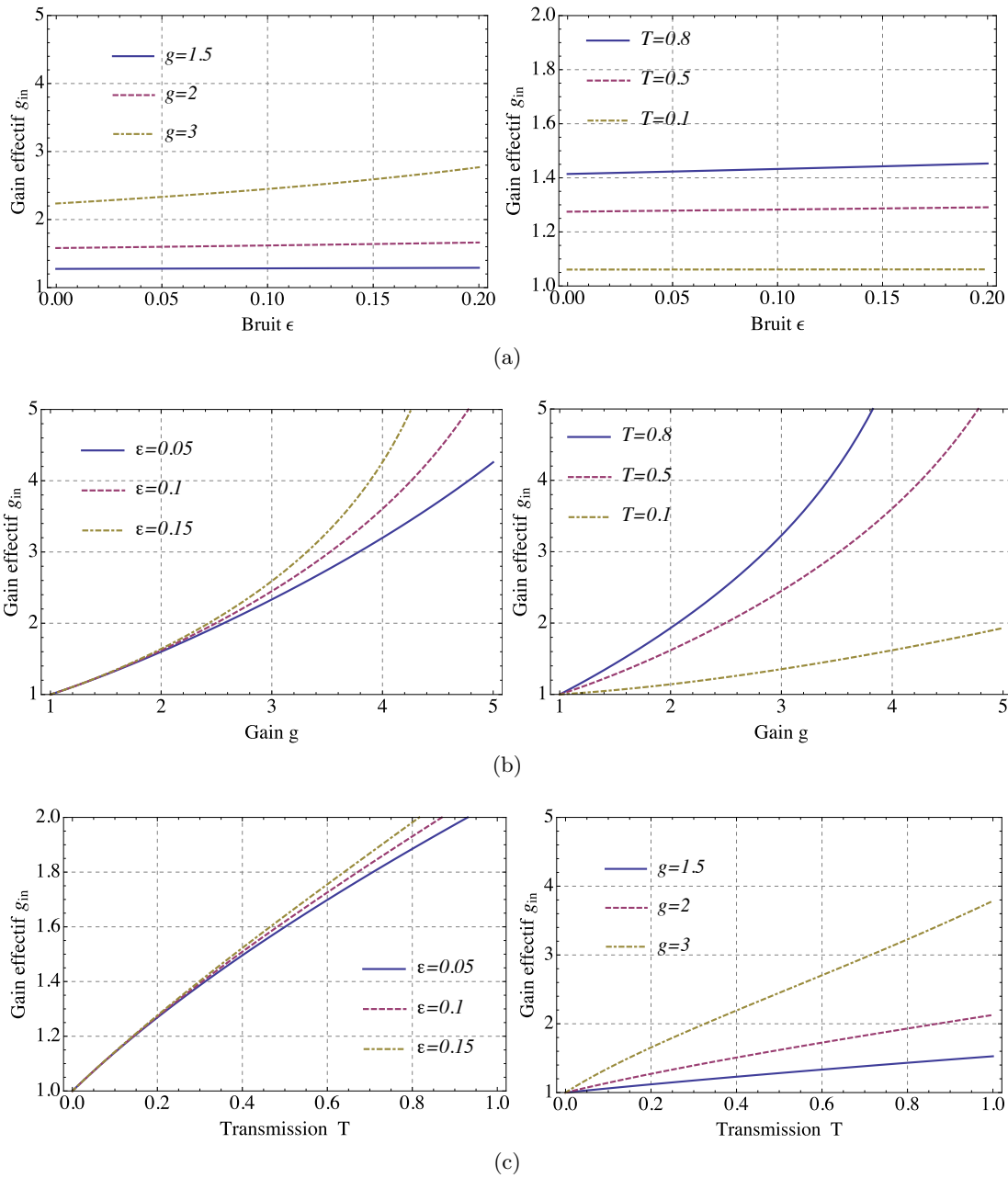
Gain effectif  $g_{in}$ 

FIGURE 9.9 – Gain effectif  $g_{in}$  : (a) en fonction du bruit  $\epsilon$ , pour  $T=0.5$  (gauche) et  $g=1.5$  (droite); (b) en fonction du gain  $g$ , pour  $T=0.5$  (gauche) et  $\epsilon=0.1$  (droite); (c) en fonction de la transmission  $T$ , pour  $g=2$  (gauche) et  $\epsilon=0.1$  (droite).

Le gain effectif est montré sur la figure 9.9. Pour de petits gains  $g$ , il dépend relativement peu du bruit  $\epsilon$ . La dépendance en  $T$  est en revanche beaucoup plus forte. On peut dire qualitativement que  $g_{in}$  est plus petit que  $g$  lorsque le bruit est petit, mais peut devenir supérieur lorsque le bruit  $\epsilon$  augmente.

### 9.4.3 Sens physique des paramètres effectifs

Les paramètres effectifs (9.51) et (9.52) ont été introduits dans le but de pouvoir utiliser un système physique équivalent sans NLA en cryptographie quantique, où le canal effectif est de transmission  $\eta \leq 1$ . Plusieurs conditions sont à respecter afin que cette interprétation soit valable et que le canal effectif puisse avoir une interprétation physique.

#### Contrainte liée au bruit du canal

Les états thermiques déplacés (9.35) ont une variance  $\frac{1+g^2\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}$ . Afin que ces états ne divergent pas, il est donc nécessaire que  $g\lambda_{\text{ch}} < 1$ . Cette condition est vérifiée si

$$g < \sqrt{1 + \frac{2}{T\epsilon}}. \quad (9.53)$$

On obtient également la même condition sur  $g$  afin que  $g_{\text{in}}$  soit réel, *i.e.* que  $(g^2-1)T\epsilon < 2$ , et également afin que  $\eta$  ne diverge pas.

#### Contrainte sur $\lambda$

Pour que l'état EPR effectif  $|\zeta\rangle$  soit physique,  $\zeta$  doit vérifier  $0 \leq \zeta < 1$ . Puisque Alice optimise sa variance de modulation, on peut en fait supposer que  $\lambda$  est toujours choisi de telle sorte que  $\zeta$  ait une valeur physique

$$0 \leq \zeta < 1 \Rightarrow 0 \leq \lambda < \frac{1}{g_{\text{in}}}, \quad (9.54)$$

ce qui est possible tant que (9.53) est vérifiée.

#### Contrainte sur $\eta$ - gain maximal $g^{\text{max}}$

Le paramètre  $\eta$  peut quant à lui être interprété comme une transmission effective si  $0 \leq \eta \leq 1$ . Cette condition est respectée tant que le gain reste inférieur à une valeur limite  $g^{\text{max}}$  que nous avons déjà évoqué :

$$g^{\text{max}}(T, \epsilon) = \sqrt{\frac{\epsilon[T(\epsilon-4)+2] + 4\sqrt{\frac{T(\epsilon-2)+2}{\epsilon}} - 2\sqrt{\epsilon[T(\epsilon-2)+2]} + 4T - 4}{T(\epsilon-2)^2}} \quad (9.55)$$

Remarquablement, la condition (9.53) est toujours vérifiée lorsque  $g \leq g^{\text{max}}$ . En effet, on peut montrer que l'on a toujours la relation

$$g^{\text{max}} < \sqrt{1 + \frac{2}{T\epsilon}}. \quad (9.56)$$

Cette relation est d'ailleurs assez intuitive : lorsque  $g \rightarrow \sqrt{1 + \frac{2}{T\epsilon}}$ ,  $\eta$  diverge vers  $+\infty$ , et est donc nécessairement supérieur à 1.  $g$  est donc supérieur à  $g^{\text{max}}$ , qui est par définition le gain tel que  $\eta=1$ .



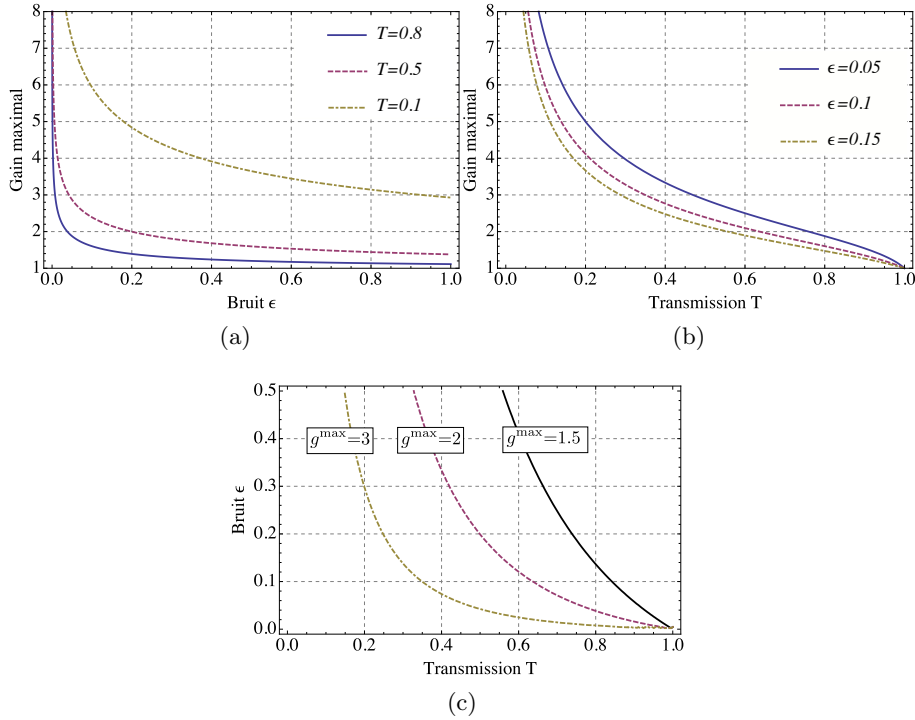


FIGURE 9.10 – Gain maximal  $g^{\max}$  : (a) en fonction du bruit  $\epsilon$  ; (b) en fonction de la transmission  $T$  ; (c) Courbes  $g^{\max}$  ayant une valeur donnée en fonction de la transmission  $T$  et du bruit  $\epsilon$ .

### Contrainte sur $\epsilon^g$ et $\Delta$

$\epsilon^g$  et  $\Delta$  doivent rester positifs pour correspondre à des bruits ajoutés. En ce qui concerne  $\epsilon^g$ , défini pour  $\eta \leq 1$ , on peut montrer que

$$g \leq g^{\max} \Rightarrow \epsilon^g \geq 0, \quad (9.57)$$

l'égalité étant atteinte seulement si  $\epsilon=0$ . Concernant  $\Delta$ , on peut poser  $G=\eta$  lorsque  $\eta \geq 1$ . On montre alors que

$$\Delta > 0 \Rightarrow g < \sqrt{1 + \frac{2}{T\epsilon}}. \quad (9.58)$$

### Intérêt de $g^{\max}$

Ainsi, en nous assurant que le gain du NLA utilisé est inférieur à  $g^{\max}$ , nous sommes certains que le canal effectif correspond bien à un canal introduisant des pertes, que le NLA effectif possède un gain réel, et plus généralement que le canal effectif ne “diverge” pas.

La figure 9.10 montre ce gain maximal en fonction de  $T$  et  $\epsilon$ . Globalement, on peut dire qu'il augmente quand  $T$  ou  $\epsilon$  diminuent. A la limite où  $\epsilon \rightarrow 0$ ,  $g^{\max} \rightarrow +\infty$  : comme nous le verrons,  $\eta$  prend une expression simple dans ce cas et garde toujours une valeur physique (et inférieure à 1) quel que soit  $g$ . Puisque  $\eta$  augmente lorsque  $\epsilon$  augmente (Fig. 9.7), on comprend intuitivement qu'une valeur de 1 sera atteinte pour un gain plus faible, et donc que le gain maximal doit diminuer.

Enfin,  $\lim_{\substack{T \rightarrow 1 \\ \epsilon > 0}} g^{\max} = 1$  : si  $T=1$  et  $\epsilon > 0$ , un gain supérieur à 1 conduit également à une valeur de  $\eta$  supérieure à 1.

### 9.4.4 Comportements limites des paramètres effectifs

#### Sans bruit thermique ( $T \leq 1, \epsilon=0$ )

Plusieurs tests permettent de s'assurer de la validité des paramètres effectifs (9.51) et (9.52). A la limite où il n'y a pas de bruit thermique ( $\epsilon=0$ ) :

$$\epsilon=0 \Rightarrow \begin{cases} \zeta = \lambda \sqrt{1 + (g^2 - 1)T} \\ \eta = \frac{g^2 T}{1 + (g^2 - 1)T} \\ \epsilon^g = 0 \end{cases} \quad (9.59)$$

On retrouve ainsi les expressions précédemment obtenues par la référence [Ralph08]. Remarquons que lorsque  $\epsilon=0$ , le gain effectif  $g_{\text{in}}=\zeta/\lambda$  est toujours inférieur à  $g$ . De même, la transmission effective est toujours supérieure à  $T$ , et inférieure à 1 (Fig. 9.11). C'est seulement lorsque  $g$  tend vers l'infini qu'elle tend vers 1. Lorsque  $\epsilon=0$  et que  $g$  est fini, le canal effectif ne peut donc pas être un canal sans pertes.

La méthode initialement proposée dans [Ralph08] permet de trouver l'expression des paramètres effectifs (9.59) en travaillant en base de Fock à partir de l'état EPR initial  $|\lambda\rangle$ . Après le passage dans le canal quantique, modélisé par une lame séparatrice, cet état est transformé en

$$|\lambda, T\rangle = \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} \sum_{k=0}^n \lambda^n \sqrt{C_n^k} T^{\frac{k}{2}} (1-T)^{\frac{n-k}{2}} |n\rangle_A |k\rangle_B |n-k\rangle_E, \quad (9.60)$$

où A,B et E correspondent respectivement aux modes d'Alice, Bob et Eve. Bob applique ensuite le NLA sur son mode :

$$|\lambda, T, g\rangle = (\mathbb{I} \otimes \hat{T} \otimes \mathbb{I}) |\lambda, T\rangle \quad (9.61a)$$

$$= \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} \sum_{k=0}^n \sqrt{C_n^k} \lambda^n T^{\frac{k}{2}} (1-T)^{\frac{n-k}{2}} g^k |n\rangle_A |k\rangle_B |n-k\rangle_E \quad (9.61b)$$

En posant ensuite  $\zeta = \lambda \sqrt{1 + (g^2 - 1)T}$  et  $\eta = \frac{g^2 T}{1 + (g^2 - 1)T}$ , on voit que

$$\zeta^n \eta^{\frac{k}{2}} (1-\eta)^{\frac{n-k}{2}} = \lambda^n T^{\frac{k}{2}} (1-T)^{\frac{n-k}{2}} g^k. \quad (9.62)$$

On retrouve bien l'équivalence entre un système effectif et le système réel avec le NLA.

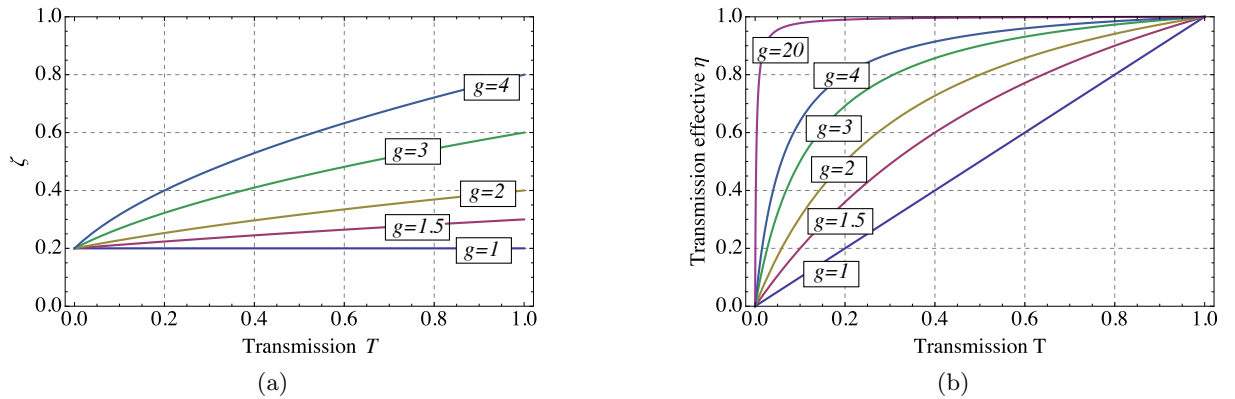


FIGURE 9.11 – Paramètres effectifs pour un canal sans bruit ajouté : (a)  $\zeta$  en fonction de la transmission  $T$ , pour  $\lambda=0.2$ ; (b)  $\eta$  en fonction de la transmission  $T$ .

### Sans NLA ( $g=1$ )

Lorsque  $g=1$ , on retrouve naturellement les paramètres initiaux :

$$g=1 \Rightarrow \begin{cases} \zeta = \lambda \\ \eta = T \\ \epsilon^g = \epsilon \end{cases} \quad (9.63)$$

### 9.4.5 Vérification numérique

La validité des paramètres (9.52) a également été vérifiée numériquement avec des simulations utilisant Matlab et la Quantum Optics Toolbox. Le principe de la simulation est de calculer directement l'état amplifié après le canal quantique  $(T, \epsilon)$ , et de le comparer à l'état obtenu en utilisant un NLA de gain  $g_{\text{in}}$  avant un canal quantique  $(\eta, \epsilon^g)$ .

On peut répéter cette procédure pour différents états initiaux, et à chaque fois l'accord avec le calcul analytique est excellent, ce qui valide l'expression des paramètres effectifs et la démarche globale du calcul.

Les étapes de la simulation sont résumées ci-dessous. On commence par calculer directement l'état amplifié après le canal quantique  $(T, \epsilon)$  :

- On fixe une dimension de l'espace de Hilbert  $N$  suffisamment grande pour négliger les effets dus à la troncature. Pour des états contenant quelques photons,  $N=25$  donne en général de bons résultats.
- On crée un état quantique quelconque  $|\psi_{\text{in}}\rangle$ . Afin de réduire la dimension du problème, il est préférable de travailler avec un vecteur plutôt qu'avec une matrice densité. En pratique, on peut par exemple prendre  $|\psi_{\text{in}}\rangle = \hat{D}(1)|0\rangle$  ou  $|\psi_{\text{in}}\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ .
- On crée un opérateur  $\hat{U}_{\text{BS}}$  associé à la lame séparatrice modélisant les pertes du canal

$$\hat{U}_{\text{BS}}(T) = e^{\text{acos} \sqrt{T}(\hat{a}^\dagger \otimes \hat{a} - \hat{a} \otimes \hat{a}^\dagger)}. \quad (9.64)$$

Le deuxième mode correspond au mode de l'environnement (Eve).

- On crée un état EPR  $|\chi_{\text{eve}}\rangle = \hat{S}_2(r_{\text{eve}})|0\rangle \otimes |0\rangle$  correspondant aux deux modes de l'état EPR utilisé pour modéliser le bruit thermique. Le paramètre  $r_{\text{eve}}$  est tel que

$$\cosh 2r_{\text{eve}} = \frac{T}{1-T} \left( \frac{1-T}{T} + \epsilon \right). \quad (9.65)$$

Un des modes est injecté dans la séparatrice modélisant les pertes du canal afin d'introduire le bruit thermique, l'autre est tracé. Là encore, on aurait pu travailler directement avec un état thermique, mais cela augmente la dimension du problème et ralentit donc les calculs. On peut éventuellement utiliser une dimension différente de  $N$  pour définir  $|\chi_{\text{eve}}\rangle$ .

- On calcule l'état pur tripartite après la séparatrice

$$|\psi_{\text{out}}\rangle = \left( \hat{U}_{\text{BS}}(T) \otimes \mathbb{I} \right) |\psi_{\text{in}}\rangle \otimes |\chi_{\text{eve}}\rangle, \quad (9.66)$$

puis on trace sur les deux modes de l'environnement

$$\hat{\rho}_{\text{out}} = \text{Tr}_{\text{eve}}\{|\psi_{\text{out}}\rangle\langle\psi_{\text{out}}|\}. \quad (9.67)$$

- On applique le NLA de gain  $g$  (naturellement tronqué à  $N-1$  photons puisque l'on travail dans un espace de dimension  $N$ )

$$\hat{\rho}_{\text{out}}^{\text{NLA}} = e^{\hat{a}^\dagger \hat{a} \ln g} \hat{\rho}_{\text{out}} e^{\hat{a}^\dagger \hat{a} \ln g}. \quad (9.68)$$

L'état  $\hat{\rho}_{\text{out}}^{\text{NLA}}$  n'est pas normalisé, mais nous pouvons le garder tel quel.

On répète ensuite cette procédure, mais en appliquant d'abord le NLA effectif avant le canal effectif :

- On part du même état initial  $|\psi_{\text{in}}\rangle$ .
- On applique le NLA effectif de gain  $g_{\text{in}}$  sur l'état initial,

$$|\psi_{\text{NLA}}\rangle = e^{\hat{a}^\dagger \hat{a} \ln g_{\text{in}}} |\psi_{\text{in}}\rangle. \quad (9.69)$$

- On crée un opérateur  $\hat{U}_{\text{BS}}$  modélisant les pertes du canal effectif

$$\hat{U}_{\text{BS}}(\eta) = e^{\text{acos } \sqrt{\eta}(\hat{a}^\dagger \otimes \hat{a} - \hat{a} \otimes \hat{a}^\dagger)}. \quad (9.70)$$

- On crée un état EPR  $|\chi_{\text{eve}}^{\text{NLA}}\rangle = \hat{S}_2(r_{\text{eve}}^{\text{NLA}})|0\rangle \otimes |0\rangle$  correspondant aux deux modes de l'environnement. Le paramètre  $r_{\text{eve}}^{\text{NLA}}$  est tel que

$$\cosh 2r_{\text{eve}}^{\text{NLA}} = \frac{\eta}{1-\eta} \left( \frac{1-\eta}{\eta} + \epsilon^g \right). \quad (9.71)$$

- On calcule l'état pur tripartite après la lame séparatrice

$$|\psi_{\text{out}}^{\text{eff}}\rangle = \left( \hat{U}_{\text{BS}}(\eta) \otimes \mathbb{I} \right) |\psi_{\text{in}}^{\text{NLA}}\rangle \otimes |\chi_{\text{eve}}^{\text{NLA}}\rangle, \quad (9.72)$$

puis on trace sur les deux modes de l'environnement

$$\hat{\rho}_{\text{out}}^{\text{eff}} = \text{Tr}_{\text{eve}} \{ |\psi_{\text{out}}^{\text{eff}}\rangle \langle \psi_{\text{out}}^{\text{eff}}| \}. \quad (9.73)$$

On compare ensuite les deux états obtenus  $\hat{\rho}_{\text{out}}^{\text{NLA}}$  et  $\hat{\rho}_{\text{out}}^{\text{eff}}$ , en tenant compte du coefficient de normalisation  $\mu$  (9.45). Tant que la valeur  $N$  est suffisamment grande, ils sont parfaitement identiques. Nous avons également effectué ce test en utilisant des mélanges statistiques à la place de  $|\psi_{\text{in}}\rangle$ , ou encore des opérateurs du type  $|n\rangle\langle m|$ , avec  $m \neq n$ .

#### 9.4.6 Prise en compte d'une troncature

Dans ce chapitre, ainsi que dans le chapitre 10, nous considérons une version idéale du NLA décrite par  $\hat{T} = g^{\hat{n}}$ . La restriction à un NLA tronqué tel que  $\hat{M}_{\text{suc}}^N$  défini par (8.11) peut en fait se faire assez simplement. En effet, on remarque que

$$\hat{M}_{\text{suc}}^N = \frac{1}{g^N} \sum_{n=0}^N g^n |n\rangle\langle n| \quad (9.74a)$$

$$= \frac{1}{g^N} \left[ \sum_{k=0}^N |k\rangle\langle k| \right] \hat{T}. \quad (9.74b)$$

Voyons maintenant comment utiliser cette propriété : supposons que l'on envoie un état  $\hat{\rho}_{\text{in}}$  dans le canal  $(T, \epsilon)$ , produisant un état  $\hat{\rho}_{\text{out}}$  (9.21). On applique ensuite soit le NLA décrit par  $\hat{T}$ , soit le NLA décrit par  $\hat{M}_{\text{suc}}^N$ . Dans le premier cas, on obtient l'état  $\hat{\rho}_{\text{out}}^{\text{NLA}}$  (9.36), et on peut utiliser l'équivalence avec le système effectif que nous avons montré. Dans le second cas, on obtient un état non gaussien  $\hat{\rho}_{\text{out}}^{\text{M}}$ , qui peut être calculé à partir de  $\hat{\rho}_{\text{out}}^{\text{NLA}}$  :

$$\hat{\rho}_{\text{out}}^{\text{M}} = \frac{\frac{1}{g^{2N}} \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}} \hat{\Pi}}{\frac{1}{g^{2N}} \text{Tr} \{ \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}} \hat{\Pi} \}} = \frac{\hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}} \hat{\Pi}}{\text{Tr} \{ \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}} \}} \quad (9.75)$$

avec  $\hat{\Pi} = \sum_{k=0}^N |n\rangle\langle n|$ , et en utilisant le fait que  $\hat{\Pi}^2 = \hat{\Pi}$ . Ainsi, on peut utiliser le système effectif, et simplement projeter l'état en sortie du canal effectif sur le sous-espace contenant  $N$  photons au maximum.

Si l'on cherche à calculer l'état gaussien de même matrice de covariance que  $\hat{\rho}_{\text{out}}^{\text{NLA}}$ , cette opération n'est en fait même pas nécessaire. En effet, la valeur moyenne d'un opérateur  $\hat{A}$  est donnée par :

$$\text{Tr}\{\hat{A} \hat{\rho}_{\text{out}}^{\text{M}}\} = \frac{1}{\text{Tr}\{\hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}}\}} \text{Tr}\{\hat{A} \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}} \hat{\Pi}\} \quad (9.76a)$$

$$= \frac{1}{\text{Tr}\{\hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}}\}} \text{Tr}\{\hat{\Pi} \hat{A} \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}}\} \quad (9.76b)$$

Ainsi, la valeur moyenne de  $\hat{A}$  pour l'état  $\hat{\rho}_{\text{out}}^{\text{M}}$  est donc égale (au facteur  $1/\text{Tr}\{\hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}}\}$  près) à la valeur moyenne de  $\hat{\Pi} \hat{A} \hat{\Pi}$  pour l'état  $\hat{\rho}_{\text{out}}^{\text{NLA}}$ . Celle-ci peut se calculer en utilisant les fonctions de Wigner et la formule (2.73) :

$$\text{Tr}\{\hat{\Pi} \hat{A} \hat{\Pi} \hat{\rho}_{\text{out}}^{\text{NLA}}\} = 4\pi \int dx dp W_{\hat{\Pi} \hat{A} \hat{\Pi}}(x, p) W_{\hat{\rho}_{\text{out}}^{\text{NLA}}}(x, p) \quad (9.77)$$

De cette manière, on peut donc calculer les valeurs moyennes et les moments des quadratures de  $\hat{\rho}_{\text{out}}^{\text{M}}$ , en remplaçant  $\hat{A}$  par  $\hat{X}$ ,  $\hat{P}$ ,  $\hat{X}^2$  ou  $\hat{P}^2$ .

## 9.5 Application aux communications quantiques

### 9.5.1 Suppression des pertes et atténuateur sans bruit

#### Principe

M. Mićuda *et al.* ont introduit le concept d'*atténuateur sans bruit* afin de réduire les pertes effectives d'un canal [Mićuda12]. Cet atténuateur est un amplificateur sans bruit de gain  $\nu < 1$ , qui peut être simplement implémenté en envoyant l'état à atténuer sur une lame séparatrice de transmission en amplitude  $\nu$ , dont l'autre mode d'entrée est vide, et en conditionnant à la mesure du vide sur la sortie du mode réfléchi.

Le principe de leur protocole est le suivant : l'état initial est d'abord atténué d'un facteur  $\nu$ . Puis il est envoyé dans le canal quantique, supposé sans bruit, de transmission en amplitude  $\sqrt{T}$ . Enfin, il est amplifié avec un NLA de gain  $g = \frac{1}{\nu\sqrt{T}}$ . Les auteurs parlent de "suppression des pertes" : leur combinaison d'atténuateur et d'amplificateur réduit efficacement les dégradations dues aux pertes, en produisant un état final qui est un mélange de l'état initial non dégradé et de termes d'autant plus faibles que  $\nu$  est faible. A la limite où  $\nu \rightarrow 0$ , tout se passe comme si le canal n'introduisait plus de pertes, et se comportait comme l'opérateur identité, sans modifier l'état initial d'aucune façon [Mićuda12].

Outre le fait que cette limite correspond à une probabilité de succès nulle, le canal effectif de [Mićuda12] ne prend pas une forme simple lorsque  $\nu \neq 0$ . De plus, le cas d'un canal bruité semble délicat à généraliser car les calculs en base de Fock deviennent vite assez complexes.

En utilisant les résultats de la section 9.3, la généralisation de cette notion de suppression des pertes est en revanche immédiate. Nous avons montré qu'un NLA de gain  $g$  placé après le canal est équivalent à un NLA effectif de gain  $g_{\text{in}}$  placé avant le canal effectif. Dès lors, il suffit d'utiliser un atténuateur de gain  $1/g_{\text{in}}$  pour compenser l'effet du NLA effectif, puisque

$$(1/g_{\text{in}})^{\hat{n}} g_{\text{in}}^{\hat{n}} = \mathbb{I}. \quad (9.78)$$

Il ne reste alors plus que le canal effectif, comme illustré sur la figure 9.12.

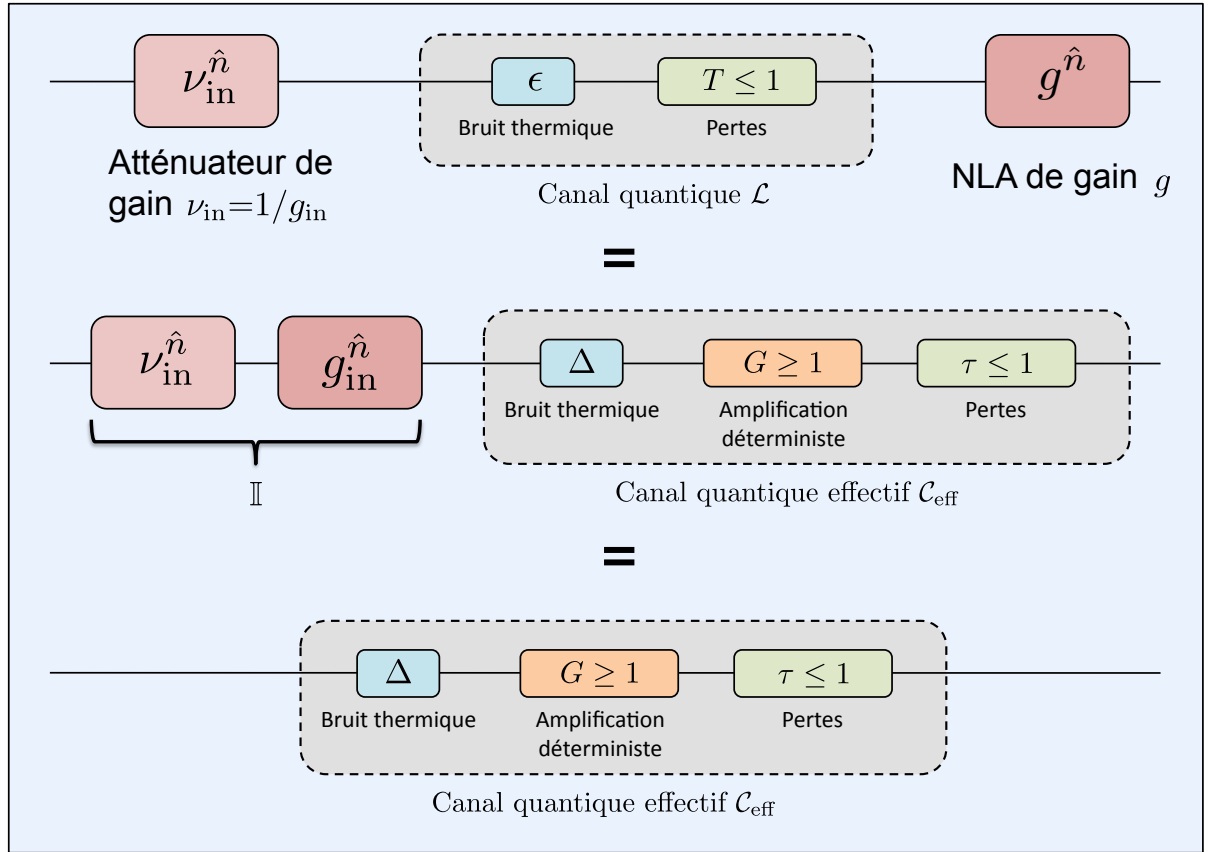


FIGURE 9.12 – Canal équivalent en utilisant un atténuateur de gain  $\nu=1/g_{\text{in}}$ . Le facteur de normalisation  $\mu$  n'est pas représenté.

### Canal sans bruit

Lorsque le canal n'est pas bruité ( $\epsilon=0$ ), les paramètres effectifs prennent la forme (9.59). On peut donc rigoureusement obtenir un canal effectif de transmission  $\eta$  en utilisant un atténuateur de gain

$$\nu = \frac{1}{g_{\text{in}}} = \frac{1}{\sqrt{1+(g^2-1)T}}. \quad (9.79)$$

Comme nous l'avons vu précédemment, pour un gain fini, le canal effectif a une transmission  $\eta$  supérieure à  $T$ , mais qui reste inférieure à 1. Pour un gain  $g \gg 1$ , (9.79) devient

$$\nu \simeq \frac{1}{g\sqrt{T}}. \quad (9.80)$$

Nous retombons alors sur les valeurs d'atténuation et de gain considérées dans [Mičuda12]. A la limite où  $g \rightarrow +\infty$ , la transmission effective tend vers 1 et on obtient un canal sans pertes. L'avantage de notre méthode est de pouvoir rigoureusement obtenir un canal de transmission  $\eta$ , quelle que soit la valeur de  $g$ .

Ainsi, une amplification sans bruit de faible gain ne permettra pas de complètement supprimer les pertes, mais pourra quand même les réduire, lorsqu'elle est associée à un atténuateur de gain adapté.

### Suppression totale des pertes d'un canal bruité, avec un gain fini

La même méthode nous permet d'obtenir le canal quantique effectif  $\mathcal{C}_{\text{eff}}$  dans le cas général d'un canal avec un bruit thermique  $\epsilon$ . L'étude des paramètres effectifs dans la section 9.3 a montré que, sous l'effet du bruit,  $\eta$  peut devenir égal à 1 pour une valeur de gain  $g^{\text{max}}$  qui est cette fois finie.

Nous pouvons donc en tirer deux conclusions : la première est qu'il est possible de supprimer totalement les pertes d'un canal bruité avec un NLA de gain fini et un atténuateur de gain  $1/g_{\text{in}}$  non nul. Le gain nécessaire est d'autant plus faible que le canal est bruité (cf. Fig. 9.10 (c)).

La seconde conclusion est que l'on peut ajouter volontairement du bruit thermique à l'entrée d'un canal afin d'augmenter la transmission effective, pour un NLA de gain donné. De ce fait, un canal sans bruit introduisant des pertes  $T$  peut être converti en canal sans pertes, introduisant un bruit  $\Delta$  (lorsque l'on ajoute un bruit  $\epsilon$ ).

### Un atténuateur est-il toujours bénéfique ?

Il existe au moins deux types d'états pour lesquels un atténuateur sans bruit n'est pas forcément nécessaire : les états de Fock et les états cohérents. Tous deux ne sont pas "déformés" par un NLA, et acquièrent un coefficient global qui disparaît lors de la normalisation. Ainsi, un état de Fock  $|n\rangle$  ne sera simplement pas affecté par le NLA effectif  $g_{\text{in}}^{\hat{n}}$ , ni par un éventuel atténuateur sans bruit. Pour ce type d'état, l'amplification de gain  $g$  après le canal  $(T, \epsilon)$  est donc directement équivalente au canal effectif  $\mathcal{C}_{\text{eff}}$ , sans avoir besoin d'un atténuateur.

Un état cohérent  $|\alpha\rangle$  sera quand à lui amplifié par le NLA effectif  $g_{\text{in}}^{\hat{n}}$ , mais puisque cette amplification est sans bruit, elle est bénéfique ! L'amplitude moyenne de l'état en sortie du canal effectif étant égale à  $g_{\text{in}}\sqrt{\eta}\alpha$ , et sa variance à  $1+\eta\Delta=1+g_{\text{in}}^2\eta(\frac{\Delta}{g_{\text{in}}^2})$ , on peut définir une "transmission effective équivalente"  $\tilde{\eta}=g_{\text{in}}^2\eta$  supérieure à  $\eta$ , et un "bruit effectif équivalent"  $\tilde{\Delta}=\Delta/g_{\text{in}}^2$  inférieur à  $\Delta$ .

Bien sûr, lorsque l'état initial est composé d'une superposition de ces états, l'atténuateur sans bruit retrouve son intérêt afin de ne pas biaiser la superposition sous l'action de  $g_{\text{in}}^{\hat{n}}$ .

#### 9.5.2 "Concentration de phase"

Nous avons vu qu'un état cohérent  $|\alpha\rangle$  envoyé dans le canal  $(T, \epsilon)$  et amplifié avec un NLA de gain  $g$  est transformé en un état thermique déplacé, de valeur moyenne  $\tilde{g}\sqrt{T}\alpha$ , et de variance  $\frac{1+g^2\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2}$ . Rappelons les expressions de  $\tilde{g}$  et  $\lambda_{\text{ch}}$ , respectivement données par (9.34) et (9.24) :

$$\tilde{g} = g \frac{1-\lambda_{\text{ch}}^2}{1-g^2\lambda_{\text{ch}}^2} \quad \text{et} \quad \lambda_{\text{ch}}^2 = \frac{T\epsilon}{2+T\epsilon} \quad (9.81)$$

Ce résultat permet de calculer simplement l'amplification d'un état cohérent auquel on aurait ajouté volontairement du bruit  $\epsilon$  : il suffit de poser  $T=1$ . Comme nous l'avons déjà indiqué,  $\tilde{g}$  est supérieur à  $g$ , et ce d'autant plus que le canal est bruité, *i.e.* que le bruit ajouté  $\epsilon$  est grand. Ce phénomène semble en tout point similaire à celui de la "concentration de phase" [Marek10a, Usuga10], présentée au chapitre 8 : le bruit ajouté est équivalent à un déplacement aléatoire autour de  $\alpha$ . L'amplificateur sans bruit favorise ensuite les états cohérents  $|\beta\rangle\langle\beta|$  de grandes amplitudes dans le mélange statistique, grâce au facteur  $\exp[(g^2-1)|\beta|^2]$  qu'il introduit. Ce facteur dépend exponentiellement de  $\beta$ , alors qu'une ou plusieurs soustractions de photons introduisent un coefficient polynomial : on s'attend donc à ce qu'un petit ajout de bruit avant le NLA puisse conduire à une grande "concentration de phase".

