



THÈSE
présentée à
L'UNIVERSITÉ DE NICE–SOPHIA ANTIPOLIS
École Doctorale : Sciences Fondamentales et Appliquées

pour obtenir le grade de
DOCTEUR EN SCIENCES
dans la spécialité : **PHYSIQUE**

par
Olivier Alibart

**SOURCE DE PHOTONS UNIQUES ANNONCÉS
À 1550 nm EN OPTIQUE GUIDÉE
POUR LES COMMUNICATIONS QUANTIQUES**

Soutenue le 13 décembre 2004 devant la commission composée de :

Izo Abram	Directeur de recherche, CNRS (Paris)	<i>Président</i>
Pascal Baldi	Maître de Conférence, Université de Nice	<i>Directeur</i>
Giuseppe Leo	Professeur, Université Paris VII	<i>Rapporteur</i>
Jean-Philippe Poizat	Chargé de recherche, CNRS (Grenoble)	<i>Rapporteur</i>
John G. Rarity	Professeur, Université de Bristol (Angleterre)	
Sébastien Tanzilli	Assistant professeur, Université de Genève (Suisse)	

à 14h00 au Laboratoire de Physique de la Matière Condensée



THÈSE
présentée à
L'UNIVERSITÉ DE NICE–SOPHIA ANTIPOLIS
École Doctorale : Sciences Fondamentales et Appliquées

pour obtenir le grade de
DOCTEUR EN SCIENCES
dans la spécialité : **PHYSIQUE**

par
Olivier Alibart

**SOURCE DE PHOTONS UNIQUES ANNONCÉS
À 1550 nm EN OPTIQUE GUIDÉE
POUR LES COMMUNICATIONS QUANTIQUES**

Soutenue le 13 décembre 2004 devant la commission composée de :

Izo Abram	Directeur de recherche, CNRS (Paris)	<i>Président</i>
Pascal Baldi	Maître de Conférence, Université de Nice	<i>Directeur</i>
Giuseppe Leo	Professeur, Université Paris VII	<i>Rapporteur</i>
Jean-Philippe Poizat	Chargé de recherche, CNRS (Grenoble)	<i>Rapporteur</i>
John G. Rarity	Professeur, Université de Bristol (Angleterre)	
Sébastien Tanzilli	Assistant professeur, Université de Genève (Suisse)	

à 14h00 au Laboratoire de Physique de la Matière Condensée

Remerciements

Une fois n'est pas coutume, je souhaiterais commencer ces remerciements en ne remerciant justement pas :

Les photons qui sous prétexte d'une quelconque mouvance dite « poissonnienne » n'en font qu'à leur tête et s'arrangent pour arriver de façon complètement aléatoire et désordonnée à la sortie de la source et ceci malgré trois années de dressage intensif...

Les photodiodes *InGaAs* qui n'auront cessé de jouer avec leurs copines les molécules d'eau alors que je les en avais défendues... J'aurais sans cesse dû les séparer dès que je les laissais quelques jours sans surveillance, afin de pouvoir travailler sereinement.

Mais les trois années de thèse ne se résument pas à ces petits embêtements et je garde de très bons souvenirs de cette période. Il faut savoir que beaucoup de personnes au LPMC y sont pour quelque chose et j'espère que tous se reconnaîtront dans ce qui m'a le plus marqué et que je retiendrai :

Pour commencer, je tiens à remercier *Jean-Pierre Romagnan*, directeur du LPMC à mon arrivée, pour m'avoir accueilli au sein de son laboratoire avec la gentillesse et l'attention qui le caractérisent¹. Lui a succédé à ce poste, *Gérard Monnom* que je remercie chaleureusement d'être venu au mois d'Août (bravant le soleil, la chaleur... et les touristes) prendre des nouvelles et me remonter le moral en phase finale de rédaction à raison de deux à trois fois par semaine.

Ces lignes n'y suffiront pas, mais je tiens à remercier *Pascal Baldi*... mon premier directeur de thèse et moi son premier thésard (snif!). De ces trois années de thèse sous sa direction, je me rappellerai nos longues discussions sur la compréhension des phénomènes physiques cachés derrière notre « petite source de photons uniques » se finissant souvent autour d'un solo de

¹Les premiers mots qui me viennent à l'esprit pour qualifier Jean-Pierre sont sa gentillesse et son attention et, ne pouvant faire une plus belle phrase que celle de S. Tanzilli, je m'associe à lui en lui empruntant ses mots.

Metallica, des Smashing Pumkins ou DU blues d'AC/DC. Je garderai de son encadrement une plus grande rigueur et je le remercie pour la confiance, le soutien et la grande liberté d'orientation qu'il m'a accordé au prix de quelques lentilles, wafers et diode laser !

Je tiens à associer *Dan Ostrowsky* aux gens qui ont marqué cette thèse et je le remercie pour les nombreuses questions pertinentes qu'il a soulevées durant ces trois ans qui nous ont permis de mieux appréhender les tenants et les aboutissants de ce sujet. Aussi, je me souviendrai toujours de la conférence à Lausanne pour les soirées en sa compagnie et surtout l'inoubliable restaurant gastronomique où il nous a emmené.

Enfin, je souhaite aussi remercier *Marc De Micheli*, sans qui je serais peut être encore entrain de chercher la recette du poling et du guide d'onde parfait. Marc s'est énormément investi sur la partie technique de mon travail et c'est au travers de nombreuses discussions que le protocole de fabrication d'un échantillon complet a abouti.

Une partie de ce travail a été réalisée au GAP-Optique de Genève et je tiens à témoigner toute ma reconnaissance à *Nicolas Gisin* de m'avoir proposé de venir dans son laboratoire réaliser l'expérience et bénéficier ainsi des connaissances et du matériel présents au 12, rue de l'école de médecine.

Je souhaiterais remercier *Giuseppe Leo* et *Jean-Philippe Poizat* d'avoir accepté la charge d'être rapporteur de ce manuscrit et *Izo Abram* d'avoir trouvé du temps libre dans son agenda pour participer à cette soutenance.

Je tiens également à adresser ma sincère reconnaissance à *John G. rarity* pour avoir accepté de faire partie de ce jury et m'avoir proposé un post-doc au sein de son groupe. Je profite de ces lignes pour lui rappeler qu'il a lui aussi participé à rendre cette conférence à Lausanne inoubliable. . .

En parlant de Lausanne, je n'aurais jamais imaginé ce cas de figure possible au début de ma thèse, mais je suis extrêmement heureux que *Sébastien Tanzilli* ait fait partie de ce jury. Je veux qu'il sache combien son aide et ses conseils ont été cruciaux et déterminants durant ces trois années et que son aide sur la relecture du manuscrit a dépassé toutes mes attentes. Je ne le remercierai jamais assez pour ça et je me rend compte que nous nous sommes toujours croisés ces dernières années entre Nice, Genève et Bristol. . . mais je ne suis pas inquiet on se retrouvera un jour sur Nice.

Bon, je vais casser un mythe : « Non ! Un étudiant ne travaille pas 24 h sur 24 h, il mange de temps en temps ». Dans cette optique, je souhaiterais que les personnes avec qui j'ai partagé mes repas sachent combien ces moments ont été agréables. J'en garde un d'autant bon souvenir que c'est avec eux que j'ai découvert la joie de se faire fragger à *Counter-strike*® ; nos soirées réseaux me manquent déjà :

Le Oire, mon freerider préféré, celui avec qui une banale sortie VTT se transforme en une sortie *Joe Bar Team*® à la « *si ça passait, c'était beau* ». Ces sorties vélo me manqueront à Bristol mais je reviendrai... Greg ne change pas ! En attendant y'a toujours le sumo !

En revenant, je souhaite aussi revoir *Le Barth* alias le « le roi de la technologie », je garde un excellent souvenir des soirées pizzas au labo et de notre après-midi *TOEIC*. Ton mariage était vraiment génial et j'espère bien vous revoir tous les deux... peut être à la Foux d'Allos ?

Et bien sûr, *Le Gayp* et sa chère et tendre *Mimi*. Si j'ai un seul conseil à donner aux nouveaux, c'est : « En cas de problème avec *PhotoShop*®, ne demandez pas au Gayp... Désinstallez ! Ce sera plus simple ». En tout cas, merci à vous deux pour les chocolats et pour cette soirée fantastique tout là-haut à St-Blaise haut-lieu de l'activisme anti-industrie du disque² et de l'ADSL pour tous.

Le Jé, le fragueur du bureau du bas, merci pour ton oscillo, ta disponibilité et surtout pour ton aide quand j'étais en pleine galère de branchement.

Gilles Llorca le luttor à ski avec qui j'ai partagé des sessions de ski/snow-board inoubliables mais qui malheureusement restera un bipède skieur... Enfin, il en faut bien pour damer la piste en chasse-neige !

Super Babou, ta playlist magique, tes jeux de mots pourris et la *2.9 power* vont me manquer mais avec *l'ipod* et *MSN* c'est un bout de toi qui part aussi pour Bristol (euh pour trackmania il faut encore que je trouve un moyen...). Embrasse bien fort *Céline* de ma part et bon courage pour votre dernière ligne droite et rappelez-vous : « on a jamais été aussi près » (ou prêt, c'est selon...). Aussi, je m'associe une dernière fois à Basile pour clamer bien haut que malgré les apparences, nous ne sommes pour rien dans la disparition de nos ordinateurs !