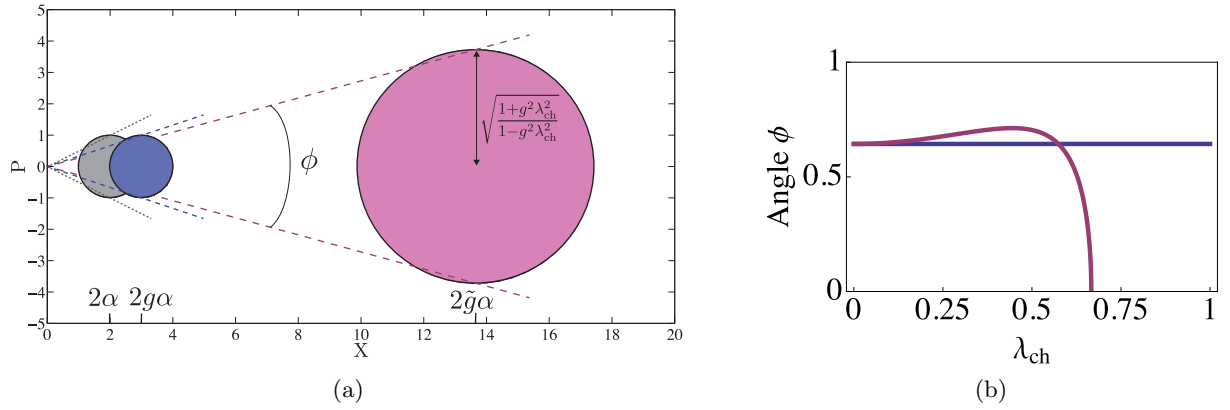


FIGURE 9.13 – Concentration de phase en ajoutant du bruit avant un NLA de gain donné : (a) fonction de wigner (représentée avec un écart-type de rayon) de l'état initial  $\alpha=1$  (en gris), de l'état amplifié avec un NLA de gain  $g=1.5$  (en bleu), et de l'état amplifié après ajout d'un bruit  $\epsilon$  (en rose) ; (b) angle  $\phi$  de l'état cohérent  $\alpha=1$  amplifié avec un NLA de gain  $g=1.5$  fixé, sans ajout de bruit (en bleu), et en ajoutant un bruit  $\epsilon$  vérifiant  $\lambda_{ch}^2 = \frac{\epsilon}{2+\epsilon}$  (en rose).

La figure 9.13 montre qu'il semble être effectivement possible d'utiliser cet ajout de bruit afin d'améliorer l'estimation de la phase d'un état cohérent, pour un NLA de gain donné. Le critère utilisé est cependant assez qualitatif : on définit l'incertitude de phase  $\phi$  par<sup>4</sup>

$$\tan \frac{\phi}{2} = \frac{\text{Ecart-type}}{\text{Amplitude moyenne}} = \frac{\sqrt{\frac{1+g^2\lambda_{ch}^2}{1-g^2\lambda_{ch}^2}}}{2\tilde{g}\alpha}, \quad (9.82)$$

comme montré sur la figure 9.13 (a). La figure 9.13 (b) compare l'angle  $\phi$  obtenu avec et sans ajout de bruit : on constate clairement une amélioration lorsque l'ajout de bruit est suffisamment important. Soulignons que la divergence observée survient lorsque le bruit ajouté ne permet plus de satisfaire la condition  $g < \sqrt{1 + \frac{2}{T\epsilon}}$  (avec  $T=1$ , et  $1+\epsilon = \frac{1+\lambda_{ch}^2}{1-\lambda_{ch}^2}$ ).

Cette amélioration sur l'estimation de la phase serait à comparer plus précisément avec les autres protocoles tels que [Marek10a, Usuga10] utilisant une soustraction de photon, afin de voir si l'utilisation d'un NLA est réellement avantageuse compte tenu des difficultés expérimentales pour l'implémenter. Il faudrait également comparer l'intérêt expérimental de cet ajout de bruit, par rapport à simplement augmenter le gain du NLA.

## 9.6 Conclusion

Ce chapitre nous a permis d'étudier plusieurs propriétés générales d'un amplificateur sans bruit. Nous avons d'abord montré que la probabilité de succès peut être bornée supérieurement par une valeur constante égale à  $1/g^2$ , lorsque le NLA est utilisé avec des états thermiques.

Nous avons ensuite montré l'équivalence entre un canal introduisant des pertes et du bruit thermique suivi d'un amplificateur sans bruit, et un amplificateur sans bruit effectif placé en amont d'un canal effectif. Outre les applications potentielles que nous avons présenté, dont la suppression des pertes d'un canal, cette équivalence nous permettra d'étudier l'utilisation d'un NLA en cryptographie quantique, qui fait l'objet du chapitre suivant.

4. Le facteur 2 pour l'amplitude moyenne provient de la convention  $N_0=1$ .





## Chapitre 10

# L'amplificateur sans bruit non déterministe en cryptographie quantique

### Sommaire

---

|                                                                                      |            |
|--------------------------------------------------------------------------------------|------------|
| <b>10.1 Introduction</b>                                                             | <b>195</b> |
| <b>10.2 Calcul des taux secrets</b>                                                  | <b>196</b> |
| 10.2.1 Taux secrets sans le NLA                                                      | 196        |
| 10.2.2 Taux secrets avec le NLA                                                      | 197        |
| <b>10.3 Amélioration des performances - attaques collectives</b>                     | <b>197</b> |
| 10.3.1 Considérations préliminaires                                                  | 198        |
| 10.3.2 Une distance de transmission augmentée, et une plus grande tolérance au bruit | 200        |
| 10.3.3 Que se passe t'il quand le gain augmente trop ?                               | 206        |
| 10.3.4 Augmentation arbitraire de la distance maximale de transmission               | 207        |
| <b>10.4 Attaques individuelles et non amélioration des performances avec le NLA</b>  | <b>211</b> |
| 10.4.1 Démonstration pour un canal ajoutant du bruit                                 | 211        |
| 10.4.2 Démonstration exacte pour tout $T$ , pour un canal sans bruit                 | 214        |
| <b>10.5 Conclusion</b>                                                               | <b>216</b> |

---

### 10.1 Introduction

Les protocoles de cryptographie quantique ne permettent d'échanger une clé secrète que si les conditions ne sont pas trop favorables pour Eve : toutes les imperfections du canal de transmission sont à son avantage, si bien que si les pertes ou le bruit sont trop importants, elle finit par obtenir plus d'information qu'Alice sur les données de Bob. Le taux secret décroît donc à mesure que les pertes ou le bruit augmentent, et peut devenir nul pour une certaine distance de transmission.

Puisqu'un amplificateur sans bruit non déterministe permet de réduire les pertes effectives d'un canal, il apparaît alors comme un outil potentiel pour améliorer les performances d'un protocole de QKD à variables continues.

Ce chapitre est consacré à cette application de l'amplificateur sans bruit, en se focalisant sur le protocole GG02. Nous détaillerons explicitement les améliorations apportées pour un protocole sécurisé contre les attaques collectives, en montrant qu'il est possible d'améliorer la distance maximale de transmission, ainsi que la tolérance face au bruit ajouté. Puis nous montrerons que de manière peut être un peu surprenante, le NLA n'améliore jamais le taux secret contre les attaques individuelles.

Ce chapitre est en partie basé sur les résultats qui ont donné lieu à la publication [Blandino12c].

## 10.2 Calcul des taux secrets

Le calcul des taux secrets contre les attaques individuelles et collectives a été présenté dans la section 7.4.4. Rappelons en ici les principaux éléments : le canal quantique est modélisé par une transmission  $T$ , et un bruit ajouté équivalent en entrée  $\epsilon$ . Nous nous plaçons dans le cas général d'une détection homodyne imparfaite, avec une efficacité  $\nu$  et un bruit électronique (sur les mesures)  $\kappa$ .  $V = \frac{1+\lambda^2}{1-\lambda^2} = V_A + 1$  est la variance de l'état thermique d'Alice<sup>1</sup>,  $\chi_{\text{line}} = \frac{1-T}{T} + \epsilon$  est le bruit ajouté par le canal ramené à l'entrée incluant l'effet des pertes,  $\chi_{\text{hom}} = \frac{1-\nu}{\nu} + \frac{\kappa}{\nu}$  est le bruit ajouté par la détection homodyne ramené à l'entrée de celle-ci, et  $\chi_{\text{tot}} = \chi_{\text{line}} + \frac{\chi_{\text{hom}}}{T}$  est le bruit total ramené à l'entrée du canal.

### Information mutuelle entre Alice et Bob

L'information mutuelle entre Alice et Bob est :

$$I_{AB} = \frac{1}{2} \log_2 \left[ \frac{V + \chi_{\text{tot}}}{1 + \chi_{\text{tot}}} \right] \quad (10.1)$$

### Information mutuelle entre Bob et Eve - attaques individuelles

Lorsque Eve est restreinte à des attaques individuelles, l'information mutuelle avec Bob est donnée par :

$$I_{BE} = \frac{1}{2} \log_2 \left[ \frac{T^2 (V + \chi_{\text{tot}}) \left( \frac{1}{V} + \chi_{\text{line}} \right)}{1 + T \chi_{\text{hom}} \left( \frac{1}{V} + \chi_{\text{line}} \right)} \right] \quad (10.2)$$

### Information mutuelle entre Bob et Eve - attaques collectives

Lorsque Eve peut implémenter des attaques collectives, l'information mutuelle avec Bob est bornée par la borne de Holevo  $\chi_{BE}$ , définie par (7.9).

#### 10.2.1 Taux secrets sans le NLA

Le taux secret contre les attaques individuelles, sans le NLA, est donné par :

$$\Delta I_{\text{Ind}}(\lambda, T, \epsilon) = \beta I_{AB}(\lambda, T, \epsilon) - I_{BE}(\lambda, T, \epsilon) \quad (10.3)$$

Le taux secret contre les attaques collectives, sans le NLA, est donné par :

$$\Delta I_{\text{H}}(\lambda, T, \epsilon) = \beta I_{AB}(\lambda, T, \epsilon) - \chi_{BE}(\lambda, T, \epsilon) \quad (10.4)$$

---

1. Dans ce chapitre, nous utilisons la convention  $N_0=1$ .

La dépendance par rapport à  $\lambda$ ,  $T$  et  $\epsilon$  est soulignée car ces formules seront utilisées avec les paramètres effectifs pour calculer le taux secret avec le NLA. En revanche, les autres paramètres ne sont pas transformés, et leur dépendance est implicite.

### 10.2.2 Taux secrets avec le NLA

Avec le NLA, le taux est obtenu en utilisant (10.3) ou (10.4) avec les paramètres effectifs, pondéré par la probabilité de succès  $P_{\text{suc}}$  de l'amplification :

$$\Delta I_{\text{Ind}}^{\text{NLA}}(\lambda, T, \epsilon) = P_{\text{suc}} \Delta I_{\text{Ind}}(\zeta, \eta, \epsilon^g) \quad (10.5)$$

$$\Delta I_{\text{H}}^{\text{NLA}}(\lambda, T, \epsilon) = P_{\text{suc}} \Delta I_{\text{H}}(\zeta, \eta, \epsilon^g) \quad (10.6)$$

Puisque  $\zeta$  dépend linéairement de  $\lambda$  et que le canal quantique effectif en est indépendant, nous pouvons le considérer comme un paramètre libre, optimisé pour maximiser le taux secret. Ceci fixe ensuite la valeur de  $\lambda = \zeta / g_{\text{in}}$ . Cette modulation de l'état EPR envoyé par Alice est en fait équivalente à l'utilisation d'un atténuateur quantique présenté dans la section 9.5 : un atténuateur de gain  $1/g_{\text{in}}$  transforme simplement un état EPR de paramètre  $\zeta$  en un état EPR de paramètre  $\zeta/g_{\text{in}}$ . Ainsi, avec un atténuateur avant le canal, Alice enverrait un état EPR  $|\zeta\rangle$ , alors que sans atténuateur elle envoie un état  $|\zeta/g_{\text{in}}\rangle$ .

Nous attribuerons à  $P_{\text{suc}}$  la borne  $1/g^2$  obtenue dans la section 9.2, en gardant à l'esprit que les valeurs des taux secrets seront très optimistes, mais que cela n'a pas d'incidence sur l'amélioration des performances que nous allons détailler : seul le taux secret sera diminué, mais les améliorations de la distance maximale de transmission et des autres performances seront inchangées. La probabilité de succès agit simplement comme un facteur de proportionnalité qui ne change pas le fait qu'un taux secret soit positif ou négatif. Avec une probabilité de succès trop optimiste, le seul inconvénient est de ne pas pouvoir donner une estimation numérique précise du taux secret.

## 10.3 Amélioration des performances - attaques collectives

Nous avons maintenant tous les outils pour attaquer le cœur du problème, et étudier l'utilité du NLA en cryptographie quantique. Notre étude combinera autant que possible des résultats numériques exacts, et des développements perturbatifs au premier ordre en  $T$ , valables pour de fortes pertes. Ces développements nous permettront d'obtenir des résultats analytiques qui seront fort utiles pour comprendre les modifications apportées par le NLA, et qui sont de plus en très bon accord avec les résultats numériques. Nous attacherons d'ailleurs une importance certaine à étudier et vérifier leurs conditions de validité.

Nous commencerons par une étude du développement perturbatif du taux secret avec le NLA  $\Delta I_{\text{H}}^{\text{NLA}}$ , ce qui nous permettra d'obtenir une formule analytique de la distance maximale de transmission permettant l'échange d'une clé secrète. En la comparant avec la distance maximale de transmission sans NLA, nous montrerons un résultat important : les pertes admissibles sont augmentées de  $20 \log_{10} g$  dB avec le NLA. Nous montrerons également que le NLA améliore la résistance au bruit, en permettant un échange de clé pour un canal davantage bruité. Nous illustrerons ces résultats par plusieurs exemples. Puis nous verrons comment prolonger arbitrairement la distance de transmission en adaptant le gain en fonction des pertes.

### 10.3.1 Considérations préliminaires

#### Développement perturbatif du taux secret avec le NLA

La complexité de la formule (10.4) rend l'obtention de résultats analytiques extrêmement difficile sans faire d'approximations. Heureusement, dans la plupart des cas le taux secret est encore positif pour de petites valeurs de transmission  $T$ , permettant ainsi d'effectuer un développement au premier ordre. Le calcul détaillé est effectué en annexe G, dans le cas général d'un canal bruité et d'une détection homodyne imparfaite :

$$\Delta I_H^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 T}{(1-\lambda^2)^2 (1+\kappa) 2 \ln 2} \left\{ (1-\lambda^2)^2 (1+\kappa) \epsilon \ln \frac{\epsilon}{2} + (1-\lambda^2)^2 \nu \epsilon \ln [g^2 T] + \nu \lambda^4 \ln \lambda^4 + \right. \\ \left. (\lambda^2 - 1) \left( \nu (\epsilon - \lambda^2 (2\beta + \epsilon)) + (\lambda^2 - 1) \epsilon (\nu - \kappa - 1) \ln \left[ \frac{\epsilon(1+\kappa-\nu)}{2(1+\kappa)} \right] \right) \right\} \quad (10.7)$$

La plupart du temps, nous supposons que la détection homodyne est parfaite ( $\nu=1, \kappa=0$ ), afin de simplifier les calculs. Dans ce cas, (10.7) devient :

$$\Delta I_H^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 T}{(1-\lambda^2)^2 2 \ln 2} \left\{ \lambda^4 \ln \lambda^4 + (\lambda^2 - 1) \left[ -\lambda^2 (2\beta + \epsilon) + (\lambda^2 - 1) \epsilon \left( \ln \frac{\epsilon}{2} + \ln [g^2 T] \right) + \epsilon \right] \right\} \quad (10.8)$$

Cette formule, bien qu'écrite sous une forme différente, est identique à celle présentée dans [Blandino12c]. Enfin, il pourra également être utile de considérer le cas d'un canal n'ajoutant pas de bruit ( $\epsilon=0$ ), afin de montrer que dans ce cas le NLA n'apporte pas d'améliorations. Toujours en supposant la détection homodyne parfaite, (10.8) devient :

$$\Delta I_H^{\text{NLA}} \simeq P_{\text{suc}} \frac{\lambda^2 g^2 T}{(1-\lambda^2)^2 \ln 2} [\beta (1-\lambda^2) + \lambda^2 \ln \lambda^2] \quad (10.9)$$

#### Optimisation de la variance de modulation

Puisque nous autorisons toujours Alice à optimiser sa variance de modulation, et donc  $\lambda$ , on peut chercher la valeur optimale  $\lambda_{\text{opt}}$  qui maximise le taux secret. On montre sans difficulté que la dérivée de (10.7) par rapport à  $\lambda$  est annulée lorsque la condition suivante est vérifiée (Fig. 10.1) :

$$\frac{\lambda_{\text{opt}}^2}{1-\lambda_{\text{opt}}^2} (\lambda_{\text{opt}}^2 - 4 \ln \lambda_{\text{opt}} - 1) = \beta \quad (10.10)$$

La variance optimale ne dépend donc que de l'efficacité de réconciliation  $\beta$ , quel que soit le bruit ajouté, la transmission et les imperfections de la détection homodyne – lorsque le développement au premier ordre en  $T$  est valide –.

Cette optimisation permet de supprimer une variable de (10.7). En pratique, plutôt que de chercher  $\lambda_{\text{opt}}$  pour une valeur de  $\beta$  donnée, il est plus simple de faire la démarche inverse en considérant que (10.10) donne la valeur optimale de  $\beta$  pour une valeur  $\lambda_{\text{opt}}$  donnée. On peut ainsi directement utiliser cette expression de  $\beta$  dans les formules des développements pour obtenir les taux secrets maximisés sur  $\lambda$ , en se souvenant que  $\lambda_{\text{opt}}$  n'est plus un paramètre libre, mais prend une valeur qui dépend de  $\beta$ . Afin de simplifier les notations, nous continuerons d'utiliser la notation  $\lambda$  pour lorsque la variance de modulation est optimisée, étant sous-entendu qu'il s'agit de  $\lambda_{\text{opt}}$ .

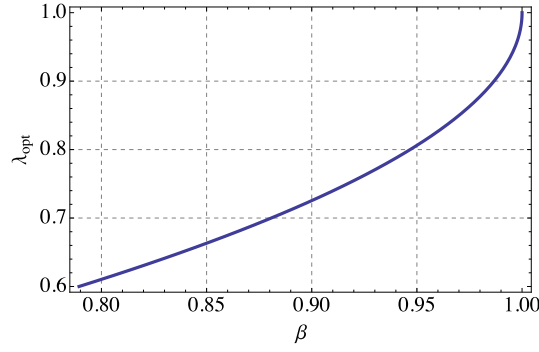


FIGURE 10.1 –  $\lambda_{\text{opt}}$  en fonction de  $\beta$ , donné par l'équation 10.10.  $\beta=0.95$  correspond à  $\lambda_{\text{opt}} \simeq 0.8065$ .

Pour une détection homodyne imparfaite, la maximisation de (10.7) donne :

$$\Delta I_{\text{H}}^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 T}{(1-\lambda^2)^2 (1+\kappa) 2 \ln 2} \left\{ (1+\kappa) (1-\lambda^2)^2 \epsilon \ln \frac{\epsilon}{2} + (1-\lambda^2)^2 \nu \epsilon \ln [g^2 T] + \nu \lambda^4 \ln \lambda^4 + \right. \\ \left. \nu [(\lambda^2-1) (2\lambda^4 - \lambda^2 \epsilon + \epsilon) - 2\lambda^4 \ln \lambda^4] + (1-\lambda^2)^2 \epsilon (\nu - \kappa - 1) \ln \left[ \frac{\epsilon(1+\kappa-\nu)}{2(1+\kappa)} \right] \right\} \quad (10.11)$$

Pour une détection homodyne parfaite, ce taux devient :

$$\Delta I_{\text{H}}^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 T}{(1-\lambda^2)^2 2 \ln 2} \left\{ (1-\lambda^2)^2 \epsilon \left( \ln \frac{\epsilon}{2} + \ln [g^2 T] \right) - \lambda^4 \ln \lambda^4 + (\lambda^2-1) (2\lambda^4 - \lambda^2 \epsilon + \epsilon) \right\} \quad (10.12)$$

Enfin, pour une détection homodyne parfaite et en l'absence de bruit :

$$\Delta I_{\text{H}}^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 \lambda^4 T}{(1-\lambda^2)^2 2 \ln 2} (\lambda^2 - \ln \lambda^2 - 1) \quad (10.13)$$

Puisque que  $\lambda^2 - \ln \lambda^2 - 1 > 0$  pour toute valeur de  $\lambda$  comprise entre 0 et 1, on voit tout de suite qu'il n'y a pas de distance maximale de transmission lorsqu'il n'y a pas de bruit, même pour  $\beta < 1$ . Nous reviendrons sur ce point dans le paragraphe suivant.

Par la suite, le développement perturbatif du taux secret sera toujours utilisé avec la variance optimale, si bien que sauf mention contraire nous n'utiliserons plus que les expressions (10.11), (10.12) et (10.13) pour les calculs analytiques. Ces formules permettent aussi d'obtenir le développement du taux secret sans NLA en prenant simplement  $g=1$  et  $P_{\text{suc}}=1$ .

### Distance maximale de transmission

La relative simplicité des formules précédentes permet de calculer analytiquement la transmission limite  $T_{\text{lim}}^{\text{NLA}}$  du canal réel en dessous de laquelle le taux secret s'annule. Dans le cas général avec une variance de modulation non optimisée, on trouve à partir de (10.7) que :

$$T_{\text{lim}}^{\text{NLA}} = \frac{1}{g^2} \frac{2}{\epsilon} (1+\kappa)^{\frac{\nu-\kappa-1}{\nu}} (1+\kappa-\nu)^{\frac{1+\kappa}{\nu}-1} \lambda^{-\frac{4\lambda^4}{(1-\lambda^2)^2}} e^{\frac{\lambda^2(2\beta+\epsilon)-\epsilon}{(\lambda^2-1)\epsilon}} \quad (10.14)$$

Cette formule est donnée à titre informatif, puisque nous utiliserons plutôt les taux maximisés. Pour une détection homodyne imparfaite avec la variance de modulation optimale, (10.11) donne :

$$T_{\text{lim}}^{\text{NLA}} = \frac{1}{g^2} \frac{2}{\epsilon} (1+\kappa)^{\frac{\nu-\kappa-1}{\nu}} (1+\kappa-\nu)^{\frac{1+\kappa}{\nu}-1} \lambda^{\frac{4\lambda^4}{(1-\lambda^2)^2\epsilon}} e^{\frac{2\lambda^4-\lambda^2\epsilon+\epsilon}{\epsilon-\lambda^2\epsilon}} \quad (10.15)$$

Et enfin, pour une détection homodyne parfaite avec la variance optimale, (10.12) donne :

$$T_{\text{lim}}^{\text{NLA}} = \frac{1}{g^2} \frac{2}{\epsilon} \exp \left[ \frac{2\lambda^4-\lambda^2\epsilon+\epsilon}{\epsilon(1-\lambda^2)} + \frac{4\lambda^4}{\epsilon(1-\lambda^2)^2} \ln \lambda \right] \quad (10.16)$$

Ces trois expressions tendent toutes vers 0 lorsque le bruit tend vers 0,

$$\lim_{\epsilon \rightarrow 0} T_{\text{lim}}^{\text{NLA}} = 0 \quad (10.17)$$

ce qui montre qu'il n'y a pas de distance maximale en l'absence de bruit, même si  $\beta < 1$ . En revanche la distance de transmission est toujours limitée en présence de bruit, même lorsque  $\beta = 1$ . Le fait que la transmission intervienne à l'intérieur de l'expression (10.11) explique qu'il puisse y avoir une distance maximale, puisque si elle n'est qu'en facteur global, elle a pour seul effet de diminuer le taux mais sans jamais strictement l'annuler. Comme on pouvait intuitivement s'en douter, ces transmissions limites ne dépendent pas de la probabilité de succès  $P_{\text{suc}}$ .

Remarquons enfin que (10.14), (10.15) et (10.16) sont en général en excellent accord avec les courbes numériques, mais il peut arriver que cela ne soit plus le cas si  $\epsilon$  ou  $\beta$  (et donc  $\lambda_{\text{opt}}$ ) sont trop grands. Il convient donc de toujours vérifier leur validité en les comparant numériquement à (10.6).

### 10.3.2 Une distance de transmission augmentée, et une plus grande tolérance au bruit

#### Augmentation de la distance maximale de transmission

Dans toutes les expressions de  $T_{\text{lim}}^{\text{NLA}}$ , le gain du NLA n'intervient qu'à travers le facteur  $\frac{1}{g^2}$ . Puisque la transmission limite  $T_{\text{lim}}$  sans NLA est obtenue en prenant simplement  $g=1$ , on peut relier les transmissions limites avec et sans NLA par la relation suivante :

$$T_{\text{lim}}^{\text{NLA}} = \frac{1}{g^2} T_{\text{lim}} \quad (10.18)$$

Le NLA permet donc de diminuer la transmission limite, ce qui revient à augmenter la distance de transmission<sup>2</sup>. En exprimant les pertes en dB, une transmission diminuée d'un facteur  $g^2$  correspond à une augmentation des pertes de :

$$\Delta \mathcal{P} = 20 \log_{10} g \text{ dB} \quad (10.19)$$

En d'autres termes, le protocole de QKD fonctionne pour un canal ayant  $\Delta \mathcal{P}$  dB de pertes supplémentaires, en utilisant un NLA. Soulignons encore une fois que cette amélioration ne

---

2. Nous parlerons selon les cas d'une augmentation de la distance maximale de transmission, d'une diminution de la transmission minimale, ou d'une augmentation des pertes maximales admissibles. Ces trois formulations sont bien sûr équivalentes.

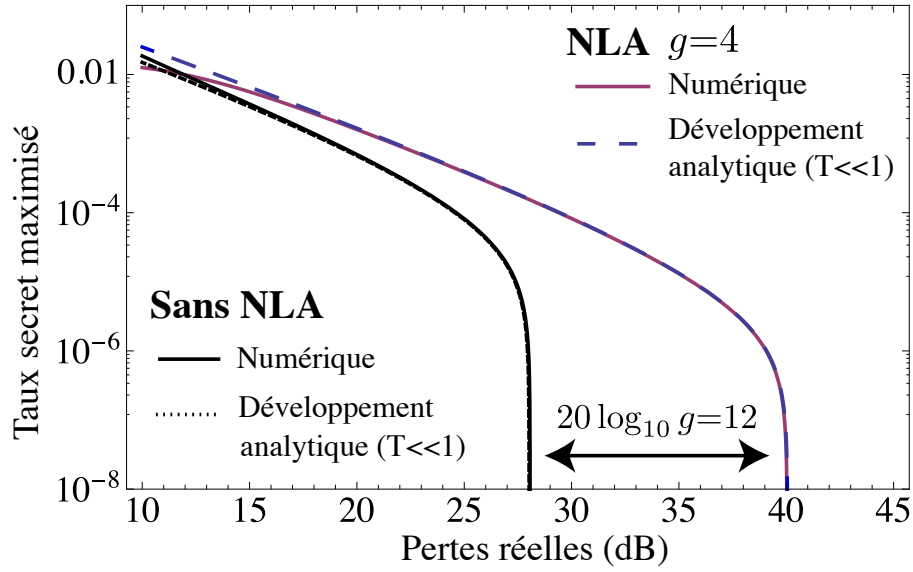


FIGURE 10.2 – Comparaison de  $\Delta I_H$  et  $\Delta I_H^{\text{NLA}}$  et de leurs développements perturbatifs, en fonction des pertes en dB, pour  $\beta=0.95$  et  $\epsilon=0.05$ . Les courbes numériques sont obtenues avec les formules (10.4) et (10.6) en utilisant  $P_{\text{suc}}=1/g^2$ , maximisées numériquement sur  $\lambda$  ou  $\zeta$  pour chaque valeur de pertes. Les développements analytiques correspondent à la formule (10.12) (avec  $P_{\text{suc}}=1$  et  $g=1$  sans NLA), en utilisant  $\lambda_{\text{opt}}=0.8065$ . La détection homodyne est supposée parfaite.

dépend pas des autres paramètres, *i.e.* de  $\beta$ , de la variance de modulation, du bruit thermique, ou des imperfections de la détection homodyne.

On peut comprendre l'équation (10.18) de manière intuitive, en regardant le développement des paramètres effectifs au premier ordre en  $T$  :

$$\eta \simeq g^2 T \quad \epsilon^g = \epsilon + \frac{1}{2}(g^2 - 1)(2 - \epsilon)\epsilon T \quad \zeta \simeq \lambda + \frac{1}{2}(g^2 - 1)\lambda T \quad (10.20)$$

Il ne peut pas y avoir de termes d'ordre 0 en  $T$  dans le développement du taux secret, car sinon ce dernier ne serait pas nul pour une transmission nulle. Il doit donc forcément y avoir  $T$  en facteur global. Pour que l'expression totale soit du premier ordre en  $T$ , les développements de  $\epsilon^g$  et de  $\zeta$  doivent donc être restreints à l'ordre 0. Tout se passe donc comme si le NLA transformait  $T$  en  $\eta=g^2 T$ , sans modifier les autres paramètres. La transmission limite  $T_{\text{lim}}^{\text{NLA}}$  sera donc telle que  $g^2 T_{\text{lim}}^{\text{NLA}} = T_{\text{lim}}$ , ce qui redonne la formule (10.18).

Par souci de simplicité, nous supposons maintenant que la détection homodyne est parfaite. Les résultats précédents sont illustrés sur la figure 10.2, où l'on compare d'une part le taux obtenu avec un NLA de gain  $g=4$  et  $P_{\text{suc}}=1/g^2$  au taux obtenu sans NLA, et d'autre part l'accord entre les développements perturbatifs et les expressions numériques. On voit tout d'abord que ce dernier est très bon lorsque les pertes deviennent suffisamment importantes. Les valeurs des pertes maximales données par l'équation (10.16), pour des gains  $g=1$  et  $g=4$  sont respectivement de 28.1 et 40.1 dB, ce qui correspond parfaitement aux pertes maximales admissibles observées numériquement. Ceci confirme une tolérance à  $20 \log_{10} g=12$  dB de pertes supplémentaires avec le NLA.

On vérifie sur la figure 10.3 que  $g=4$  est inférieur au gain maximal autorisé  $g^{\text{max}}$  lorsque les pertes sont supérieures à environ 5 dB, pour  $\epsilon=0.05$ . La figure 10.2, avec des pertes commençant



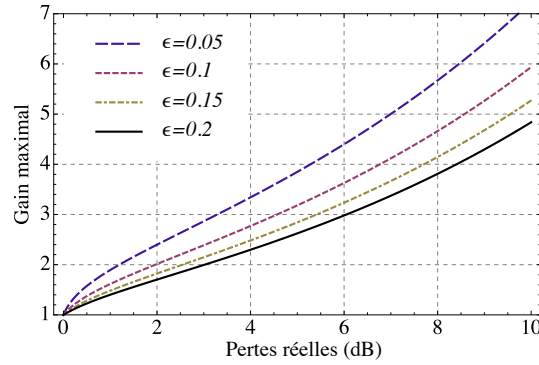


FIGURE 10.3 – Gain maximal  $g^{\max}$  donné par (9.55) en fonction des pertes en dB, pour plusieurs valeurs de  $\epsilon$ .

à 10 dB, correspond donc à un régime autorisé garantissant que  $\eta < 1$ .

### Amélioration de la tolérance au bruit

La distance maximale de transmission dépend du bruit ajouté par le canal. Plus celui-ci est important, et plus la distance maximale est faible. Pour des pertes données, on peut définir le bruit maximal  $\epsilon_{\max}$  tolérable au delà duquel le taux secret devient nul. La figure 10.4 montre sa dépendance en fonction de la transmission  $T$ , sans utiliser de NLA. On voit que  $\epsilon_{\max}$  dépend de  $\beta$  de manière cruciale : lorsque  $\beta$  diminue, la forme de la courbe reste globalement inchangée, mais diminuée d'un facteur supérieur à 2 lorsque  $\beta$  passe de 1 à 0.85.

Cette figure peut en fait être lue de deux manières : elle donne soit la transmission minimale  $T_{\lim}$  pour une valeur de  $\epsilon$ , soit le bruit maximal  $\epsilon_{\max}$  pour une valeur de  $T$ . Comme toutes les autres figures de ce type que nous utiliserons, elle est obtenue en calculant numériquement le taux secret pour une table de valeurs de  $T$  et  $\epsilon$  s'étendant sur une plage assez grande, pour lequel le taux prend des valeurs positives et négatives. Ces valeurs sont ensuite interpolées par une fonction polynomiale  $f(T, \epsilon)$  à deux paramètres, beaucoup plus facile à manipuler numériquement que (10.4). On peut en particulier tracer la courbe  $f(T, \epsilon) = 0$  en fonction de  $T$  et  $\epsilon$ , afin d'obtenir la figure 10.4.

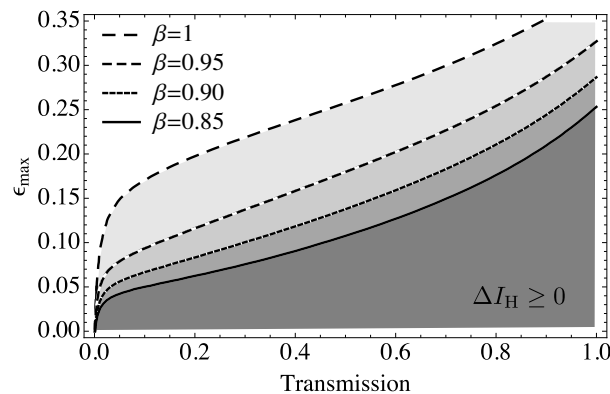


FIGURE 10.4 – Bruit  $\epsilon^{\max}$  à partir duquel le taux secret sans NLA s'annule, en fonction de la transmission  $T$ . Le taux est calculé avec (10.4) et maximisé numériquement sur  $\lambda$  pour chaque valeur de  $T$ .

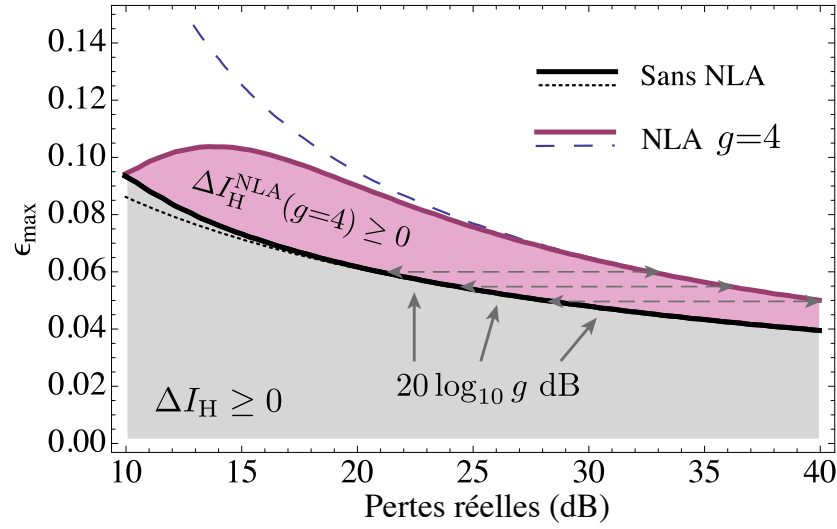


FIGURE 10.5 – Bruit à partir duquel le taux secret est nul, en fonction des pertes en dB, sans NLA et pour un NLA de gain  $g=4$ . Les tirets correspondent aux développements perturbatifs, et les traits pleins aux expressions numériques. Les courbes numériques sont obtenues avec les formules (10.4) et (10.6), maximisées numériquement sur  $\lambda$  ou  $\zeta$  pour chaque valeur de pertes. Les développements analytiques correspondent à la formule (10.12) (avec  $P_{\text{suc}}=1$  et  $g=1$  sans NLA), en utilisant  $\lambda_{\text{opt}}=0.8065$ . La détection homodyne est supposée parfaite, et  $\beta=0.95$ .

Notons également que même pour  $\beta=1$ , la distance de transmission est toujours finie dès que  $\epsilon>0$ , comme indiqué par (10.16).

La figure 10.5 montre  $\epsilon_{\text{max}}$  en fonction des pertes réelles en dB, sans NLA et pour un NLA de gain  $g=4$ . Par *réelles* nous entendons pertes du canal physiquement utilisé, dues à la transmission  $T$ . Cette précision est importante, comme nous le verrons par la suite. On compare également l'accord entre les expressions numériques et les développements perturbatifs. Ce dernier est excellent à partir d'environ 25 dB de pertes. En revanche, la tolérance au bruit est très optimiste pour le développement perturbatif avec le NLA pour de faibles pertes.

Pour une valeur de bruit fixée, on retrouve l'amélioration de  $20 \log_{10} g$  dB de pertes tolérables en présence du NLA dès que le développement perturbatif est valide. On voit clairement que cette amélioration ne dépend pas des pertes, comme nous l'avons vu analytiquement.

Pour une valeur de pertes fixées, on voit que le NLA augmente la tolérance au bruit (zone en violet). Cette amélioration n'est pas constante et dépend des pertes, mais elle peut être significative. Nous verrons en revanche qu'elle est limitée : il n'est pas toujours possible d'obtenir un taux secret positif lorsque le bruit est trop important, même en augmentant le gain.

La forme particulière de la courbe de  $\epsilon_{\text{max}}$  avec le NLA peut être comprise en s'aidant de la figure 10.6. Sur cette figure, la courbe  $\epsilon_{\text{max}}$  correspond au bruit maximal tolérable en fonction de la transmission réelle  $T$  (c'est une des courbes de la figure 10.4). Les courbes en tirets sont des courbes paramétriques du couple  $(\eta[T, \epsilon], \epsilon^g[T, \epsilon])$ , obtenues en faisant varier  $T$  et en gardant  $\epsilon$  constant : l'axe des abscisses correspond à la transmission effective  $\eta[T, \epsilon]$ , alors que celui des ordonnées au bruit effectif  $\epsilon^g[T, \epsilon]$ . Pour cette raison, les axes portent l'appellation “réelle” ou “effective”, selon la courbe qui est regardée. En vertu de

$$\Delta I_{\text{H}}^{\text{NLA}}(\lambda, T, \epsilon) \geq 0 \iff \Delta I_{\text{H}}(\zeta, \eta, \epsilon^g) \geq 0, \quad (10.21)$$

on peut en déduire graphiquement les zones pour lesquelles le taux secret sera positif avec le

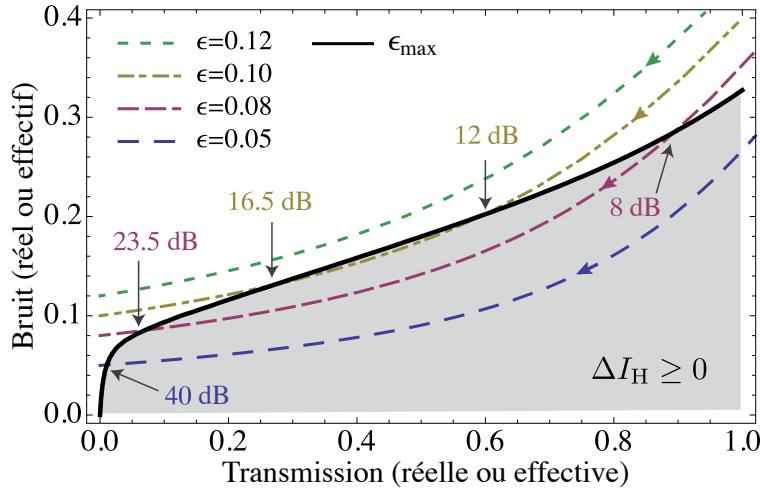


FIGURE 10.6 – Courbes paramétriques du couple  $(\eta[T, \epsilon], \epsilon^g[T, \epsilon])$  en faisant varier les pertes réelles, et en fixant le bruit réel  $\epsilon$ , pour un NLA de gain  $g=4$ . Les pertes réelles varient de la valeur correspondant à  $g^{\max}=4$ , jusqu'à 50 dB. Les flèches indiquent leur sens d'augmentation. Les valeurs indiquées en dB correspondent aux pertes réelles donnant la valeur des paramètres effectifs pointée par la flèche. Le trait plein correspond au bruit maximal  $\epsilon_{\max}$  sans NLA, pour la transmission réelle et pour  $\beta=0.95$ .

NLA : il suffit que  $\epsilon^g$  soit inférieur au  $\epsilon_{\max}$  sans NLA correspondant à  $\eta$ . En d'autres termes, il faut que la courbe  $(\eta, \epsilon^g)$  soit en dessous de  $\epsilon_{\max}$ .

Remarquons que  $\epsilon_{\max}$  est obtenu en optimisant  $\lambda$ , or c'est  $\zeta$  qui intervient dans (10.21). En fait, puisque l'on peut considérer que  $\zeta$  est un paramètre libre, il est identique d'optimiser  $\lambda$  afin de maximiser  $\Delta I_H^{\text{NLA}}(\lambda, T, \epsilon)$ , ou bien d'optimiser  $\zeta$  afin de maximiser  $\Delta I_H(\zeta, \eta, \epsilon^g)$ . En effet, seul  $\zeta$  dépend de  $\lambda$ , et en vertu de (10.6) et (9.51),

$$\max_{\lambda} \Delta I_H^{\text{NLA}}(\lambda, T, \epsilon) = P_{\text{suc}} \max_{\lambda} \Delta I_H(\zeta = g_{\text{in}} \lambda, \eta, \epsilon^g) \quad (10.22a)$$

$$= P_{\text{suc}} \max_{\zeta} \Delta I_H(\zeta, \eta, \epsilon^g). \quad (10.22b)$$

On peut donc bien directement comparer la courbe  $(\eta, \epsilon^g)$  à la courbe  $(\eta, \epsilon_{\max}[\eta])$  pour déterminer si le taux secret est positif avec le NLA. Lorsque  $\epsilon^g > \epsilon_{\max}$ , le bruit effectif est trop important, compte tenu de la transmission effective, pour donner un taux positif.

Voyons maintenant ce qu'il se passe en fonction du bruit  $\epsilon$  : si  $\epsilon$  est trop important, on voit que quelles que soient les pertes, les paramètres effectifs décrivent une trajectoire qui ne passe jamais dans une zone de taux positif (voir par exemple la courbe correspondant à  $\epsilon=0.12$ ). Pour des valeurs de bruit plus faibles (par exemple  $\epsilon=0.10$ ), il peut arriver que les paramètres effectifs rentrent dans une zone de taux positif à partir d'une certaine valeur de pertes (12 dB), puis en ressortent (16.5 dB), ce qui correspond exactement au comportement observé sur la figure 10.5. Enfin, pour un faible bruit (par exemple  $\epsilon=0.05$ ), les paramètres effectifs peuvent déjà se trouver dans une zone de taux positifs pour de faibles pertes. Dans tous les cas, ils finissent toujours par en ressortir : la valeur des pertes réelles correspondant à l'intersection avec  $\epsilon_{\max}$  correspond alors à  $T_{\text{lim}}^{\text{NLA}}$ . Par exemple pour  $\epsilon=0.05$ , les courbes s'intersectent pour environ 40 dB de pertes réelles, comme on s'y attendait selon la figure 10.2.

Les pertes réelles indiquées sur la figure 10.6 peuvent être obtenues facilement en utilisant la figure 10.7 : les pertes réelles  $y$  sont représentées en fonction de la transmission effective,

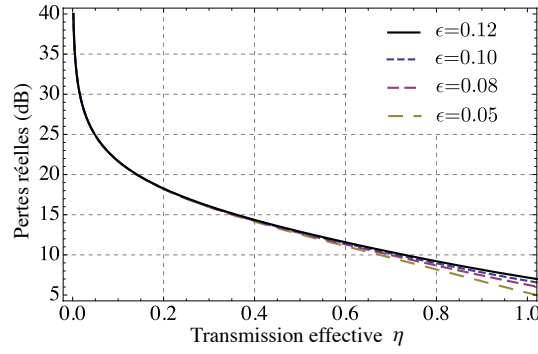


FIGURE 10.7 – Pertes réelles en dB permettant d’obtenir la transmission effective  $\eta$ , pour  $g=4$  et pour plusieurs valeurs de bruit.

pour plusieurs valeurs de bruit. Il suffit donc de repérer la transmission effective correspondant à l’intersection entre la courbe paramétrique et  $\epsilon_{\max}$  sur la figure 10.6, pour en déduire les pertes réelles correspondantes. Notons que le bruit n’a que très peu d’effet sur cette courbe, car on est très vite dans un régime où  $\eta \simeq g^2 T$ .

### Validité du développement perturbatif

Revenons sur les conditions de validité du développement perturbatif. La formule (10.18) n’impose pas de condition sur la valeur du gain, et permet donc en théorie d’augmenter arbitrairement la distance maximale de transmission. Bien que cette formule soit valable uniquement dans le cadre du développement au premier ordre en  $T$ , l’accord est en fait toujours vérifié avec la valeur numérique, même pour de très grands gains (Fig. 10.8). Un grand gain n’est de toute façon utilisable que pour des pertes suffisamment importantes de manière à assurer qu’il soit inférieur  $g^{\max}$ , et de ce fait la transmission est donc toujours suffisamment faible pour garder le développement perturbatif valide. La figure 10.8 montre que la distance maximale de transmission donnée par le développement perturbatif est toujours valable, même pour un gain très important. En revanche, il y a une zone de “transition” pour laquelle l’expression numérique et le développement ne sont pas en bon accord. En dessous d’environ 55 dB, le fait que le bruit effectif est trop important et ne permet pas d’avoir un taux positif pour ces valeurs de paramètres n’apparaît pas dans le développement perturbatif.

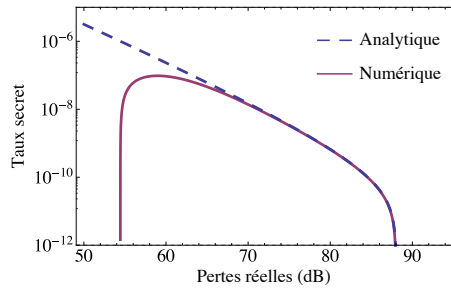


FIGURE 10.8 – Comparaison entre le développement perturbatif et l’expression numérique du taux secret, pour un gain  $g=1000$  et  $\epsilon=0.05$ . Ce gain est autorisé car  $g^{\max}=1000$  à partir d’environ 52.6 dB, et est supérieur à 1000 pour des pertes plus importantes. Sans NLA  $\lambda=0.8065$ , et avec le NLA  $\zeta=0.8065$ .

Enfin, la figure 10.9 montre le type de désaccord qu'il peut y avoir entre le développement perturbatif et l'expression numérique lorsque les pertes limites sont trop faibles. Elles sont dans ce cas plus petites que les pertes à partir desquelles le développement devient valide.

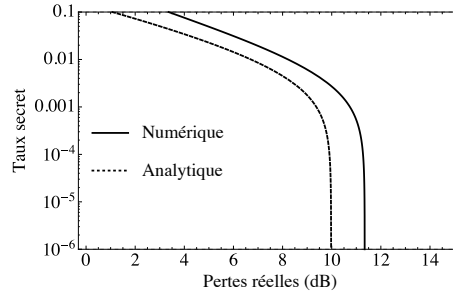


FIGURE 10.9 – Limite de validité du développement perturbatif, sans NLA. Pour  $\lambda=0.8065$ ,  $\epsilon=0.0865$  et  $\beta=0.95$ .

### 10.3.3 Que se passe t'il quand le gain augmente trop ?

Que se passe t'il lorsque, pour des pertes et un bruit fixés, on augmente le gain ? Le taux secret augmente t'il indéfiniment ? La réponse est négative. Une augmentation du gain conduit toujours à une décroissance du taux secret qui finit par le rendre négatif – sauf éventuellement pour certains cas où  $g^{\max}$  est atteint avant –.

Considérons par exemple le cas de la figure 10.10, pour lequel les pertes et le bruit sont fixés. Pour  $g=1$ , le taux secret est négatif. Compte tenu des conclusions des paragraphes précédents, si le bruit n'est pas trop important, un gain  $g$  prolonge les pertes admissibles de  $20 \log_{10} g$  dB. On s'attend donc à obtenir un taux secret positif à partir d'une certaine valeur du gain, ici égale à 1.62. En revanche, lorsque le gain est supérieur à 28, le taux redevient négatif.

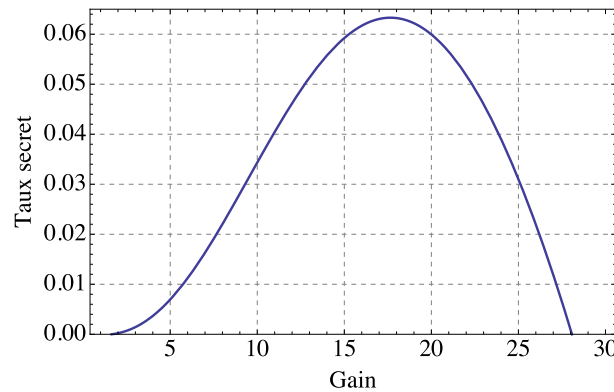


FIGURE 10.10 – Non augmentation arbitraire du taux secret en augmentant le gain. Le taux est calculé numériquement avec (10.6) en prenant  $\zeta=0.8065$ , 25 dB de pertes,  $\epsilon=0.06$ , et  $\beta=0.95$ . Une maximisation sur  $\lambda$  changerait simplement l'allure de la courbe et les valeurs de gain pour lesquels le taux secret est positif, mais l'interprétation physique est identique. La probabilité de succès n'est pas prise en compte pour ne pas modifier l'allure de la courbe par le facteur  $\frac{1}{g^2}$ , mais de toute façon cela ne changerait pas le comportement pour  $g=1.67$  et  $g=28$ .

Un coup d'oeil à la figure 10.11 nous fournit l'explication. Cette figure représente la courbe paramétrique du couple  $(\eta, \epsilon^g)$  en fonction du gain, ainsi que  $\epsilon_{\max}$  pour  $\beta=0.95$ . Comme pour la

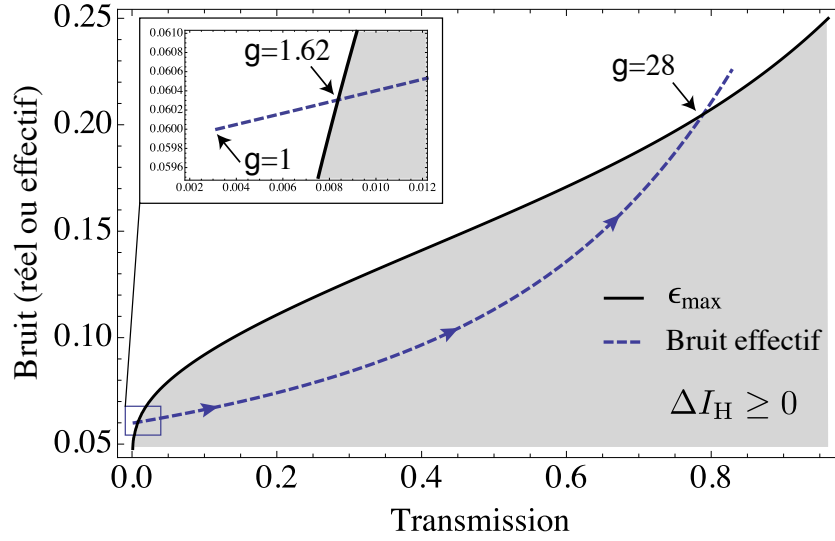


FIGURE 10.11 – Courbe paramétrique du couple  $(\eta, \epsilon^g)$  en faisant varier le gain, et en fixant le bruit réel  $\epsilon=0.06$  et les pertes réelles à 25 dB. Les flèches indiquent le sens d'augmentation du gain.  $\epsilon_{\max}$  est obtenu pour  $\lambda=0.8065$  et  $\beta=0.95$ .

figure 10.6, le taux secret avec le NLA est positif si les paramètres effectifs se trouvent en dessous de la courbe de  $\epsilon_{\max}$ . Pour  $g=1$ , l'encadré nous montre que cela n'est pas le cas, comme observé sur la figure 10.10. Puis les paramètres effectifs intersectent  $\epsilon_{\max}$  pour  $g=1.26$ , donnant un taux secret positif. Ils restent en dessous de  $\epsilon_{\max}$  pour une certaine plage de gains, mais intersectent à nouveau sa courbe pour  $g=28$ .

Autrement dit, lorsque le gain est trop grand, le bruit effectif ne permet pas d'avoir un taux secret positif pour la transmission effective correspondante. Les valeurs précises des paramètres effectifs sont données dans la table 10.1. Notons que ce phénomène n'apparaît pas avec les développements perturbatifs.

|              | $g=1.62$ | $g=28$ |
|--------------|----------|--------|
| $\eta$       | 0.00826  | 0.787  |
| $\epsilon^g$ | 0.0603   | 0.204  |

TABLE 10.1 – Gain et paramètres effectifs associés donnant un taux secret nul, pour la figure 10.10.

### 10.3.4 Augmentation arbitraire de la distance maximale de transmission

Nous avons vu qu'il est toujours possible de prolonger la distance maximale de transmission, à la condition toutefois que le bruit ne soit pas trop important. Nous avons vu également que si le gain est trop important, le bruit effectif peut être trop grand pour avoir un taux secret positif. Le gain à utiliser dépend donc des pertes réelles : nous allons maintenant voir comment l'adapter simplement afin d'éliminer ce problème. Ceci nous permettra en plus d'avoir une démarche indépendante de (10.18).

Commençons par étudier la forme des paramètres effectifs en supposant que  $g \gg 1$  et  $\epsilon \ll 1$ .

Pour la transmission effective  $\eta$  :

$$\eta = \frac{g^2 T}{1 + (g^2 - 1) T [\frac{1}{4} (g^2 - 1) (\epsilon - 2) \epsilon T - \epsilon + 1]} \simeq \frac{g^2 T}{1 + g^2 T [1 - \frac{1}{2} g^2 T \epsilon]} \quad (10.23)$$

Si l'on suppose que  $g^2 T \simeq 1$ , on peut négliger le terme  $g^2 T \epsilon / 2$  devant 1, et on obtient  $\eta \simeq \frac{1}{2}$ . En procédant de même pour  $\epsilon^g$ , et en supposant de plus que  $\epsilon \ll 1$  :

$$\epsilon^g = \epsilon + \frac{1}{2} (g^2 - 1) (2 - \epsilon) \epsilon T \simeq \epsilon (1 + g^2 T) \simeq 2\epsilon \quad (10.24)$$

En conclusion, nous voyons, de manière certes très qualitative, que lorsque  $g^2 T \simeq 1$  la transmission effective tend vers 1/2, alors que le bruit effectif tend vers deux fois le bruit initial. Et ceci, quelles que soient les pertes : il peut donc ne plus y avoir de distance maximale de transmission si le bruit effectif n'est pas trop important.

Raisonnons maintenant de manière plus précise. On utilise un gain  $g = \frac{1}{\sqrt{T}}$  (Fig. 10.12) dans les expressions complètes de  $\eta$  et de  $\epsilon^g$ , puis on fait tendre  $T$  vers 0 pour avoir une limite valable pour de fortes pertes :

$$\begin{aligned} \eta^* &= \lim_{g=\frac{1}{\sqrt{T}}, T \rightarrow 0} \eta = \frac{1}{2} \frac{1}{1 - \frac{3}{4}\epsilon + \frac{1}{8}\epsilon^2} \\ \epsilon^{g*} &= \lim_{g=\frac{1}{\sqrt{T}}, T \rightarrow 0} \epsilon^g = \left(2 - \frac{\epsilon}{2}\right) \epsilon \end{aligned} \quad (10.25)$$

Ces expressions ne supposent pas de condition particulière sur le bruit  $\epsilon$ . Dans le cas où il est petit devant 1, on retrouve bien sûr les résultats obtenus qualitativement. La transmission effective tend donc vers une valeur proche de 1/2, alors que le bruit effectif est proche du bruit initial multiplié par 2. Si, pour une transmission  $\eta^*$  et un bruit  $\epsilon^{g*}$ , le taux secret sans NLA est positif, il n'y a alors pas de distance maximale de transmission.

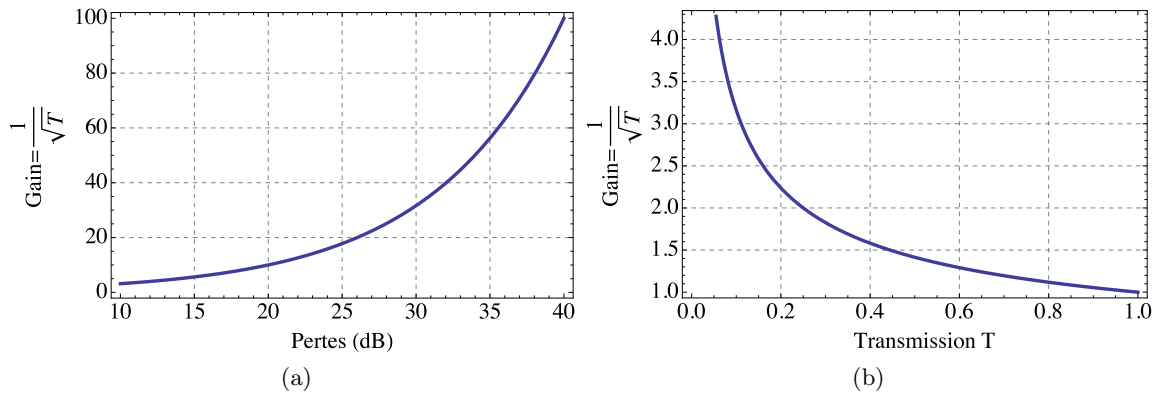


FIGURE 10.12 – Gain variable  $g = \frac{1}{\sqrt{T}}$  en fonction des pertes en dB (a) et de la transmission  $T$  (b).

La figure 10.13 illustre ce résultat. Elle représente le taux secret obtenu numériquement sans NLA, avec un NLA de gain 4, et avec un NLA de gain variable. On y voit clairement une disparition de la distance maximale de transmission. Le taux secret décroît néanmoins avec le gain variable, en raison de la probabilité de succès  $1/g^2 = T$  de plus en plus faible à mesure que les pertes augmentent.

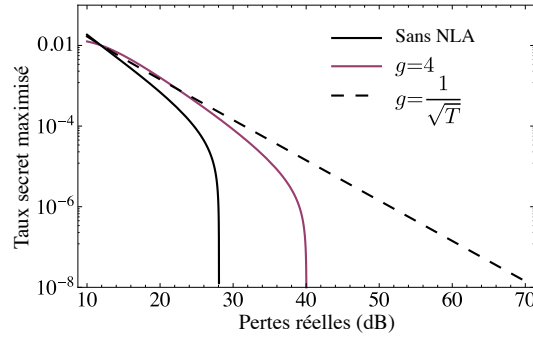


FIGURE 10.13 – Taux secret calculé numériquement avec (10.4) et (10.6), maximisé numériquement sur  $\lambda$  ou  $\zeta$ .  $\epsilon=0.05$  et  $\beta=0.95$ . La détection homodyne est supposée parfaite.

La figure 10.14 montre  $\epsilon_{\max}$  pour le gain variable, en reprenant également les courbes numériques sans NLA et pour un NLA de gain  $g=4$  de la figure 10.5. Avec le gain variable, le taux secret reste positif quelles que soient les pertes si  $\epsilon \leq 0.0965$ . Le gain variable n'est pas forcément le gain maximisant la tolérance au bruit, comme on le voit pour des pertes d'environ 15 dB, pour lesquelles  $g=4$  donne de meilleurs résultats. En revanche, il a l'avantage de donner des résultats facilement compréhensibles, et il est de plus toujours inférieur à  $g^{\max}$  pour les valeurs de bruit considérées (Fig. 10.15).

Le bruit maximal admissible avec le NLA de gain variable est facilement obtenu en s'aidant de la figure 10.16. Le taux secret est positif tant que  $\epsilon^{g^*}$  est inférieur à  $\epsilon_{\max}$ . On trouve la transmission effective limite  $\eta^*$  correspondant à l'intersection de  $\epsilon^{g^*}$  et  $\epsilon_{\max}$ , puis le bruit réel  $\epsilon$  correspondant. Pour  $\beta=0.95$ , le bruit effectif maximal est  $\epsilon^{g^*}=0.188$ , ce qui correspond à  $\epsilon=0.0965$  (en utilisant (10.25) , comme observé sur la figure 10.14).

Notons pour finir que le développement perturbatif ne donne pas de bons résultats lorsque l'on utilise le gain variable. Pour une raison quelque peu fortuite, le taux secret avec le gain variable correspond toujours à la “zone de transition” où les pertes sont trop faibles pour assurer une bonne validité du développement (Fig. 10.17).

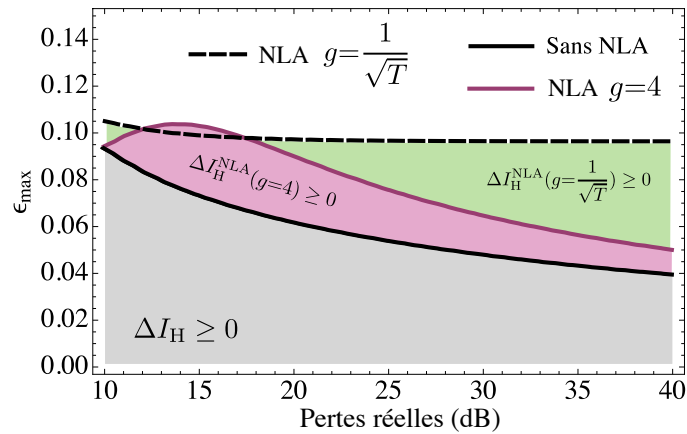


FIGURE 10.14 – Bruit à partir duquel le taux secret est nul, en fonction des pertes en dB, pour  $\beta=0.95$ . Les courbes numériques sont obtenues avec les formules (10.4) et (10.6), maximisées numériquement sur  $\lambda$  ou  $\zeta$  pour chaque valeur de pertes. La détection homodyne est supposée parfaite.



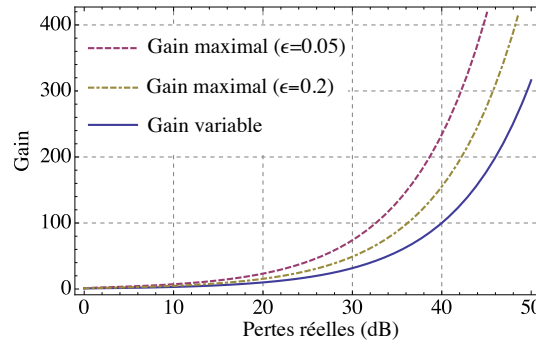


FIGURE 10.15 – Comparaison entre le gain variable  $g=1/\sqrt{T}$  et le gain maximal  $g^{\max}$  pour différentes valeurs de  $\epsilon$ .

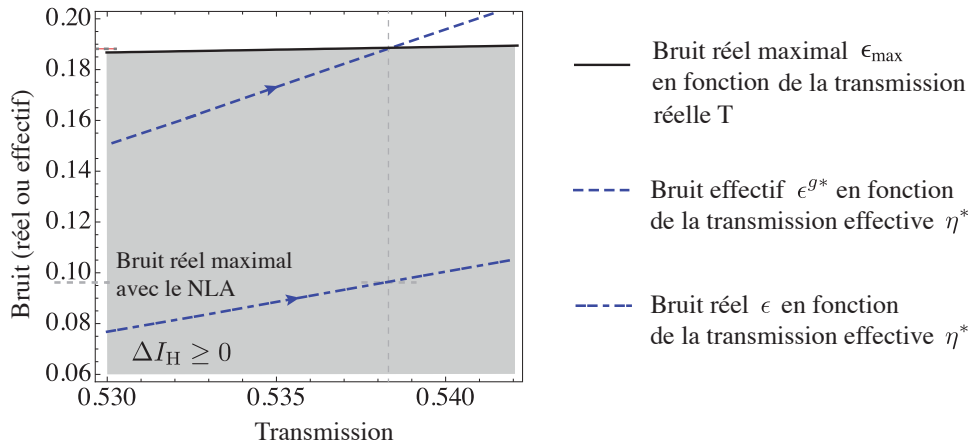


FIGURE 10.16 – Bruit maximal avec un NLA de gain variable : le bruit effectif maximal correspond à l'intersection de la courbe paramétrique du couple  $(\eta^*, \epsilon^{g*})$  en fonction de  $\epsilon$  et de  $\epsilon_{\max}$  (empruntée à la figure 10.4 pour  $\beta=0.95$ ). La courbe paramétrique du couple  $(\eta^*, \epsilon)$  en fonction de  $\epsilon$  donne alors le bruit réel  $\epsilon$  correspondant. Les flèches donnent le sens d'augmentation de  $\epsilon$ .

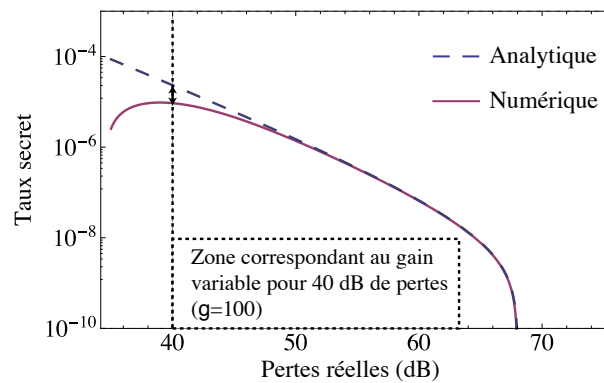


FIGURE 10.17 – Comparaison entre l'expression numérique et le développement perturbatif du taux secret dans la zone correspondant au gain variable (pour 40 dB de pertes). Pour  $\epsilon=0.05$  et  $\beta=0.95$ .

## 10.4 Attaques individuelles et non amélioration des performances avec le NLA

Notre étude de l'utilisation du NLA en cryptographie quantique avait initialement commencé en considérant le cas le plus simple : le taux secret contre les attaques individuelles, pour un canal qui n'ajoute pas de bruit. L'objectif était de voir dans un premier temps si le NLA pouvait, théoriquement, augmenter le taux secret. Les conclusions se sont vite révélées négatives : non seulement il n'y a pas de distance maximale de transmission dans ces conditions, et donc pas d'amélioration possible de ce côté là, mais en plus le NLA donne *toujours* un taux secret inférieur, même en prenant en compte la probabilité de succès maximale  $P_{\text{suc}}=1/g^2$ .

La conclusion se révèle identique lorsque le canal introduit du bruit : le NLA ne permet pas d'augmenter le taux secret contre les attaques individuelles. Cette section présente ces résultats, de manière numérique pour un canal avec bruit, et de manière analytique sans développement perturbatif pour un canal sans bruit.

### 10.4.1 Démonstration pour un canal ajoutant du bruit

#### Régime de fortes pertes

Un développement perturbatif du taux secret (10.5) au premier ordre en  $T$  donne :

$$\Delta I_{\text{Ind}}^{\text{NLA}} \simeq P_{\text{suc}} g^2 T \nu \frac{\beta \lambda^2 + (-2 + \beta + \epsilon) \lambda^4 - \epsilon}{(1 + \kappa)(1 - \lambda^4) \ln 2} \quad (10.26)$$

La variance de modulation optimale (*i.e.* le paramètre  $\lambda$  optimal) est obtenue en cherchant la valeur  $\lambda_{\text{opt}}^{\text{Ind}}$  qui annule la dérivée de (10.26) par rapport à  $\lambda$  :

$$\lambda_{\text{opt}}^{\text{Ind}} = \sqrt{\frac{2 - \beta - 2\sqrt{1 - \beta}}{\beta}} = \tanh \left[ \frac{1}{2} \text{sech}^{-1} \left[ \sqrt{1 - \beta} \right] \right] \quad (10.27)$$

Comme pour les attaques collectives,  $\lambda_{\text{opt}}^{\text{Ind}}$  ne dépend que de  $\beta$ . Il est représenté sur la figure 10.18.

En injectant  $\lambda_{\text{opt}}^{\text{Ind}}$  dans le développement du taux secret (10.26), on obtient la formule analytique du développement optimisé :

$$\Delta I_{\text{Ind}} \simeq \frac{g^2 \nu T}{(1 + \kappa) 2 \ln 2} \left[ 2 \left( 1 - \epsilon - \sqrt{1 - \beta} \right) - \beta \right] \quad (10.28)$$

Ce taux reste positif tant que  $2(1 - \epsilon - \sqrt{1 - \beta}) - \beta > 0$ , ce qui est satisfait pour :

$$\epsilon < \epsilon_{\text{max}}^{\text{Ind}} = 1 - \sqrt{1 - \beta} - \frac{\beta}{2} \quad (10.29)$$

Ce bruit maximal est représenté sur la figure 10.19. Il vaut environ 0.3 pour  $\beta=0.95$ .

En conclusion, cette étude perturbative nous montre que d'une part il n'y a pas de distance maximale de transmission lorsque  $\epsilon < \epsilon_{\text{max}}^{\text{Ind}}$ , même lorsque la détection homodyne est imparfaite, et d'autre part que le taux secret est simplement multiplié par  $g^2$  en utilisant le NLA. Une probabilité de succès inférieure à  $1/g^2$  ne pourra donc pas conduire à une augmentation du taux. Cette étude préliminaire n'étant valable que pour des fortes pertes, il reste à étudier l'effet du NLA pour des pertes plus faibles. Nous verrons que les conclusions seront toutefois identiques.

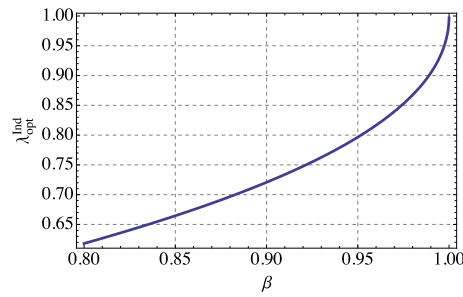


FIGURE 10.18 –  $\lambda_{\text{opt}}^{\text{Ind}}$  en fonction de  $\beta$ , donné par (10.27).  $\beta=0.95$  correspond à  $\lambda_{\text{opt}} \simeq 0.7965$ .

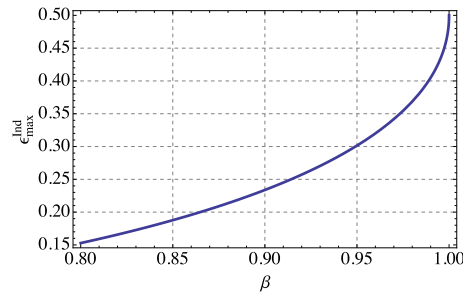


FIGURE 10.19 – Bruit maximal admissible  $\epsilon_{\text{max}}$ , donné par (10.29).

### Pertes quelconques

Un bruit supérieur à  $\epsilon_{\text{max}}^{\text{Ind}}$  induit une distance maximale de transmission. Dès lors, il serait envisageable que le NLA soit bénéfique dans ce régime. Comme observé sur la figure 10.20, il n'en est rien : contrairement au cas des attaques collectives, le NLA induit soit une baisse de la tolérance au bruit, soit au mieux aucune amélioration. Plus le gain est faible, et plus la courbe de  $\epsilon_{\text{max}}^{\text{Ind}}$  avec le NLA “suivra” celle de  $\epsilon_{\text{max}}^{\text{Ind}}$  sans le NLA.

On peut comprendre ce phénomène en s'aidant de la figure 10.21, qui se base sur un raisonnement similaire à celui de la figure 10.11. Des courbes paramétriques du couple  $(\eta, \epsilon^g)$  y sont représentées en fonction du gain, pour un exemple de 7.5 dB de pertes réelles et plusieurs valeurs de bruit. Pour  $g=1$ , la valeur maximale admissible de  $\epsilon$  est égale à  $\epsilon_{\text{max}}^{\text{Ind}} \simeq 0.35$  sans NLA. Pour un bruit inférieur, le taux secret est positif, comme observé sur la figure 10.20 pour 7.5 dB de pertes, sur la courbe sans NLA. On distingue ensuite deux comportements lorsque le gain augmente : si  $\epsilon$  était inférieur à  $\epsilon_{\text{max}}^{\text{Ind}}$ , une augmentation du gain finit par rendre  $\epsilon^g$  supérieur à  $\epsilon_{\text{max}}^{\text{Ind}}$ . La valeur du gain correspondant peut être obtenue en s'aidant de la figure 10.22. Dans ce cas, le NLA n'est non seulement d'aucune utilité, mais il est même désavantageux. Si  $\epsilon$  était supérieur à  $\epsilon_{\text{max}}^{\text{Ind}}$ , la courbe paramétrique de  $(\eta, \epsilon^g)$  n'intersecte jamais celle de  $\epsilon_{\text{max}}^{\text{Ind}}$  : le NLA ne permet donc pas de se ramener dans une zone de taux positif. Nous en concluons que le NLA n'améliore pas la résistance au bruit, quelles que soient les pertes, pour le cas des attaques individuelles.

On peut trouver le gain correspondant à un bruit maximal donné en s'aidant de la figure 10.22. Supposons par exemple que l'on veuille connaître le gain pour lequel  $\epsilon_{\text{max}}^{\text{Ind}}=0.252$ . Sur la figure 10.21, on trouve la transmission effective  $\eta=0.73$  correspondant à l'intersection de la courbe paramétrique pour  $\epsilon=0.252$  et la courbe  $\epsilon_{\text{max}}^{\text{Ind}}$ . On regarde ensuite le gain correspondant à cette transmission effective sur la figure 10.22, et on trouve  $g=2.6$ . Ceci est bien cohérent avec la figure 10.20, puisque  $\epsilon_{\text{max}}^{\text{Ind}}=0.252$  pour 7.5 dB de pertes et  $g=2.6$ .

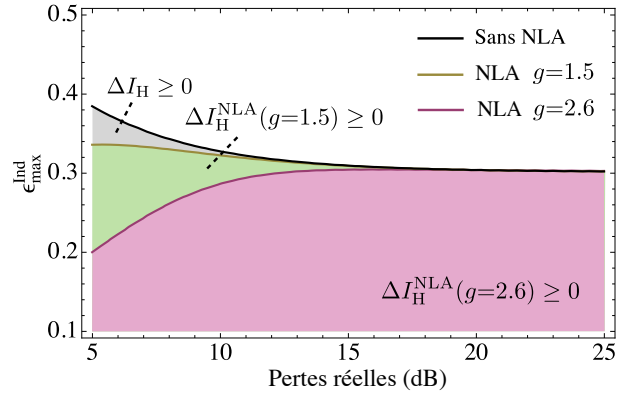


FIGURE 10.20 – Bruit à partir duquel le taux secret est nul, en fonction des pertes en dB, sans NLA, pour un NLA de gain  $g=1.5$ , et pour un NLA de gain  $g=4$ . Les courbes numériques sont obtenues avec (10.3) et (10.5), et maximisées numériquement sur  $\lambda$  ou  $\zeta$  pour chaque valeur de pertes. La détection homodyne est supposée parfaite, et  $\beta=0.95$ .

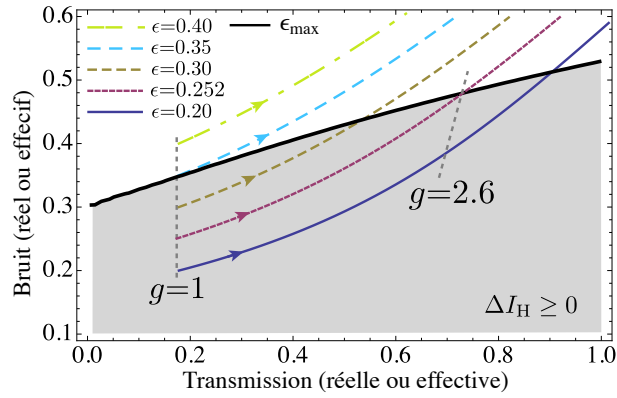


FIGURE 10.21 – Courbe paramétrique du couple  $(\eta, \epsilon^g)$  en fonction de  $g$  pour 7.5 dB de pertes et différentes valeurs de  $\epsilon$ , et comparaison avec  $\epsilon_{\max}^{\text{Ind}}$  sans NLA. Les flèches donnent le sens d'augmentation de  $g$ .  $\epsilon_{\max}^{\text{Ind}}$  est obtenu en maximisant numériquement sur  $\lambda$  pour chaque valeur de  $T$ , avec  $\beta=0.95$ .