Ceci me permet d'enchaîner sur les personnes qui étaient, sont passés ou sont actuellement dans le bureau 2.9, lieu magique où il fait bon travailler, bien qu'un peu bruyant et souvent agrémenté de musique de merde :

Tout avait bien commencé, par un message de bienvenue « Ah, non ça va pas être possible on est déjà trois dans ce bureau » mais j'ai découvert dans ce bureau *Poufiñho*, *Marie Angiñho* et *Tanziñho*. Sachez qu'en guise d'hommage, le téléphone sans fil porte toujours vos noms et en attendant, je souhaite bonne route à *Pierre*, *Delphine* et *Louise* tandis que je souhaite un bon retour à Nice à *Dima*.

Gillou, mon cher collègue, mon cher boulet de salle blanche, rien est plus pareil là-bas depuis que tu as tourné du côté obscur du professorat. Bonne

²<http://www.audionautes.net>

route avec ta Magalie et bon courage pour le CAPES.

Guillaume, en fait, au début de ce paragraphe, je ne pensais pas vraiment que notre musique était pourrie mais peut-être que si finalement... car tu n'auras tenu qu'un an dans ce bureau. En tout cas bon courage pour tes deux dernières années tu as du boulot en perspective.

Pierre, il va falloir imposer ta musique de merde à toi si tu veux pas trop entendre les merdes des autres. Bon courage avec tes copains *GHZ* et pour les trois années à venir dans ce bureau 2.9...

Dernier venu, *Sorin* qui nous a appris que les Roumains (et Roumaines !) avaient des origines latines et a su imposer à maintes occasions son coté latin-lover. On est bien d'accord : « Ce qu'il y a de plus mieux dans le sport, ce sont les sportives ! »

Je ne veux pas oublier les autres « bureaux étudiants » du LPMC :

D'abord les « vieux », *Katia* et *Loïc* « mi-homme, mi-vis en titane » précédents détenteurs de l'art ancestral de la fabrication de guides PPLN, ils ont fait de moi un des gardiens de ce secret. Gardien, je l'étais pour 3 ans, maintenant je l'ai confié à *Sorin* et malheureusement pour lui, ça n'aide pas avec les filles ! Merci pour votre aide et bonne continuation à la petite famille *Chanvillard*.

Laurent, merci pour ta gentillesse, ton aide pour le neophyte que je suis en L^AT_EX et pour ton super tableau que je t'ai piqué. Bonne chance à Lyon et courage y'en a que pour un an d'être loin de *Béatrice* !

Oliver, a mi me gusto español y pienso que veremonos ¿porqué no en Venezuela ? Voila tous mes souvenirs d'espagnol... Penses-tu que ce sera suffisant pour venir te voir chez toi ? En tout cas bon courage pour la dernière ligne droite ¡Hasta luego !

Pour finir avec les « bureaux étudiants », le bureau 2.19 qui a été renouvelé complètement cette année et se présente comme un concurrent sérieux au bureau 2.9 pour la douceur de vivre. Alors bon courage à *Davide*, *Nadir* et *David*, profitez bien de ce bureau.

Pierre de Valérie et *Valérie de Pierre de Mouans* que je remercie pour les superbes affiches d'annonce de la thèse qu'elle a faite et pour les bons moments passés au mariage du Barth.

Les collègues de foot en salle : *Franck* alias « Patator », *Wilfried* que je remercie pour ses commentaires avisés lors de la préparation de mon exposé oral et *Fabrice* mon alter-ego au foot doté d'un style de jeu original mêlant « adresse » et « physique ».

Eric Picholle n'a pas fait de foot en salle, mais je tiens aussi à le remercier chaleureusement pour son regard pertinent sur mon travail et ses précieux

conseils pour la préparation de l'oral :

Je remercie *Jean-Claude* « le roi du vide et de la micro-électronique » pour son soutien et son aide lors des problèmes rencontrés durant cette thèse.

Sportivement parlant, je ne remercierai jamais assez *Dédé* et *Christophe* sponsors officiels de ma carrière bmxistique. Merci pour les pièces, pour les coups de main et pour votre gentillesse (d'ailleurs le Dragonfly, le Jackflash et le Kona s'associent à moi et vous remercient pour les superbes pièces tuning que vous leur avez faites).

Enfin, je tiens à remercier *Denise Siedler*, *Christine Ubaldi*, *Martine Cecchetti*, *Nicole Fruchart* et *Annette Chiche* sans qui la vie d'un étudiant qui doit remplir des ordres de missions pour voyager, acheter du matériel, etc... ne serait pas aussi facile.

Durant ce travail, j'ai passé quelques temps au *CRHEA* et au *GAP-Optique*, je souhaite saluer les personnes que j'ai rencontré là bas :

Helge Haas à qui je dois beaucoup en salle blanche, pour ta gentillesse légendaire, merci !

Pour avoir partagé les repas au restaurant France Telecom et quelques parties de foot, merci à *Sebastien Pezzagna* et aussi *Sébastien Chenot* qui prend la relève de *Helge* en salle blanche.

Je garde un excellent souvenir d'avoir travaillé avec *Sylvain Fasel* pendant 15 jours. Je le remercie pour les discussions enrichissantes que nous avons eu et pour m'avoir fait découvrir *Couleur 3* et *Accuradio*, je lui souhaite bon courage pour la suite avec ses copains les plasmons.

Évidement *Alexios Beveratos* faisait partie de ces discussions mais je tiens à le remercier pour les quinze jours que nous avons passé en collocation à la découverte de Genève.

Enfin, je tiens à remercier *Ivan*, *Hugues*, *Matthaeus* « j'emballe les filles en 3s chrono », *Sofian* « le roi du sushi » (titre partagé avec *Hugo*) qui m'ont accueilli à bras ouverts durant mon séjour Genevois et m'ont permis de passer d'agréables moments en dehors de la salle de manip'. Merci encore !

Je souhaite adresser ma profonde reconnaissance à *Carole Fernandes* et *Jean-Paul Dalban* qui m'ont fait confiance en me permettant d'enseigner la physique au DU IRI. Sachez que j'en garde un très bon souvenir et que cela a éveillé en moi le goût de l'enseignement pour les années à venir.

Les semaines de travail ne faisant pas sept jours ! Je tiens à saluer les *colleriders* avec qui j'ai passé des week-end exceptionnels, je pense à *la famille bringer* qui m'a fait découvrir mon (faible) potentiel de « racer », les *Salvadori brothers*, *Mat* et *Phil* qui m'ont sorti de la race pour le free, le snowboard et le surf.

Je salue aussi *Thierry* qui est quelque part à l'origine de mon départ pour le sud et je pense aux prochaines soirées gastronomiques qu'on se fera quand il rentrera des États-Unis avec plein de bonnes recettes !

Pour finir, je souhaiterais remercier *mon frère* et *Olivier Durand-Drouhin* qui m'ont fait une très agréable surprise en descendant ensemble pour la soutenance de thèse et qui se sont sacrifiés pour assister à l'ultime répétition et aussi goûter le kir.

Merci aussi à *Marie-Pierre*, *Evelyne* et *Charles* qui ont organisé un buffet qui restera dans les annales pour longtemps d'un point de vue aussi bien culinaire que décoratif.

Évidement, je ne serais pas là avec autant de joie à faire de la physique sans *mes parents* qui m'ont toujours permis de faire ce qui m'attirais et soutenu dans mes choix qu'ils soient scolaires ou récréatifs — je pense aux rédactions de français, aux sorties en tandem VTT par temps pourris ou encore aux voyages chez le vélociste ou à l'hôpital, qui par le plus grand des hasards étaient voisins... —

Et pour finir, j'ai une pensée pour celle qui a le plus enduré ces trois années de thèse en partageant ma vie au quotidien : je pense bien sûr à *Magalie* et j'ai un petit peu honte de la remercier en partant une année loin d'elle, mais je veux qu'elle sache à quel point je lui suis reconnaissant pour sa patience et pour toutes les fois où elle a fait semblant de trouver mes photons « uniques » pour me remonter le moral.

Table des matières

INTRODUCTION GÉNÉRALE	15
PROLOGUE	21
1 La conversion paramétrique comme source de photons uniques	33
1.1 Une source de paires de photons	34
1.1.1 Introduction à l'interaction lumière-matière	34
1.1.2 L'interaction lumière-matière mise en équation	36
1.2 Une source de photons asynchrone	44
1.3 Une source de photons uniques !	47
1.3.1 Définition des paramètres	47
1.3.2 Calcul de la probabilité P_2	51
1.3.3 Calcul de la probabilité P_1	51
1.3.4 Calcul de la probabilité P_0	52
1.3.5 La fonction d'autocorrélation d'ordre deux : $g^{(2)}(0)$	52
1.3.6 Étude de l'influence des paramètres de la source sur ses performances	55
1.3.7 Résumé de l'influence des paramètres	62
1.3.8 Choix des paramètres de fonctionnement	62
1.4 Conclusion du chapitre 1	63
2 La technologie de l'optique guidée	65
2.1 Des guides d'ondes sur PPLN	66
2.1.1 Les guides d'ondes SPE	66
2.1.2 Le Quasi Accord de Phase	68
2.1.3 Caractérisation de la fluorescence paramétrique	77
2.1.4 Caractérisation spatiale transverse des modes	87
2.2 Les composants fibrés de la source	93
2.2.1 La fibre optique	94
2.2.2 Les connecteurs optiques	96
2.2.3 Le filtre de pompe	97
2.2.4 Le WDM 1310/1550 nm	98

2.2.5	Pertes des composants assemblés	99
2.2.6	La photodiode à avalanche en Germanium	99
2.2.7	L'électronique de traitement du signal	101
2.3	Détermination du coefficient de couplage guide-fibre	102
2.3.1	Le montage et les paramètres utilisés	102
2.3.2	Les résultats	105
2.4	Conclusion du chapitre 2	108
3	Mesures expérimentales des performances de la source	111
3.1	Les mesures de P_0 , P_1 et P_2 en mode asynchrone	112
3.1.1	Le montage	113
3.1.2	Le Modèle d'analyse des données brutes	115
3.1.3	Les résultats expérimentaux	121
3.1.4	Comparaisons avec les prévisions du chapitre 1	123
3.1.5	Conclusion	129
3.2	Mesure directe de $g^{(2)}(0)$ en mode asynchrone	131
3.2.1	Rappel sur la mesure expérimentale de $g^{(2)}(\tau)$	131
3.2.2	Le principe de la mesure asynchrone de $g^{(2)}(0)$	137
3.2.3	Le montage expérimental	137
3.2.4	Le programme d'acquisition et de traitement des données	141
3.2.5	L'analyse théorique des données	143
3.2.6	La mesure expérimentale et l'évaluation du bruit	145
3.2.7	Les résultats	148
3.3	Conclusion du chapitre 3	151
4	Perspectives	153
4.1	Le pompage en régime impulsionnel	154
4.2	L'utilisation d'une APD– <i>Silicium</i> comme trigger	155
4.3	Réduction de la largeur spectrale des photons	158
4.4	Conclusion du chapitre 4	162
	CONCLUSION GÉNÉRALE	165
	ANNEXES	169
A	Sécurité de la distribution quantique de clé	171
A.1	Qu'est ce le taux de bits utilisables ?	171
A.1.1	L'amplification de confidentialité	171
A.1.2	La correction d'erreurs	172
A.2	La distance maximale de l'échange quantique de clé	173
A.3	Évaluation de p^{net}	174

A.4	Le taux de bits utilisables par impulsion	174
A.5	Calcul du taux de bits utilisables à partir de données expérimentales	175
B	Caractéristiques des APDs utilisées	177
B.1	Les photodiodes à avalanches	177
B.2	L'APD <i>Germanium</i>	180
B.2.1	Mesure de l'efficacité de détection en mode passif . . .	180
B.2.2	Courbe d'étalonnage de l'APD- <i>Ge</i>	182
B.3	L'APD <i>InGaAs</i>	182
B.3.1	Mesure de l'efficacité de détection en mode déclenché .	182
B.3.2	Courbe d'étalonnage des APD- <i>InGaAs</i>	185
C	Intervalle de temps séparant deux paires	187
C.1	Rappels théoriques	187
C.2	La mesure expérimentale	189
D	La fabrication des guides PPLN	193
D.1	Le substrat de niobate de lithium	193
D.2	Le retournement des domaines	196
D.2.1	Le dispositif expérimental	196
D.2.2	Le contrôle des domaines	199
D.3	Les guides d'ondes <i>SPE</i>	202
D.3.1	Le protocole expérimental	203
E	Liste des symboles utilisés	207

Introduction générale

Apparue au milieu du siècle dernier, la physique quantique est souvent associée à des « expériences de pensée » ou à des expériences fondamentales comme le test des inégalités de Bell [1] qui montrent que la classe des théories locales ne peut fournir une description complète des corrélations qu'il existe entre des particules intriquées [2, 3, 4]. Toutefois, dans notre « société de l'information », les scientifiques ont perçu l'intérêt de la théorie quantique appliquée à celle de l'information [5, 6, 7, 8, 9]. Nous sommes au début des années 90, l'*information quantique* est un nouveau champ de recherche dont l'objectif est de tirer parti des possibilités offertes par la mécanique quantique pour traiter l'information d'une manière plus efficace [10]. Les deux composantes principales sont les *communications quantiques*, qui apportent une sécurité accrue par rapport aux systèmes de cryptographie classique et d'autre part le *calcul quantique*, pour lequel de nouveaux algorithmes basés sur les principes de la mécanique quantique permettent de diminuer radicalement les temps de calculs nécessaires pour résoudre certains problèmes [11], grâce à la superposition cohérente d'états possible pour les systèmes quantiques. En effet, ces derniers offrent une infinité de possibilités : l'état 1, l'état 0 et toutes les superpositions des deux. On ne parle plus alors de bits, mais plutôt de « *qu-bits* » (quantum bits).

L'objet quantique le plus simple à maîtriser est le photon, qui est utilisé comme moyen de transport du *qu-bit*³. Le terme « communication quantique » regroupe alors des thèmes de recherche tels que l'échange quantique de clé, la téléportation d'état ou encore la permutation d'enchevêtrement à l'aide de photons. Expérimentalement et chronologiquement, deux axes de recherche se sont dessinés :

Le premier concerne les communications à l'air libre [12, 13, 14, 15, 16], utilisant des photons dont la longueur d'onde est située dans le visible. Ces communications sont limitées à plusieurs dizaines de kilomètres, ce qui est

³Il est couramment appelé « *flying qu-bit* » en opposition aux atomes qui peuvent aussi stocker le qu-bit mais ne voyagent pas.

amplement suffisant pour communiquer d'un immeuble à l'autre ou, plus ambitieux, pour établir une communication terre-satellite.

Le deuxième axe de recherche a pour objectif de tirer parti de l'optique guidée pour transmettre l'information plus loin. En effet, l'industrie des télécommunications a su tirer parti du très faible coût de fabrication des fibres optiques en silice et adapter la fréquence de ses signaux à la plage où les pertes sont les plus faibles : le proche infrarouge. Il existe donc aujourd'hui deux plages de longueurs d'ondes pour les communications dans les fibres optiques : 1310 nm et 1550 nm qui correspondent respectivement à *la première et seconde fenêtre télécom*. À ces longueurs d'ondes, il est possible de s'échanger des informations sur plusieurs centaines de km sans ré-amplifier le signal optique ; distance qu'il faut comparer (à débit constant) à la limite de l'ordre du km associée au traditionnel « fil en cuivre » ou de la dizaine de km associée à une communication à « l'air libre ». On estime qu'aujourd'hui plus de 80% des communications longues distances sont transportées le long de plus de 25 millions de kilomètres de câbles optiques partout dans le monde. Il aura fallu peu de temps aux communications quantiques pour profiter des formidables possibilités qu'offrent les fibres optiques pour échanger « quantiquement » de l'information à l'aide de photons uniques. Nous sommes maintenant en 1995 et l'équipe du GAP-Optique échange une clé secrète à travers 23 km de fibre optique installée sous le lac Léman [17]. Pourtant, aujourd'hui encore, il n'existe pas « d'outils » réellement *adaptés* aux communications quantiques dans les fibres optiques. Nous pensons, entre autres, à des sources capables de produire des photons uniques, qui constituent l'un des ingrédients de base des communications quantiques. La source de photons uniques « idéale » se définit ainsi :

- Compacte, elle doit pouvoir s'insérer de façon simple, stable et efficace dans les systèmes de télécommunication à fibres disponibles actuellement. En effet, l'insertion des photons (ou qu-bits) générés doit se faire avec le moins de pertes possibles et leur longueur d'onde doit être centrée sur l'une des fenêtres préférentielles, à savoir 1310 ou 1550 nm .
- Elle doit posséder une probabilité d'émettre un photon unique proche de un.
- Enfin, le taux d'impulsion contenant deux photons doit être le plus faible possible.

Pourquoi chercher à faire une source de photons uniques présentant la plus forte probabilité d'émettre un photon unique et la plus faible d'en émettre simultanément deux, alors qu'il a été démontré qu'un laser cohérent atténué suffit à simuler approximativement une source de photons uniques ?

En effet, si nous atténuons suffisamment un laser pour qu'il envoie en moyenne 0,1 photon par impulsion, nous obtenons une source capable d'émettre un photon unique pour dix impulsions émises ($P_1 = 0,1$), tandis qu'elle émettra deux photons simultanés pour deux cents impulsions émises ($P_2 = 0,005$). Ce montage fournit donc une relativement bonne approximation d'une source de photons uniques, avec toutefois des défauts importants :

- Les protocoles, comme notamment la téléportation ou l'échange quantique de clé secrète, sont actuellement limités par cette trop faible probabilité P_1 qui engendre des temps de mesures très longs ou des débits d'échange très faibles.
- La grande quantité d'impulsions vides ($P_0 = 0,9$) est d'autant plus gênante qu'à l'heure actuelle il n'existe pas de détecteurs parfaits et que, dans ce cas, la proportion du bruit lors de la mesure prend une place prépondérante.
- Enfin, la quantité d'impulsions contenant deux photons limite considérablement la distance de transmission d'une clé sécurisée⁴ ou réduit la qualité des interférences dans une expérience de téléportation quantique.