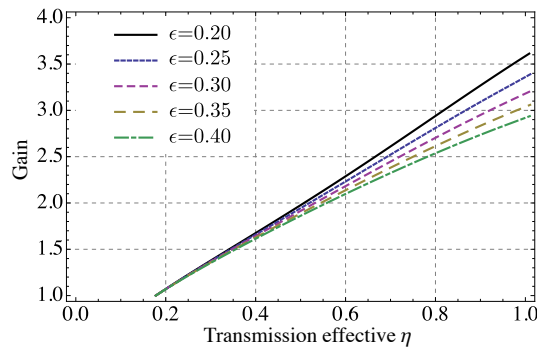


FIGURE 10.22 – Gain permettant d'obtenir la transmission effective  $\eta$ , pour 7.5 dB de pertes réelles et plusieurs valeurs de bruit.

### 10.4.2 Démonstration exacte pour tout $T$ , pour un canal sans bruit

Même si le NLA n'améliore pas la résistance au bruit, pourrait-il augmenter la valeur du taux secret, pour un canal sans bruit ? Ou plutôt, y serait-il théoriquement autorisé ? Car il est clair qu'une implémentation expérimentale aura sûrement une probabilité de succès inférieure à la borne maximale  $1/g^2$ . Comme on peut s'en douter compte tenu des conclusions précédentes, la réponse est également négative.

Dans ce qui suit, nous allons démontrer analytiquement, sans approximation ni développement perturbatif, que cette borne  $1/g^2$  garantit que le taux secret avec une variance de modulation optimale ne puisse pas être augmenté, quel que soit le gain et quelles que soient les pertes.

#### Maximisation du taux secret

Commençons par réécrire le taux secret sans NLA, donné par (10.3) pour  $\epsilon=0$  et une détection homodyne parfaite, en gardant cette fois-ci la variance de modulation  $V$  plutôt que le paramètre  $\lambda$  :

$$\Delta I_{\text{Ind}} = \frac{1}{2} \log_2 \left[ \frac{(\frac{V+B}{1+B})^\beta}{T^2(V+B)(\frac{1}{V}+B)} \right] = \frac{1}{2} \log_2 \left[ \frac{V[1+(V-1)T]^{\beta-1}}{V+T(1-V)} \right] \quad (10.30)$$

avec  $B = \frac{1-T}{T}$ . Dérivons maintenant l'expression dans le logarithme par rapport à  $V$ , afin de trouver la variance qui la maximise :

$$\frac{\partial}{\partial V} \left[ \frac{V[1+(V-1)T]^{\beta-1}}{V+T(1-V)} \right] = - \frac{T[(V-1)T+1]^{\beta-2} [V^2(\beta-1)(T-1)+T(1-\beta V)-1]}{[V+T(1-V)]^2} \quad (10.31)$$

Cette dérivée s'annule lorsque  $V^2(\beta-1)(T-1)+T(1-\beta V)-1=0$ . On trouve les solutions de ce polynôme sans difficulté :

$$V_{\pm} = \frac{\beta T \pm \sqrt{\beta^2 T^2 - 4(T-1)(\beta T - \beta - T + 1)}}{2(\beta T - \beta - T + 1)} \quad (10.32)$$

Il est facile de vérifier que  $V_-$  peut prendre des valeurs négatives, et ne correspond pas à une solution physique. Notons  $V_{\text{max}} = V_+ = \frac{1+\lambda_{\text{max}}^2}{1-\lambda_{\text{max}}^2}$  la variance de modulation optimale.

On peut également écrire  $\lambda_{\text{max}}$  en fonction de  $V_{\text{max}}$ , et faire un développement à l'ordre le plus faible en  $T$ ,

$$\lambda_{\text{max}}^{\text{DL}} \simeq \tanh \left[ \frac{1}{2} \text{sech}^{-1} \left[ \sqrt{1-\beta} \right] \right], \quad (10.33)$$

ce qui correspond exactement à (10.27), et confirme donc la validité de  $V_+$ . On note  $\Delta I_{\text{Ind}}^{\text{max}}(T)$  le taux secret maximisé sans NLA :

$$\Delta I_{\text{Ind}}^{\text{max}}(T) = \max_{\lambda} \Delta I_{\text{Ind}}(\lambda, T) = \Delta I_{\text{Ind}}(\lambda_{\text{max}}, T) \quad (10.34)$$

La dépendance de  $\lambda_{\text{max}}$  en  $T$  et  $\beta$  est sous entendue, et nous écrirons  $\lambda_{\text{max}}[T]$  lorsque cela sera nécessaire.

Le taux secret avec le NLA est maximisé lorsque  $\zeta = \lambda_{\max}[\eta]$ . Compte tenu de l'expression (9.59) selon laquelle  $\zeta = \lambda \sqrt{1 + (g^2 - 1)T}$ , le paramètre  $\lambda_{\max}^{\text{NLA}}$  est donc

$$\lambda_{\max}^{\text{NLA}}[T, g] = \frac{\lambda_{\max}[\eta]}{\sqrt{1 + (g^2 - 1)T}}. \quad (10.35)$$

Le taux secret maximisé avec le NLA s'écrit alors

$$\Delta I_{\text{Ind}}^{\text{NLA}, \max}(T) = \Delta I_{\text{Ind}}^{\text{NLA}}(\lambda_{\max}^{\text{NLA}}, T) = P_{\text{suc}} \Delta I_{\text{Ind}}(\lambda_{\max}[\eta], \eta). \quad (10.36)$$

### Comparaison des taux secrets

Il reste enfin à comparer les taux secrets avec et sans NLA. Nous allons montrer que (10.34) est toujours supérieur à (10.36), même en utilisant la probabilité de succès la plus optimiste  $P_{\text{suc}} = 1/g^2$ . On pose pour cela

$$\Xi = \frac{1}{g^2} \Delta I_{\text{Ind}}(\lambda_{\max}[\eta], \eta) - \Delta I_{\text{Ind}}(\lambda_{\max}[T], T), \quad (10.37)$$

afin de montrer que  $\Xi \leq 0$ . Pour ce faire, on commence par calculer la dérivée seconde de  $\Xi$  par rapport à  $\beta$  :

$$\frac{\partial^2}{\partial \beta^2} \Xi = -\frac{T}{(1-\beta) \ln 4} \left( \frac{1}{b} - \frac{1}{a} \right) \quad (10.38)$$

avec  $a = \sqrt{-4\beta(T-1)^2 + \beta^2 g^4 T^2 + 4(T-1)^2}$  et  $b = \sqrt{(\beta-2)^2 T^2 + 8(\beta-1)T - 4\beta + 4}$ . On vérifie ensuite que cette dérivée seconde est toujours négative : comme  $\beta < 1$ , son signe dépend uniquement de  $a - b$ . Puisque  $a^2 - b^2 = \beta^2 T^2 (g^4 - 1) \geq 0$ , on en déduit que la dérivée seconde (10.38) est donc toujours négative.

Ceci implique que la dérivée première de  $\Xi$  est décroissante avec  $\beta$ . Cette dérivée s'écrit :

$$\begin{aligned} \frac{\partial}{\partial \beta} \Xi = & \frac{1}{g^2 \ln 4} \ln \left[ \frac{T \left( \beta [(g^4 - 2)T + 2] + g^2 \sqrt{8(\beta-1)T - 4\beta + T^2 (\beta(\beta g^4 - 4) + 4) + 4 + 2T - 2} \right) + 2(\beta-1)(T-1)}{2([g^2 - 1]T + 1)(\beta-1)(T-1)} \right] \\ & - \frac{1}{g^2 \ln 4} g^2 \ln \left[ \frac{T \left( \sqrt{(\beta-2)^2 T^2 + 8(\beta-1)T - 4\beta + 4} + \beta T \right)}{2(\beta-1)(T-1)} + 1 - T \right] \end{aligned} \quad (10.39)$$

On montre sans difficulté que  $\lim_{\beta \rightarrow 0} \frac{\partial}{\partial \beta} \Xi = 0$ . Or comme la dérivée première est décroissante, on en conclut qu'elle est toujours négative. Finalement, on peut conclure que  $\Xi$  est une fonction monotone et décroissante de  $\beta$ . Comme  $\lim_{\beta \rightarrow 0} \Xi = 0$ , on en conclut que  $\Xi \leq 0$ , pour toutes valeurs de  $\beta$ ,  $g$ , et  $T$ .

Le taux secret contre les attaques individuelles pour un canal sans bruit est donc *toujours* inférieur en utilisant le NLA, même dans les conditions les plus optimistes.

## 10.5 Conclusion

Nous avons montré qu'un amplificateur sans bruit permet d'améliorer deux critères importants en cryptographie quantique : la distance maximale de transmission, et la tolérance au bruit introduit par le canal. Notre étude s'est concentrée sur le protocole GG02, un des plus simples à mettre en œuvre expérimentalement. Puisque tous les protocoles gaussiens sont unifiés dans une même représentation à intrication virtuelle où seules les mesures effectuées par Alice et Bob diffèrent [Weedbrook12], les paramètres effectifs sont bien sûr toujours les mêmes. Il en résulte alors des conclusions similaires concernant l'amélioration de la tolérance face aux pertes et au bruit.

Nous avons ici considéré un amplificateur parfait. Bien qu'impossible à réaliser strictement, il est quand même possible de s'en rapprocher autant que voulu. Notre étude constitue donc un cas limite à ce que l'on peut obtenir avec un amplificateur sans bruit. Toute implémentation plus réaliste donnera des améliorations dépendant du degré de ressemblance avec l'amplificateur parfait. Pour les types de modulations utilisées en cryptographie quantique, les ressources nécessaires à une bonne approximation ne semblent toutefois pas si déraisonnables : des simulations numériques montrent que, avec le théorème d'optimalité gaussienne et seulement 7 ou 8 étages dans le protocole à ciseaux quantiques, on peut obtenir un taux secret presque égal à celui obtenu avec l'amplificateur parfait.

Une implémentation *physique* de l'amplificateur ne semble d'ailleurs pas indispensable pour les protocoles à états cohérents, comme récemment introduit dans les références [Fiurášek12] et [Walk13]. Ces deux articles montrent que l'on peut reproduire l'action de l'amplificateur sans bruit *virtuellement* par post-sélection, en attribuant un certain poids aux données mesurées par Bob. Dans la référence [Fiurášek12], J. Fiurášek et N. Cerf se sont intéressés au cas où Bob utilise une détection hétérodyne. Le NLA est virtuellement reproduit en pondérant les mesures par une fonction gaussienne due au facteur  $\exp[(g^2-1)|\alpha|^2/2]$  lors de l'amplification d'un état cohérent. Dans la référence [Walk13], N. Walk *et al.* ont considéré le cas où Bob utilise une détection homodyne. Leur méthode de post-sélection est légèrement différente : le principe est de modifier la distribution (gaussienne) des mesures homodynes, afin d'obtenir une variance différente. Cette post-sélection permet de reproduire l'action d'un amplificateur sans bruit, mais contrairement à [Fiurášek12], il doit être suivi d'une amplification déterministe indépendante de la phase, et d'un ajout de bruit thermique avec une lame séparatrice : cette configuration est extrêmement proche de l'équivalence avec un NLA effectif suivi du canal effectif que nous avons développé au chapitre 9. Il semble donc que nos résultats puissent être utiles pour ce type d'implémentation virtuelle, et c'est sans doute une direction qui mérite une étude plus approfondie. Pour les deux implémentations [Fiurášek12] et [Walk13], une amélioration des performances similaire à celle montrée dans ce chapitre a été observée.

L'avantage de ces implémentations virtuelles est évident en terme de possibilité d'implémentation. Notons toutefois que [Fiurášek12] et [Walk13] considèrent une détection hétérodyne ou homodyne parfaite. Afin de pouvoir mettre en œuvre concrètement ces implémentations virtuelles, il serait nécessaire de savoir comment les imperfections modifient, le cas échéant, la fonction de filtrage à utiliser. Enfin, ces méthodes reproduisent des distributions gaussiennes à la limite où aucune mesure n'est conservée. Une implémentation réaliste introduira une non gaussianité, qui même faible, devra être prise en compte dans les preuves de sécurité.

# Chapitre 11

## Conclusion et perspectives

### Sommaire

|                                   |     |
|-----------------------------------|-----|
| <a href="#">11.1 Conclusion</a>   | 217 |
| <a href="#">11.2 Perspectives</a> | 218 |

### 11.1 Conclusion

Au cours de cette thèse, nous avons pu aborder plusieurs domaines variés de l'information quantique avec des variables continues.

Nous avons d'abord étudié la discordance quantique, qui permet de quantifier les corrélations de nature quantique dans un système. Nous avons démontré qu'il était possible d'obtenir une estimation proche des limites fondamentales données par les bornes de Cramér-Rao, en utilisant des mesures homodynes, qui sont parmi les plus simples que nous ayons à notre disposition.

Nous nous sommes ensuite intéressés au calcul quantique, en réalisant et en caractérisant une porte de phase. Nous avons proposé une méthode de caractérisation avant tout axée vers une faisabilité expérimentale, en combinant efficacement un modèle analytique et des tests expérimentaux. Nous avons ainsi pu souligner l'importance de la qualité des états utilisés en tant que qubits.

Enfin, nous avons mené une étude théorique sur les potentialités offertes par l'amplificateur sans bruit en information quantique. Nous avons d'abord démontré une équivalence avec un système effectif ouvrant la voie à de nombreuses applications, en n'étant pas limité au cadre des variables continues. Nous avons pu en mettre certaines en lumière, pour lesquelles une étude plus approfondie serait intéressante.

Nous avons ensuite utilisé nos résultats afin d'étudier l'intérêt de cet amplificateur en cryptographie quantique. Nous avons pu montrer qu'il permet d'accroître la distance de transmission, et d'extraire une clé secrète dans des conditions qui ne le permettraient pas sans lui. Ces améliorations ont quand même un coût : la probabilité de succès de l'amplification, qui limite drastiquement le taux secret que l'on peut obtenir. Cette probabilité de succès est même tellement faible que l'utilisation de l'amplificateur sans bruit n'est pas avantageuse lorsque l'on peut s'en passer pour avoir un taux secret positif.

Le taux secret est d'ailleurs un des freins au développement de la cryptographie quantique, qui peine à séduire en proposant une sécurité inconditionnelle à des débits de plusieurs ordres



de grandeurs inférieurs à ceux auxquels nous sommes habitués pour transférer des données dont le cryptage est certes, en théorie, déchiffable, mais qui procure quand même une très bonne sécurité dans l'état actuel de la technologie. Pour contourner ce problème tout en profitant des avantages de pouvoir créer une clé secrète à distance, des utilisations hybrides se développent : la clé n'est plus utilisée pour un codage parfaitement sûr, mais dans des algorithmes de chiffrement symétrique. Le taux secret devient alors moins crucial, puisque la taille des clés nécessaires est beaucoup plus faible. En revanche, la distance maximale de transmission reste toujours un problème. Et c'est justement ce critère que l'amplificateur sans bruit permet d'améliorer.

## 11.2 Perspectives

Le dispositif expérimental est actuellement en cours de modification afin d'améliorer la qualité des états produits, et de pouvoir les utiliser pour des opérations plus complexes. Un amplificateur optique – déterministe – utilisant un laser de pompe est à l'étude, et pourra sans doute être utilisé afin d'augmenter la puissance de nos impulsions, ce qui permettra de diminuer les focalisations dans les cristaux et d'améliorer la pureté des états.

Une autre perspective concerne bien sûr l'amplificateur sans bruit. De nombreuses applications restent sans doute à découvrir, tant son potentiel apparaît important. Peut-être permettra-t-il d'améliorer d'autres protocoles, voire de réaliser certaines expériences pour la première fois, comme un test des inégalités de Bell sans échappatoire. Plusieurs protocoles ont été proposés à cette fin, et c'est sans doute une piste prometteuse.

Un autre axe de recherche concerne la faisabilité expérimentale des diverses applications de l'amplificateur sans bruit. Même s'il peut être utile dans sa version idéale, il est nécessaire d'étudier si ces bénéfices sont toujours présents avec une version "expérimentale", tronquée à quelques étages pour l'implémentation avec les ciseaux quantiques, ou limitée à quelques additions et soustractions pour l'approximation polynomiale, et en prenant en compte les diverses imperfections.

Peu d'expériences ont été faites dans le cadre d'une utilisation du NLA. La plupart des expériences ont été des démonstrations de sa faisabilité, qui étaient nécessaires étant donné ses propriétés surprenantes. L'implémentation virtuelle semble tout à fait réalisable puisqu'elle ne nécessite qu'un traitement des données adéquat, mais à condition de savoir comment prendre en compte les imperfections de la détection homodyne ou hétérodyne.

D'autres expériences simples, telles que la suppression des pertes en ajoutant du bruit thermique comme nous l'avons proposé dans ce manuscrit, semblent tout à fait réalisables avec notre dispositif expérimental, et pourraient également constituer un axe de recherche intéressant.

## Quatrième partie

### Annexes



## Annexe A

# Quelques propriétés opératoriels bien utiles

### A.1 Evolution d'une fonction d'opérateurs

Soit  $\hat{A}$  un opérateur quelconque et  $\hat{S}$  un opérateur inversible, qui ne soit pas forcément unitaire. Pour une fonction d'opérateur  $f$  décomposable en série, on a la relation<sup>1</sup> [Puri01] :

$$\boxed{\hat{S}f[\hat{A}]\hat{S}^{-1} = f[\hat{S}\hat{A}\hat{S}^{-1}]} \quad (\text{A.1})$$

Cette relation est fondamentale. Elle permet par exemple de calculer la transformation des opérateurs couramment utilisés (déplacement, squeezing,...) sous l'action d'une évolution libre ou d'une lame séparatrice.

La démonstration est relativement simple. En décomposant  $f[A]$  en série,  $\hat{S}f[\hat{A}]\hat{S}^{-1}$  sera composé de termes du type  $\hat{S}\hat{A}^m\hat{S}^{-1}$ . En introduisant l'identité  $\mathbb{I}=\hat{S}\hat{S}^{-1}$  entre chaque produit de  $\hat{A}$ , on obtient  $\hat{S}\hat{A}^m\hat{S}^{-1}=[\hat{S}\hat{A}\hat{S}^{-1}]^m$ . La série donne ensuite  $f[\hat{S}\hat{A}\hat{S}^{-1}]$ .

### A.2 Evolution des opérateurs $\hat{a}$ et $\hat{a}^\dagger$ sous l'action d'un hamiltonien quadratique

Afin de démontrer simplement (2.123a), on peut utiliser une technique inspirée de la référence<sup>2</sup> [Puri01]. Définissons l'opérateur  $\hat{F}=\alpha_1\hat{a}+\alpha_2\hat{a}^\dagger+\alpha_3\hat{a}^\dagger\hat{a}$ , et cherchons l'opérateur  $\hat{a}(\theta)$  tel que

$$\hat{a}(\theta) = \exp(-\theta\hat{F})\hat{a}\exp(\theta\hat{F}). \quad (\text{A.2})$$

Pour cela, dérivons l'expression précédente par rapport à  $\theta$  :

$$\frac{d}{d\theta}\hat{a}(\theta) = -\exp(-\theta\hat{F})\hat{F}\hat{a}\exp(\theta\hat{F}) + \exp(-\theta\hat{F})\hat{a}\hat{F}\exp(\theta\hat{F}) \quad (\text{A.3a})$$

$$= \exp(-\theta\hat{F})[\hat{a}, \hat{F}]\exp(\theta\hat{F}) \quad (\text{A.3b})$$

$$= \exp(-\theta\hat{F})(\alpha_2 + \alpha_3\hat{a})\exp(\theta\hat{F}) \quad (\text{A.3c})$$

$$= \alpha_2 + \alpha_3\hat{a}(\theta) \quad (\text{A.3d})$$

---

1. Voir en particulier page 39.

2. Voir en particulier page 41 .

La solution de cette équation, compte tenu de  $\hat{\mathbf{a}}(\theta=0)=\hat{\mathbf{a}}$ , est :

$$\hat{\mathbf{a}}(\theta) = \hat{\mathbf{a}} \exp(\theta\alpha_3) + \frac{\alpha_2}{\alpha_3} [\exp(\alpha_3\theta) - 1] \quad (\text{A.4})$$

On pourrait être tenté de prendre directement l'hermitique conjugué de (A.4) afin d'obtenir l'évolution de  $\hat{\mathbf{a}}^\dagger$ . Ceci n'est malheureusement pas possible, car dans le cas général  $[\exp(\theta\hat{\mathbf{F}})]^\dagger$  n'est pas forcément égal à  $[\exp(\theta\hat{\mathbf{F}})]^{-1} = \exp(-\theta\hat{\mathbf{F}})$ . En appliquant la même méthode, on montre que pour

$$\hat{\mathbf{a}}^\dagger(\theta) = \exp(-\theta\hat{\mathbf{F}})\hat{\mathbf{a}}^\dagger\exp(\theta\hat{\mathbf{F}}), \quad (\text{A.5})$$

on a

$$\frac{d}{d\theta}\hat{\mathbf{a}}^\dagger(\theta) = -\alpha_1 - \alpha_3\hat{\mathbf{a}}^\dagger(\theta), \quad (\text{A.6})$$

dont la solution est donnée par

$$\hat{\mathbf{a}}^\dagger(\theta) = \hat{\mathbf{a}}^\dagger \exp(-\theta\alpha_3) + \frac{\alpha_1}{\alpha_3} [\exp(-\alpha_3\theta) - 1]. \quad (\text{A.7})$$

### Application à une rotation

Pour une évolution libre ou une rotation d'un angle  $\phi$ , on a  $\hat{\mathbf{F}} = -i\phi\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}}$ . En posant  $\theta=1, \alpha_3=-i\phi$  et  $\alpha_1=\alpha_2=0$ , on obtient :

$$e^{i\phi\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}}} \hat{\mathbf{a}} e^{-i\phi\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}}} = \hat{\mathbf{a}} e^{-i\phi} \quad (\text{A.8a})$$

$$e^{i\phi\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}}} \hat{\mathbf{a}}^\dagger e^{-i\phi\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}}} = \hat{\mathbf{a}}^\dagger e^{+i\phi} \quad (\text{A.8b})$$

### Application à un déplacement

Pour un déplacement de  $\alpha$ , on a  $\hat{\mathbf{F}} = \alpha\hat{\mathbf{a}}^\dagger - \alpha^*\hat{\mathbf{a}}$ . En posant  $\theta=1, \alpha_1 = -\alpha^*, \alpha_2=\alpha$  et  $\alpha_3=0$ , on obtient (2.123a) et (2.123b) :

$$\hat{\mathbf{D}}^\dagger(\alpha)\hat{\mathbf{a}}\hat{\mathbf{D}}(\alpha) = \hat{\mathbf{a}} + \alpha \quad (\text{A.9a})$$

$$\hat{\mathbf{D}}^\dagger(\alpha)\hat{\mathbf{a}}^\dagger\hat{\mathbf{D}}(\alpha) = \hat{\mathbf{a}}^\dagger + \alpha^* \quad (\text{A.9b})$$

## A.3 “Désintriquer” une exponentielle

### A.3.1 Formule de Baker-Hausdorff

Lorsque deux opérateurs  $\hat{\mathbf{A}}$  et  $\hat{\mathbf{B}}$  commutent avec leur commutateur  $\hat{\mathbf{C}}=[\hat{\mathbf{A}},\hat{\mathbf{B}}]$ , on a la relation [Cohen-Tannoudji97c] :

$$\exp(\theta[\hat{\mathbf{A}} + \hat{\mathbf{B}}]) = \exp(\theta\hat{\mathbf{B}}) \exp(\theta\hat{\mathbf{A}}) \exp(\theta^2\hat{\mathbf{C}}/2) \quad (\text{A.10a})$$

$$= \exp(\theta\hat{\mathbf{A}}) \exp(\theta\hat{\mathbf{B}}) \exp(-\theta^2\hat{\mathbf{C}}/2) \quad (\text{A.10b})$$

### Application à l'opérateur de déplacement

Pour l'opérateur de déplacement  $\hat{D}(\alpha) = \exp[\alpha\hat{a}^\dagger - \alpha^*\hat{a}]$ , on identifie  $\theta=1$ ,  $\hat{A}=\alpha\hat{a}^\dagger$ ,  $\hat{B}=-\alpha^*\hat{a}$ , et  $\hat{C}=|\alpha|^2$ . La formule (A.10) donne donc :

$$\hat{D}(\alpha) = e^{-\frac{1}{2}|\alpha|^2} e^{\alpha\hat{a}^\dagger} e^{-\alpha^*\hat{a}} \quad (\text{A.11a})$$

$$= e^{+\frac{1}{2}|\alpha|^2} e^{-\alpha^*\hat{a}} e^{\alpha\hat{a}^\dagger} \quad (\text{A.11b})$$

### A.3.2 Relations de commutation de SU(2) et lame séparatrice

Soient  $\hat{J}_+$ ,  $\hat{J}_-$  et  $\hat{J}_3$  trois opérateurs qui vérifient les relations de commutation de SU(2), correspondant à un moment cinétique :

$$[\hat{J}_3, \hat{J}_\pm] = \pm \hat{J}_\pm \quad (\text{A.12})$$

$$[\hat{J}_+, \hat{J}_-] = 2\hat{J}_3 \quad (\text{A.13})$$

On a alors la relation [Barnett03, Truax85] :

$$\exp[\lambda_+ \hat{J}_+ + \lambda_- \hat{J}_- + \lambda_3 \hat{J}_3] = \exp[\Lambda_+ \hat{J}_+] \exp[\ln(\Lambda_3) \hat{J}_3] \exp[\Lambda_- \hat{J}_-] \quad (\text{A.14a})$$

$$= \exp[\Lambda_- \hat{J}_-] \exp[-\ln(\Lambda_3) \hat{J}_3] \exp[\Lambda_+ \hat{J}_+] \quad (\text{A.14b})$$

avec

$$\Lambda_3 = \left( \cosh \alpha - \frac{\lambda_3}{2\alpha} \sinh \alpha \right)^{-2} \quad (\text{A.15})$$

$$\Lambda_\pm = \frac{2\lambda_\pm \sinh \alpha}{2\alpha \cosh \alpha - \lambda_3 \sinh \alpha} \quad (\text{A.16})$$

$$\alpha^2 = \frac{1}{4} \lambda_3^2 + \lambda_+ \lambda_- \quad (\text{A.17})$$

### Application à la lame séparatrice

L'opérateur  $\hat{U}_{BS}(\theta) = \exp[\theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger)]$  associé à une lame séparatrice peut être décomposé en utilisant ces relations, en identifiant  $\hat{J}_+ = \hat{a}^\dagger \hat{b}$ ,  $\hat{J}_- = \hat{a} \hat{b}^\dagger$  et  $\hat{J}_3 = \frac{1}{2}(\hat{a}^\dagger \hat{a} - \hat{b}^\dagger \hat{b})$ . On a  $\lambda_3=0$ ,  $\lambda_\pm = \pm \theta$ , et donc  $\alpha = \pm i\theta$ ,  $\Lambda_3 = \cos^{-2} \theta$  et  $\Lambda_\pm = \pm \tan \theta$ . On peut donc écrire :

$$\exp[\theta(\hat{a}^\dagger \hat{b} - \hat{a} \hat{b}^\dagger)] = \exp[\hat{a}^\dagger \hat{b} \tan \theta] \exp[-(\hat{a}^\dagger \hat{a} - \hat{b}^\dagger \hat{b}) \ln(\cos \theta)] \exp[-\hat{a} \hat{b}^\dagger \tan \theta] \quad (\text{A.18a})$$

$$= \exp[-\hat{a} \hat{b}^\dagger \tan \theta] \exp[(\hat{a}^\dagger \hat{a} - \hat{b}^\dagger \hat{b}) \ln(\cos \theta)] \exp[\hat{a}^\dagger \hat{b} \tan \theta] \quad (\text{A.18b})$$

### A.3.3 Relations de commutation de SU(1,1) et opérateurs de squeezing

Soient  $\hat{K}_+$ ,  $\hat{K}_-$  et  $\hat{K}_3$  trois opérateurs qui satisfont les relations de commutation de SU(1,1) :

$$[\hat{K}_3, \hat{K}_\pm] = \pm \hat{K}_\pm \quad (\text{A.19})$$

$$[\hat{K}_+, \hat{K}_-] = -2\hat{K}_3 \quad (\text{A.20})$$

On a alors la relation [Barnett03, Truax85] :

$$\exp \left[ \gamma_+ \hat{\mathbf{K}}_+ + \gamma_- \hat{\mathbf{K}}_- + \gamma_3 \hat{\mathbf{K}}_3 \right] = \exp \left[ \Gamma_+ \hat{\mathbf{K}}_+ \right] \exp \left[ \ln(\Gamma_3) \hat{\mathbf{K}}_3 \right] \exp \left[ \Gamma_- \hat{\mathbf{K}}_- \right] \quad (\text{A.21a})$$

$$= \exp \left[ \Gamma_- \hat{\mathbf{K}}_- \right] \exp \left[ -\ln(\Gamma_3) \hat{\mathbf{K}}_3 \right] \exp \left[ \Gamma_+ \hat{\mathbf{K}}_+ \right] \quad (\text{A.21b})$$

avec

$$\Gamma_3 = \left( \cosh \beta - \frac{\gamma_3}{2\beta} \sinh \beta \right)^{-2} \quad (\text{A.22})$$

$$\Gamma_{\pm} = \frac{2\gamma_{\pm} \sinh \beta}{2\beta \cosh \beta - \gamma_3 \sinh \beta} \quad (\text{A.23})$$

$$\beta^2 = \frac{1}{4} \gamma_3^2 - \gamma_+ \gamma_- \quad (\text{A.24})$$

Ces relations permettent de décomposer les opérateurs de squeezing monomode et bimode, qui sont tous les deux constitués d'opérateurs vérifiant les relations (A.19) et (A.20).

### Application au squeezing monomode

Pour l'opérateur de squeezing monomode  $\hat{\mathbf{S}}(r) = \exp \left[ \frac{r}{2} \hat{\mathbf{a}}^2 - \frac{r}{2} (\hat{\mathbf{a}}^\dagger)^2 \right]$ , on identifie  $\hat{\mathbf{K}}_+ = \frac{1}{2} (\hat{\mathbf{a}}^\dagger)^2$ ,  $\hat{\mathbf{K}}_- = \hat{\mathbf{K}}_+^\dagger$ , et  $\hat{\mathbf{K}}_3 = \frac{1}{4} (2\hat{\mathbf{a}}^\dagger \hat{\mathbf{a}} + 1)$ . On a de plus  $\gamma_3 = 0, \gamma_{\pm} = \mp r$ , et donc  $\beta = r, \Gamma_3 = (\cosh r)^{-2}$ , et  $\Gamma_{\pm} = \mp \tanh r$ . On peut donc écrire :

$$\exp \left[ \frac{r}{2} \hat{\mathbf{a}}^2 - \frac{r}{2} (\hat{\mathbf{a}}^\dagger)^2 \right] = \exp \left[ -\frac{1}{2} (\hat{\mathbf{a}}^\dagger)^2 \tanh r \right] \exp \left[ -\frac{1}{2} (2\hat{\mathbf{a}}^\dagger \hat{\mathbf{a}} + 1) \ln(\cosh r) \right] \exp \left[ +\frac{1}{2} \hat{\mathbf{a}}^2 \tanh r \right] \quad (\text{A.25})$$

L'application sur le vide permet plusieurs simplifications, car  $\exp \left[ +\frac{1}{2} (\tanh r) \hat{\mathbf{a}}^2 \right] |0\rangle = |0\rangle$  et  $\exp \left[ -\frac{1}{2} \ln(\cosh r) (2\hat{\mathbf{a}}^\dagger \hat{\mathbf{a}} + 1) \right] |0\rangle = \frac{1}{\sqrt{\cosh r}} |0\rangle$ . On obtient finalement la décomposition du vide comprimé  $\hat{\mathbf{S}}(r)|0\rangle$  en base de Fock :

$$\exp \left[ \frac{r}{2} \hat{\mathbf{a}}^2 - \frac{r}{2} (\hat{\mathbf{a}}^\dagger)^2 \right] |0\rangle = \frac{1}{\sqrt{\cosh r}} \exp \left[ -\frac{1}{2} \tanh r (\hat{\mathbf{a}}^\dagger)^2 \right] |0\rangle \quad (\text{A.26a})$$

$$= \frac{1}{\sqrt{\cosh r}} \sum_{n=0}^{\infty} \frac{1}{n!} \sqrt{(2n)!} \left( -\frac{1}{2} \tanh r \right)^n |2n\rangle \quad (\text{A.26b})$$

### Application au squeezing bimode

Pour l'opérateur de squeezing bimode  $\hat{\mathbf{S}}_2(r) = \exp \left[ r \hat{\mathbf{a}} \hat{\mathbf{b}} - r \hat{\mathbf{a}}^\dagger \hat{\mathbf{b}}^\dagger \right]$ , on identifie  $\hat{\mathbf{K}}_+ = \hat{\mathbf{a}}^\dagger \hat{\mathbf{b}}^\dagger, \hat{\mathbf{K}}_- = \hat{\mathbf{K}}_+^\dagger$ , et  $\hat{\mathbf{K}}_3 = \frac{1}{2} (\hat{\mathbf{a}}^\dagger \hat{\mathbf{a}} + \hat{\mathbf{b}}^\dagger \hat{\mathbf{b}} + 1)$ . On a également  $\gamma_3 = 0, \gamma_{\pm} = \mp r$ , et donc  $\beta = r, \Gamma_3 = (\cosh r)^{-2}$ , et  $\Gamma_{\pm} = \mp \tanh r$ .

On peut donc écrire :

$$\exp \left[ r \hat{\mathbf{a}} \hat{\mathbf{b}} - r \hat{\mathbf{a}}^\dagger \hat{\mathbf{b}}^\dagger \right] = \exp \left[ -\hat{\mathbf{a}}^\dagger \hat{\mathbf{b}}^\dagger \tanh r \right] \exp \left[ -(\hat{\mathbf{a}}^\dagger \hat{\mathbf{a}} + \hat{\mathbf{b}}^\dagger \hat{\mathbf{b}} + 1) \ln(\cosh r) \right] \exp \left[ +\hat{\mathbf{a}} \hat{\mathbf{b}} \tanh r \right] \quad (\text{A.27})$$

On obtient facilement la décomposition du vide comprimé bimode en base de Fock, compte tenu de  $\exp \left[ +\hat{\mathbf{a}}\hat{\mathbf{b}} \tanh r \right] |0\rangle|0\rangle = |0\rangle|0\rangle$ , et  $\exp \left[ -(\hat{\mathbf{a}}^\dagger\hat{\mathbf{a}} + \hat{\mathbf{b}}^\dagger\hat{\mathbf{b}} + 1) \ln(\cosh r) \right] |0\rangle|0\rangle = \frac{1}{\cosh r} |0\rangle|0\rangle$  :

$$\exp \left[ r\hat{\mathbf{a}}\hat{\mathbf{b}} - r\hat{\mathbf{a}}^\dagger\hat{\mathbf{b}}^\dagger \right] |0\rangle|0\rangle = \frac{1}{\cosh r} \exp \left[ -\hat{\mathbf{a}}^\dagger\hat{\mathbf{b}}^\dagger \tanh r \right] |0\rangle|0\rangle \quad (\text{A.28a})$$

$$= \frac{1}{\cosh r} \sum_{n=0}^{\infty} (-\tanh r)^n |n\rangle|n\rangle \quad (\text{A.28b})$$

$$= \sqrt{1-\lambda^2} \sum_{n=0}^{\infty} (-\lambda)^n |n\rangle|n\rangle \quad (\text{A.28c})$$

avec  $\lambda = \tanh r$ .





## Annexe B

# Modèle multimode simplifié de l'OPA

### B.1 Modélisation multimode et lien avec le modèle empirique

Le modèle empirique de l'OPA imparfait se justifie bien d'un point de vue expérimental, puisqu'il donne un très bon accord avec les mesures. Afin de comprendre l'origine physique de ses paramètres, nous développons dans cette annexe un traitement multimode simplifié de l'OPA, que nous relierons ensuite au modèle empirique.

Reprenons l'hamiltonien d'interaction multimode (3.6), en ne considérant qu'un traitement fréquentiel en configuration dégénérée pour simplifier, et en ne prenant en compte qu'un seul mode spatial transverse. On obtient dans ce cas [Rohde07] :

$$\hat{\mathcal{W}} = -i\hbar \int d\omega_1 d\omega_2 \psi(\omega_1, \omega_2) \hat{\mathbf{a}}^\dagger(\omega_1) \hat{\mathbf{a}}^\dagger(\omega_2) + \text{h.c.} \quad (\text{B.1})$$

Sous cette forme, le squeezing n'apparaît pas clairement à cause des corrélations entre les termes à différentes fréquences. Une réécriture très utile fait appel à la décomposition de Schmidt [Parker00, Law00, Lamata05]. On montre en effet que toute fonction de deux variables  $\psi(\omega_1, \omega_2)$  suffisamment régulière peut se décomposer sur deux bases de fonctions orthonormales à une variable :

$$\psi(\omega_1, \omega_2) = \sum_k c_k \phi_k(\omega_1) \varphi_k(\omega_2) \quad (\text{B.2})$$

Pour une configuration dégénérée, les deux bases de fonctions sont identiques,  $\psi(\omega_1, \omega_2) = \sum_k c_k \phi_k(\omega_1) \phi_k(\omega_2)$  [Christ11]. On peut donc réécrire<sup>1</sup>

$$\lambda \int d\omega_1 d\omega_2 \psi(\omega_1, \omega_2) \hat{\mathbf{a}}^\dagger(\omega_1) \hat{\mathbf{a}}^\dagger(\omega_2) = \sum_k \frac{1}{2} \zeta_k^* \hat{\mathbf{A}}_k^{\dagger 2}, \quad (\text{B.3})$$

avec  $\hat{\mathbf{A}}_k^\dagger = \int d\omega \phi_k(\omega) \hat{\mathbf{a}}^\dagger(\omega)$ , et  $\lambda$  la durée d'interaction. Puisque les fonctions  $\{\phi_k\}$  sont orthogonales, les modes  $\{\hat{\mathbf{A}}_k\}$  le sont également et vérifient  $[\hat{\mathbf{A}}_k, \hat{\mathbf{A}}_m^\dagger] = \delta_{km}$ . On peut donc réécrire l'évolution due à  $\hat{\mathcal{W}}$  comme un produit tensoriel de squeezers monomodes [Wasilewski06, Christ11,

---

1. Jusqu'à maintenant, nous n'avons considéré qu'un paramètre de compression réel. Un facteur de phase correspond simplement à une rotation dans l'espace des phases. Par exemple, pour  $\xi = r e^{-2i\theta}$ , on reconnaît une rotation d'un angle  $\theta$  (puisque  $\hat{\mathbf{a}}$  évolue en  $\hat{\mathbf{a}} e^{-i\theta}$ ). L'état obtenu est donc comprimé selon  $\hat{\mathbf{X}}_\theta$ .