Ce travail de thèse s'inscrit dans la continuité de celui de S. Tanzilli, afin de démontrer l'apport de la *technologie de l'optique intégrée* à cette discipline en pleine expansion que sont les *communications quantiques* [18]. En suivant l'idée de Mandel d'utiliser des paires de photons [19], nous avons pour objectif de montrer la faisabilité d'une source de photons uniques annoncés compacte et performante basée sur la génération paramétrique dans un guide d'ondes réalisé par échange protonique sur un substrat de niobate de lithium polarisé périodiquement (PPLN) et régie par le quasi-accord de phase. Plus encore, ce qui nous intéresse, c'est de véritablement conférer à notre source les caractéristiques qui lui permettront d'être l'un des composants clés des futures expériences de communications quantiques qui commencent à sortir des laboratoires afin d'échanger des qu-bits sur grandes distances via les fibres optiques [20, 21, 22, 23, 24].

Les guides d'ondes PPLN que nous fabriquons depuis quelques années au laboratoire s'adaptent parfaitement aux critères de la source idéale que nous avons énoncés plus haut. En effet, sachant que la configuration guidée offre un confinement des ondes sur une longueur pouvant atteindre plusieurs centimètres, il est possible d'obtenir des probabilités de conversion des photons de

⁴voir annexe A

pompe en paires de photons encore jamais atteintes à l'aide d'un cristal massif [25, 26]. Ceci permet, par ailleurs, de réduire considérablement la puissance de pompe requise et permet l'utilisation de simple diodes laser n'émettant pas plus de quelques micro-Watts. Aussi, le processus de génération paramétrique spontanée en quasi-accord de phase permet d'obtenir naturellement des paires de photons à pratiquement toutes les longueurs d'ondes désirées et plus précisément à celles qui nous intéressent (1310 nm et 1550 nm). Enfin, la récolte des paires, à la sortie du guide, se fait à l'aide d'une simple fibre optique télécom qui, par nature, s'insère parfaitement dans n'importe quel réseau de communication optique. Parmi les critères de la source idéale figurait une haute probabilité d'émettre un photon unique ($P_1 \approx 1$) ainsi qu'une très faible probabilité d'en émettre deux ($P_2 \approx 0$). La vérification de ces deux derniers critères représente le travail fait à Nice (et à Genève) et qui a consisté à développer deux nouvelles méthodes de caractérisation des performances d'une source de photons uniques asynchrone.

Contenu du manuscrit

PROLOGUE

Avant de nous lancer dans le cœur du manuscrit, nous avons jugé utile de regrouper dans un prologue les différentes manières de réaliser une source de photons uniques. Nous définirons tout d'abord les termes de *sources de photons à la demande* et de *sources de photons annoncés*, puis nous donnerons quelques exemples représentatifs appartenant à ces deux communautés, pour que le lecteur, peu familier avec les sources de photons uniques, y trouve les bases des différentes techniques utilisées. Notons que le lecteur y trouvera également une introduction approfondie sur l'utilisation des paires de photons pour réaliser une source de photons uniques annoncés.

CHAPITRE 1

*Pour réaliser une source de photons uniques annoncés, il faut tout d'abord disposer d'une source de paires de photons, puis les séparer pour que la détection de l'un annonce la présence de l'autre*⁵. Théoriquement, la probabilité d'avoir un photon annoncé est ainsi égale à 1. . . pourtant pour de nombreuses raisons expérimentales⁶, cette probabilité P_1 sera inférieure à un. En introduisant au préalable le concept de *source asynchrone* et de *communication asynchrone* pour justifier l'origine des événements à deux photons, ce chapitre sera l'occasion de rappeler le formalisme associé à la production de

⁵C'est ce que dit la recette!

⁶Avec en chef de file, les pertes suivies par les erreurs de détections. . .

paires de photons en régime continu et de s'en servir pour établir théoriquement les probabilités P_1 et P_2 d'avoir un ou deux photons émis par la source. Ces calculs prendront en compte, par exemple, les pertes vues par les paires ou la faible efficacité du détecteur servant à annoncer l'émission des photons uniques. Les formules développées ici ne sont pas liées à notre configuration guidée mais sont applicables à n'importe quelle source de photons uniques annoncés asynchrone fonctionnant en régime continu.

CHAPITRE 2

Fort de cela, nous avons baptisé le chapitre deux « technologie de l'optique guidée ». Nous y développerons les particularités et avantages expérimentaux que procure l'utilisation de guides d'ondes pour la production de paires de photons. Ce sera aussi l'occasion de décrire et de caractériser les composants optiques qui interviennent dans la réalisation complète de la source et nous finirons ce chapitre par une mesure des pertes totales, rencontrées par les photons annoncés au sein de la source.

CHAPITRE 3

Il existe une méthode bien connue pour estimer le caractère « unique » des photons émis par une source de photons. Pourtant, de part son caractère synchrone, cette méthode ne convient pas à notre source et nous avons donc développé à Nice et à Genève deux nouvelles méthodes expérimentales permettant de mesurer les performances d'une source de photons uniques asynchrone. Après le recoupement de ces résultats expérimentaux avec les prédictions théoriques du chapitre 1, nous avons souhaité réunir dans un tableau les performances des principales réalisations expérimentales exposées dans le prologue.

CHAPITRE 4

Après avoir conclu quant aux performances de notre source, nous finirons ce manuscrit en étudiant les améliorations qu'il est possible de lui apporter à très court terme pour accroître encore ses performances. Il se trouve que le détecteur qui sert à annoncer l'arrivée des photons uniques constitue le principal point faible de la source et nous proposons alors une estimation réaliste du gain qu'apporterait un système de détection plus performant. Nous profiterons aussi de ce chapitre pour étudier la réduction de la largeur spectrale de nos photons annoncés.

Prologue sur les sources de photons uniques

Nous allons tacher dans ce prologue d'introduire les principaux concepts qui permettent d'obtenir des photons uniques et d'expliquer succinctement le principe de fonctionnement qui leur est propre.

Nous pouvons d'ores et déjà classer les principales réalisations en deux grandes familles :

Les sources de photons uniques à la demande qui reposent sur une conversion *efficace et contrôlée* de l'énergie⁷ apportée au système, sous forme de photons. Nous entendons par là que la probabilité d'émettre un photon après avoir reçu une quantité « d'énergie initiale » doit être très élevée et qu'en même temps la quantité de photons émis doit être extrêmement bien contrôlée pour toujours être égale à un. En résumé, ce type de source émet un photon unique avec une probabilité $P_1 \approx 1$ juste après un apport d'énergie de la part de l'utilisateur. On dit communément que ces *sources fournissent des photons uniques à la demande*. Nous pouvons classer dans cette famille, les émetteurs uniques comme une molécule ou un centre coloré dans le diamant qui sont des systèmes quantiques à deux niveaux qui, une fois excités, ne peuvent émettre qu'un seul photon à la fois.

Les sources de photons uniques annoncés qui reposent uniquement sur une très forte corrélation entre l'émission d'un photon et un autre événement. En contrepartie, la conversion « énergie-photon » est très mal contrôlée (voire aléatoire). Le principe retenu est celui de la *post-sélection* qui consiste à ne pas savoir quand va être émis le photon, mais de connaître à posteriori quand il a été émis grâce à l'observation de l'événement qui lui est corrélé. Dans ce cas l'utilisateur ne contrôle

⁷quelque soit l'origine et la forme de cette énergie

pas l'émission du photon, qui peut intervenir n'importe quand, mais dispose d'un indicateur qui le prévient quand le photon a été émis. Les paires de photons issues d'atomes froids et de la conversion paramétrique correspondent à ce type de source.

Mais, nous sommes en droit de nous demander ce qui permet de qualifier un système de « *source de photons uniques* ». Nous verrons dans le chapitre 1 que la fonction d'autocorrélation d'ordre deux, $g^{(2)}(0)$, permet de quantifier le caractère « unique » des photons émis. Pour faire simple, cette fonction quantifie le taux d'impulsions à deux photons par rapport à celles n'en contenant qu'un seul. Ainsi, une « vraie » source de photons uniques doit présenter un $g^{(2)}(0)$ le plus proche possible de zéro qui correspond à une source parfaite n'émettant jamais deux photons simultanément, tandis qu'une source poissonnienne présente un $g^{(2)}(0)$ égal à un. Toutes les sources présentant un $g^{(2)}(0)$ compris entre zéro et un sont dites « *sub-poissonniennes* ».

Ainsi, le laser atténué, qui suit une probabilité poissonnienne d'émettre un photon unique, présente un $g^{(2)}(0)$ égal à un et ne mérite donc pas l'appellation de source de photons uniques. Il n'en reste pas moins qu'une assez bonne approximation, car il réalise seulement un « espacement » des photons mais ne permet pas d'observer des photons ayant un comportement global dit « solitaire ».

Intéressons nous maintenant aux principales méthodes qui ont permis de réaliser de « vraies » sources de photons uniques. Historiquement, la première réalisation, en 1977, d'une source de photons uniques à la demande a consisté à utiliser un jet atomique de très faible densité [27]. Cette expérience a permis d'observer un effet de dégroupement mais pas de véritable comportement de photons uniques, car la statistique des atomes présents dans la zone d'observation était elle-même poissonnienne. L'expérience a été reprise en 1987 avec un ion unique piégé [28], ce qui a permis d'observer un véritable comportement de source de photons uniques ($g^{(2)}(0) = 0,04$). Entre temps, Mandel a introduit, en 1986, le concept des photons uniques annoncés grâce à l'utilisation de paires de photons [19] issues de l'interaction paramétrique optique. L'expérience a montré qu'il était possible d'obtenir un réel comportement de type « source de photons uniques » par post-sélection des photons ayant pourtant à l'origine une statistique poissonnienne. Avec le développement de la technologie des semi-conducteurs, d'autres moyens d'obtenir des photons uniques sont apparus. En contrôlant la croissance d'un empilement de semi-conducteurs, il est possible d'obtenir des « boîtes quantiques » capables d'émettre un photon unique en réponse à une excitation électrique,

tandis qu'une autre méthode consiste à contrôler les ondes acoustiques de surface dans ces matériaux pour n'émettre qu'un seul photon à la fois. Plus récemment, les progrès accomplis dans la manipulation des atomes en cavité des nuages d'atomes froids a permis de trouver de nouvelles techniques pour créer des paires de photons et suivre le concept introduit par Mandel. Nous verrons toutefois que ces techniques ne sont pas de simple copies de la technique de Mandel mais apportent de réelles nouveautés.

Après ce bref historique, il est temps de passer en revue les principales sources de photons uniques. Nous avons choisi de classer les sources en fonction de leur appartenance aux deux familles qui sont les *sources de photons uniques à la demande* et les *sources de photons uniques annoncés*.

1 – Sources de photons à la demande

Dans le cas des molécules et des centres colorés dans le diamant, le concept repose sur un système quantique à deux niveaux qui, préalablement excité, revient à son état initial en émettant un seul photon.

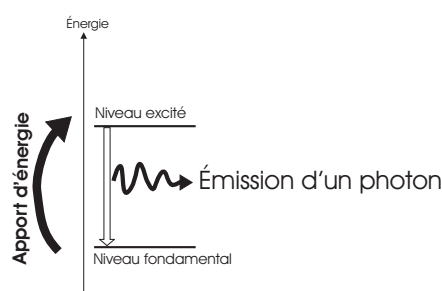


FIG. 1 – Schéma d'un système quantique à deux niveaux.

Toute l'art de réaliser une source de photons uniques, à partir de ces émetteurs, consiste à choisir un système quantique qui possède une conversion « apport d'énergie – émission d'un photon » la plus proche de un et à isoler un seul de ces émetteurs. En règle générale, la largeur spectrale des photons émis et le taux de répétition de la source sont déterminés⁸ par la durée de vie du niveau excité du système. Enfin l'inconvénient lié à ce type d'émetteur est l'angle solide d'émission des photons qui est généralement égal à 4π stéradians rendant peu efficace la collection des photons produits.

⁸En partie seulement... car il faut tenir compte, par exemple, des *phonons* dans le cristal jouent un rôle sur la durée de vie du niveau excité.

a—Molécules

Au premier abord, les molécules semblent être les candidats idéaux. Leurs caractéristiques chimiques sont généralement très bien connues ce qui permet de prévoir la position des niveaux atomiques, donc le spectre des photons émis ainsi que le schéma d'excitation. De plus, il est assez simple de préparer des échantillons ne contenant qu'une seule molécule sur plusieurs mm^2 . Leur grand désavantage est le manque de photo-stabilité. Les molécules sous excitation optique subissent une transformation chimique irréversible, après avoir émis un certain nombre de photons (de l'ordre de 10^6). Ce phénomène est connu sous le nom de photo-blanchiment et intervient généralement après quelques centaines de ms d'excitation. Mais bien que les molécules ne soient pas stables à température ambiante, elles le sont à très basse température [29]. Une molécule, refroidie à $1,5 K$ a été utilisée en 1999 pour réaliser une source de photons uniques [30]. Dans ce cas précis l'excitation s'effectue à l'aide d'un laser et l'émission du photon est déclenchée par l'adjonction d'un champ électrique. Il faut cependant noter la réalisation de sources de photons uniques à température ambiante, excitées par un simple laser impulsif reportées dans les références [31, 32].

b—Centres colorés dans le diamant

À l'inverse des joailliers, il faut chercher un diamant avec un défaut unique. Le cristal de diamant est naturellement transparent à tout apport d'énergie, mais l'existence d'un défaut (ou centre coloré, voir figure 2) va définir dans le cristal un système quantique à deux niveaux capable d'absorber l'énergie et de la ré-émettre sous la forme d'un photon unique. À l'instar des molécules, l'émission des photons se fait dans un angle solide de 4π stéradians mais son avantage est sa photostabilité intrinsèque à température ambiante. Cependant la largeur spectrale des photons émis est telle qu'ils nécessitent un filtrage en longueur d'onde, qui réduit au final la probabilité d'avoir un photon. Notons qu'une réalisation expérimentale de distribution quantique de clé a pu être réalisée à l'aide de ce type de source [33]. Le lecteur pourra trouver des informations supplémentaires sur les centres colorés dans le diamant en se tournant vers les références [34, 35, 36].

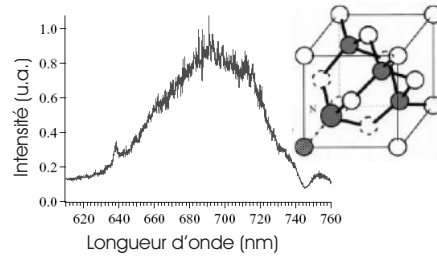


FIG. 2 – Centre coloré dans le diamant qui correspond à un atome d'azote à la place d'un atome de carbone. Nous pouvons observer le spectre d'émission assez large du système à deux niveau ainsi constitué [36].

c–Semi-conducteurs

Une amélioration évidente des sources précédentes serait de réduire la largeur spectrale des photons émis et de réussir à confiner un émetteur dans une structure guidante pour augmenter le taux de collection des photons⁹. Cette idée n'a pas encore été expérimentalement réalisée pour les molécules ou les centres colorés, mais la communauté des semi-conducteurs a développé une solution qui répond à ces attentes.

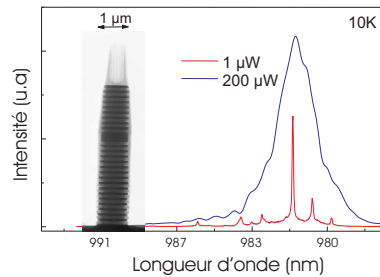


FIG. 3 – Microcavité autour d'une boîte quantique et son spectre d'émission [37].

Les boîtes quantiques semi-conductrices *InAs* présentent de bonnes caractéristiques à basse température (5 K). Ces émetteurs sont en effet photostables et le niveau excité a une durée de vie de l'ordre de quelques *ns* ce qui procure une très faible largeur spectrale aux photons uniques. Ces boîtes quantiques sont en fait des îlots semi-conducteurs d'une hauteur de quelques

⁹qui reste à l'heure actuelle très faible, de l'ordre de quelques %.

nm , pour un diamètre de quelques dizaines de nm . De par cette taille réduite, les niveaux électroniques à l'intérieur de la boîte sont quantifiés et chaque désexcitation d'une paire *électron-trou* injectée dans la boîte donne naissance à un seul photon. Mais le principal intérêt des boîtes quantiques est d'être assez « simple¹⁰ » à insérer directement à l'intérieur d'une *microcavité* qui va permettre d'augmenter la probabilité que le photon soit émis suivant l'axe de cette cavité. Cette microcavité en forme de pilier (voir figure 3) va permettre de guider les photons uniques émis par la boîte quantique jusqu'à son sommet où ceux-ci seront facilement récoltés. Les références [38, 37] ont montré un très fort couplage entre la cavité et la boîte quantique et ont obtenu des coefficients de couplage, du photon unique dans la structure guidante de la cavité, pouvant aller jusqu'à 80%. Pour l'instant, le principal inconvénient des boîtes quantiques se situe au niveau leur température de fonctionnement. Celle-ci nécessite l'utilisation d'un cryostat qui ajoute des pertes sur le parcours des photons entre le micropilier et l'utilisateur diminuant au final la probabilité de récolter un photon unique utilisable. Notons toutefois que ce type de source a permis de réaliser une expérience de téléportation [39].