Blow90]

$$\exp\left[-\frac{i}{\hbar}\lambda\hat{\mathcal{W}}\right] = \bigotimes_k \exp\left[\frac{1}{2}\zeta_k\hat{A}_k^2 - \frac{1}{2}\zeta_k^*\hat{A}_k^{\dagger 2}\right] \quad (\text{B.4})$$

transformant les modes  $\{\hat{A}_k\}$  selon (2.139)

$$\hat{A}'_k = \hat{A}_k \cosh \zeta_k - \hat{A}_k^\dagger \sinh \zeta_k. \quad (\text{B.5})$$

Nous avons vu dans la section 3.4.1 qu'une détection homodyne ne voit que la composante du champ dans le mode de l'oscillateur local  $\hat{\mathbf{d}} = \int d\omega \alpha_{\text{OL}}^*(\omega) \hat{\mathbf{a}}(\omega)$ . On peut décomposer ce mode dans la base  $\{\phi_k\} : \alpha_{\text{OL}}(\omega) = \sum_k M_k e^{-i\theta_k} \phi_k(\omega)$ , avec  $M_k$  réel, ce qui permet d'écrire

$$\hat{\mathbf{d}} = \sum_k M_k e^{i\theta_k} \hat{A}_k. \quad (\text{B.6})$$

### B.1.1 Efficacité homodyne parfaite

Raisonnons maintenant en suivant une démarche similaire à [Tualle-Brouiri09]. Sous l'action de la transformation (B.4), le mode  $\hat{\mathbf{d}}$  se transforme en

$$\hat{\mathbf{d}}' = \sum_k M_k e^{i\theta_k} \left( \hat{A}_k \cosh \zeta_k - \hat{A}_k^\dagger \sinh \zeta_k \right). \quad (\text{B.7})$$

Afin de pouvoir justifier le modèle empirique présenté dans le chapitre 3, on introduit un mode  $\hat{\mathbf{a}}$ , d'abord amplifié par un OPA dégénéré de paramètre  $R$ , puis par un OPA non dégénéré de paramètre  $\Gamma R$ . On montre sans difficulté que ces deux amplifications transforment le mode  $\hat{\mathbf{a}}$  en

$$\hat{\mathbf{a}}' = \hat{\mathbf{a}} \cosh \Gamma R \cosh R - \hat{\mathbf{a}}^\dagger \cosh \Gamma R \sinh R - \hat{\mathbf{c}}^\dagger \sinh \Gamma R \quad (\text{B.8})$$

où  $\hat{\mathbf{c}}$  est le mode fictif de l'OPA non dégénéré. En comparant les expressions (B.7) et (B.8), on voit que l'on peut identifier  $\hat{\mathbf{a}}'$  et  $\hat{\mathbf{d}}'$  à condition que :

$$\hat{\mathbf{a}} \cosh \Gamma R \cosh R = \sum_k M_k e^{i\theta_k} \hat{A}_k \cosh \zeta_k \quad (\text{B.9a})$$

$$\hat{\mathbf{a}}^\dagger \cosh \Gamma R \sinh R + \hat{\mathbf{c}}^\dagger \sinh \Gamma R = \sum_k M_k e^{i\theta_k} \hat{A}_k^\dagger \sinh \zeta_k \quad (\text{B.9b})$$

On trouve les valeurs de  $R$  et  $\Gamma$  permettant de satisfaire ces conditions sans problème : puisque  $[\hat{\mathbf{a}}, \hat{\mathbf{a}}^\dagger] = 1$  et  $[\hat{A}_k, \hat{A}_m^\dagger] = \delta_{km}$ , il suffit de poser

$$\cosh \Gamma R \cosh R = \sqrt{\sum_k M_k^2 \cosh^2 \zeta_k}. \quad (\text{B.10})$$

Pour la deuxième condition, on calcule le commutateur entre (B.9a) et (B.9b) :

$$\cosh^2 \Gamma R \cosh R \sinh R = \sum_k M_k^2 e^{2i\theta_k} \cosh \zeta_k \sinh \zeta_k \quad (\text{B.11})$$

En divisant cette équation par le carré de (B.10), on obtient la valeur<sup>2</sup> de  $R$  :

$$R = \operatorname{arctanh} \left[ \frac{\sum_k M_k^2 e^{2i\theta_k} \cosh \zeta_k \sinh \zeta_k}{\sum_k M_k^2 \cosh^2 \zeta_k} \right] \quad (\text{B.12})$$

On obtient ensuite la valeur de  $\Gamma$  en utilisant (B.10) :

$$\Gamma = \frac{1}{R} \operatorname{arcosh} \left[ \frac{1}{\cosh R} \sqrt{\sum_k M_k^2 \cosh^2 \zeta_k} \right] \quad (\text{B.13})$$

Enfin, on peut vérifier que le coefficient de  $\hat{\mathbf{c}}^\dagger$  est cohérent avec ces expressions en utilisant le fait que  $[\hat{\mathbf{a}}', \hat{\mathbf{a}}'^\dagger] = [\hat{\mathbf{d}}', \hat{\mathbf{d}}'^\dagger] = 1$ . Naturellement, lorsque seul un coefficient  $M_k$  est non nul, le mode de la détection homodyne correspond à un mode de Schmidt. On vérifie alors que  $R = \zeta_k$  et  $\Gamma = 0$ , comme on pouvait s'y attendre.

### B.1.2 Prise en compte de l'efficacité homodyne

Les paramètres  $R$  et  $\Gamma$  introduits dans cette annexe permettent de modéliser les données mesurées par une détection homodyne parfaite. Afin de tenir compte de l'efficacité  $\eta < 1$ , on peut introduire deux autres paramètres  $r$  et  $\gamma$ , correspondant à ceux présentés dans le chapitre 3 et dans le reste de ce manuscrit, tels que :

$$\eta \left( \frac{h}{s} + h - 1 \right) + 1 - \eta = \frac{H}{S} + H \quad (\text{B.14a})$$

$$\eta (hs + h - 1) + 1 - \eta = HS + H - 1 \quad (\text{B.14b})$$

---

2. On peut supposer que  $R$  est réel. Sinon, on peut modifier les équations (B.9a) et (B.9b) pour tenir compte du facteur de phase.



## Annexe C

# Eléments de photodétection

L'expression du champ quantique (2.40), obtenue dans le chapitre 2 a été obtenue sans approximation. Malheureusement, cette forme présente deux difficultés qui compliquent fortement les calculs : le facteur  $\mathfrak{E}(\mathbf{k}) = \sqrt{\hbar c \|\mathbf{k}\|} / (16\epsilon_0 \pi^3)$  et le vecteur polarisation  $\boldsymbol{\epsilon}_\lambda(\mathbf{k})$  dépendent de la variable intégrée  $\mathbf{k}$ . Une approximation fort utile consiste donc à supposer que ces termes peuvent être sortis de l'intégrale : c'est heureusement une approximation qui ne sera pas trop illégitime dans notre cas. Sortir le facteur  $\mathfrak{E}(\mathbf{k})$  consiste à supposer que le mode du champ excité est de faible largeur spectrale par rapport à la fréquence centrale, ce qui sera le cas pour nos impulsions femto-secondes. Enfin, le vecteur polarisation peut être considéré constant dans une certaine mesure sous l'approximation paraxiale, que nous ferons également ici.

Sous ces approximations et pour une polarisation donnée, écrivons donc l'opérateur champ scalaire, en incluant l'évolution libre :

$$\hat{\mathbf{E}}(\mathbf{r}, t) = \mathfrak{E}(\mathbf{k}_0) \int d^3\mathbf{k} i \left( \hat{\mathbf{a}}(\mathbf{k}) e^{i\mathbf{k}\cdot\mathbf{r} - ic\|\mathbf{k}\|t} - \hat{\mathbf{a}}^\dagger(\mathbf{k}) e^{-i\mathbf{k}\cdot\mathbf{r} + ic\|\mathbf{k}\|t} \right) \quad (\text{C.1})$$

avec  $\mathbf{k}_0$  est le vecteur d'onde correspondant à la fréquence centrale  $\omega_0$ . Remarquons que le facteur  $e^{-ic\|\mathbf{k}\|t}$  ne peut pas être sorti de l'intégrale car il oscille rapidement en fonction du temps.

Le signal de photodétection est proportionnel à l'opérateur  $\hat{\mathbf{E}}^{(-)} \hat{\mathbf{E}}^{(+)}$  [Cohen-Tannoudji01]. Pour une photodiode placée en  $z$ , de surface  $D$  mesurant pendant un temps  $T$ , le courant produit sera proportionnel à

$$\hat{i} = \int_S d^2S \int_0^T dt \hat{\mathbf{E}}^{(-)}(\mathbf{r}, t) \hat{\mathbf{E}}^{(+)}(\mathbf{r}, t) \quad (\text{C.2a})$$

$$\simeq \mathfrak{E}^2(\mathbf{k}_0) \int_S \int_0^T d^2S dt \int d^3\mathbf{k} d^3\mathbf{q} \hat{\mathbf{a}}^\dagger(\mathbf{q}) \hat{\mathbf{a}}(\mathbf{k}) e^{i(\mathbf{k}-\mathbf{q})\cdot\mathbf{r} - ic(\|\mathbf{k}\| - \|\mathbf{q}\|)t} \quad (\text{C.2b})$$

Si la surface de la photodiode est largement supérieure à l'étendue spatiale des modes du champ excités, et que le temps d'intégration est lui aussi très supérieur à la durée des impulsions, on peut raisonnablement étendre les domaines d'intégration à l'infini. On peut alors utiliser la

relation  $\int_{-\infty}^{+\infty} dx e^{-ixy} = 2\pi\delta(y)$  :

$$\hat{i} \simeq \mathfrak{E}^2(\mathbf{k}_0) \int d^2S dt \int d^3\mathbf{k} d^3\mathbf{q} \hat{a}^\dagger(\mathbf{q}) \hat{a}(\mathbf{k}) e^{i(\mathbf{k}-\mathbf{q}) \cdot \mathbf{r} - ic(\|\mathbf{k}\| - \|\mathbf{q}\|)t} \quad (\text{C.3a})$$

$$= \mathfrak{E}^2(\mathbf{k}_0) \int d^3\mathbf{k} d^3\mathbf{q} \hat{a}^\dagger(\mathbf{q}) \hat{a}(\mathbf{k}) \underbrace{\int d^2S dt e^{i(\mathbf{k}-\mathbf{q}) \cdot \mathbf{r} - ic(\|\mathbf{k}\| - \|\mathbf{q}\|)t}}_{\propto \delta(\mathbf{k}-\mathbf{q})} \quad (\text{C.3b})$$

$$= \mathcal{N}(\mathbf{k}_0) \int d^3\mathbf{k} \hat{a}^\dagger(\mathbf{k}) \hat{a}(\mathbf{k}) \quad (\text{C.3c})$$

où  $\mathcal{N}(\mathbf{k}_0)$  inclut tous les facteurs de proportionnalité que nous n'avons pas besoin d'explicitier.

Le signal de photodétection est donc proportionnel au nombre total de photons.

## Annexe D

# Modèle simplifié de l'origine de $\xi$

On considère que l'état en sortie de l'OPA peut être décomposé sur deux sous-espaces de Hilbert, l'un correspondant au mode spatio-temporel vu par la détection homodyne (dh), et l'autre à un mode spatio-temporel orthogonal (N), peuplé par du bruit qui pourra déclencher l'APD. Une matrice densité  $\hat{\rho}_{\text{in}}$  peut donc s'écrire comme :

$$\hat{\rho}_{\text{in}} = \hat{\rho}_{\text{dh}} \otimes \hat{\rho}_{\text{N}} \quad (\text{D.1})$$

L'hypothèse faite ici est que les deux modes sont peuplés de façon indépendante, et que l'état initial n'est pas dans une superposition cohérente de ces deux modes. On peut aussi considérer que l'on a laissé tomber la partie orthogonale filtrée par l'APD, car elle ne contribuera pas au conditionnement.

Appelons  $\hat{U}_{\text{BS}} = \hat{U}_{\text{dh}} \otimes \hat{U}_{\text{N}}$  l'opérateur de la lame séparatrice utilisé pour la soustraction de photon, et agissant sur  $\hat{\rho}_{\text{in}}$ . La matrice densité après cet opérateur est donc

$$\hat{\hat{\rho}} = \hat{U}_{\text{BS}} (\hat{\rho}_{\text{in}} \otimes |0_c, 0_d\rangle\langle 0_c, 0_d|) \hat{U}_{\text{BS}}^\dagger \quad (\text{D.2a})$$

$$= \hat{\hat{\rho}}_{\text{dh,c}} \otimes \hat{\hat{\rho}}_{\text{N,d}} \quad (\text{D.2b})$$

avec  $c$  et  $d$  les modes vides de la lame séparatrice correspondant respectivement aux modes dh et N. Considérons maintenant l'action de l'APD sur les modes  $c$  et  $d$ . On appelle  $C$  l'évènement "détection de au moins un photon dans le mode  $c$ ", et  $D$  "détection de au moins un photon dans le mode  $d$ ". On considère que le conditionnement est réussi si l'évènement  $C \cup D$  est réalisé, le  $\cup$  pouvant être exclusif ou non. Nous le considérerons exclusif, ce qui signifie que l'on considère que l'APD détecte dans un mode ou dans l'autre, mais pas dans les deux à la fois. L'évènement associé est donc  $C \cap \bar{D} + \bar{C} \cap D$ .

L'opérateur associé à l'évènement  $C$ , agissant sur le mode  $c$  de  $\hat{\hat{\rho}}$ , est  $\mathbb{I}_c - |0\rangle\langle 0|_c$ . L'opérateur associé à l'évènement  $C \cap \bar{D}$ , agissant sur  $\hat{\hat{\rho}}$ , s'écrit alors :

$$\hat{\Pi}_C = (\mathbb{I}_c - |0\rangle\langle 0|_c) \otimes |0\rangle\langle 0|_d \otimes \mathbb{I}_{\text{dh}} \otimes \mathbb{I}_{\text{N}} \quad (\text{D.3})$$

De même, l'opérateur associé à l'évènement  $\bar{C} \cap D$  s'écrit :

$$\hat{\Pi}_D = |0\rangle\langle 0|_c \otimes (\mathbb{I}_d - |0\rangle\langle 0|_d) \otimes \mathbb{I}_{\text{dh}} \otimes \mathbb{I}_{\text{N}} \quad (\text{D.4})$$

L'opérateur associé au conditionnement s'écrit alors :

$$\hat{\Pi} = \hat{\Pi}_C + \hat{\Pi}_D \quad (\text{D.5})$$



L'état final non normalisé produit lorsque l'APD donne un clic est donc :

$$\hat{\rho}_{\text{out}} = \text{Tr}_{c,d}\{\hat{\Pi}\hat{\rho}\hat{\Pi}^\dagger\} \quad (\text{D.6a})$$

$$= \text{Tr}_{c,d}\{\hat{\Pi}_C\hat{\rho}\hat{\Pi}_C^\dagger + \hat{\Pi}_D\hat{\rho}\hat{\Pi}_D^\dagger + \hat{\Pi}_C\hat{\rho}\hat{\Pi}_D^\dagger + \hat{\Pi}_D\hat{\rho}\hat{\Pi}_C^\dagger\} \quad (\text{D.6b})$$

$$= \text{Tr}_{c,d}\{\hat{\Pi}_C\hat{\rho}\hat{\Pi}_C^\dagger\} + \text{Tr}_{c,d}\{\hat{\Pi}_D\hat{\rho}\hat{\Pi}_D^\dagger\} \quad (\text{D.6c})$$

$$= \hat{\rho}_{\text{out},C} + \hat{\rho}_{\text{out},D} \quad (\text{D.6d})$$

En effet, voit sans difficulté que  $\text{Tr}_{c,d}\{\hat{\Pi}_C\hat{\rho}\hat{\Pi}_D^\dagger\} = \text{Tr}_{c,d}\{\hat{\Pi}_D\hat{\rho}\hat{\Pi}_C^\dagger\} = 0$  car pour les deux modes c et d, on projette sur “au moins un photon” d'un côté et “0 photon” de l'autre.

En posant enfin

$$\xi = \frac{\text{Tr}\{\hat{\rho}_{\text{out},C}\}}{\text{Tr}\{\hat{\rho}_{\text{out},C} + \hat{\rho}_{\text{out},D}\}} \quad \text{et} \quad 1 - \xi = \frac{\text{Tr}\{\hat{\rho}_{\text{out},D}\}}{\text{Tr}\{\hat{\rho}_{\text{out},C} + \hat{\rho}_{\text{out},D}\}}, \quad (\text{D.7})$$

et en se souvenant que la trace d'un produit tensoriel est égale au produit des traces, l'état final “physique” normalisé  $\hat{\rho}_{\text{out}} = \hat{\rho}_{\text{out}} / \text{Tr}\{\hat{\rho}_{\text{out}}\}$  peut s'écrire :

$$\hat{\rho}_{\text{out}} = \frac{\hat{\rho}_{\text{out},C} + \hat{\rho}_{\text{out},D}}{\text{Tr}\{\hat{\rho}_{\text{out},C} + \hat{\rho}_{\text{out},D}\}} \quad (\text{D.8a})$$

$$= \xi \frac{\hat{\rho}_{\text{out},C}}{\text{Tr}\{\hat{\rho}_{\text{out},C}\}} + (1-\xi) \frac{\hat{\rho}_{\text{out},D}}{\text{Tr}\{\hat{\rho}_{\text{out},D}\}} \quad (\text{D.8b})$$

L'équation (D.8a) peut être interprétée comme la somme de deux processus  $\tilde{\mathcal{F}}^\vee$  et  $\tilde{\mathcal{F}}^\times$  probabilistes et linéaires, qui est ensuite normalisée :

$$\hat{\rho}_{\text{out}} = \frac{\tilde{\mathcal{F}}^\vee(\hat{\rho}_{\text{in}}) + \tilde{\mathcal{F}}^\times(\hat{\rho}_{\text{in}})}{\text{Tr}\{\tilde{\mathcal{F}}^\vee(\hat{\rho}_{\text{in}}) + \tilde{\mathcal{F}}^\times(\hat{\rho}_{\text{in}})\}} \quad (\text{D.9})$$

avec  $\tilde{\mathcal{F}}^\vee(\hat{\rho}_{\text{in}}) = \hat{\rho}_{\text{out},C}$  et  $\tilde{\mathcal{F}}^\times(\hat{\rho}_{\text{in}}) = \hat{\rho}_{\text{out},D}$ . Une tomographie quantique multimode de  $\tilde{\mathcal{F}}^\vee + \tilde{\mathcal{F}}^\times$  serait en théorie possible, mais expérimentalement très délicate à mettre en œuvre, puisqu'elle devrait inclure les modes du bruit comme expliqué dans la section 6.4.2.

A aucune étape de la transformation, les modes dh et N ne sont couplés entre eux. On peut donc écrire

$$\tilde{\mathcal{F}}^\vee(\hat{\rho}_{\text{in}}) = \tilde{\mathcal{F}}_{\text{dh}}^\vee(\hat{\rho}_{\text{dh}}) \otimes \tilde{\mathcal{F}}_{\text{N}}^\vee(\hat{\rho}_{\text{N}}), \quad \text{et} \quad \tilde{\mathcal{F}}^\times(\hat{\rho}_{\text{in}}) = \tilde{\mathcal{F}}_{\text{dh}}^\times(\hat{\rho}_{\text{dh}}) \otimes \tilde{\mathcal{F}}_{\text{N}}^\times(\hat{\rho}_{\text{N}}). \quad (\text{D.10})$$

En notant  $c_{\text{N}}^\vee = \text{Tr}\{\tilde{\mathcal{F}}_{\text{N}}^\vee(\hat{\rho}_{\text{N}})\}$  et  $c_{\text{N}}^\times = \text{Tr}\{\tilde{\mathcal{F}}_{\text{N}}^\times(\hat{\rho}_{\text{N}})\}$ , après trace partielle sur le mode du bruit, l'état dans le mode de la détection homodyne s'écrit

$$\hat{\rho}_{\text{out,dh}} = \frac{c_{\text{N}}^\vee \tilde{\mathcal{F}}_{\text{dh}}^\vee(\hat{\rho}_{\text{dh}}) + c_{\text{N}}^\times \tilde{\mathcal{F}}_{\text{dh}}^\times(\hat{\rho}_{\text{dh}})}{\text{Tr}\{c_{\text{N}}^\vee \tilde{\mathcal{F}}_{\text{dh}}^\vee(\hat{\rho}_{\text{dh}}) + c_{\text{N}}^\times \tilde{\mathcal{F}}_{\text{dh}}^\times(\hat{\rho}_{\text{dh}})\}} \quad (\text{D.11a})$$

$$= \xi \mathcal{E}^\vee(\hat{\rho}_{\text{in}}) + (1-\xi) \mathcal{E}^\times(\hat{\rho}_{\text{in}}). \quad (\text{D.11b})$$

Puisque la probabilité de succès dépend de l'état initial du bruit  $\hat{\rho}_{\text{N}}$  de part les coefficients  $c_{\text{N}}^\vee$  et  $c_{\text{N}}^\times$ , on ne peut faire de tomographie de processus quantique en considérant uniquement le mode de la détection homodyne. Les transformations  $\mathcal{E}^\vee$  et  $\mathcal{E}^\times$  produisent des états normalisés, et sont non linéaires.

## Annexe E

# Une autre implémentation approchée du NLA

Cette annexe s'intéresse à la transcription en variable continue d'une implémentation "physique" de l'amplificateur sans bruit, du type  $\hat{T}_N = \sum_{k=0}^N g^k |k\rangle \langle k|$ , telle que modélisée par (8.11).

### E.1 Résultat préliminaire sur les états de Fock

Commençons par démontrer un résultat préliminaire. Soit  $\hat{U}(\phi) = \exp[-i\phi\hat{n}]$  l'opérateur de rotation. Définissons l'opérateur non unitaire  $\hat{\Pi}_n$  créant une "superposition circulaire" :

$$\hat{\Pi}_n = \frac{1}{2\pi} \int_0^{2\pi} d\phi \hat{U}(\phi) e^{i\phi n} \quad (\text{E.1})$$

Regardons maintenant l'action de  $\hat{\Pi}_n$  sur un état  $|\psi\rangle = \sum_k c_k |k\rangle$  quelconque :

$$\hat{\Pi}_n |\psi\rangle = \left[ \frac{1}{2\pi} \int_0^{2\pi} d\phi \hat{U}(\phi) e^{i\phi n} \right] \left[ \sum_k c_k |k\rangle \right] \quad (\text{E.2a})$$

$$= \sum_k c_k \underbrace{\frac{1}{2\pi} \int_0^{2\pi} d\phi e^{i\phi(n-k)}}_{\delta_{n,k}} |k\rangle \quad (\text{E.2b})$$

$$= c_n |n\rangle \quad (\text{E.2c})$$

On en conclut que  $\hat{\Pi}_n$  est égal au projecteur sur  $|n\rangle$  :

$$\boxed{\frac{1}{2\pi} \int_0^{2\pi} d\phi \hat{U}(\phi) e^{i\phi n} = |n\rangle \langle n|} \quad (\text{E.3})$$

En particulier, la superposition circulaire pour  $n=0$  d'un état quelconque produit toujours le vide. Cette conclusion peut sembler quelque peu curieuse. Les chats de Schrödinger pairs sont en effet obtenus avec une version discrète de  $\hat{\Pi}_{n=0}$ , en superposant l'état initial et un déphasage de  $\pi$  :

$$|\alpha\rangle + |-\alpha\rangle = \left( \hat{U}(0) + \hat{U}(\pi) \right) |\alpha\rangle \quad (\text{E.4})$$

De même, un “chat à N pattes” s’obtient en superposant N phases :

$$\sum_{k=0}^{N-1} |\alpha e^{2i\frac{k\pi}{N}}\rangle = \left[ \sum_{k=0}^{N-1} \hat{U}(2i\frac{k\pi}{N}) \right] |\alpha\rangle \quad (\text{E.5})$$

On pourrait donc s’attendre à ce que l’action de  $\hat{\Pi}_{n=0}$  produise également un état hautement non classique. Il n’en est rien, puisque l’état obtenu sera toujours le vide :

$$\hat{\Pi}_{n=0}|\alpha\rangle = (\langle 0|\alpha\rangle)|0\rangle \quad (\text{E.6})$$

## E.2 Décomposition de $\hat{T}_N$ sur la base des déplacements

### E.2.1 Sans coupure d’intégration

Les opérateurs de déplacements constituent une base sur laquelle on peut décomposer n’importe quel opérateur  $\hat{O}$  [Ferraro05] :

$$\hat{O} = \frac{1}{\pi} \int d^2\alpha \text{Tr}\{\hat{O}\hat{D}(\alpha)\}\hat{D}^\dagger(\alpha) \quad (\text{E.7})$$

Pour  $\hat{T}$  qui est non borné, une telle décomposition risque fortement de ne pas converger. Par contre, il ne doit pas y avoir de problème pour  $\hat{T}_N$  qui est un opérateur bien physique. Calculons les coefficients de la décomposition :

$$\text{Tr}\{\hat{T}_N\hat{D}(\alpha)\} = \sum_{k=0}^{\infty} \langle k| \left[ \sum_{m=0}^N g^m |m\rangle \langle m| \right] \hat{D}(\alpha) |k\rangle \quad (\text{E.8a})$$

$$= \sum_{k=0}^N g^k \langle k| \hat{D}(\alpha) |k\rangle \quad (\text{E.8b})$$

$$= \sum_{k=0}^N g^k e^{-\frac{1}{2}|\alpha|^2} L_k(|\alpha|^2) \quad (\text{E.8c})$$

où  $L_k$  est le k-ième polynôme de Laguerre. On a donc :

$$\hat{T}_N = \frac{1}{\pi} \sum_{k=0}^N g^k \int d^2\alpha e^{-\frac{1}{2}|\alpha|^2} L_k(|\alpha|^2) \hat{D}^\dagger(\alpha) \quad (\text{E.9})$$

Remarquons que l’on peut faire le changement de variable  $\alpha \rightarrow -\alpha$  et donc remplacer  $\hat{D}^\dagger(\alpha)$  par  $\hat{D}(\alpha)$ . Nous n’avons pour l’instant pas fait d’approximation : notre objectif sera de montrer que l’on peut obtenir  $\hat{T}_N$  avec une bonne précision en limitant le domaine d’intégration, en intégrant jusqu’à un rayon maximal  $\alpha_m$ .

Passons maintenant en coordonnées polaires en posant  $\alpha = \gamma e^{-i\phi}$  avec  $\gamma$  réel, et appliquons  $\hat{T}_N$  sur un état de Fock  $|n\rangle$  :

$$\hat{T}_N |n\rangle = \frac{1}{\pi} \sum_{k=0}^N g^k \int_0^\infty \int_0^{2\pi} \gamma d\gamma d\phi e^{-\frac{1}{2}\gamma^2} L_k(\gamma^2) \hat{D}(\gamma e^{-i\phi}) |n\rangle \quad (\text{E.10})$$

Puisque  $\hat{D}(\gamma e^{-i\phi}) = \hat{U}(\phi) \hat{D}(\gamma) \hat{U}(-\phi)$ , on a  $\hat{D}(\gamma e^{-i\phi})|n\rangle = e^{i\phi n} \hat{U}(\phi) \hat{D}(\gamma)|n\rangle$ , et donc :

$$\hat{T}_N|n\rangle = \frac{1}{\pi} \sum_{k=0}^N g^k \int_0^\infty \gamma d\gamma e^{-\frac{1}{2}\gamma^2} L_k(\gamma^2) \left[ \int_0^{2\pi} d\phi e^{i\phi n} \hat{U}(\phi) \right] \hat{D}(\gamma)|n\rangle \quad (\text{E.11})$$

On utilise maintenant le résultat préliminaire (E.3) pour remplacer le terme entre crochets :

$$\hat{T}_N|n\rangle = \frac{1}{\pi} \sum_{k=0}^N g^k \int_0^\infty \gamma d\gamma e^{-\frac{1}{2}\gamma^2} L_k(\gamma^2) 2\pi |n\rangle \langle n| \hat{D}(\gamma)|n\rangle \quad (\text{E.12a})$$

$$= \sum_{k=0}^N g^k \int_0^\infty 2\gamma d\gamma e^{-\gamma^2} L_k(\gamma^2) L_n(\gamma^2) |n\rangle \quad (\text{E.12b})$$

En posant  $z = \gamma^2$ , on obtient finalement :

$$\hat{T}_N|n\rangle = \left[ \sum_{k=0}^N g^k \int_0^\infty dz e^{-z} L_k(z) L_n(z) \right] |n\rangle \quad (\text{E.13})$$

Puisque les polynômes de Laguerre sont orthogonaux avec la métrique  $e^{-x} dx$ ,

$$\int_0^\infty dz e^{-z} L_k(z) L_n(z) = \delta_{k,n}, \quad (\text{E.14})$$

on retrouve finalement que  $\hat{T}_N|n\rangle = \sum_{k=0}^N g^k \delta_{k,n} |n\rangle = g^n |n\rangle$  lorsque  $n \leq N$ .

## E.2.2 Avec une coupure d'intégration

Le calcul précédant nous confirme en premier lieu que (E.9) est correcte. Ecrit de la sorte, il peut sembler n'être qu'un détour peu commode pour retrouver l'action de  $\hat{T}$  en base de Fock. Cependant, tout l'intérêt provient de l'interprétation que l'on peut en faire : si l'on était capable d'implémenter expérimentalement une superposition de déplacements quelconques  $\hat{D}(\alpha)$ , il serait alors possible d'implémenter  $\hat{T}_N$  en pondérant ces déplacements avec des coefficients donnés par (E.8).

L'amplitude maximale des déplacements sera toutefois forcément limitée par les ressources expérimentales. En appelant  $\alpha_m$  le rayon d'intégration maximal, on définit

$$\hat{T}_N^{\alpha_m} = \frac{1}{\pi} \sum_{k=0}^N g^k \int_{|\alpha| \leq \alpha_m} d^2\alpha e^{-\frac{1}{2}|\alpha|^2} L_k(|\alpha|^2) \hat{D}^\dagger(\alpha) \quad (\text{E.15})$$

L'expression (E.13) nous permet de déterminer précisément l'effet de cette coupure, en remplaçant la borne infinie par  $\alpha_m^2$  (ne pas oublier qu'il y a eu un changement de variable). L'action de  $\hat{T}_N^{\alpha_m}$  sur un état de Fock s'écrit alors :

$$\hat{T}_N^{\alpha_m}|n\rangle = \left[ \sum_{k=0}^N g^k \int_0^{\alpha_m^2} dz e^{-z} L_k(z) L_n(z) \right] |n\rangle \quad (\text{E.16})$$

L'opérateur  $\hat{T}_N^{\alpha_m}$  peut donc également s'écrire comme

$$\hat{T}_N^{\alpha_m} = \sum_{n=0}^N \left[ \sum_{k=0}^N g^k \int_0^{\alpha_m^2} dz e^{-z} L_k(z) L_n(z) \right] |n\rangle \langle n|. \quad (\text{E.17})$$

Cette expression permet de définir les éléments d'un POVM associé à cette implémentation approchée du NLA, définis de manière similaire à (8.11) :

$$\hat{M}_{\text{suc}}^{N,\alpha_m} = \frac{1}{g^N} \hat{T}_N^{\alpha_m} \quad (\text{E.18})$$

L'opérateur  $\hat{M}_{\text{suc}}^{N,\alpha_m}$  est défini de telle sorte que  $\hat{M}_{\text{suc}}^{N,\alpha_m\dagger} \hat{M}_{\text{suc}}^{N,\alpha_m} + \hat{M}_{\text{fail}}^{N,\alpha_m\dagger} \hat{M}_{\text{fail}}^{N,\alpha_m} = \mathbb{I}$ .

Finalement, cette coupure aura pour effet de ne pas strictement annuler l'intégrale de recouvrement E.14 pour  $k \neq n$ , et de ne pas la rendre strictement égale à 1 pour  $k = n$ . On peut donc comparer le coefficient total obtenu devant chaque état  $|n\rangle$  à  $g^n$  afin de déterminer une valeur de coupure qui lui procure la précision voulue. Un exemple est donné dans la table E.1, où une amplitude de  $\alpha_m = 6$  permet d'obtenir des coefficients précis à un minimum d'environ 2% pour  $N=8$ , avec une très bonne précision pour les premiers états de Fock.

| Etat de Fock                | n=0   | n=1   | n=2   | n=3   | n=4   | n=5   | n=6   | n=7  | n=8  |
|-----------------------------|-------|-------|-------|-------|-------|-------|-------|------|------|
| $g^n$                       | 1     | 3     | 9     | 27    | 81    | 243   | 729   | 2187 | 6561 |
| Coefficient avec la coupure | 0.999 | 3.000 | 8.992 | 27.08 | 80.36 | 246.6 | 713.5 | 2239 | 6425 |
| Ecart relatif absolu (%)    | 0.001 | 0.015 | 0.086 | 0.3   | 0.8   | 1.5   | 2.1   | 2.4  | 2.1  |

TABLE E.1 – Comparaison des coefficients de  $\hat{T}_N$  et  $\hat{T}_N^{\alpha_m}$  pour les états de Fock  $|n\rangle$ , pour  $N=8$ ,  $g=3$ , et  $\alpha_m=6$ .

L'amplitude  $\alpha_m$  donnant une bonne précision augmente d'autant plus que  $N$  augmente, principalement en raison du facteur  $g^k$  dans (E.16) qui augmente plus vite avec  $k$  que l'intégrale de recouvrement ne décroît, pour un  $n$  fixé.

### E.3 Implémentation “physique” d'une superposition de déplacements

Dans cette section, nous esquissons les grandes lignes d'un raisonnement qui semble permettre d'implémenter une superposition de déplacements quelconque  $\int d^2\alpha f(\alpha) \hat{D}(\alpha)$ , et qui semble être relié à [Clausen99]. Ce raisonnement est basé sur l'implémentation d'un seul déplacement à partir d'une lame séparatrice et d'un état cohérent intense. On mélange un état  $|\phi\rangle$  servant de ressource avec l'état initial à déplacer sur une séparatrice de faible réflectivité en amplitude  $r$ . On fait ensuite une mesure  $\langle\beta|$  sur le deuxième mode de sortie du BS qui va agir comme un conditionnement.  $\langle\beta|$  est associé à une observable quelconque mais devant être choisie avant de préparer  $|\phi\rangle$ .

L'état de ressource  $|\phi\rangle$  est du type :

$$|\phi\rangle = \int d^2\lambda f(\lambda) \frac{1}{\langle\beta|\lambda/r\rangle} |\lambda/r\rangle, \quad (\text{E.19})$$

où  $|\lambda\rangle$  est un état cohérent. Le coefficient  $\frac{1}{\langle\beta|\lambda/r\rangle}$  sert à ce que la superposition finale ne soit pas modifiée par le conditionnement sur  $\langle\beta|$ .

Quelques simulations numériques avec Matlab pour une superposition discrète de déplacements ont donné des résultats encourageants, mais qui doivent encore être approfondis, notamment pour voir les conditions de convergence vers la superposition de déplacements souhaitée. Remarquons que l'on peut choisir  $\langle\beta|$  de manière à maximiser la probabilité de succès du conditionnement.

## Annexe F

# Compléments sur le NLA

Dans cette annexe, nous revenons sur la définition (9.2) de  $\mathcal{T}$ , afin de voir sous quelles conditions elle correspond bien à la définition d'une opération quantique, c'est-à-dire une opération linéaire, complètement positive, et pour laquelle la trace de l'état de sortie est inférieure ou égale à 1. Il suffit pour cela qu'elle admette une décomposition de Kraus [Nielsen00]

$$\mathcal{T}[\hat{\rho}] = \sum_n \hat{E}_n \hat{\rho} \hat{E}_n^\dagger, \quad (\text{F.1})$$

avec  $\sum_n \hat{E}_n^\dagger \hat{E}_n \leq \mathbb{I}$ . Il ne semble pas évident que la définition (9.2) vérifie ces propriétés, puisque la normalisation avec le dénominateur  $\text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}$  semble en particulier la rendre non linéaire. Commençons par remarquer que chaque terme de la décomposition (F.1) peut également s'écrire :

$$\hat{E}_n \hat{\rho} \hat{E}_n^\dagger = \underbrace{\text{Tr}\{\hat{E}_n \hat{\rho} \hat{E}_n^\dagger\}}_{:=P_n} \frac{\hat{E}_n \hat{\rho} \hat{E}_n^\dagger}{\underbrace{\text{Tr}\{\hat{E}_n \hat{\rho} \hat{E}_n^\dagger\}}_{:=\hat{\rho}_n}} \quad (\text{F.2})$$

L'état  $\hat{\rho}_n$  est un état physique de trace 1, obtenu avec la probabilité  $P_n$ . L'état total  $\sum_n \hat{E}_n \hat{\rho} \hat{E}_n^\dagger$  peut être interprété comme un mélange statistique d'états  $\hat{\rho}_n$  avec des probabilités  $P_n$ . Puisque  $\hat{T}$  augmente la trace, il n'est pas possible de l'interpréter comme un opérateur  $\hat{E}$  de la décomposition de Kraus, ni d'interpréter  $\text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}$  comme une probabilité de succès.