Il existe une autre famille de semi-conducteurs, que l'on peut manipuler à température ambiante : ce sont les nanocristaux de *CdSe*. Ces sources présentent les caractéristiques d'une source de photons uniques [40, 41], mais l'émission d'un nanocristal de *CdSe* présente aléatoirement des périodes de fonctionnement entrecoupées de période d'inactivité (de quelques minutes à plusieurs heures) ne permettant pas, pour l'instant, de les utiliser comme une source de photons uniques.

d—Ondes acoustiques de surface (SAW)

Le principe est assez proche du fonctionnement d'une boîte quantique, mais diffère dans la manière d'apporter les paires *électron-trou* au semi-conducteur. Il a été montré qu'une onde acoustique de surface se propageant à travers le semi-conducteur donne naissance à des paires *électron-trou* qui sont localisées dans les minimums de l'onde et qui voyagent à la vitesse de cette dernière (voir figure 4). Comme la quantité de paires contenues dans un minimum de l'onde acoustique de surface peut être extrêmement bien contrôlée pour valoir un, la production d'états contenant un photon unique est alors envisageable.

À notre connaissance, il n'existe aucune réalisation expérimentale de source de photons uniques à partir d'ondes acoustiques de surface, mais de

¹⁰Simple n'est sans doute pas le terme approprié, car d'un point de vue technologique, insérer une boîte quantique dans une microcavité est loin d'être une formalité.

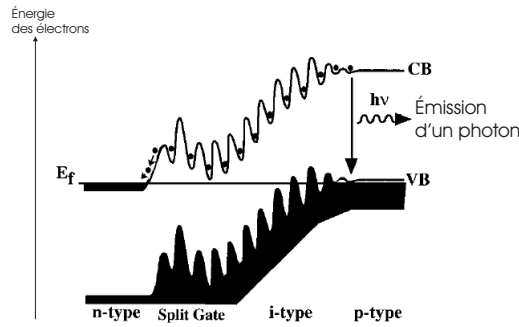


FIG. 4 – Utilisation des ondes acoustiques de surface pour la génération de photons uniques [42]. Les niveaux CB , VB et E_f représentent respectivement les énergies de la bande de Conduction, de la bande de Valence et du niveau de Fermi.

nombreuses références [42, 43] laissent présager de la faisabilité de ce concept.

2 – Sources de photons annoncés

Dans le cas des *sources de photons annoncés* (de l'anglais « *Heralded Photons Sources* ») nous n'allons pas chercher à contrôler l'émission d'un photon à partir d'un apport d'énergie, mais à trouver des processus qui vont, avant tout, émettre des photons uniquement par paires. Il est important que les deux photons d'une paire soient extrêmement bien corrélés, c'est-à-dire que la présence de l'un garantisse toujours la présence de l'autre. La génération paramétrique de paires de photons est l'un des tout premiers candidats et nous allons voir que de nouvelles méthodes originales permettent d'obtenir aussi des paires de photons.

a–Génération paramétrique de paires de photons

Imaginons une source ayant une statistique poissonnienne avec un nombre moyen de photons par impulsion $\bar{n} = 0,1$. La probabilité d'avoir i photons s'écrit alors :

$$P_i = \frac{(\bar{n})^i}{i!} e^{-\bar{n}}$$

La quantité d'impulsions contenant deux photons n'est pas nulle et seulement une impulsion sur dix contient effectivement un photon unique. Il y a donc deux inconvénients : nous ne savons pas dans quelle impulsion est situé le photon unique et ce type de source présente un $g^{(2)}(0) \approx \frac{2P_2}{P_1^2} \approx 1$.

Ici l'objectif est fort simple. Nous n'allons pas réduire le nombre d'impulsions contenant deux photons mais chercher à obtenir une information sur celles qui contenaient au moins un photon puis de ne garder que celles là. Bien évidemment la probabilité d'avoir un photon par impulsion est toujours égale à 0,1 en moyenne, mais la probabilité d'avoir au moins un photon par *impulsion annoncée* est artificiellement égale à un. Cette technique consiste en une *post-sélection* des impulsions dans lesquelles nous sommes sûr qu'il y avait au moins un photon.

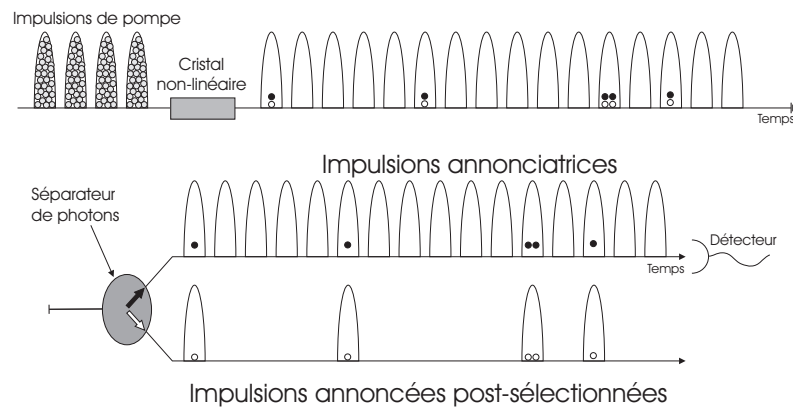


FIG. 5 – Principe de la post-sélection des impulsions contenant au moins un photon. Sur la figure du bas les impulsions « vides » ont été éliminées par post-sélection.

Maintenant, à l'aide de la figure 5, prenons un cristal non-linéaire qui va convertir des impulsions lumineuses incidentes, dites de pompe, en impulsions lumineuses contenant des paires de photons noirs et blancs sur le schéma. Sans entrer dans les détails, acceptons que le processus soit aléatoire et que son efficacité soit faible. Alors la quantité de paires de photons contenues dans une impulsion suit une statistique poissonnienne. À la sortie du cristal, la plupart des impulsions sont vides et seulement quelques unes contiennent une paire. Remarquons toutefois, sur la figure 5, que la probabilité qu'il y ait deux paires dans une impulsion n'est pas nulle. En supposant que nous disposions d'un moyen de séparer spatialement les photons d'une paire, la détection d'un photon noir sur la voie supérieure nous garantit la présence de son homologue blanc sur la voie inférieure. En ne considérant finalement que les événements qui ont été post-sélectionnés, la probabilité d'obtenir un photon unique par impulsion devient très proche¹¹ de un, tandis que

¹¹Il ne faut pas oublier de retrancher les cas où deux photons étaient présents dans l'impulsion

celle d'en avoir deux reste faible. Ce concept a été introduit par Mandel en 1986 [19] et depuis de nombreuses démonstrations expérimentales ont suivi [44, 45, 46, 47, 48, 49, 50, 51].

Si la création simultanée des photons de la paire est à la base du principe des photons annoncés, elle peut représenter un inconvénient pour l'utilisateur qui n'a pas moyen de contrôler aisément l'intervalle de temps entre la détection du photon annonceur et l'émission du photon annoncé. évidemment, en ajoutant un délai optique sur le trajet d'un des photons, l'utilisateur peut retarder l'instant d'arrivée de tous les photons noirs par rapport aux blancs, mais cette méthode ne permet pas de l'ajuster pour chaque photon émis. Le principal reproche fait à ce type de source est donc de ne pouvoir les utiliser dans des expériences où l'envoi du photon doit suivre une fréquence d'horloge fixe. À ce titre, notons toutefois que les références [52, 53] proposent un moyen judicieux d'obtenir une *source de photons à la demande* à partir de paires de photons. Simplement à l'aide d'une boucle de stockage, qui jouerait le rôle d'une « mémoire tampon » capable d'absorber un débit trop élevé ou justement de compenser un débit de photon trop faible en rendant les photons stockés, la source émettrait des photons uniques à intervalles réguliers.

Il se trouve que les sources de photons annoncés à partir d'atomes en cavité et d'atomes froids permettent de gérer « simplement » cette durée.

b—Atomes en cavité

Le principe repose sur la présence d'un atome de *Rubidium* au sein d'une cavité optique et sur la capacité qu'a le système *atome-cavité* d'émettre un photon unique lors de sa désexcitation. La présence d'un atome unique au sein de la cavité suit une statistique poissonnienne. Cette occurrence a donc une faible probabilité d'exister, par contre une fois dans la cavité l'atome y séjourne environ une centaine de microsecondes ce qui laisse largement le temps à l'utilisateur de détecter sa présence et d'obtenir un photon émis par le système *atome-cavité*. Prenons un peu le temps d'expliquer le fonctionnement du système à l'aide de la figure 6.

Des atomes de Rubidium sont préparés dans l'état $|u\rangle$ dans un piège magnéto-optique (MOT). Situé juste au dessus d'une cavité, le piège laisse s'échapper quelques atomes qui tombent à l'intérieur de la cavité. Vu les dimensions de la cavité et la vitesse de chute des atomes, ces derniers y séjournent environ $500\ \mu\text{s}$. La cavité étant résonnante avec la transition atomique $|e\rangle \rightsquigarrow |g\rangle$ du *Rubidium*, il ne faut plus considérer un atome à l'intérieur de la cavité mais un système indissociable *atome-cavité*. Dans cette configuration, les états notés $|0\rangle$ et $|1\rangle$ représentent la quantité de photons présents à l'intérieur de la cavité.