En revanche, toute implémentation *physique* de  $\hat{T}$  produisant un état  $\hat{\rho}_{\text{out}}$  avec une probabilité de succès  $\text{Tr}\{\hat{\rho}_{\text{out}}\}$  doit nécessairement pouvoir s'écrire comme résultant d'une évolution du système avec un environnement extérieur, sur lequel on fait une mesure correspondant au conditionnement [Nielsen00] :

$$\hat{\rho}_{\text{out}} = \langle 1_E | \hat{U} (\hat{\rho} \otimes |0_E\rangle \langle 0_E|) \hat{U}^\dagger | 1_E \rangle \quad (\text{F.3a})$$

$$= \hat{E}_{\text{NLA}} \hat{\rho} \hat{E}_{\text{NLA}}^\dagger \quad (\text{F.3b})$$

avec  $\hat{E}_{\text{NLA}} := \langle 1_E | \hat{U} | 0_E \rangle$ . Bien entendu, l'état initial de l'environnement  $|0_E\rangle$ , son état de conditionnement  $|1_E\rangle$ , et l'évolution unitaire  $\hat{U}$  dépendent de la façon dont est implémenté le NLA, mais nous n'avons pas besoin de connaître leurs expressions exactes ici. Avec ce formalisme, la probabilité de succès est :

$$P(\hat{\rho}) = \text{Tr}\{\langle 1_E | \hat{U} (\hat{\rho} \otimes |0_E\rangle \langle 0_E|) \hat{U}^\dagger | 1_E \rangle\} \quad (\text{F.4a})$$

$$= \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger \hat{E}_{\text{NLA}} \hat{\rho}\} \quad (\text{F.4b})$$

Nous faisons ensuite l'hypothèse que l'implémentation physique est suffisamment bonne pour pouvoir assimiler l'état physique *normalisé*  $\hat{\rho}_{\text{out}}/\text{Tr}\{\hat{\rho}_{\text{out}}\}$  à l'état produit par le NLA parfait  $\hat{T}$  :

$$\frac{\hat{E}_{\text{NLA}}\hat{\rho}\hat{E}_{\text{NLA}}^\dagger}{\text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\}} \simeq \frac{\hat{T}\hat{\rho}\hat{T}}{\text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}} \quad (\text{F.5})$$

Nous pouvons donc réécrire (9.2) sous la forme :

$$\mathcal{T}[\hat{\rho}] = \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\} \frac{\hat{T}\hat{\rho}\hat{T}}{\text{Tr}\{\hat{T}\hat{\rho}\hat{T}\}} + (1 - \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\})|0\rangle\langle 0| \quad (\text{F.6a})$$

$$= \hat{E}_{\text{NLA}}\hat{\rho}\hat{E}_{\text{NLA}}^\dagger + (1 - \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\})|0\rangle\langle 0| \quad (\text{F.6b})$$

Sous cette forme, on voit sans difficulté que  $\mathcal{T}$  est bien linéaire :

$$\mathcal{T}[p_1\hat{\rho}_1 + p_2\hat{\rho}_2] = p_1\mathcal{T}[\hat{\rho}_1] + p_2\mathcal{T}[\hat{\rho}_2] \quad (\text{F.7})$$

Afin d'obtenir une décomposition de Kraus (F.1) il reste à trouver un opérateur  $\hat{B}$  tels que :

$$(1 - \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\})|0\rangle\langle 0| = \hat{B}\hat{\rho}\hat{B}^\dagger \quad (\text{F.8})$$

$$\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}} + \hat{B}^\dagger\hat{B} = \mathbb{I} \quad (\text{F.9})$$

En combinant ces deux équations, on peut alors trouver une relation que doit satisfaire  $\hat{B}$  :

$$(1 - \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\})|0\rangle\langle 0| = (\text{Tr}\{\hat{\rho}\} - \text{Tr}\{\hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}\hat{\rho}\})|0\rangle\langle 0| \quad (\text{F.10a})$$

$$= (\text{Tr}\{[\mathbb{I} - \hat{E}_{\text{NLA}}^\dagger\hat{E}_{\text{NLA}}]\hat{\rho}\})|0\rangle\langle 0| \quad (\text{F.10b})$$

$$= \text{Tr}\{\hat{B}^\dagger\hat{B}\hat{\rho}\}|0\rangle\langle 0| \quad (\text{F.10c})$$

$$= |0\rangle\text{Tr}\{\hat{B}\hat{\rho}\hat{B}^\dagger\}\langle 0| \quad (\text{F.10d})$$

$$= \sum_m |0\rangle\langle m|\hat{B}\hat{\rho}\hat{B}^\dagger|m\rangle\langle 0| \quad (\text{F.10e})$$

$$= \hat{B}\hat{\rho}\hat{B}^\dagger \quad (\text{F.10f})$$

Lorsque la condition (F.9) est vérifiée, (F.8) est équivalente à

$$\sum_m |0\rangle\langle m|\hat{B} = \hat{B}. \quad (\text{F.11})$$

Un opérateur  $\hat{B}$  qui satisfait ces conditions peut s'écrire  $\hat{B} = \sum_m c_m |0\rangle\langle m|$ , où  $c_m$  est défini à partir de (F.9).

## Annexe G

# Développement perturbatif des taux secrets

Dans cette annexe, nous détaillons les calculs pour obtenir les développements perturbatifs des taux secrets donnés dans la section 7.4.4, et utilisés dans la section 10.3. On considère directement l'utilisation d'un NLA, en utilisant les paramètres effectifs.

### G.1 Développement de l'information de Holevo

#### Fonctions $\mu$

La première étape consiste à développer les fonctions  $\mu$  définis par (7.11b) au premier ordre en  $T$ , en utilisant les paramètres effectifs :

$$\mu_1 \simeq \frac{1-\lambda^2 (2g^2T-1)}{1-\lambda^2} \quad (\text{G.1a})$$

$$\mu_2 \simeq 1+g^2T\epsilon \quad (\text{G.1b})$$

$$\mu_3 \simeq \frac{(1-\lambda^2)}{(\lambda^2-1)^3 (1+\kappa)} [(\lambda^4-1)(\kappa+1)+2\lambda^2 g^2T (\lambda^2(\nu-\kappa-1)+\kappa+1)] \quad (\text{G.1c})$$

$$\mu_4 \simeq \frac{g^2T\epsilon(1-\nu+\kappa)+\kappa+1}{1+\kappa} \quad (\text{G.1d})$$

#### Fonctions $G$

On injecte ensuite ces développements dans la fonction  $G$  (7.11a), puis on fait à nouveau un développement au premier ordre en  $T$  :

$$G \left[ \frac{\mu_1-1}{2} \right] \simeq \frac{1}{\ln 2} \left( -\frac{\lambda^2 \ln [\lambda^2]}{\lambda^2-1} + \ln \left[ \frac{1}{1-\lambda^2} \right] + \frac{\lambda^2 g^2T \ln [\lambda^2]}{1-\lambda^2} \right) \quad (\text{G.2a})$$

$$G \left[ \frac{\mu_2-1}{2} \right] \simeq \frac{g^2T\epsilon}{2 \ln 2} \left( 1 - \ln [g^2T] - \ln \left[ \frac{\epsilon}{2} \right] \right) \quad (\text{G.2b})$$

$$G \left[ \frac{\mu_3-1}{2} \right] \simeq \frac{1}{\ln 2} \left( -\frac{\lambda^2 \ln [\lambda^2]}{\lambda^2-1} + \ln \left[ \frac{1}{1-\lambda^2} \right] + \frac{\lambda^2 g^2T \ln [\lambda^2] (\lambda^2(\nu-\kappa-1)+\kappa+1)}{(\lambda^2-1)^2 (\kappa+1)} \right) \quad (\text{G.2c})$$

$$G \left[ \frac{\mu_4-1}{2} \right] \simeq \frac{g^2T\epsilon(1-\nu+\kappa)}{(1+\kappa)2 \ln 2} \left( 1 - \ln [g^2T] - \ln \left[ \frac{\epsilon(1-\nu+\kappa)}{2(1+\kappa)} \right] \right) \quad (\text{G.2d})$$



Il ne reste plus qu'à sommer ou soustraire ces fonctions selon la formule (7.10) pour obtenir le développement de l'information de Holevo.

## G.2 Développement de $I_{AB}$

L'information mutuelle entre Alice et Bob est développée sans difficulté à partir de (7.5) :

$$I_{AB} \simeq \frac{g^2 T \nu \lambda^2}{(1-\lambda^2)(1+\kappa) \ln 2} \quad (\text{G.3})$$

## G.3 Développement du taux secret contre les attaques collectives

En utilisant les développements précédents, on obtient le développement du taux secret au premier ordre en  $T$ , pour un canal bruité et une détection homodyne imparfaite :

$$\begin{aligned} \Delta I_{\text{H}}^{\text{NLA}} \simeq P_{\text{suc}} \frac{g^2 T}{(1-\lambda^2)^2 (1+\kappa) 2 \ln 2} \left\{ (1-\lambda^2)^2 (1+\kappa) \epsilon \ln \frac{\epsilon}{2} + (1-\lambda^2)^2 \nu \epsilon \ln [g^2 T] + \nu \lambda^4 \ln \lambda^4 + \right. \\ \left. (\lambda^2 - 1) \left( \nu (\epsilon - \lambda^2 (2\beta + \epsilon)) + (\lambda^2 - 1) \epsilon (\nu - \kappa - 1) \ln \left[ \frac{\epsilon(1+\kappa-\nu)}{2(1+\kappa)} \right] \right) \right\} \quad (\text{G.4}) \end{aligned}$$

## G.4 Développement de $I_{BE}$

On peut également développer l'information mutuelle entre Eve et Bob dans les cas des attaques individuelles, à partir de (7.6) :

$$I_{BE} \simeq \frac{g^2 T \nu [\lambda^4 (\epsilon - 2) - \epsilon]}{(\lambda^4 - 1)(\kappa + 1) \ln 2} \quad (\text{G.5})$$

## Annexe H

# Autre dérivation des paramètres effectifs

La première méthode que nous avons utilisé pour dériver les paramètres effectifs était uniquement basée sur la matrice de covariance entre Alice et Bob [Blandino12c]. Le but était de trouver des paramètres tels que la matrice de covariance de l'état amplifié puisse être retrouvée en envoyant un état EPR de paramètre  $\zeta$  à travers un canal quantique de transmission  $\eta$  et d'excès de bruit  $\epsilon^g$  (Fig. H.1). Nous allons présenter cette méthode, qui, bien que moins générale, peut quand même être intéressante car elle utilise des arguments différents. Nous verrons que les paramètres effectifs  $\eta$  et  $\epsilon^g$  sont identiques à ceux trouvés dans la section 9.3, et celui de l'état EPR correspond à l'amplification de l'état EPR initial avec un NLA de gain  $g_{\text{in}} : \zeta = g_{\text{in}} \lambda$ .

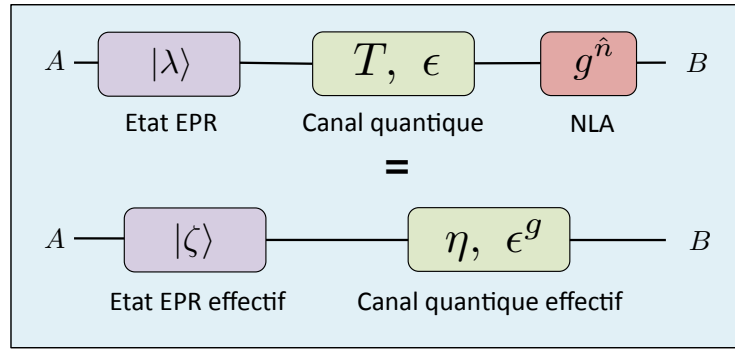


FIGURE H.1 – Système effectif équivalent pour le cas d'un état EPR.

## H.1 Obtention des paramètres effectifs

### H.1.1 Conditions en utilisant un état thermique déplacé

La matrice de covariance Alice Bob sans amplification est donnée par :

$$\Gamma_{AB} = \begin{pmatrix} V(\lambda)\mathbb{I} & \sqrt{T[V(\lambda)^2-1]}\mathbb{Z} \\ \sqrt{T[V(\lambda)^2-1]}\mathbb{Z} & T[V(\lambda) + \frac{1-T}{T} + \epsilon]\mathbb{I} \end{pmatrix} \quad (\text{H.1})$$

avec  $V(\lambda)=V_A+1=\cosh(2r)=\frac{1+\lambda^2}{1-\lambda^2}$ , et  $\epsilon$  le bruit thermique ajouté par le canal, ramené à l'entrée. On note également<sup>1</sup> que  $V_B = 1+TV_A+T\epsilon$ . Nous cherchons donc des paramètres effectifs de manière à exprimer la matrice de covariance de l'état amplifié comme :

$$\Gamma_{AB}^g = \begin{pmatrix} V(\zeta)\mathbb{I} & \sqrt{\eta[V(\zeta)^2-1]}\mathbb{Z} \\ \sqrt{\eta[V(\zeta)^2-1]}\mathbb{Z} & \eta[V(\zeta)+\frac{1-\eta}{\eta}+\epsilon^g]\mathbb{I} \end{pmatrix} \quad (\text{H.2})$$

La première étape de la résolution est identique à celle présentée dans la section 9.3 : il s'agit de considérer l'action du NLA sur un état thermique déplacé. Lorsque la mesure hétérodyne d'Alice est  $\alpha_A$ , le deuxième mode de l'état EPR est projeté sur un état cohérent d'amplitude proportionnelle à  $\lambda\alpha_A$  (annexe I).

Cet état cohérent est ensuite envoyé dans le canal bruité de transmission  $T$ , ce qui le transforme en un état thermique déplacé, de moyenne  $\sqrt{T}\lambda\alpha_A$  et de variance

$$\frac{1+\lambda_{\text{ch}}^2}{1-\lambda_{\text{ch}}^2} = 1 + T\epsilon. \quad (\text{H.3})$$

L'amplification de cet état est similaire à (9.35), en ne tenant pas compte du facteur de normalisation :

$$\hat{T} \left[ \hat{D}(\sqrt{T}\lambda\alpha_A) \hat{\rho}_{\text{th}}(\lambda_{\text{ch}}) \hat{D}^\dagger(\sqrt{T}\lambda\alpha_A) \right] \hat{T} \propto \hat{D}(\tilde{g}\sqrt{T}\lambda\alpha_A) \hat{\rho}_{\text{th}}(g\lambda_{\text{ch}}) \hat{D}^\dagger(\tilde{g}\sqrt{T}\lambda\alpha_A) \quad (\text{H.4})$$

Ceci implique donc la transformation :

$$\begin{aligned} \sqrt{T}\lambda\alpha_A &\xrightarrow{\text{NLA}} \tilde{g}\sqrt{T}\lambda\alpha_A \\ \lambda_{\text{ch}} &\xrightarrow{\text{NLA}} g\lambda_{\text{ch}} \end{aligned} \quad (\text{H.5})$$

avec  $\tilde{g}=g\frac{1-\lambda_{\text{ch}}^2}{1-g\lambda_{\text{ch}}^2}$ .

### H.1.2 Conditions en utilisant un état thermique non déplacé

Lorsque Bob n'a aucune information sur l'état envoyé par Alice (ou de manière équivalente, avant qu'Alice ne fasse sa mesure), il reçoit un thermique non déplacé  $\hat{\rho}_B$ , dont la variance est donnée par (H.1). On peut donc introduire un paramètre  $\lambda^*$ , tel que

$$\hat{\rho}_B = (1-\lambda^{*2}) \sum_{n=0}^{\infty} (\lambda^*)^{2n} |n\rangle\langle n|. \quad (\text{H.6})$$

Le paramètre  $\lambda^*$  doit par définition vérifier  $\frac{1+\lambda^{*2}}{1-\lambda^{*2}}=1+TV_A+T\epsilon$ , ce qui donne

$$\lambda^{*2} = \frac{T[\lambda^2(2-\epsilon)+\epsilon]}{2-\lambda^2[2+T(\epsilon-2)]+T\epsilon}. \quad (\text{H.7})$$

Le NLA transforme ensuite cet état en un autre état thermique de paramètre  $g\lambda^*$ , et de variance  $\frac{1+(g\lambda^*)^2}{1-(g\lambda^*)^2}$ . Compte tenu de (H.7), nous obtenons donc la transformation :

$$\frac{T[\lambda^2(2-\epsilon)+\epsilon]}{2-\lambda^2[2+T(\epsilon-2)]+T\epsilon} \xrightarrow{\text{NLA}} g^2 \frac{T[\lambda^2(2-\epsilon)+\epsilon]}{2-\lambda^2[2+T(\epsilon-2)]+T\epsilon} \quad (\text{H.8})$$

---

1. On utilise la convention  $N_0=1$ .

### H.1.3 Résolution du système

Nous pouvons maintenant chercher les paramètres effectifs qui permettent d'exprimer la matrice de covariance sous la forme (H.2), en utilisant les équations (H.5) et (H.8) :

$$\sqrt{\eta}\zeta = \tilde{g}\sqrt{T}\lambda \quad (\text{H.9a})$$

$$\lambda_{\text{ch}}^g = g\lambda_{\text{ch}} \quad (\text{H.9b})$$

$$\frac{\eta[\zeta^2(2-\epsilon^g) + \epsilon^g]}{2-\zeta^2[2+\eta(\epsilon^g-2)]+\eta\epsilon^g} = g^2 \frac{T[\lambda^2(2-\epsilon) + \epsilon]}{2-\lambda^2[2+T(\epsilon-2)]+T\epsilon} \quad (\text{H.9c})$$

Ici  $\lambda_{\text{ch}}^g$  est tel que  $\frac{1+(\lambda_{\text{ch}}^g)^2}{1-(\lambda_{\text{ch}}^g)^2} = 1 + \eta\epsilon^g$ . Les deux premières équations sont identiques à (9.43a) et (9.43b), avec  $G=1$  et  $\tau=\eta$ . Il suffit ensuite de résoudre le système pour trouver  $\zeta$ ,  $\eta$  et  $\epsilon^g$  en fonction de  $\lambda$ ,  $T$ ,  $\epsilon$  et  $g$  :

$$\boxed{\begin{aligned} \zeta &= \lambda \sqrt{\frac{2 + (g^2-1)(2-\epsilon)T}{2 - (g^2-1)\epsilon T}} \\ \eta &= \frac{g^2 T}{1 + (g^2-1)T[\frac{1}{4}(g^2-1)(\epsilon-2)\epsilon T - \epsilon + 1]} \\ \epsilon^g &= \epsilon + \frac{1}{2}(g^2-1)(2-\epsilon)\epsilon T \end{aligned}} \quad (\text{H.10})$$

Nous obtenons donc bien les mêmes paramètres effectifs que dans la section 9.4, pour le cas particulier où  $G=1$  et  $\eta \leq 1$ .

## H.2 Lien avec les paramètres effectifs dans le cas général

Si cette méthode de résolution n'a une interprétation physique que lorsque  $\eta \leq 1$ , les paramètres (H.10) sont en réalité toujours valables même lorsque  $\eta > 1$ , et sont très simplement reliés aux paramètres (9.44) obtenus dans le cas général.

Considérons le cas général où le canal effectif  $\mathcal{C}_{\text{eff}}$  est constitué d'un amplificateur de gain  $G \geq 1$  et de pertes  $\tau \leq 1$ . Nous avons montré (cf. (9.48)) que l'on peut définir

$$\chi_{\text{ch}} = \frac{G-1}{G} + \frac{1-\tau}{\tau G} = \frac{\tau(G-2)+1}{\tau G} \quad (\text{H.11})$$

comme étant le bruit ajouté par ce canal ramené à l'entrée. On remarque ensuite deux points : d'une part  $\chi_{\text{ch}}$  peut également s'écrire sous la forme

$$\chi_{\text{ch}} = \frac{1-\tau G}{\tau G} + \frac{2(G-1)}{G}, \quad (\text{H.12})$$

et d'autre part, l'égalité suivante est toujours vérifiée :

$$\epsilon^g = \Delta + \frac{2(G-1)}{G} \quad (\text{H.13})$$

où  $\Delta$  est l'excès de bruit du canal effectif dans le cas général, défini par (9.44).

Ainsi, le bruit total ramené à l'entrée  $\chi_{\text{tot}}$  (cf. (9.47)) est égal à :

$$\chi_{\text{tot}} = \Delta + \frac{G-1}{G} + \frac{1-\tau}{\tau G} \quad (\text{H.14a})$$

$$= \Delta + \frac{1-\tau G}{\tau G} + \frac{2(G-1)}{G} \quad (\text{H.14b})$$

$$= \epsilon^g + \frac{1-\eta}{\eta} \quad (\text{H.14c})$$

avec  $\eta = \tau G$ . De ce fait, un état cohérent est donc toujours transformé en un état thermique déplacé, de variance

$$\tau G \left( 1 + \chi_{\text{tot}} \right) = \eta \left( 1 + \epsilon^g + \frac{1-\eta}{\eta} \right) \quad (\text{H.15a})$$

$$= 1 + \eta \epsilon^g. \quad (\text{H.15b})$$

Avec cette définition de  $\epsilon^g$ , tout ce passe donc comme si le canal effectif introduisait des pertes, même si  $\eta > 1$ . Dans ce cas,  $\epsilon^g$  contient une contribution  $\frac{2(G-1)}{G}$  qui tient compte du bruit ajouté par le fait que  $\eta > 1$ .

Finalement, l'expression des paramètres obtenus avec la méthode présentée dans cette annexe est donc valable quel que soit  $\eta$ . Lorsque  $\eta > 1$ ,  $\epsilon^g$  ne correspond pas simplement à l'excès de bruit du canal, mais contient une correction afin de tenir compte du bruit de l'amplification déterministe.

# Annexe I

## Mesure hétérodyne d'Alice

### I.1 Mesure hétérodyne d'un mode d'un état EPR

Dans ce manuscrit, nous avons souvent utilisé le fait qu'une mesure hétérodyne sur un mode d'un état EPR projette l'autre mode sur un état cohérent. Cette annexe montre ce résultat d'une manière simple sans passer par les matrices de covariances. Une autre présentation détaillée peut être également trouvée dans [Grosshans03a].

Une mesure hétérodyne est modélisée par un projecteur [Weedbrook12]

$$\hat{\mathbf{E}}(\alpha) = \frac{1}{\sqrt{\pi}} |\alpha\rangle\langle\alpha|, \quad (\text{I.1})$$

et une mesure  $\alpha$  projette un état  $\hat{\rho}$  sur l'état

$$\hat{\rho}_\alpha = \frac{1}{p_\alpha} \hat{\mathbf{E}}(\alpha) \hat{\rho} \hat{\mathbf{E}}^\dagger(\alpha), \quad (\text{I.2})$$

où  $p_\alpha = \text{Tr}\{\hat{\mathbf{E}}(\alpha) \hat{\rho} \hat{\mathbf{E}}^\dagger(\alpha)\}$  est égale à la probabilité de succès. Puisque nous considérons ici un état EPR  $|\lambda\rangle$  pur, nous pouvons travailler avec des kets :

$$\hat{\mathbf{E}}(\alpha)|\lambda\rangle = \frac{1}{\sqrt{\pi}} \sqrt{1-\lambda^2} |\alpha\rangle \otimes \left( \sum_{n=0}^{\infty} \lambda^n \langle\alpha|n\rangle |n\rangle \right) \quad (\text{I.3a})$$

$$= \frac{1}{\sqrt{\pi}} \sqrt{1-\lambda^2} |\alpha\rangle \otimes e^{-\frac{1}{2}|\alpha|^2} \left( \sum_{n=0}^{\infty} \frac{(\lambda\alpha)^n}{\sqrt{n!}} |n\rangle \right) \quad (\text{I.3b})$$

$$= \frac{1}{\sqrt{\pi}} \sqrt{1-\lambda^2} e^{\frac{1}{2}|\alpha|^2(\lambda^2-1)} |\alpha\rangle \otimes |\lambda\alpha\rangle \quad (\text{I.3c})$$

Le deuxième mode est donc projeté sur l'état cohérent  $|\lambda\alpha\rangle$ . Le premier mode, toujours présent dans le calcul, est absorbé par la détection d'Alice.

La probabilité de succès est donnée par :

$$p_\alpha = \frac{1-\lambda^2}{\pi} e^{|\alpha|^2(\lambda^2-1)} \quad (\text{I.4})$$

On vérifie que l'on a bien  $\int d^2\alpha p_\alpha = 1$  :

$$\int d^2\alpha p_\alpha = \frac{1-\lambda^2}{\pi} \left( \int dx e^{x^2(\lambda^2-1)} \right)^2 = 1 \quad (\text{I.5})$$

Le vecteur déplacement est quand à lui égal à  $2\sqrt{N_0}\lambda(\alpha_x, \alpha_y)$ . On vérifie enfin que l'on retrouve la variance de modulation  $V_A N_0$ , pour la valeur moyenne de la quadrature  $\hat{\mathbf{X}}$  par exemple. :

$$\int d\alpha_x d\alpha_y (2\sqrt{N_0}\lambda\alpha_x)^2 p_\alpha = \sqrt{\frac{1-\lambda^2}{\pi}} \int d\alpha_x (2\sqrt{N_0}\lambda\alpha_x)^2 e^{\alpha_x^2(\lambda^2-1)} \quad (\text{I.6a})$$

$$= \frac{2\lambda^2}{1-\lambda^2} N_0 \quad (\text{I.6b})$$

$$= \left( \frac{1+\lambda^2}{1-\lambda^2} - 1 \right) N_0 \quad (\text{I.6c})$$

$$= V_A N_0 \quad (\text{I.6d})$$

**Remarque :**  $V_A N_0$  correspond à la variance de modulation de l'amplitude des quadratures, et non à la variance de modulation de l'amplitude de l'état cohérent. Il faut donc prendre en compte le facteur  $2\sqrt{N_0}$  dans (I.6a).

## I.2 Mesure hétérodyne d'un état cohérent

Montrons maintenant qu'une mesure hétérodyne introduit une unité de bruit de photon supplémentaire. Soit  $|\beta\rangle$  un état cohérent que l'on souhaite mesurer. On suppose, sans perte de généralité, que  $\beta$  est réel. On fait une mesure hétérodyne avec le projecteur (I.1). La probabilité de mesurer un état  $|\alpha\rangle$  est alors donnée par

$$p_\alpha^{\text{coh}} = \text{Tr}\{\hat{\mathbf{E}}(\alpha)|\beta\rangle\langle\beta|\hat{\mathbf{E}}^\dagger(\alpha)\} = e^{-\frac{1}{2}|\alpha-\beta|^2} = e^{-\frac{1}{2}(\alpha_x-\beta)^2 - \frac{1}{2}\alpha_y^2}. \quad (\text{I.7})$$

La valeur moyenne de l'amplitude mesurée est bien égale à  $2\sqrt{N_0}\beta$  :

$$\int d^2\alpha \frac{1}{\pi} 2\sqrt{N_0}(\alpha_x + i\alpha_y) e^{-\frac{1}{2}(\alpha_x-\beta)^2 - \frac{1}{2}\alpha_y^2} = 2\sqrt{N_0}\beta \quad (\text{I.8})$$

En revanche, la variance de la quadrature  $\hat{\mathbf{X}}$  est égale à :

$$\int d^2\alpha \frac{1}{\pi} \left(2\sqrt{N_0}\alpha_x\right)^2 e^{-\frac{1}{2}(\alpha_x-\beta)^2 - \frac{1}{2}\alpha_y^2} - \left(2\sqrt{N_0}\beta\right)^2 = 2N_0 \quad (\text{I.9})$$

De même pour la variance de la quadrature  $\hat{\mathbf{P}}$  :

$$\int d^2\alpha \frac{1}{\pi} \left(2\sqrt{N_0}\alpha_y\right)^2 e^{-\frac{1}{2}(\alpha_x-\beta)^2 - \frac{1}{2}\alpha_y^2} = 2N_0 \quad (\text{I.10})$$

La détection hétérodyne introduit bien une unité  $N_0$  de bruit supplémentaire sur la variance des quadratures.

# Bibliographie

- [Acín06] A. ACÍN, S. MASSAR, & S. PIRONIO, *Efficient quantum key distribution secure against no-signalling eavesdroppers*, New Journal of Physics **8**, 126 (2006), doi:[10.1088/1367-2630/8/8/126](https://doi.org/10.1088/1367-2630/8/8/126).
- [Acín07] A. ACÍN, N. BRUNNER, N. Gisin, S. MASSAR, S. PIRONIO, & V. SCARANI, *Device independent security of quantum cryptography against collective attacks*, Physical Review Letters **98**, 230501 (2007), doi:[10.1103/PhysRevLett.98.230501](https://doi.org/10.1103/PhysRevLett.98.230501).
- [Adami97] C. ADAMI & N. J. CERF, *von neumann capacity of noisy quantum channels*, Physical Review A **56**, 3470 (1997), doi:[10.1103/PhysRevA.56.3470](https://doi.org/10.1103/PhysRevA.56.3470).
- [Adesso10] G. ADESSO & A. DATTA, *Quantum versus classical correlations in gaussian states*, Physical Review Letters **105**, 030501 (2010), doi:[10.1103/PhysRevLett.105.030501](https://doi.org/10.1103/PhysRevLett.105.030501).
- [Aiello05] A. AIELLO & J. P. WOERDMAN, *Exact quantization of a paraxial electromagnetic field*, Physical Review A **72**, 060101 (2005), doi:[10.1103/PhysRevA.72.060101](https://doi.org/10.1103/PhysRevA.72.060101).
- [Anis12] A. ANIS & A. I. LVOVSKY, *Maximum-likelihood coherent-state quantum process tomography*, New Journal of Physics **14**, 105021 (2012), doi:[10.1088/1367-2630/14/10/105021](https://doi.org/10.1088/1367-2630/14/10/105021).
- [Appel07] J. APPEL, D. HOFFMAN, E. FIGUEROA, & A. I. LVOVSKY, *Electronic noise in optical homodyne tomography*, Physical Review A **75**, 035802 (2007), doi:[10.1103/PhysRevA.75.035802](https://doi.org/10.1103/PhysRevA.75.035802).
- [Appel08] W. APPEL, *Mathématiques pour la physique et les physiciens !*, H K, 4e édition edn. (2008), ISBN 2351410394.
- [Babichev03] S. A. BABICHEV, J. RIES, & A. I. LVOVSKY, *Quantum scissors : Teleportation of single-mode optical states by means of a nonlocal single photon*, Europhysics Letters (EPL) **64**, 1 (2003), doi:[10.1209/epl/i2003-00504-y](https://doi.org/10.1209/epl/i2003-00504-y).
- [Barbieri10] M. BARBIERI, N. SPAGNOLO, M. G. GENONI, F. FERREYROL, R. BLANDINO, M. G. A. PARIS, P. GRANGIER, & R. TUALLE-BROURI, *Non gaussianity of quantum states : An experimental test on single-photon-added coherent states*, Physical Review A **82**, 063833 (2010), doi:[10.1103/PhysRevA.82.063833](https://doi.org/10.1103/PhysRevA.82.063833).
- [Barbieri11] M. BARBIERI, F. FERREYROL, R. BLANDINO, R. TUALLE-BROURI, & P. GRANGIER, *Nondeterministic noiseless amplification of optical signals : a review of recent experiments*, Laser Physics Letters **8**, 411 (2011), doi:[10.1002/lapl.201010143](https://doi.org/10.1002/lapl.201010143).



- [Barnett89] S. M. BARNETT & S. J. D. PHOENIX, *Entropy as a measure of quantum optical correlation*, Physical Review A **40**, 2404 (1989), doi:[10.1103/PhysRevA.40.2404](#).
- [Barnett91] S. M. BARNETT & S. J. D. PHOENIX, *Information theory, squeezing, and quantum correlations*, Physical Review A **44**, 535 (1991), doi:[10.1103/PhysRevA.44.535](#).
- [Barnett03] S. M. BARNETT & P. M. RADMORE, *Methods in Theoretical Quantum Optics*, Oxford University Press (2003), ISBN 9780198563617.
- [Barnett09] S. M. BARNETT & S. CROKE, *Quantum state discrimination*, Advances in Optics and Photonics **1**, 238 (2009), doi:[10.1364/AOP.1.000238](#).
- [Bartlett02] S. D. BARTLETT, B. C. SANDERS, S. L. BRAUNSTEIN, & K. NEMOTO, *Efficient classical simulation of continuous variable quantum information processes*, Physical Review Letters **88**, 097904 (2002), doi:[10.1103/PhysRevLett.88.097904](#).
- [Bellomo12] B. BELLOMO, G. L. GIORGI, F. GALVE, R. LO FRANCO, G. COMPAGNO, & R. ZAMBRINI, *Unified view of correlations using the square-norm distance*, Physical Review A **85**, 032104 (2012), doi:[10.1103/PhysRevA.85.032104](#).
- [Bennett84] C. H. BENNETT & G. BRASSARD, *Quantum cryptography : public-key distribution and coin tossing*, Proc. IEEE Conf. on Comp., Sys. and Sig., Bangalore, India 175 (1984).
- [Bennett92] C. H. BENNETT, *Quantum cryptography using any two nonorthogonal states*, Physical Review Letters **68**, 3121 (1992), doi:[10.1103/PhysRevLett.68.3121](#).
- [Bennett93] C. H. BENNETT, G. BRASSARD, C. CRÉPEAU, R. JOZSA, A. PERES, & W. K. WOOTTERS, *Teleporting an unknown quantum state via dual classical and einstein-podolsky-rosen channels*, Physical Review Letters **70**, 1895 (1993), doi:[10.1103/PhysRevLett.70.1895](#).
- [Bennett01] C. H. BENNETT, D. P. DiVINCENZO, P. W. SHOR, J. A. SMOLIN, B. M. TERHAL, & W. K. WOOTTERS, *Remote state preparation*, Physical Review Letters **87**, 077902 (2001), doi:[10.1103/PhysRevLett.87.077902](#).
- [Blandino12a] R. BLANDINO, F. FERREYROL, M. BARBIERI, P. GRANGIER, & R. TUALLE-BROURI, *Characterization of a  $\pi$ -phase shift quantum gate for coherent-state qubits*, New Journal of Physics **14**, 013017 (2012), doi:[10.1088/1367-2630/14/1/013017](#).
- [Blandino12b] R. BLANDINO, M. G. GENONI, J. ETESSE, M. BARBIERI, M. G. A. PARIS, P. GRANGIER, & R. TUALLE-BROURI, *Homodyne estimation of gaussian quantum discord*, Physical Review Letters **109**, 180402 (2012), doi:[10.1103/PhysRevLett.109.180402](#).
- [Blandino12c] R. BLANDINO, A. LEVERRIER, M. BARBIERI, J. ETESSE, P. GRANGIER, & R. TUALLE-BROURI, *Improving the maximum transmission distance of continuous-variable quantum key distribution using a noiseless amplifier*, Physical Review A **86**, 012327 (2012), doi:[10.1103/PhysRevA.86.012327](#).

- [Bloch06] M. BLOCH, A. THANGARAJ, S. W. McLAUGHLIN, & J. M. MEROLLA, *LDPC-based gaussian key reconciliation*, in *IEEE Information Theory Workshop, 2006. ITW '06 Punta del Este*, 116–120, IEEE (2006), ISBN 1-4244-0035-X, doi:[10.1109/ITW.2006.1633793](https://doi.org/10.1109/ITW.2006.1633793).
- [Blow90] K. J. BLOW, R. LOUDON, S. J. D. PHOENIX, & T. J. SHEPHERD, *Continuum fields in quantum optics*, *Physical Review A* **42**, 4102 (1990), doi:[10.1103/PhysRevA.42.4102](https://doi.org/10.1103/PhysRevA.42.4102).
- [Bongioanni10] I. BONGIOANNI, L. SANSONI, F. SCIARRINO, G. VALLONE, & P. MATALONI, *Experimental quantum process tomography of non-trace-preserving maps*, *Physical Review A* **82**, 042307 (2010), doi:[10.1103/PhysRevA.82.042307](https://doi.org/10.1103/PhysRevA.82.042307).
- [Bouwmeester97] D. BOUWMEESTER, J.-W. PAN, K. MATTLE, M. EIBL, H. WEINFURTER, & A. ZEILINGER, *Experimental quantum teleportation*, *Nature* **390**, 575 (1997), doi:[10.1038/37539](https://doi.org/10.1038/37539).
- [Bowen03] W. P. BOWEN, R. SCHNABEL, P. K. LAM, & T. C. RALPH, *Experimental investigation of criteria for continuous variable entanglement*, *Physical Review Letters* **90**, 043601 (2003), doi:[10.1103/PhysRevLett.90.043601](https://doi.org/10.1103/PhysRevLett.90.043601).
- [Bowen04] W. P. BOWEN, R. SCHNABEL, P. K. LAM, & T. C. RALPH, *Experimental characterization of continuous-variable entanglement*, *Physical Review A* **69**, 012304 (2004), doi:[10.1103/PhysRevA.69.012304](https://doi.org/10.1103/PhysRevA.69.012304).
- [Brask12] J. B. BRASK, N. BRUNNER, D. CAVALCANTI, & A. LEVERRIER, *Bell tests for continuous-variable systems using hybrid measurements and heralded amplifiers*, *Physical Review A* **85**, 042116 (2012), doi:[10.1103/PhysRevA.85.042116](https://doi.org/10.1103/PhysRevA.85.042116).
- [Braunstein94] S. L. BRAUNSTEIN & C. M. CAVES, *Statistical distance and the geometry of quantum states*, *Physical Review Letters* **72**, 3439 (1994), doi:[10.1103/PhysRevLett.72.3439](https://doi.org/10.1103/PhysRevLett.72.3439).
- [Braunstein98] S. L. BRAUNSTEIN & H. J. KIMBLE, *Teleportation of continuous quantum variables*, *Physical Review Letters* **80**, 869 (1998), doi:[10.1103/PhysRevLett.80.869](https://doi.org/10.1103/PhysRevLett.80.869).
- [Braunstein05] S. L. BRAUNSTEIN & P. VAN LOOCK, *Quantum information with continuous variables*, *Reviews of Modern Physics* **77**, 513 (2005), doi:[10.1103/RevModPhys.77.513](https://doi.org/10.1103/RevModPhys.77.513).
- [Brida10] G. BRIDA, I. P. DEGIOVANNI, A. FLORIO, M. GENOVESE, P. GIORDA, A. MEDA, M. G. A. PARIS, & A. SHURUPOV, *Experimental estimation of entanglement at the quantum limit*, *Physical Review Letters* **104**, 100501 (2010), doi:[10.1103/PhysRevLett.104.100501](https://doi.org/10.1103/PhysRevLett.104.100501).
- [Broducth12] A. BRODUCTH, A. DATTA, K. MODI, Á. RIVAS, & C. A. RODRÍGUEZ-ROSARIO, *Vanishing quantum discord is not necessary for completely positive maps*, arXiv :1212.4387 (2012).
- [Brouiri00] R. BROURI, A. BEVERATOS, J.-P. POIZAT, & P. GRANGIER, *Photon antibunching in the fluorescence of individual color centers in diamond*, *Optics Letters* **25**, 1294 (2000), doi:[10.1364/OL.25.001294](https://doi.org/10.1364/OL.25.001294).