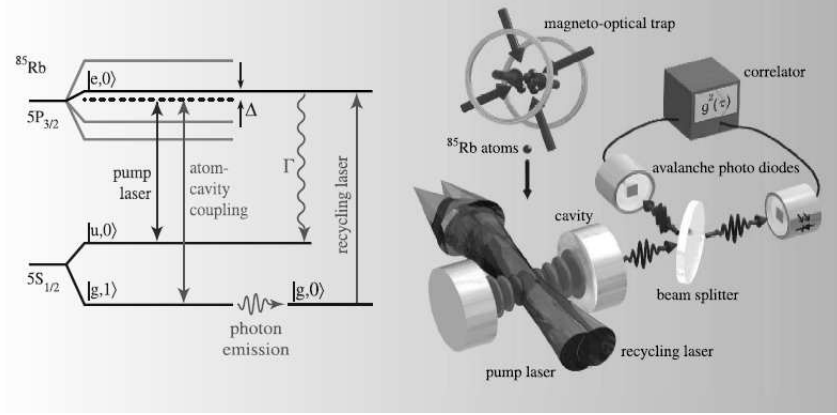


FIG. 6 – Atomes en cavité pour la génération de photons uniques [54].

Une première impulsion laser va donc permettre de tester la présence d'un atome dans la cavité, en transférant le système *atome-cavité* de l'état $|u, 0\rangle$ vers l'état $|e, 0\rangle$ et en mesurant l'émission d'un photon lorsque le système se relaxe vers l'état $|g, 1\rangle$. La probabilité qu'un atome soit présent dans la cavité est très faible. L'utilisateur refait donc cette opération jusqu'à ce qu'il détecte un photon émis hors de la cavité et, dans ce cas, il envoie une seconde impulsion laser dans la cavité qui va *re-préparer* l'atome dans l'état $|u, 0\rangle$. À ce stade, l'utilisateur est sûr qu'il y a un atome dans la cavité et le système est prêt à émettre un photon unique quand l'utilisateur enverra pour la troisième fois une impulsion laser dans la cavité.

L'intérêt du système est qu'une fois validée la présence d'un atome dans la cavité, nous disposons d'environ $500 \mu\text{s}$ pour avancer ou retarder l'émission du photon qui peut être alors synchronisé avec une horloge de travail fixe.

Par soucis de pédagogie, nous avons simplifié le protocole expérimental, mais nous invitons le lecteur à lire la référence [54] pour en apprendre plus.

c–Condensat d'atomes froids

Les atomes froids permettent aussi ce genre de manipulation. Considérons un ensemble d'atomes froids de *Césium* maintenu dans un piège magnéto-optique (*MOT*).

Une impulsion laser initiale, dite « *d'écriture* », permet de faire passer le système de l'état initial $|a\rangle$ vers un état dit *d'excitation collective* $|b\rangle$. Cette transition passe par un niveau virtuel $|e\rangle$ et s'accompagne de l'émission de

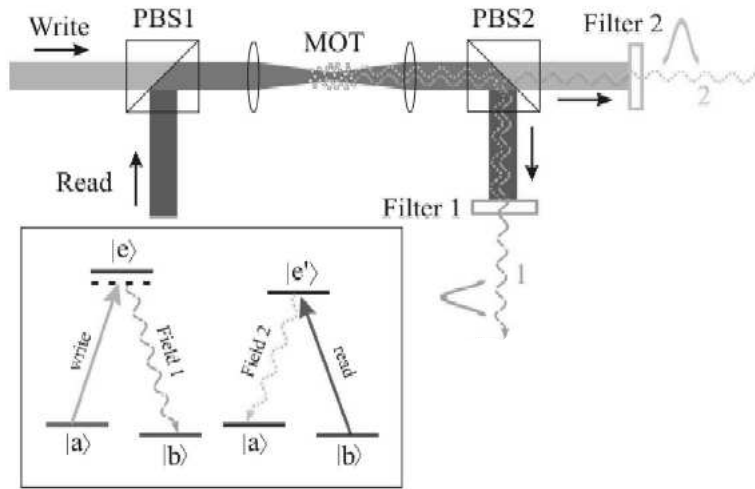


FIG. 7 – Utilisation d'un condensat d'atomes froids pour la génération de photons uniques.

photons par *diffusion Stokes*. Tout l'art d'obtenir un photon unique consiste à disposer d'un temps d'interaction très court entre l'impulsion d'écriture et le nuage atomique pour que le nombre moyen de photon émis par diffusion Stokes soit plus petit que un.

L'utilisateur va donc envoyer des impulsions lumineuses jusqu'à ce que « l'écriture » soit confirmée par l'émission d'un photon à travers le filtre 1 de la figure 7. À partir de cet instant, l'ensemble d'atomes est dans l'état excité $|b\rangle$ qui ne peut revenir à son état initial qu'en émettant un photon unique. Il se trouve que cet état excité est métastable¹², l'utilisateur dispose donc d'un temps $\mathcal{A}$ pour envoyer une seconde impulsion dite de « lecture » qui forcera le système atomique à retourner dans son état fondamental $|a\rangle$. Cette désexcitation s'accompagne aussi de l'émission d'un photon unique au travers du filtre 2.

Cette explication est tirée de la référence [55] et nous invitons le lecteur à lire la référence [56] pour découvrir les multiples possibilités¹³ qu'offrent les ensembles d'atomes froids pour les communications quantiques.

¹²Expérimentalement, il est stable 300 ns et peut être étendu à plusieurs secondes selon la référence [55].

¹³Il se trouve que cet état métastable $|b\rangle$ permet de stocker l'information portée par un *Q-bit*.

Ce prologue sur les différentes réalisations de sources de photons uniques n'a pas pour objectif d'être exhaustif et il prétend simplement introduire au lecteur, peu familier avec les sources de photons uniques, les principaux protagonistes de ce domaine. Nous espérons avoir fait un tour d'horizon des principaux concepts utilisés, avoir montré combien ils étaient différents et quelles étaient leurs limites respectives. À ce sujet, nous souhaitons attirer l'attention du lecteur sur la revue scientifique *New Journal of Physics*¹⁴ qui propose justement un numéro spécial qui regroupe la plupart des articles concernant des sources de photons uniques.

Dans le cadre ce prologue, nous n'avons pas jugé utile de donner au lecteur les performances expérimentales des différentes sources rencontrées. Toutefois, pour quantifier les performances de chacunes¹⁵ et pour pouvoir les départager entre-elles, il est judicieux de définir des figures de mérites pertinentes. La première de toutes est évidemment la probabilité d'obtenir le photon unique, couramment écrite P_1 ; vient ensuite la probabilité d'avoir deux photons qui contrairement à P_1 doit être la plus faible possible. Pour quantifier P_2 , nous utiliserons plus couramment le paramètre $g^{(2)}(0)$ dont la valeur fournit une indication représentative de la probabilité d'avoir deux photons divisée par celle d'en avoir un seul. Un $g^{(2)}(0)$ proche de zéro décrira alors une source qui présente un réel comportement de source de photons uniques.

¹⁴En accès libre par internet à l'adresse : <http://www.iop.org/EJ/journal/NJP>

¹⁵Notons que les performances de ces sources sont disponibles, via un tableau, au chapitre 3 (page 128).

Chapitre 1

La conversion paramétrique comme source de paires de photons ... et de photons uniques !

Nous avons vu dans le prologue sur les sources de photons uniques comment utiliser des paires de photons pour réaliser une source de photons annoncés. Cette section avait pour but d'expliquer le concept de base introduit par Mandel en 1986 et nous avons supposé alors que la génération des paires de photons était peu efficace et suivait une statistique poissonnienne. Ce premier chapitre sera l'occasion d'expliquer plus précisément l'origine des paires de photons. Nous allons donc décrire les interactions paramétriques optiques dans un matériau non-linéaire et nous montrerons comment celles-ci sont en mesure de nous fournir des paires de photons et quelles en sont leurs caractéristiques.

La première partie de ce chapitre traitera donc de la génération paramétrique dans le cas le plus général possible. Contrairement au prologue où nous avons laissé entendre que l'utilisateur obtenait des impulsions lumineuses contenant une¹ paire de photons, nous décrirons dans la seconde partie un mode de fonctionnement continu où il n'est plus possible de parler en terme d'impulsions lumineuses. Dans notre configuration, l'instant de création des paires n'est plus défini précisément et nous tâcherons donc d'expliquer comment réussir à distinguer et à isoler les paires les unes des autres pour obtenir quand même une source de photons uniques annoncés. Ce régime de fonctionnement est dit *asynchrone* par identification au protocole de communication *ATM* : « *Asynchronous Transfert Mode* » largement utilisé

¹ou plusieurs. . .

pour les télécommunications.

Une fois ces points éclaircis, nous prendrons le temps de définir les figures de mérites associées aux sources de photons uniques et nous utiliserons ce que nous avons appris de l'interaction paramétrique et de l'*ATM* pour essayer de prévoir théoriquement les performances que nous sommes en droit d'attendre de notre source. Par performances, nous entendons la probabilité P_1 d'avoir un photon et $g^{(2)}(0)$ qui quantifie le caractère unique des photons émis.