- [Calvo05] G. F. CALVO, A. PICON, & E. BAGAN, *A quantum field theory twist to photon angular momentum*, arXiv :quant-ph/0509040 (2005).
- [Cavalcanti11] D. CAVALCANTI, L. AOLITA, S. BOIXO, K. MODI, M. PIANI, & A. WINTER, *Operational interpretations of quantum discord*, Physical Review A **83**, 032324 (2011), doi:[10.1103/PhysRevA.83.032324](https://doi.org/10.1103/PhysRevA.83.032324).
- [Caves82] C. M. CAVES, *Quantum limits on noise in linear amplifiers*, Physical Review D **26**, 1817 (1982), doi:[10.1103/PhysRevD.26.1817](https://doi.org/10.1103/PhysRevD.26.1817).
- [Caves12] C. M. CAVES, J. COMBES, Z. JIANG, & S. PANDEY, *Quantum limits on phase-preserving linear amplifiers*, Physical Review A **86**, 063802 (2012), doi:[10.1103/PhysRevA.86.063802](https://doi.org/10.1103/PhysRevA.86.063802).
- [Cerf96] N. J. CERF & C. ADAMI, *Accessible information in quantum measurement*, arXiv :quant-ph/9611032 (1996).
- [Cerf97] N. J. CERF & C. ADAMI, *Negative entropy and information in quantum mechanics*, Physical Review Letters **79**, 5194 (1997), doi:[10.1103/PhysRevLett.79.5194](https://doi.org/10.1103/PhysRevLett.79.5194).
- [Cerf01] N. J. CERF, M. LÉVY, & G. V. ASSCHE, *Quantum distribution of gaussian keys using squeezed states*, Physical Review A **63**, 052311 (2001), doi:[10.1103/PhysRevA.63.052311](https://doi.org/10.1103/PhysRevA.63.052311).
- [Cerf07] N. J. CERF, G. LEUCHS, & E. S. POLZIK, *Quantum Information With Continuous Variables of Atoms and Light*, Imperial College Press (2007), ISBN 1860947603.
- [Cheffles98] A. CHEFLES, *Unambiguous discrimination between linearly independent quantum states*, Physics Letters A **239**, 339 (1998), doi:[10.1016/S0375-9601\(98\)00064-4](https://doi.org/10.1016/S0375-9601(98)00064-4).
- [Chiao08] R. CHIAO & J. GARRISON, *Quantum Optics*, Oxford University Press, USA (2008), ISBN 9780198508861.
- [Chiuri11] A. CHIURI, G. VALLONE, M. PATERNOSTRO, & P. MATALONI, *Extremal quantum correlations : Experimental study with two-qubit states*, Physical Review A **84**, 020304 (2011), doi:[10.1103/PhysRevA.84.020304](https://doi.org/10.1103/PhysRevA.84.020304).
- [Christ11] A. CHRIST, K. LAIHO, A. ECKSTEIN, K. N. CASSEMIRO, & C. SILBERHORN, *Probing multimode squeezing with correlation functions*, New Journal of Physics **13**, 033027 (2011), doi:[10.1088/1367-2630/13/3/033027](https://doi.org/10.1088/1367-2630/13/3/033027).
- [Chuang97] I. L. CHUANG & M. A. NIELSEN, *Prescription for experimental determination of the dynamics of a quantum black box*, Journal of Modern Optics **44**, 2455 (1997), doi:[10.1080/09500349708231894](https://doi.org/10.1080/09500349708231894).
- [Ciccarello12a] F. CICCARELLO & V. GIOVANNETTI, *Creating quantum correlations through local nonunitary memoryless channels*, Physical Review A **85**, 010102 (2012), doi:[10.1103/PhysRevA.85.010102](https://doi.org/10.1103/PhysRevA.85.010102).
- [Ciccarello12b] F. CICCARELLO & V. GIOVANNETTI, *Local-channel-induced rise of quantum correlations in continuous-variable systems*, Physical Review A **85**, 022108 (2012), doi:[10.1103/PhysRevA.85.022108](https://doi.org/10.1103/PhysRevA.85.022108).
- [Clausen99] J. CLAUSEN, M. DAKNA, L. KNÖLL, & D. G. WELSCH, *Conditional quantum-state transformation at a beam splitter*, Journal of Optics

- B : Quantum and Semiclassical Optics **1**, 332 (1999), doi:[10.1088/1464-4266/1/3/306](https://doi.org/10.1088/1464-4266/1/3/306).
- [Clerk10] A. A. CLERK, M. H. DEVORET, S. M. GIRVIN, F. MARQUARDT, & R. J. SCHOELKOPF, *Introduction to quantum noise, measurement, and amplification*, Reviews of Modern Physics **82**, 1155 (2010), doi:[10.1103/RevModPhys.82.1155](https://doi.org/10.1103/RevModPhys.82.1155).
- [Cleve98] R. CLEVE, A. EKERT, C. MACCHIAVELLO, & M. MOSCA, *Quantum algorithms revisited*, Proceedings of the Royal Society of London. Series A : Mathematical, Physical and Engineering Sciences **454**, 339 (1998), doi:[10.1098/rspa.1998.0164](https://doi.org/10.1098/rspa.1998.0164).
- [Cochrane99] P. T. COCHRANE, G. J. MILBURN, & W. J. MUNRO, *Macroscopically distinct quantum-superposition states as a bosonic code for amplitude damping*, Physical Review A **59**, 2631 (1999), doi:[10.1103/PhysRevA.59.2631](https://doi.org/10.1103/PhysRevA.59.2631).
- [Cohen-Tannoudji96] C. COHEN-TANNOUDJI, J. DUPONT-ROC, & G. GRYNBERG, *Processus d'interaction entre photons et atomes*, EDP Sciences (1996), ISBN 9782868833587.
- [Cohen-Tannoudji97a] C. COHEN-TANNOUDJI, B. DIU, & F. LALOË, *Complément-EII*, in *Mécanique quantique I*, Hermann (1997), ISBN 2705660747.
- [Cohen-Tannoudji97b] C. COHEN-TANNOUDJI, B. DIU, & F. LALOË, *Complément-GV*, in *Mécanique quantique I*, Hermann (1997), ISBN 2705660747.
- [Cohen-Tannoudji97c] C. COHEN-TANNOUDJI, B. DIU, & F. LALOË, *Mécanique quantique I*, Hermann (1997), ISBN 2705660747.
- [Cohen-Tannoudji01] C. COHEN-TANNOUDJI, J. DUPONT-ROC, & G. GRYNBERG, *Photons et atomes : introduction à l'électrodynamique quantique*, EDP Sciences (2001), ISBN 9782868835352.
- [Dakić10] B. DAKIĆ, V. VEDRAL, & Č. BRUKNER, *Necessary and sufficient condition for nonzero quantum discord*, Physical Review Letters **105**, 190502 (2010), doi:[10.1103/PhysRevLett.105.190502](https://doi.org/10.1103/PhysRevLett.105.190502).
- [Dakić12] B. DAKIĆ, Y. O. LIPP, X. MA, M. RINGBAUER, S. KROPATSCHEK, S. BARZ, T. PATEREK, V. VEDRAL, A. ZEILINGER, *et al.*, *Quantum discord as resource for remote state preparation*, Nature Physics **8**, 666 (2012), doi:[10.1038/nphys2377](https://doi.org/10.1038/nphys2377).
- [Darquié05] B. DARQUIÉ, M. P. A. JONES, J. DINGJAN, J. BEUGNON, S. BERGAMINI, Y. SORTAIS, G. MESSIN, A. BROWAEYS, & P. GRANGIER, *Controlled single-photon emission from a single trapped two-level atom*, Science **309**, 454 (2005), doi:[10.1126/science.1113394](https://doi.org/10.1126/science.1113394).
- [Datta05] A. DATTA, S. T. FLAMMIA, & C. M. CAVES, *Entanglement and the power of one qubit*, Physical Review A **72**, 042316 (2005), doi:[10.1103/PhysRevA.72.042316](https://doi.org/10.1103/PhysRevA.72.042316).
- [Datta07] A. DATTA & G. VIDAL, *Role of entanglement and correlations in mixed-state quantum computation*, Physical Review A **75**, 042310 (2007), doi:[10.1103/PhysRevA.75.042310](https://doi.org/10.1103/PhysRevA.75.042310).

- [Datta08] A. DATTA, A. SHAJI, & C. M. CAVES, *Quantum discord and the power of one qubit*, Physical Review Letters **100**, 050502 (2008), doi:[10.1103/PhysRevLett.100.050502](https://doi.org/10.1103/PhysRevLett.100.050502).
- [Dawson06] C. M. DAWSON, H. L. HASELGROVE, & M. A. NIELSEN, *Noise thresholds for optical quantum computers*, Physical Review Letters **96**, 020501 (2006), doi:[10.1103/PhysRevLett.96.020501](https://doi.org/10.1103/PhysRevLett.96.020501).
- [Deutsch92] D. DEUTSCH & R. JOZSA, *Rapid solution of problems by quantum computation*, Proceedings of the Royal Society of London. Series A : Mathematical and Physical Sciences **439**, 553 (1992), doi:[10.1098/rspa.1992.0167](https://doi.org/10.1098/rspa.1992.0167).
- [Devetak05] I. DEVETAK & A. WINTER, *Distillation of secret key and entanglement from quantum states*, Proceedings of the Royal Society A : Mathematical, Physical and Engineering Science **461**, 207 (2005), doi:[10.1098/rspa.2004.1372](https://doi.org/10.1098/rspa.2004.1372).
- [Devitt09] S. J. DEVITT, K. NEMOTO, & W. J. MUNRO, *The idiots guide to quantum error correction*, arXiv :0905.2794 (2009).
- [DiVincenzo95] D. DIVINCENZO, *Quantum computation*, Science **270**, 255 (1995), doi:[10.1126/science.270.5234.255](https://doi.org/10.1126/science.270.5234.255).
- [Duan98] L.-M. DUAN & G.-C. GUO, *Probabilistic cloning and identification of linearly independent quantum states*, Physical Review Letters **80**, 4999 (1998), doi:[10.1103/PhysRevLett.80.4999](https://doi.org/10.1103/PhysRevLett.80.4999).
- [Dunjko12] V. DUNJKO & E. ANDERSSON, *Truly noiseless probabilistic amplification*, Physical Review A **86**, 042322 (2012), doi:[10.1103/PhysRevA.86.042322](https://doi.org/10.1103/PhysRevA.86.042322).
- [Einstein05] A. EINSTEIN, *Über einen die erzeugung und verwandlung des lichtes betreffenden heuristischen gesichtspunkt*, Annalen der Physik **322**, 132 (1905), doi:[10.1002/andp.19053220607](https://doi.org/10.1002/andp.19053220607).
- [Einstein35] A. EINSTEIN, B. PODOLSKY, & N. ROSEN, *Can quantum-mechanical description of physical reality be considered complete ?*, Physical Review **47**, 777 (1935), doi:[10.1103/PhysRev.47.777](https://doi.org/10.1103/PhysRev.47.777).
- [Eisert02] J. EISERT, S. SCHEEL, & M. B. PLENIO, *Distilling gaussian states with gaussian operations is impossible*, Physical Review Letters **89**, 137903 (2002), doi:[10.1103/PhysRevLett.89.137903](https://doi.org/10.1103/PhysRevLett.89.137903).
- [Ekert91] A. K. EKERT, *Quantum cryptography based on bell's theorem*, Physical Review Letters **67**, 661 (1991), doi:[10.1103/PhysRevLett.67.661](https://doi.org/10.1103/PhysRevLett.67.661).
- [Ferraro05] A. FERRARO, S. OLIVARES, & M. G. A. PARIS, *Gaussian states in continuous variable quantum information*, arXiv :quant-ph/0503237 (2005), (Bibliopolis, Napoli, 2005) ISBN 88-7088-483-X.
- [Ferraro10] A. FERRARO, L. AOLITA, D. CAVALCANTI, F. M. CUCCHIETTI, & A. ACÍN, *Almost all quantum states have nonclassical correlations*, Physical Review A **81**, 052318 (2010), doi:[10.1103/PhysRevA.81.052318](https://doi.org/10.1103/PhysRevA.81.052318).
- [Ferreyrol10] F. FERREYROL, M. BARBIERI, R. BLANDINO, S. FOSSIER, R. TUALLE-BROURI, & P. GRANGIER, *Implementation of a nondeterministic optical noiseless amplifier*, Physical Review Letters **104**, 123603 (2010), doi:[10.1103/PhysRevLett.104.123603](https://doi.org/10.1103/PhysRevLett.104.123603).



- [Ferreyrol11a] F. FERREYROL, *Manipulation de champs quantiques mésoscopiques*, Ph.D. thesis, Université Paris-Sud 11 (2011).
- [Ferreyrol11b] F. FERREYROL, R. BLANDINO, M. BARBIERI, R. TUALLE-BROURI, & P. GRANGIER, *Experimental realization of a nondeterministic optical noiseless amplifier*, Physical Review A **83**, 063801 (2011), doi:[10.1103/PhysRevA.83.063801](https://doi.org/10.1103/PhysRevA.83.063801).
- [Filip05] R. FILIP, P. MAREK, & U. L. ANDERSEN, *Measurement-induced continuous-variable quantum interactions*, Physical Review A **71**, 042308 (2005), doi:[10.1103/PhysRevA.71.042308](https://doi.org/10.1103/PhysRevA.71.042308).
- [Fiurášek01] J. FIURÁŠEK, *Optical implementation of continuous-variable quantum cloning machines*, Physical Review Letters **86**, 4942 (2001), doi:[10.1103/PhysRevLett.86.4942](https://doi.org/10.1103/PhysRevLett.86.4942).
- [Fiurášek02] J. FIURÁŠEK, *Gaussian transformations and distillation of entangled gaussian states*, Physical Review Letters **89**, 137904 (2002), doi:[10.1103/PhysRevLett.89.137904](https://doi.org/10.1103/PhysRevLett.89.137904).
- [Fiurášek04] J. FIURÁŠEK, *Optimal probabilistic cloning and purification of quantum states*, Physical Review A **70**, 032308 (2004), doi:[10.1103/PhysRevA.70.032308](https://doi.org/10.1103/PhysRevA.70.032308).
- [Fiurášek09] J. FIURÁŠEK, *Engineering quantum operations on traveling light beams by multiple photon addition and subtraction*, Physical Review A **80**, 053822 (2009), doi:[10.1103/PhysRevA.80.053822](https://doi.org/10.1103/PhysRevA.80.053822).
- [Fiurášek12] J. FIURÁŠEK & N. J. CERF, *Gaussian postselection and virtual noiseless amplification in continuous-variable quantum key distribution*, Physical Review A **86**, 060302 (2012), doi:[10.1103/PhysRevA.86.060302](https://doi.org/10.1103/PhysRevA.86.060302).
- [Fossier09a] S. FOSSIER, E. DIAMANTI, T. DEBUISCHERT, R. TUALLE-BROURI, & P. GRANGIER, *Improvement of continuous-variable quantum key distribution systems by using optical preamplifiers*, Journal of Physics B : Atomic, Molecular and Optical Physics **42**, 114014 (2009), doi:[10.1088/0953-4075/42/11/114014](https://doi.org/10.1088/0953-4075/42/11/114014).
- [Fossier09b] S. FOSSIER, E. DIAMANTI, T. DEBUISCHERT, A. VILLING, R. TUALLE-BROURI, & P. GRANGIER, *Field test of a continuous-variable quantum key distribution prototype*, New Journal of Physics **11**, 045023 (2009), doi:[10.1088/1367-2630/11/4/045023](https://doi.org/10.1088/1367-2630/11/4/045023).
- [Furusawa98] A. FURUSAWA, J. L. SØRENSEN, S. L. BRAUNSTEIN, C. A. FUCHS, H. J. KIMBLE, & E. S. POLZIK, *Unconditional quantum teleportation*, Science **282**, 706 (1998), doi:[10.1126/science.282.5389.706](https://doi.org/10.1126/science.282.5389.706).
- [Gagatsos12] C. N. GAGATSOS, E. KARPOV, & N. J. CERF, *Probabilistic phase-insensitive optical squeezer in compliance with causality*, Physical Review A **86**, 012324 (2012), doi:[10.1103/PhysRevA.86.012324](https://doi.org/10.1103/PhysRevA.86.012324).
- [García-Patrón06] R. GARCÍA-PATRÓN & N. J. CERF, *Unconditional optimality of gaussian attacks against continuous-variable quantum key distribution*, Physical Review Letters **97**, 190503 (2006), doi:[10.1103/PhysRevLett.97.190503](https://doi.org/10.1103/PhysRevLett.97.190503).
- [García-Patrón07] R. GARCÍA-PATRÓN, *Quantum information with optical continuous variables : from bell tests to key Distribution/Information quantique avec*

- variables continues optiques : des tests de bell à la distribution de clé*, Ph.D. thesis, ULB (2007).
- [García-Patrón09] R. GARCÍA-PATRÓN & N. J. CERF, *Continuous-variable quantum key distribution protocols over noisy channels*, Physical Review Letters **102**, 130501 (2009), doi:[10.1103/PhysRevLett.102.130501](https://doi.org/10.1103/PhysRevLett.102.130501).
- [Genoni08] M. G. GENONI, P. GIORDA, & M. G. A. PARIS, *Optimal estimation of entanglement*, Physical Review A **78**, 032303 (2008), doi:[10.1103/PhysRevA.78.032303](https://doi.org/10.1103/PhysRevA.78.032303).
- [Gerhardt11] I. GERHARDT, Q. LIU, A. LAMAS-LINARES, J. SKAAR, C. KURTSIEFER, & V. MAKAROV, *Full-field implementation of a perfect eavesdropper on a quantum cryptography system*, Nature Communications **2**, 349 (2011), doi:[10.1038/ncomms1348](https://doi.org/10.1038/ncomms1348).
- [Gerrits10] T. GERRITS, S. GLANCY, T. S. CLEMENT, B. CALKINS, A. E. LITA, A. J. MILLER, A. L. MIGDALL, S. W. NAM, R. P. MIRIN, *et al.*, *Generation of optical coherent-state superpositions by number-resolved photon subtraction from the squeezed vacuum*, Physical Review A **82**, 031802 (2010), doi:[10.1103/PhysRevA.82.031802](https://doi.org/10.1103/PhysRevA.82.031802).
- [Ghosh86] R. GHOSH, C. K. HONG, Z. Y. OU, & L. MANDEL, *Interference of two photons in parametric down conversion*, Physical Review A **34**, 3962 (1986), doi:[10.1103/PhysRevA.34.3962](https://doi.org/10.1103/PhysRevA.34.3962).
- [Giedke02] G. GIEDKE & J. IGNACIO CIRAC, *Characterization of gaussian operations and distillation of gaussian states*, Physical Review A **66**, 032316 (2002), doi:[10.1103/PhysRevA.66.032316](https://doi.org/10.1103/PhysRevA.66.032316).
- [Gilchrist04] A. GILCHRIST, K. NEMOTO, W. J. MUNRO, T. C. RALPH, S. GLANCY, S. L. BRAUNSTEIN, & G. J. MILBURN, *Schrodinger cats and their power for quantum information processing*, Journal of Optics B : Quantum and Semiclassical Optics **6**, S828 (2004), doi:[10.1088/1464-4266/6/8/032](https://doi.org/10.1088/1464-4266/6/8/032).
- [Gilchrist05] A. GILCHRIST, N. K. LANGFORD, & M. A. NIELSEN, *Distance measures to compare real and ideal quantum processes*, Physical Review A **71**, 062310 (2005), doi:[10.1103/PhysRevA.71.062310](https://doi.org/10.1103/PhysRevA.71.062310).
- [Giorda10] P. GIORDA & M. G. A. PARIS, *Gaussian quantum discord*, Physical Review Letters **105**, 020503 (2010), doi:[10.1103/PhysRevLett.105.020503](https://doi.org/10.1103/PhysRevLett.105.020503).
- [Giorda12] P. GIORDA, M. ALLEGRA, & M. G. A. PARIS, *Quantum discord for gaussian states with non-gaussian measurements*, Physical Review A **86**, 052328 (2012), doi:[10.1103/PhysRevA.86.052328](https://doi.org/10.1103/PhysRevA.86.052328).
- [Gisin02] N. GISIN, G. RIBORDY, W. TITTEL, & H. ZBINDEN, *Quantum cryptography*, Reviews of Modern Physics **74**, 145 (2002), doi:[10.1103/RevModPhys.74.145](https://doi.org/10.1103/RevModPhys.74.145).
- [Gisin07] N. GISIN & R. THEW, *Quantum communication*, Nature Photonics **1**, 165 (2007), doi:[10.1038/nphoton.2007.22](https://doi.org/10.1038/nphoton.2007.22).
- [Gisin10] N. GISIN, S. PIRONIO, & N. SANGOUARD, *Proposal for implementing device-independent quantum key distribution based on a heralded qubit amplifier*, Physical Review Letters **105**, 070501 (2010), doi:[10.1103/PhysRevLett.105.070501](https://doi.org/10.1103/PhysRevLett.105.070501).

- [Glancy04] S. GLANCY, H. M. VASCONCELOS, & T. C. RALPH, *Transmission of optical coherent-state qubits*, Physical Review A **70**, 022317 (2004), doi:[10.1103/PhysRevA.70.022317](#).
- [Gottesman99] D. GOTTESMAN & I. L. CHUANG, *Demonstrating the viability of universal quantum computation using teleportation and single-qubit operations*, Nature **402**, 390 (1999), doi:[10.1038/46503](#).
- [Gottesman01] D. GOTTESMAN, A. KITAEV, & J. PRESKILL, *Encoding a qubit in an oscillator*, Physical Review A **64**, 012310 (2001), doi:[10.1103/PhysRevA.64.012310](#).
- [Grangier98] P. GRANGIER, J. A. LEVENSON, & J.-P. POIZAT, *Quantum non-demolition measurements in optics*, Nature **396**, 537 (1998), doi:[10.1038/25059](#).
- [Grangier04] P. GRANGIER, B. SANDERS, & J. VUCKOVIC, *Focus on single photons on demand*, New Journal of Physics **6** (2004), doi:[10.1088/1367-2630/6/1/E04](#).
- [Greiner96] W. GREINER & J. REINHARDT, *Field Quantization*, Springer (1996), ISBN 9783540591795.
- [Grosshans01] F. GROSSHANS & P. GRANGIER, *Effective quantum efficiency in the pulsed homodyne detection of a  $n$ -photon state*, The European Physical Journal D - Atomic, Molecular, Optical and Plasma Physics **14**, 119 (2001), doi:[10.1007/s100530170243](#).
- [Grosshans02a] F. GROSSHANS, *Communication et cryptographie quantiques avec des variables continues*, Ph.D. thesis, Université Paris Sud - Paris XI (2002).
- [Grosshans02b] F. GROSSHANS & P. GRANGIER, *Continuous variable quantum cryptography using coherent states*, Physical Review Letters **88**, 057902 (2002), doi:[10.1103/PhysRevLett.88.057902](#).
- [Grosshans03a] F. GROSSHANS, N. J. CERF, J. WENGER, R. TUALLE-BROURI, & P. GRANGIER, *Virtual entanglement and reconciliation protocols for quantum cryptography with continuous variables*, Quantum Info. Comput. **3**, 535 (2003).
- [Grosshans03b] F. GROSSHANS, G. VAN ASSCHE, J. WENGER, R. BROURI, N. J. CERF, & P. GRANGIER, *Quantum key distribution using gaussian-modulated coherent states*, Nature **421**, 238 (2003), doi:[10.1038/nature01289](#).
- [Grover96] L. K. GROVER, *A fast quantum mechanical algorithm for database search*, in *Proceedings of the twenty-eighth annual ACM symposium on Theory of computing*, STOC '96, 212–219, ACM, New York, NY, USA (1996), ISBN 0-89791-785-5, doi:[10.1145/237814.237866](#).
- [Grynberg10] G. GRYNBERG, A. ASPECT, & C. FABRE, *Introduction to Quantum Optics : From the Semi-classical Approach to Quantized Light*, Cambridge University Press (2010), ISBN 0521551129.
- [Gu12] M. GU, H. M. CHRZANOWSKI, S. M. ASSAD, T. SYMUL, K. MODI, T. C. RALPH, V. VEDRAL, & P. K. LAM, *Observing the operational significance of discord consumption*, Nature Physics **8**, 671 (2012), doi:[10.1038/nphys2376](#).



- [Hayashi06] M. HAYASHI, *Quantum Information : An Introduction*, Springer, 1 edn. (2006), ISBN 3540302654.
- [Helstrom76] HELSTROM, *Quantum Detection and Estimation Theory*, Elsevier Science (1976), ISBN 9780123400505.
- [Henderson01] L. HENDERSON & V. VEDRAL, *Classical, quantum and total correlations*, Journal of Physics A : Mathematical and General **34**, 6899 (2001), doi:[10.1088/0305-4470/34/35/315](https://doi.org/10.1088/0305-4470/34/35/315).
- [Hijlkema07] M. HIJLKEMA, B. WEBER, H. P. SPECHT, S. C. WEBSTER, A. KUHN, & G. REMPE, *A single-photon server with just one atom*, Nature Physics **3**, 253 (2007), doi:[10.1038/nphys569](https://doi.org/10.1038/nphys569).
- [Hillery00] M. HILLERY, *Quantum cryptography with squeezed states*, Physical Review A **61**, 022309 (2000), doi:[10.1103/PhysRevA.61.022309](https://doi.org/10.1103/PhysRevA.61.022309).
- [Hong85] C. K. HONG & L. MANDEL, *Theory of parametric frequency down conversion of light*, Physical Review A **31**, 2409 (1985), doi:[10.1103/PhysRevA.31.2409](https://doi.org/10.1103/PhysRevA.31.2409).
- [Horodecki05] M. HORODECKI, J. OPPENHEIM, & A. WINTER, *Partial quantum information*, Nature **436**, 673 (2005), doi:[10.1038/nature03909](https://doi.org/10.1038/nature03909).
- [Horodecki09] R. HORODECKI, P. HORODECKI, M. HORODECKI, & K. HORODECKI, *Quantum entanglement*, Reviews of Modern Physics **81**, 865 (2009), doi:[10.1103/RevModPhys.81.865](https://doi.org/10.1103/RevModPhys.81.865).
- [Hudson74] R. HUDSON, *When is the wigner quasi-probability density non-negative ?*, Reports on Mathematical Physics **6**, 249 (1974), doi:[10.1016/0034-4877\(74\)90007-X](https://doi.org/10.1016/0034-4877(74)90007-X).
- [Inoue02] K. INOUE, E. WAKS, & Y. YAMAMOTO, *Differential phase shift quantum key distribution*, Physical Review Letters **89**, 037902 (2002), doi:[10.1103/PhysRevLett.89.037902](https://doi.org/10.1103/PhysRevLett.89.037902).
- [Jackson98] J. D. JACKSON, *Classical Electrodynamics, third edition*, Wiley (1998), ISBN 047130932X.
- [Jamiolkowski72] A. JAMIOLKOWSKI, *Linear transformations which preserve trace and positive semidefiniteness of operators*, Reports on Mathematical Physics **3**, 275 (1972), doi:[10.1016/0034-4877\(72\)90011-0](https://doi.org/10.1016/0034-4877(72)90011-0).
- [Jeffers11] J. JEFFERS, *Optical amplifier-powered quantum optical amplification*, Physical Review A **83**, 053818 (2011), doi:[10.1103/PhysRevA.83.053818](https://doi.org/10.1103/PhysRevA.83.053818).
- [Jeong02] H. JEONG & M. S. KIM, *Efficient quantum computation using coherent states*, Physical Review A **65**, 042305 (2002), doi:[10.1103/PhysRevA.65.042305](https://doi.org/10.1103/PhysRevA.65.042305).
- [Jiang12] Z. JIANG, M. PIANI, & C. M. CAVES, *Ancilla models for quantum operations : for what unitaries does the ancilla state have to be physical ?*, Quantum Information Processing 1–19 (2012), doi:[10.1007/s11128-012-0500-x](https://doi.org/10.1007/s11128-012-0500-x).
- [Josse06] V. JOSSE, M. SABUNCU, N. J. CERF, G. LEUCHS, & U. L. ANDERSEN, *Universal optical amplification without nonlinearity*, Physical Review Letters **96**, 163602 (2006), doi:[10.1103/PhysRevLett.96.163602](https://doi.org/10.1103/PhysRevLett.96.163602).

- [Jouguet11] P. JOUGUET, S. KUNZ-JACQUES, & A. LEVERRIER, *Long distance continuous-variable quantum key distribution with a gaussian modulation*, Physical Review A **84**, 062317 (2011), doi:[10.1103/PhysRevA.84.062317](https://doi.org/10.1103/PhysRevA.84.062317).
- [Jouguet12a] P. JOUGUET & S. KUNZ-JACQUES, *High performance error correction for quantum key distribution using polar codes*, arXiv :1204.5882 (2012).
- [Jouguet12b] P. JOUGUET, S. KUNZ-JACQUES, T. DEBUSSCHERT, S. FOSSIER, E. DIAMANTI, R. ALLÉAUME, R. TUALLE-BROURI, P. GRANGIER, A. LEVERRIER, *et al.*, *Field test of classical symmetric encryption with continuous variables quantum key distribution*, Optics Express **20**, 14030 (2012), doi:[10.1364/OE.20.014030](https://doi.org/10.1364/OE.20.014030).
- [Jouguet12c] P. JOUGUET, S. KUNZ-JACQUES, A. LEVERRIER, P. GRANGIER, & E. DIAMANTI, *Experimental demonstration of long-distance continuous-variable quantum key distribution*, arXiv :1210.6216 (2012).
- [Jozsa94] R. JOZSA, *Fidelity for mixed quantum states*, Journal of Modern Optics **41**, 2315 (1994), doi:[10.1080/09500349414552171](https://doi.org/10.1080/09500349414552171).
- [Kim08] M. S. KIM, *Recent developments in photon-level operations on travelling light fields*, Journal of Physics B : Atomic, Molecular and Optical Physics **41**, 133001 (2008), doi:[10.1088/0953-4075/41/13/133001](https://doi.org/10.1088/0953-4075/41/13/133001).
- [Knight99] K. KNIGHT, *Mathematical Statistics*, Chapman and Hall/CRC, 1 edn. (1999), ISBN 158488178X.
- [Knill98] E. KNILL & R. LAFLAMME, *Power of one bit of quantum information*, Physical Review Letters **81**, 5672 (1998), doi:[10.1103/PhysRevLett.81.5672](https://doi.org/10.1103/PhysRevLett.81.5672).
- [Knill01] E. KNILL, R. LAFLAMME, & G. J. MILBURN, *A scheme for efficient quantum computation with linear optics*, Nature **409**, 46 (2001), doi:[10.1038/35051009](https://doi.org/10.1038/35051009).
- [Kocsis13] S. KOCSIS, G. Y. XIANG, T. C. RALPH, & G. J. PRYDE, *Heralded noiseless amplification of a photon polarization qubit*, Nature Physics **9**, 23 (2013), doi:[10.1038/nphys2469](https://doi.org/10.1038/nphys2469).
- [Kok07] P. KOK, W. J. MUNRO, K. NEMOTO, T. C. RALPH, J. P. DOWLING, & G. J. MILBURN, *Linear optical quantum computing with photonic qubits*, Reviews of Modern Physics **79**, 135 (2007), doi:[10.1103/RevModPhys.79.135](https://doi.org/10.1103/RevModPhys.79.135).
- [Kumar12] R. KUMAR, E. BARRIOS, C. KUPCHAK, & A. I. LVOVSKY, *Experimental characterization of bosonic photon creation and annihilation operators*, arXiv :1210.1150 (2012).
- [Kurtsiefer00] C. KURTSIEFER, S. MAYER, P. ZARDA, & H. WEINFURTER, *Stable solid-state source of single photons*, Physical Review Letters **85**, 290 (2000), doi:[10.1103/PhysRevLett.85.290](https://doi.org/10.1103/PhysRevLett.85.290).
- [La Porta91] A. LA PORTA & R. E. SLUSHER, *Squeezing limits at high parametric gains*, Physical Review A **44**, 2013 (1991), doi:[10.1103/PhysRevA.44.2013](https://doi.org/10.1103/PhysRevA.44.2013).

- [Lamata05] L. LAMATA & J. LEÓN, *Dealing with entanglement of continuous variables : Schmidt decomposition with discrete sets of orthogonal functions*, Journal of Optics B : Quantum and Semiclassical Optics **7**, 224 (2005), doi:[10.1088/1464-4266/7/8/004](https://doi.org/10.1088/1464-4266/7/8/004).
- [Lance05] A. M. LANCE, T. SYMUL, V. SHARMA, C. WEEDBROOK, T. C. RALPH, & P. K. LAM, *No-switching quantum key distribution using broadband modulated coherent light*, Physical Review Letters **95**, 180503 (2005), doi:[10.1103/PhysRevLett.95.180503](https://doi.org/10.1103/PhysRevLett.95.180503).
- [Landau12] L. LANDAU & E. LIFCHITZ, *Physique Théorique Mécanique*, Ellipses Marketing (2012), ISBN 2729894020.
- [Lanyon08] B. P. LANYON, M. BARBIERI, M. P. ALMEIDA, & A. G. WHITE, *Experimental quantum computing without entanglement*, Physical Review Letters **101**, 200501 (2008), doi:[10.1103/PhysRevLett.101.200501](https://doi.org/10.1103/PhysRevLett.101.200501).
- [Law00] C. K. LAW, I. A. WALMSLEY, & J. H. EBERLY, *Continuous frequency entanglement : Effective finite hilbert space and entropy control*, Physical Review Letters **84**, 5304 (2000), doi:[10.1103/PhysRevLett.84.5304](https://doi.org/10.1103/PhysRevLett.84.5304).
- [Leonhardt97] U. LEONHARDT, *Measuring the Quantum State of Light*, Cambridge University Press (1997), ISBN 9780521497305.
- [Leverrier08] A. LEVERRIER, R. ALLÉAUME, J. BOUTROS, G. ZÉMOR, & P. GRANGIER, *Multidimensional reconciliation for a continuous-variable quantum key distribution*, Physical Review A **77**, 042325 (2008), doi:[10.1103/PhysRevA.77.042325](https://doi.org/10.1103/PhysRevA.77.042325).
- [Leverrier09] A. LEVERRIER & P. GRANGIER, *Unconditional security proof of long-distance continuous-variable quantum key distribution with discrete modulation*, Physical Review Letters **102**, 180504 (2009), doi:[10.1103/PhysRevLett.102.180504](https://doi.org/10.1103/PhysRevLett.102.180504).
- [Leverrier10a] A. LEVERRIER & P. GRANGIER, *Simple proof that gaussian attacks are optimal among collective attacks against continuous-variable quantum key distribution with a gaussian modulation*, Physical Review A **81**, 062314 (2010), doi:[10.1103/PhysRevA.81.062314](https://doi.org/10.1103/PhysRevA.81.062314).
- [Leverrier10b] A. LEVERRIER, F. GROSSHANS, & P. GRANGIER, *Finite-size analysis of a continuous-variable quantum key distribution*, Physical Review A **81**, 062343 (2010), doi:[10.1103/PhysRevA.81.062343](https://doi.org/10.1103/PhysRevA.81.062343).
- [Leverrier11] A. LEVERRIER & P. GRANGIER, *Continuous variable quantum-key-distribution protocols with a non-gaussian modulation*, Physical Review A **83**, 042312 (2011), doi:[10.1103/PhysRevA.83.042312](https://doi.org/10.1103/PhysRevA.83.042312).
- [Leverrier12] A. LEVERRIER, Communication privée (2012).
- [Leverrier13] A. LEVERRIER, R. GARCÍA-PATRÓN, R. RENNER, & N. J. CERF, *Security of continuous-variable quantum key distribution against general attacks*, Physical Review Letters **110**, 030502 (2013), doi:[10.1103/PhysRevLett.110.030502](https://doi.org/10.1103/PhysRevLett.110.030502).
- [Li08] N. LI & S. LUO, *Classical states versus separable states*, Physical Review A **78**, 024303 (2008), doi:[10.1103/PhysRevA.78.024303](https://doi.org/10.1103/PhysRevA.78.024303).

- [Lloyd99] S. LLOYD & S. L. BRAUNSTEIN, *Quantum computation over continuous variables*, Physical Review Letters **82**, 1784 (1999), doi:[10.1103/PhysRevLett.82.1784](https://doi.org/10.1103/PhysRevLett.82.1784).
- [Lobino08] M. LOBINO, D. KORYSTOV, C. KUPCHAK, E. FIGUEROA, B. C. SANDERS, & A. I. LVOVSKY, *Complete characterization of quantum-optical processes*, Science **322**, 563 (2008), doi:[10.1126/science.1162086](https://doi.org/10.1126/science.1162086).
- [Lodewyck07] J. LODEWYCK, M. BLOCH, R. GARCIA-PATRON, S. FOSSIER, E. KARPOV, E. DIAMANTI, T. DEBUSSCHERT, N. J. CERF, R. TUALLE-BROURI, *et al.*, *Quantum key distribution over 25km with an all-fiber continuous-variable system*, Physical Review A **76**, 042305 (2007), doi:[10.1103/PhysRevA.76.042305](https://doi.org/10.1103/PhysRevA.76.042305).
- [Loock11] P. VAN LOOCK, *Optical hybrid approaches to quantum information*, Laser & Photonics Reviews **5**, 167–200 (2011), doi:[10.1002/lpor.201000005](https://doi.org/10.1002/lpor.201000005).
- [Loudon00] R. LOUDON, *The quantum theory of light*, Oxford University Press (2000), ISBN 9780198501770.
- [Lund08] A. P. LUND, T. C. RALPH, & H. L. HASELGROVE, *Fault-tolerant linear optical quantum computing with small-amplitude coherent states*, Physical Review Letters **100**, 030503 (2008), doi:[10.1103/PhysRevLett.100.030503](https://doi.org/10.1103/PhysRevLett.100.030503).
- [Luo08] S. LUO, *Using measurement-induced disturbance to characterize correlations as classical or quantum*, Physical Review A **77**, 022301 (2008), doi:[10.1103/PhysRevA.77.022301](https://doi.org/10.1103/PhysRevA.77.022301).
- [Madhok11] V. MADHOK & A. DATTA, *Interpreting quantum discord through quantum state merging*, Physical Review A **83**, 032323 (2011), doi:[10.1103/PhysRevA.83.032323](https://doi.org/10.1103/PhysRevA.83.032323).
- [Madsen12] L. S. MADSEN, A. BERNI, M. LASSEN, & U. L. ANDERSEN, *Experimental investigation of the evolution of gaussian quantum discord in an open system*, Physical Review Letters **109**, 030402 (2012), doi:[10.1103/PhysRevLett.109.030402](https://doi.org/10.1103/PhysRevLett.109.030402).
- [Marek09] P. MAREK & J. FIURÁŠEK, *Resources for universal quantum-state manipulation and engineering*, Physical Review A **79**, 062321 (2009), doi:[10.1103/PhysRevA.79.062321](https://doi.org/10.1103/PhysRevA.79.062321).
- [Marek10a] P. MAREK & R. FILIP, *Coherent-state phase concentration by quantum probabilistic amplification*, Physical Review A **81**, 022302 (2010), doi:[10.1103/PhysRevA.81.022302](https://doi.org/10.1103/PhysRevA.81.022302).
- [Marek10b] P. MAREK & J. FIURÁŠEK, *Elementary gates for quantum information with superposed coherent states*, Physical Review A **82**, 014304 (2010), doi:[10.1103/PhysRevA.82.014304](https://doi.org/10.1103/PhysRevA.82.014304).
- [McKeever04] J. MCKEEVER, A. BOCA, A. D. BOOZER, R. MILLER, J. R. BUCK, A. KUZMICH, & H. J. KIMBLE, *Deterministic generation of single photons from one atom trapped in a cavity*, Science **303**, 1992 (2004), doi:[10.1126/science.1095232](https://doi.org/10.1126/science.1095232).
- [Menzies09] D. MENZIES & S. CROKE, *Noiseless linear amplification via weak measurements*, arXiv :0903.4181 (2009).

- [Merali11] Z. MERALI, *Quantum computing : The power of discord*, Nature News **474**, 24 (2011), doi:[10.1038/474024a](https://doi.org/10.1038/474024a).
- [Mičuda12] M. MIČUDA, I. STRAKA, M. MIKOVA, M. DUSEK, N. J. CERF, J. FIURASEK, & M. JEZEK, *Noiseless loss suppression in quantum optical communication*, Physical Review Letters **109**, 180503 (2012), doi:[10.1103/PhysRevLett.109.180503](https://doi.org/10.1103/PhysRevLett.109.180503).
- [Mitchell03] M. W. MITCHELL, C. W. ELLENOR, S. SCHNEIDER, & A. M. STEINBERG, *Diagnosis, prescription, and prognosis of a bell-state filter by quantum process tomography*, Physical Review Letters **91**, 120402 (2003), doi:[10.1103/PhysRevLett.91.120402](https://doi.org/10.1103/PhysRevLett.91.120402).
- [Mizuno05] J. MIZUNO, K. WAKUI, A. FURUSAWA, & M. SASAKI, *Experimental demonstration of entanglement-assisted coding using a two-mode squeezed vacuum state*, Physical Review A **71**, 012304 (2005), doi:[10.1103/PhysRevA.71.012304](https://doi.org/10.1103/PhysRevA.71.012304).
- [Modi10] K. MODI, T. PATEREK, W. SON, V. VEDRAL, & M. WILLIAMSON, *Unified view of quantum and classical correlations*, Physical Review Letters **104**, 080501 (2010), doi:[10.1103/PhysRevLett.104.080501](https://doi.org/10.1103/PhysRevLett.104.080501).
- [Modi12] K. MODI, A. BRODUTCH, H. CABLE, T. PATEREK, & V. VEDRAL, *The classical-quantum boundary for correlations : Discord and related measures*, Reviews of Modern Physics **84**, 1655 (2012), doi:[10.1103/RevModPhys.84.1655](https://doi.org/10.1103/RevModPhys.84.1655).
- [Mohseni08] M. MOHSENI, A. T. REZAKHANI, & D. A. LIDAR, *Quantum-process tomography : Resource analysis of different strategies*, Physical Review A **77**, 032322 (2008), doi:[10.1103/PhysRevA.77.032322](https://doi.org/10.1103/PhysRevA.77.032322).
- [Moroder12] T. MORODER, M. CURTY, C. C. W. LIM, L. P. THINH, H. ZBINDEN, & N. Gisin, *Security of distributed-phase-reference quantum key distribution*, Physical Review Letters **109**, 260501 (2012), doi:[10.1103/PhysRevLett.109.260501](https://doi.org/10.1103/PhysRevLett.109.260501).
- [Müller12] C. R. MÜLLER, C. WITTMANN, P. MAREK, R. FILIP, C. MARQUARDT, G. LEUCHS, & U. L. ANDERSEN, *Probabilistic cloning of coherent states without a phase reference*, Physical Review A **86**, 010305 (2012), doi:[10.1103/PhysRevA.86.010305](https://doi.org/10.1103/PhysRevA.86.010305).
- [Munro05] W. J. MUNRO, K. NEMOTO, R. G. BEAUSOLEIL, & T. P. SPILLER, *High-efficiency quantum-nondemolition single-photon-number-resolving detector*, Physical Review A **71**, 033819 (2005), doi:[10.1103/PhysRevA.71.033819](https://doi.org/10.1103/PhysRevA.71.033819).
- [Navascues06] M. NAVASCUES, F. GROSSHANS, & A. ACIN, *Optimality of gaussian attacks in continuous-variable quantum cryptography*, Physical Review Letters **97**, 190502 (2006), doi:[10.1103/PhysRevLett.97.190502](https://doi.org/10.1103/PhysRevLett.97.190502).
- [Neergaard-Nielsen06] J. S. NEERGAARD-NIELSEN, B. M. NIELSEN, C. HETTICH, K. MØLMER, & E. S. POLZIK, *Generation of a superposition of odd photon number states for quantum information networks*, Physical Review Letters **97**, 083604 (2006), doi:[10.1103/PhysRevLett.97.083604](https://doi.org/10.1103/PhysRevLett.97.083604).
- [Neergaard-Nielsen10] J. S. NEERGAARD-NIELSEN, M. TAKEUCHI, K. WAKUI, H. TAKAHASHI, K. HAYASAKA, M. TAKEOKA, & M. SASAKI, *Optical*



- continuous-variable qubit*, Physical Review Letters **105**, 053602 (2010), doi:[10.1103/PhysRevLett.105.053602](https://doi.org/10.1103/PhysRevLett.105.053602).
- [Nielsen00] M. A. NIELSEN & I. L. CHUANG, *Quantum computation and quantum information*, Cambridge University Press (2000), ISBN 9780521635035.
- [Niset09] J. NISSET, J. FIURÁŠEK, & N. J. CERF, *No-go theorem for gaussian quantum error correction*, Physical Review Letters **102**, 120501 (2009), doi:[10.1103/PhysRevLett.102.120501](https://doi.org/10.1103/PhysRevLett.102.120501).
- [NIST01] NIST, *Announcing the advanced encryption standard (AES)*, Federal Information Processing Standards Publication **197** (2001).
- [O'Brien03] J. L. O'BRIEN, G. J. PRYDE, A. G. WHITE, T. C. RALPH, & D. BRANNING, *Demonstration of an all-optical quantum controlled-NOT gate*, Nature **426**, 264 (2003), doi:[10.1038/nature02054](https://doi.org/10.1038/nature02054).
- [O'Brien04] J. L. O'BRIEN, G. J. PRYDE, A. GILCHRIST, D. F. V. JAMES, N. K. LANGFORD, T. C. RALPH, & A. G. WHITE, *Quantum process tomography of a controlled-NOT gate*, Physical Review Letters **93**, 080502 (2004), doi:[10.1103/PhysRevLett.93.080502](https://doi.org/10.1103/PhysRevLett.93.080502).
- [O'Brien07] J. L. O'BRIEN, *Optical quantum computing*, Science **318**, 1567 (2007), doi:[10.1126/science.1142892](https://doi.org/10.1126/science.1142892).
- [Ollivier01] H. OLLIVIER & W. H. ZUREK, *Quantum discord : A measure of the quantumness of correlations*, Physical Review Letters **88**, 017901 (2001), doi:[10.1103/PhysRevLett.88.017901](https://doi.org/10.1103/PhysRevLett.88.017901).
- [Osorio12] C. I. OSORIO, N. BRUNO, N. SANGOUARD, H. ZBINDEN, N. GISIN, & R. T. THEW, *Heralded photon amplification for quantum communication*, Physical Review A **86**, 023815 (2012), doi:[10.1103/PhysRevA.86.023815](https://doi.org/10.1103/PhysRevA.86.023815).
- [Ou92] Z. Y. OU, S. F. PEREIRA, H. J. KIMBLE, & K. C. PENG, *Realization of the einstein-podolsky-rosen paradox for continuous variables*, Physical Review Letters **68**, 3663 (1992), doi:[10.1103/PhysRevLett.68.3663](https://doi.org/10.1103/PhysRevLett.68.3663).
- [Ou97] Z. Y. OU, *Parametric down-conversion with coherent pulse pumping and quantum interference between independent fields*, Quantum and Semiclassical Optics : Journal of the European Optical Society Part B **9**, 599 (1997), doi:[10.1088/1355-5111/9/4/009](https://doi.org/10.1088/1355-5111/9/4/009).
- [Ourjoumtsev06a] A. OURJOUTSEV, R. TUALLE-BROURI, & P. GRANGIER, *Quantum homodyne tomography of a two-photon fock state*, Physical Review Letters **96**, 213601 (2006), doi:[10.1103/PhysRevLett.96.213601](https://doi.org/10.1103/PhysRevLett.96.213601).
- [Ourjoumtsev06b] A. OURJOUTSEV, R. TUALLE-BROURI, J. LAURAT, & P. GRANGIER, *Generating optical schrödinger kittens for quantum information processing*, Science **312**, 83 (2006), doi:[10.1126/science.1122858](https://doi.org/10.1126/science.1122858).
- [Ourjoumtsev07a] A. OURJOUTSEV, *Étude théorique et expérimentale de superpositions quantiques cohérentes et d'états intriqués non-gaussiens de la lumière*, Ph.D. thesis, Université Paris Sud - Paris XI (2007).
- [Ourjoumtsev07b] A. OURJOUTSEV, A. DANTAN, R. TUALLE-BROURI, & P. GRANGIER, *Increasing entanglement between gaussian states by coherent photon subtraction*, Physical Review Letters **98**, 030502 (2007), doi:[10.1103/PhysRevLett.98.030502](https://doi.org/10.1103/PhysRevLett.98.030502).

- [Ourjoumtsev07c] A. OURJOUTSEV, H. JEONG, R. TUALLE-BROURI, & P. GRANGIER, *Generation of optical ‘Schrödinger cats’ from photon number states*, Nature **448**, 784 (2007), doi:[10.1038/nature06054](https://doi.org/10.1038/nature06054).
- [Ourjoumtsev09] A. OURJOUTSEV, F. FERREYROL, R. TUALLE-BROURI, & P. GRANGIER, *Preparation of non-local superpositions of quasi-classical light states*, Nature Physics **5**, 189 (2009), doi:[10.1038/nphys1199](https://doi.org/10.1038/nphys1199).
- [Parigi07] V. PARIGI, A. ZAVATTA, M. KIM, & M. BELLINI, *Probing quantum commutation rules by addition and subtraction of single photons to/from a light field*, Science **317**, 1890 (2007), doi:[10.1126/science.1146204](https://doi.org/10.1126/science.1146204).
- [Paris96] M. G. PARIS, *Displacement operator by beam splitter*, Physics Letters A **217**, 78 (1996), doi:[10.1016/0375-9601\(96\)00339-8](https://doi.org/10.1016/0375-9601(96)00339-8).
- [Paris09] M. G. A. PARIS, *Quantum estimation for quantum technology*, International Journal of Quantum Information **07**, 125 (2009), doi:[10.1142/S0219749909004839](https://doi.org/10.1142/S0219749909004839).
- [Paris12] M. G. A. PARIS, *The modern tools of quantum mechanics*, The European Physical Journal Special Topics **203**, 61 (2012), doi:[10.1140/epjst/e2012-01535-1](https://doi.org/10.1140/epjst/e2012-01535-1).
- [Parker00] S. PARKER, S. BOSE, & M. B. PLENIO, *Entanglement quantification and purification in continuous-variable systems*, Physical Review A **61**, 032305 (2000), doi:[10.1103/PhysRevA.61.032305](https://doi.org/10.1103/PhysRevA.61.032305).
- [Partanen12] M. PARTANEN, T. HÄYRYNEN, J. OKSANEN, & J. TULKKI, *Noiseless amplification of weak coherent fields exploiting energy fluctuations of the field*, Physical Review A **86**, 063804 (2012), doi:[10.1103/PhysRevA.86.063804](https://doi.org/10.1103/PhysRevA.86.063804).
- [Passante11] G. PASSANTE, O. MOUSSA, D. A. TROTTIER, & R. LA-FLAMME, *Experimental detection of nonclassical correlations in mixed-state quantum computation*, Physical Review A **84**, 044302 (2011), doi:[10.1103/PhysRevA.84.044302](https://doi.org/10.1103/PhysRevA.84.044302).
- [Peev09] M. PEEV, C. PACHER, R. ALLEAUME, C. BARREIRO, J. BOUDA, *et al.*, *The SECOQC quantum key distribution network in vienna*, New Journal of Physics **11**, 075001 (2009), doi:[10.1088/1367-2630/11/7/075001](https://doi.org/10.1088/1367-2630/11/7/075001).
- [Pegg98] D. T. PEGG, L. S. PHILLIPS, & S. M. BARNETT, *Optical state truncation by projection synthesis*, Physical Review Letters **81**, 1604 (1998), doi:[10.1103/PhysRevLett.81.1604](https://doi.org/10.1103/PhysRevLett.81.1604).
- [Peskin95] M. E. PESKIN & D. V. SCHROEDER, *An Introduction To Quantum Field Theory*, Westview Press (1995), ISBN 0201503972.
- [Peyronel12] T. PEYRONEL, O. FIRSTENBERG, Q.-Y. LIANG, S. HOFFERBERTH, A. V. GORSHKOV, T. POHL, M. D. LUKIN, & V. VULETIĆ, *Quantum nonlinear optics with single photons enabled by strongly interacting atoms*, Nature **488**, 57 (2012), doi:[10.1038/nature11361](https://doi.org/10.1038/nature11361).
- [Piani11] M. PIANI, S. GHARIBIAN, G. ADESSO, J. CALSAMIGLIA, P. HORODECKI, & A. WINTER, *All nonclassical correlations can be activated into distillable entanglement*, Physical Review Letters **106**, 220403 (2011), doi:[10.1103/PhysRevLett.106.220403](https://doi.org/10.1103/PhysRevLett.106.220403).

- [Pirandola08] S. PIRANDOLA, S. L. BRAUNSTEIN, & S. LLOYD, *Characterization of collective gaussian attacks and security of coherent-state quantum cryptography*, Physical Review Letters **101**, 200504 (2008), doi:[10.1103/PhysRevLett.101.200504](https://doi.org/10.1103/PhysRevLett.101.200504).
- [Plenio05] M. B. PLENIO & S. VIRMANI, *An introduction to entanglement measures*, arXiv :quant-ph/0504163 (2005), quant. Inf. Comput. **7** :1-51,2007.
- [Poyatos97] J. F. POYATOS, J. I. CIRAC, & P. ZOLLER, *Complete characterization of a quantum process : The two-bit quantum gate*, Physical Review Letters **78**, 390 (1997), doi:[10.1103/PhysRevLett.78.390](https://doi.org/10.1103/PhysRevLett.78.390).
- [Preskill98] J. PRESKILL, *Quantum information and computation*, <http://www.theory.caltech.edu/people/preskill/ph229/#lecture> (1998).
- [Puri01] R. R. PURI, *Mathematical methods of quantum optics*, Springer (2001), ISBN 9783540678021.
- [Qi10] B. QI, W. ZHU, L. QIAN, & H.-K. LO, *Feasibility of quantum key distribution through a dense wavelength division multiplexing network*, New Journal of Physics **12**, 103042 (2010), doi:[10.1088/1367-2630/12/10/103042](https://doi.org/10.1088/1367-2630/12/10/103042).
- [Rahimi-Keshari11] S. RAHIMI-KESHARI, A. SCHERER, A. MANN, A. T. REZAKHANI, A. I. LVOVSKY, & B. C. SANDERS, *Quantum process tomography with coherent states*, New Journal of Physics **13**, 013006 (2011), doi:[10.1088/1367-2630/13/1/013006](https://doi.org/10.1088/1367-2630/13/1/013006).
- [Rahimi-Keshari13] S. RAHIMI-KESHARI, C. M. CAVES, & T. C. RALPH, *Measurement-based method for verifying quantum discord*, Physical Review A **87**, 012119 (2013), doi:[10.1103/PhysRevA.87.012119](https://doi.org/10.1103/PhysRevA.87.012119).
- [Ralph99] T. C. RALPH, *Continuous variable quantum cryptography*, Physical Review A **61**, 010303 (1999), doi:[10.1103/PhysRevA.61.010303](https://doi.org/10.1103/PhysRevA.61.010303).
- [Ralph02] T. C. RALPH, W. J. MUNRO, & G. J. MILBURN, *Quantum computation based on linear optics*, vol. 4917, 1–12 (2002), doi:[10.1117/12.483016](https://doi.org/10.1117/12.483016).
- [Ralph03] T. C. RALPH, A. GILCHRIST, G. J. MILBURN, W. J. MUNRO, & S. GLANCY, *Quantum computation with optical coherent states*, Physical Review A **68**, 042319 (2003), doi:[10.1103/PhysRevA.68.042319](https://doi.org/10.1103/PhysRevA.68.042319).
- [Ralph08] T. C. RALPH & A. P. LUND, *Nondeterministic noiseless linear amplification of quantum systems*, arXiv :0809.0326 (2008), quantum Communication Measurement and Computing Proceedings of 9th International Conference, Ed. A.Lvovsky, 155-160 (AIP, New York 2009).
- [Ralph11] T. C. RALPH, *Quantum error correction of continuous-variable states against gaussian noise*, Physical Review A **84**, 022339 (2011), doi:[10.1103/PhysRevA.84.022339](https://doi.org/10.1103/PhysRevA.84.022339).
- [Reeves91] R. J. REEVES, M. G. JANI, B. JASSEMNEJAD, R. C. POWELL, G. J. MIZELL, & W. FAY, *Photorefractive properties of KNbO<sub>3</sub>*, Physical Review B **43**, 71 (1991), doi:[10.1103/PhysRevB.43.71](https://doi.org/10.1103/PhysRevB.43.71).
- [Reid00] M. D. REID, *Quantum cryptography with a predetermined key, using continuous-variable einstein-podolsky-rosen correlations*, Physical Review A **62**, 062308 (2000), doi:[10.1103/PhysRevA.62.062308](https://doi.org/10.1103/PhysRevA.62.062308).



- [Renner05a] R. RENNER, *Security of quantum key distribution*, arXiv :quant-ph/0512258 (2005).
- [Renner05b] R. RENNER & R. KÖNIG, *Universally composable privacy amplification against quantum adversaries*, in D. HUTCHISON, T. KANADE, J. KITTLER, J. M. KLEINBERG, F. MATTERN, J. C. MITCHELL, M. NAOR, O. NIERSTRASZ, C. PANDU RANGAN, *et al.*, eds., *Theory of Cryptography*, vol. 3378, 407–425, Springer Berlin Heidelberg (2005), ISBN 978-3-540-24573-5, 978-3-540-30576-7.
- [Renner09] R. RENNER & J. I. CIRAC, *de finetti representation theorem for infinite-dimensional quantum systems and applications to quantum cryptography*, Physical Review Letters **102**, 110504 (2009), doi:[10.1103/PhysRevLett.102.110504](https://doi.org/10.1103/PhysRevLett.102.110504).
- [Rivest78] R. L. RIVEST, A. SHAMIR, & L. ADLEMAN, *A method for obtaining digital signatures and public-key cryptosystems*, Commun. ACM **21**, 120–126 (1978), doi:[10.1145/359340.359342](https://doi.org/10.1145/359340.359342).
- [Rodríguez-Rosario08] C. A. RODRÍGUEZ-ROSARIO, K. MODI, A.-M. KUAH, A. SHAJI, & E. C. G. SUDARSHAN, *Completely positive maps and classical correlations*, Journal of Physics A : Mathematical and Theoretical **41**, 205301 (2008), doi:[10.1088/1751-8113/41/20/205301](https://doi.org/10.1088/1751-8113/41/20/205301).
- [Rohde07] P. P. ROHDE, W. MAUERER, & C. SILBERHORN, *Spectral structure and decompositions of optical states, and their applications*, New Journal of Physics **9**, 91 (2007), doi:[10.1088/1367-2630/9/4/091](https://doi.org/10.1088/1367-2630/9/4/091).
- [Rosenberg05] D. ROSENBERG, A. E. LITA, A. J. MILLER, & S. W. NAM, *Noise-free high-efficiency photon-number-resolving detectors*, Physical Review A **71**, 061803 (2005), doi:[10.1103/PhysRevA.71.061803](https://doi.org/10.1103/PhysRevA.71.061803).
- [Saffman10] M. SAFFMAN, T. G. WALKER, & K. MØLMER, *Quantum information with rydberg atoms*, Reviews of Modern Physics **82**, 2313 (2010), doi:[10.1103/RevModPhys.82.2313](https://doi.org/10.1103/RevModPhys.82.2313).
- [Saleh91] B. E. A. SALEH & M. C. TEICH, *Fundamentals of Photonics*, Wiley, 1st edn. (1991), ISBN 0471839655.
- [Sasaki11] M. SASAKI, M. FUJIWARA, H. ISHIZUKA, W. KLAUS, K. WAKUI, M. TAKEOKA, S. MIKI, T. YAMASHITA, Z. WANG, *et al.*, *Field test of quantum key distribution in the tokyo QKD network*, Optics Express **19**, 10387 (2011), doi:[10.1364/OE.19.010387](https://doi.org/10.1364/OE.19.010387).
- [Scarani04] V. SCARANI, A. ACÍN, G. RIBORDY, & N. Gisin, *Quantum cryptography protocols robust against photon number splitting attacks for weak laser pulse implementations*, Physical Review Letters **92**, 057901 (2004), doi:[10.1103/PhysRevLett.92.057901](https://doi.org/10.1103/PhysRevLett.92.057901).
- [Scarani05] V. SCARANI, S. IBLISDIR, N. Gisin, & A. ACÍN, *Quantum cloning*, Reviews of Modern Physics **77**, 1225 (2005), doi:[10.1103/RevModPhys.77.1225](https://doi.org/10.1103/RevModPhys.77.1225).
- [Scarani09] V. SCARANI, H. BECHMANN-PASQUINUCCI, N. J. CERF, M. DUŠEK, N. LÜTKENHAUS, & M. PEEV, *The security of practical quantum key distribution*, Reviews of Modern Physics **81**, 1301 (2009), doi:[10.1103/RevModPhys.81.1301](https://doi.org/10.1103/RevModPhys.81.1301).

- [Schrödinger35] E. SCHRÖDINGER, *Die gegenwertige situation in der quantenmechanik*, Die Naturwissenschaften **23**, 807 (1935), doi:[10.1007/BF01491891](https://doi.org/10.1007/BF01491891).
- [Shabani09] A. SHABANI & D. A. LIDAR, *Vanishing quantum discord is necessary and sufficient for completely positive maps*, Physical Review Letters **102**, 100402 (2009), doi:[10.1103/PhysRevLett.102.100402](https://doi.org/10.1103/PhysRevLett.102.100402).
- [Shannon48] C. E. SHANNON, *A mathematical theory of communication*, Bell System Technical Journal **27**, 623 (1948).
- [Shannon49] C. SHANNON, *Communication theory of secrecy systems*, Bell System Technical Journal **28**, 656 (1949).
- [Shor97] P. W. SHOR, *Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer*, SIAM Journal on Computing **26**, 1484 (1997), doi:[10.1137/S0097539795293172](https://doi.org/10.1137/S0097539795293172).
- [Simon94] R. SIMON, N. MUKUNDA, & B. DUTTA, *Quantum-noise matrix for multimode systems :  $U(n)$  invariance, squeezing, and normal forms*, Physical Review A **49**, 1567 (1994), doi:[10.1103/PhysRevA.49.1567](https://doi.org/10.1103/PhysRevA.49.1567).
- [Sivia06] D. SIVIA & J. SKILLING, *Data Analysis : A Bayesian Tutorial*, Oxford University Press, USA, 2 edn. (2006), ISBN 0198568320.
- [Slusher87] R. E. SLUSHER, P. GRANGIER, A. LAPORTA, B. YURKE, & M. J. POTASEK, *Pulsed squeezed light*, Physical Review Letters **59**, 2566 (1987), doi:[10.1103/PhysRevLett.59.2566](https://doi.org/10.1103/PhysRevLett.59.2566).
- [Stoler72] D. STOLER & S. NEWMAN, *Minimum uncertainty and density matrices*, Physics Letters A **38**, 433 (1972), doi:[10.1016/0375-9601\(72\)90240-X](https://doi.org/10.1016/0375-9601(72)90240-X).
- [Streltsov11] A. STRELTSOV, H. KAMPERMANN, & D. BRUSS, *Linking quantum discord to entanglement in a measurement*, Physical Review Letters **106**, 160401 (2011), doi:[10.1103/PhysRevLett.106.160401](https://doi.org/10.1103/PhysRevLett.106.160401).
- [Stucki05] D. STUCKI, N. BRUNNER, N. GISIN, V. SCARANI, & H. ZBINDEN, *Fast and simple one-way quantum key distribution*, Applied Physics Letters **87**, 194108 (2005), doi:[doi :10.1063/1.2126792](https://doi.org/10.1063/1.2126792).
- [Stucki09] D. STUCKI, N. WALENTA, F. VANNEL, R. T. THEW, N. GISIN, H. ZBINDEN, S. GRAY, C. R. TOWERY, & S. TEN, *High rate, long-distance quantum key distribution over 250 km of ultra low loss fibres*, New Journal of Physics **11**, 075003 (2009), doi:[10.1088/1367-2630/11/7/075003](https://doi.org/10.1088/1367-2630/11/7/075003).
- [Sych10] D. SYCH & G. LEUCHS, *Coherent state quantum key distribution with multi letter phase-shift keying*, New Journal of Physics **12**, 053019 (2010), doi:[10.1088/1367-2630/12/5/053019](https://doi.org/10.1088/1367-2630/12/5/053019).
- [Takahashi08] H. TAKAHASHI, K. WAKUI, S. SUZUKI, M. TAKEOKA, K. HAYASAKA, A. FURUSAWA, & M. SASAKI, *Generation of large-amplitude coherent-state superposition via ancilla-assisted photon subtraction*, Physical Review Letters **101**, 233605 (2008), doi:[10.1103/PhysRevLett.101.233605](https://doi.org/10.1103/PhysRevLett.101.233605).
- [Tipsmark11] A. TIPSMARK, R. DONG, A. LAGHAOUT, P. MAREK, M. JEŽEK, & U. L. ANDERSEN, *Experimental demonstration of a hadamard gate for coherent state qubits*, Physical Review A **84**, 050301 (2011), doi:[10.1103/PhysRevA.84.050301](https://doi.org/10.1103/PhysRevA.84.050301).

- [Toussaint11] U. VON TOUSSAINT, *Bayesian inference in physics*, Reviews of Modern Physics **83**, 943 (2011), doi:[10.1103/RevModPhys.83.943](https://doi.org/10.1103/RevModPhys.83.943).
- [Truax85] D. R. TRUAX, *Baker-campbell-hausdorff relations and unitarity of  $SU(2)$  and  $SU(1,1)$  squeeze operators*, Physical Review D **31**, 1988 (1985), doi:[10.1103/PhysRevD.31.1988](https://doi.org/10.1103/PhysRevD.31.1988).
- [Tualle-Brouri09] R. TUALLE-BROURI, A. OURJOUTSEV, A. DANTAN, P. GRAN-  
GIER, M. WUBS, & A. S. SØRENSEN, *Multimode model for projective  
photon-counting measurements*, Physical Review A **80**, 013806 (2009),  
doi:[10.1103/PhysRevA.80.013806](https://doi.org/10.1103/PhysRevA.80.013806).
- [Usuga10] M. A. USUGA, C. R. MÜLLER, C. WITTMANN, P. MAREK, R. FI-  
LIP, C. MARQUARDT, G. LEUCHS, & U. L. ANDERSEN, *Noise-powered  
probabilistic concentration of phase information*, Nature Physics **6**, 767  
(2010), doi:[10.1038/nphys1743](https://doi.org/10.1038/nphys1743).
- [Vedral02] V. VEDRAL, *The role of relative entropy in quantum infor-  
mation theory*, Reviews of Modern Physics **74**, 197 (2002),  
doi:[10.1103/RevModPhys.74.197](https://doi.org/10.1103/RevModPhys.74.197).
- [Vedral07] V. VEDRAL, *Introduction to Quantum Information Science*, Oxford Uni-  
versity Press, USA (2007), ISBN 0199215707.
- [Vidal99] G. VIDAL, *Entanglement of pure states for a single copy*, Physical Review  
Letters **83**, 1046 (1999), doi:[10.1103/PhysRevLett.83.1046](https://doi.org/10.1103/PhysRevLett.83.1046).
- [Waks03] E. WAKS, K. INOUE, W. OLIVER, E. DIAMANTI, & Y. YAMAMOTO,  
*High-efficiency photon-number detection for quantum information pro-  
cessing*, IEEE Journal of Selected Topics in Quantum Electronics **9**, 1502  
(2003), doi:[10.1109/JSTQE.2003.820917](https://doi.org/10.1109/JSTQE.2003.820917).
- [Walk12] N. WALK, A. P. LUND, & T. C. RALPH, *Non-deterministic noi-  
seless amplification via non-symplectic phase space transformations*,  
arXiv :1211.3794 (2012).
- [Walk13] N. WALK, T. C. RALPH, T. SYMUL, & P. K. LAM, *Security of  
continuous-variable quantum cryptography with gaussian postselection*,  
Physical Review A **87**, 020303 (2013), doi:[10.1103/PhysRevA.87.020303](https://doi.org/10.1103/PhysRevA.87.020303).
- [Wang12] S. WANG, W. CHEN, J.-F. GUO, Z.-Q. YIN, H.-W. LI, Z. ZHOU,  
G.-C. GUO, & Z.-F. HAN, *2 GHz clock quantum key distribution  
over 260 km of standard telecom fiber*, Optics Letters **37**, 1008 (2012),  
doi:[10.1364/OL.37.001008](https://doi.org/10.1364/OL.37.001008).
- [Warusfel04] A. WARUSFEL & C. DESCHAMPS, *Mathematiques "Tout-en-un", 2e  
année : Cours et exercices corrigés*, Dunod, 2e edn. (2004), ISBN  
2100075764.
- [Wasilewski06] W. WASILEWSKI, A. I. LVOVSKY, K. BANASZEK, & C. RADZEWICZ,  
*Pulsed squeezed light : Simultaneous squeezing of multiple modes*, Physi-  
cal Review A **73**, 063819 (2006), doi:[10.1103/PhysRevA.73.063819](https://doi.org/10.1103/PhysRevA.73.063819).
- [Weedbrook04] C. WEEDBROOK, A. M. LANCE, W. P. BOWEN, T. SY-  
MUL, T. C. RALPH, & P. K. LAM, *Quantum cryptography  
without switching*, Physical Review Letters **93**, 170504 (2004),  
doi:[10.1103/PhysRevLett.93.170504](https://doi.org/10.1103/PhysRevLett.93.170504).

- [Weedbrook06] C. WEEDBROOK, A. M. LANCE, W. P. BOWEN, T. SYMUL, T. C. RALPH, & P. K. LAM, *Coherent state quantum key distribution without random basis switching*, Physical Review A **73**, 022316 (2006), doi:[10.1103/PhysRevA.73.022316](#).
- [Weedbrook12] C. WEEDBROOK, S. PIRANDOLA, R. GARCÍA-PATRÓN, N. J. CERF, T. C. RALPH, J. H. SHAPIRO, & S. LLOYD, *Gaussian quantum information*, Reviews of Modern Physics **84**, 621 (2012), doi:[10.1103/RevModPhys.84.621](#).
- [Wehrl78] A. WEHRL, *General properties of entropy*, Reviews of Modern Physics **50**, 221 (1978), doi:[10.1103/RevModPhys.50.221](#).
- [Wenger04a] J. WENGER, *Dispositifs impulsionsnels pour la communication quantique à variables continues*, Ph.D. thesis, Université Paris Sud - Paris XI (2004).
- [Wenger04b] J. WENGER, R. TUALLE-BROURI, & P. GRANGIER, *Non-gaussian statistics from individual pulses of squeezed light*, Physical Review Letters **92**, 153601 (2004), doi:[10.1103/PhysRevLett.92.153601](#).
- [Wenger04c] J. WENGER, R. TUALLE-BROURI, & P. GRANGIER, *Pulsed homodyne measurements of femtosecond squeezed pulses generated by single-pass parametric deamplification*, Optics Letters **29**, 1267 (2004), doi:[10.1364/OL.29.001267](#).
- [Wootters82] W. K. WOOTTERS & W. H. ZUREK, *A single quantum cannot be cloned*, Nature **299**, 802 (1982), doi:[10.1038/299802a0](#).
- [Xiang10] G. Y. XIANG, T. C. RALPH, A. P. LUND, N. WALK, & G. J. PRYDE, *Heralded noiseless linear amplification and distillation of entanglement*, Nature Photonics **4**, 316 (2010), doi:[10.1038/nphoton.2010.35](#).
- [Yang12] S. YANG, S. ZHANG, X. ZOU, S. BI, & X. LIN, *Continuous-variable entanglement distillation with noiseless linear amplification*, Physical Review A **86**, 062321 (2012), doi:[10.1103/PhysRevA.86.062321](#).
- [Yuen83] H. P. YUEN & V. W. S. CHAN, *Noise in homodyne and heterodyne detection*, Optics Letters **8**, 177 (1983), doi:[10.1364/OL.8.000177](#).
- [Zavatta09] A. ZAVATTA, V. PARIGI, M. S. KIM, H. JEONG, & M. BELLINI, *Experimental demonstration of the bosonic commutation relation via superpositions of quantum operations on thermal light fields*, Physical Review Letters **103**, 140406 (2009), doi:[10.1103/PhysRevLett.103.140406](#).
- [Zavatta11] A. ZAVATTA, J. FIURASEK, & M. BELLINI, *A high-fidelity noiseless amplifier for quantum light states*, Nat Photon **5**, 52 (2011), doi:[10.1038/nphoton.2010.260](#).
- [Zhang00] Y. ZHANG, H. WANG, X. LI, J. JING, C. XIE, & K. PENG, *Experimental generation of bright two-mode quadrature squeezed light from a narrow-band nondegenerate optical parametric amplifier*, Physical Review A **62**, 023813 (2000), doi:[10.1103/PhysRevA.62.023813](#).
- [Zysset92] B. ZYSSET, I. BIAGGIO, & P. GÜNTHER, *Refractive indices of orthorhombic KNbO<sub>3</sub>. i. dispersion and temperature dependence*, Journal of the Optical Society of America B **9**, 380 (1992), doi:[10.1364/JOSAB.9.000380](#).

## Résumé

Cette thèse s'inscrit dans le cadre de l'information quantique avec des variables continues, en utilisant des états quantiques du champ électromagnétique. En combinant les outils propres aux variables discrètes, où la lumière est décrite en termes de photons, avec les outils des variables continues, où la lumière est décrite en termes de quadratures, nous pouvons étudier théoriquement et produire expérimentalement des états non-classiques, ainsi que des protocoles élémentaires d'information quantique. Ainsi, nous avons produit expérimentalement un état «chat de Schrödinger», superposition quantique de deux états lumineux quasi-classiques, sur lequel nous avons appliqué une porte quantique introduisant une phase dans la superposition. Nous avons ensuite analysé la qualité de cette porte en utilisant un modèle simple de notre expérience. Nous nous sommes ensuite intéressés aux corrélations quantiques, mesurées par la discordance quantique, pour une classe d'états particulièrement importants en information quantique. Nous avons quantifié la précision de nos mesures en les comparant aux bornes de Cramér-Rao classique et quantique. Enfin, nous avons étudié théoriquement l'utilisation d'un amplificateur quantique non-déterministe en cryptographie quantique. Cet amplificateur possède la propriété de pouvoir amplifier des états quantiques sans en amplifier le bruit quantique associé. Ainsi, nous avons montré qu'il permet une amélioration de la distance maximale de transmission d'une clé secrète, ainsi qu'une amélioration de la résistance au bruit introduit par le canal quantique.

**Mots-clés :** optique quantique, information quantique, variables continues, états non-gaussiens, corrélations quantiques, discordance quantique, cryptographie quantique, amplificateur sans bruit.

## Abstract

This thesis is concerned with different aspects of quantum information with the continuous variables of quantum states of light. Through the combination of the continuous and discrete descriptions, where the light is either described in terms of quadratures or photons, non-classical quantum states and elementary quantum information protocols have been theoretically studied and experimentally implemented. We have experimentally implemented a quantum superposition of two quasi-classical states of light, a “Schrödinger cat state”, which was used to feed a quantum phase gate. We have analysed the quality of this implementation by using a simple model of the experiment. We have then studied quantum correlations, as captured by the quantum discord, for an important class of states in quantum information. We have compared the precision of our measurements by using the classical and quantum Cramér-Rao bounds. Finally, we have theoretically studied the use of a non-deterministic quantum amplifier in quantum cryptography. This amplifier has the property to amplify quantum states without amplifying their quantum noise. Using this property, we have shown that it is possible to increase the maximum distance of transmission of a secret key, as well as the tolerance to the noise added by the quantum channel.

**Keywords :** quantum optics, quantum information, continuous variables, non-Gaussian states, quantum correlations, quantum discord, quantum cryptography, noiseless amplifier.