Nous clôturerons ce chapitre en expliquant l'importance relative de P_1 et de $g^{(2)}(0)$ en fonction de l'utilisation qui sera faite de la source.

1.1 Une source de paires de photons

1.1.1 L'interaction lumière-matière dans un cristal non-linéaire

Pour appréhender le fonctionnement d'un générateur de paires de photons, il est nécessaire de comprendre comment la lumière interagit avec les atomes de la matière. Imaginons-les constitués « d'électrons reliés par des ressorts aux noyaux considérés fixes ». Les photons du champ électrique optique traversant le matériau vont être absorbés et mettre en vibration les électrons qui deviennent alors, avec les noyaux, autant de dipôles oscillants susceptibles de réémettre la lumière absorbée. La réponse du matériau, ou *polarisation*, est parfaitement linéaire si les électrons vibrent à la fréquence d'excitation. À l'opposé on qualifie certains matériaux de *non-linéaires* par l'aptitude de leur polarisation à ré-émettre cette lumière non seulement à la fréquence du champ incident, mais aussi à toutes ses harmoniques. Dans le cas où plusieurs champs sont présents, toutes les combinaisons de fréquence sont possibles. Plus particulièrement, les *interactions paramétriques à trois ondes*, sont des phénomènes du 2nd ordre impliquant deux champs initiaux, dits de *pompe* et *signal*, qui sont convertis en un troisième champ *idler* à l'aide du tenseur susceptibilité d'ordre deux $\chi^{(2)}$ caractérisant les matériaux non-linéaires. On nomme $\vec{P}$ la polarisation du milieu en fonction du champ électrique [57, 58] qui sert de source aux champs électriques présents au sein du matériau par la relation $\vec{D} = \epsilon_0 \vec{E} + \vec{P}$. En plus du terme linéaire, sa composante non-linéaire, qui couple les champs de pompe et signal, s'écrit :

$$\vec{P}_{NL}(\omega_i) = \epsilon_0 [\chi^{(2)}] \vec{E}(\omega_p) \vec{E}(\omega_s) + \dots \quad (1.1)$$

Ici, la configuration qui nous intéresse de l'interaction paramétrique correspond au cas où un champ, dit de *pompe*, est envoyé dans le cristal, ainsi

qu'un champ *signal*, de fréquence plus faible, pour initier la conversion vers le champ *idler*. C'est la *conversion paramétrique de fréquences*. Les fréquences associées aux champs pompe, signal et idler vérifient la *conservation de l'énergie* :

$$\hbar\omega_p = \hbar\omega_s + \hbar\omega_i$$

Nous verrons cependant dans la suite qu'il est possible de n'envoyer dans le cristal qu'un seul champ de pompe intense et que la conversion paramétrique vers des photons, signal et idler, d'énergie plus faible peut quand même avoir lieu.

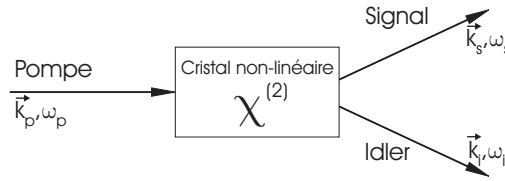


FIG. 1.1 – Illustration de l'interaction paramétrique à 3 ondes dans un cristal non-linéaire d'ordre deux, avec $\omega_{p,s,i}$ les fréquences et $\vec{k}_{p,s,i}$ les vecteurs d'onde des champs pompe, signal et idler.

À ω_{pompe} donné, il existe une infinité de couples $\{\omega_{signal}, \omega_{idler}\}$ satisfaisant la conservation de l'énergie, mais la dispersion du matériau, qui implique des vitesses de phases différentes pour les trois ondes, fait que ce qui a été créé en un point du cristal, à une fréquence donnée, est généralement détruit par interférence avec ce qui est créé, à cette même fréquence, en un autre point. C'est pourquoi les ondes signal et idler ne peuvent être créées efficacement que si, et uniquement si, une condition supplémentaire est vérifiée : la conservation de l'impulsion ou *l'accord de phase*. Cette relation traduit le fait que chaque onde se propage dans le matériau à la même vitesse que la polarisation non-linéaire qui lui sert de source. De l'infinité de couples qui satisfaisaient la conservation de l'énergie, il ne reste qu'une plage restreinte de fréquences pour les couples qui satisfont les deux équations simultanément.

Jusqu'à la fin de ce chapitre nous nous intéresserons à la conversion paramétrique dans le cas le plus général², représenté sur la figure 1.1, où pour des champs de pompe, signal et idler, monochromatiques, la condition de conservation de l'énergie et d'accord de phase seront satisfaites. Nous ne chercherons pas à donner une quelconque justification à l'accord de phase et supposons

²Cette approche est pédagogique et le lecteur désirant avoir des informations sur la conversion paramétrique en *mode guidé* ou sur la technique d'accord de phase utilisée dans notre cas se reportera directement au chapitre 2.

que les champs $\{\omega_p, \vec{k}_p\}$, $\{\omega_s, \vec{k}_s\}$ et $\{\omega_i, \vec{k}_i\}$ satisfont les équations suivantes qui correspondent respectivement à la conservation de l'énergie et à l'accord de phase :

$$\boxed{\hbar\omega_p = \hbar\omega_s + \hbar\omega_i} \quad (1.2)$$

$$\boxed{\vec{k}_p = \vec{k}_s + \vec{k}_i} \quad (1.3)$$

Pour passer à la mise en équation des interactions paramétriques optiques, nous allons considérer les champs constitués de photons et nous intéresser maintenant aux photons en tant que corpuscules. Dans le cas où seuls des photons de pompe sont présents dans le cristal non-linéaire, nous montrerons par un calcul quantique qu'ils peuvent donner naissance à des paires de photons signal et idler, bien que rien n'ait été injecté dans le cristal à ces longueurs d'onde pour initier la conversion paramétrique ; la création des paires est initiée à partir du bruit quantique. Nous reviendrons ensuite sur l'approximation monochromatique, pour déterminer la statistique des paires de photons. Une des conditions, énoncée dans le prologue, pour réaliser une source de photons annoncés repose sur la création simultanée des photons de la paire. Nous finirons cette partie, en rappelant l'expérience qui a permis de mesurer l'intervalle de temps qui sépare effectivement l'instant de création des deux photons.

1.1.2 L'interaction lumière-matière mise en équation

Écriture de l'hamiltonien d'interaction

Voyons pour l'instant comment l'interaction paramétrique peut se traduire en terme d'équations et quelles sont les solutions dans le cas général. Dans la configuration de la figure 1.1, l'Hamiltonien du système total $\hat{H}_{tot}$ peut s'écrire comme la somme d'un Hamiltonien non perturbé ($\hat{H}_0$) et d'un Hamiltonien d'interaction ($\hat{H}_{int}$) (voir le chapitre 8 de la référence [59]).

- Un hamiltonien du champ électrique contenant $n_j = \hat{a}_j^\dagger \hat{a}_j$ photons dans le mode $j = p, s, i$ correspondant respectivement aux modes de la pompe, du signal et de l'idler :

$$\hat{H}_0 = \sum_{j=p,s,i} \hbar\omega_j \left(\hat{a}_j^\dagger \hat{a}_j + \frac{1}{2} \right) \quad (1.4)$$

- Un hamiltonien de couplage des trois champs par le tenseur susceptibilité non-linéaire d'ordre deux du cristal :

$$\hat{H}_{int} = i\hbar g \left(\hat{a}_s^\dagger \hat{a}_i^\dagger \hat{a}_p - \hat{a}_s \hat{a}_i \hat{a}_p^\dagger \right) \quad (1.5)$$

où $\hat{a}_j^\dagger$ et $\hat{a}_j$ sont respectivement les opérateurs création et annihilation de photon et g le terme de couplage contenant la susceptibilité d'ordre deux $\chi^{(2)}$ caractérisant les matériaux non-linéaires (nous montrerons dans le chapitre 2 comment estimer la valeur de g).

Dans la représentation de Heisenberg, les opérateurs évoluent en fonction du temps tandis que les fonctions d'ondes restent constantes. L'équation d'évolution des opérateurs en fonction du temps s'écrit :

$$\frac{d\hat{A}}{dt} = \frac{1}{i\hbar} [\hat{A}, \hat{H}_{tot}] \quad (1.6)$$

Sachant que nous avons pour les opérateurs $\hat{a}_j^\dagger$ et $\hat{a}_j$ les relations de commutation suivantes : $[\hat{a}_m, \hat{a}_n] = [\hat{a}_m^\dagger, \hat{a}_n^\dagger] = 0$ et $[\hat{a}_m, \hat{a}_n^\dagger] = \delta_{mn}$, on peut montrer que $[\hat{n}_s + \hat{n}_i + 2\hat{n}_p, \hat{H}_{tot}] = 0$. Cette relation traduit le fait que l'opérateur $\hat{n}_s + \hat{n}_i + 2\hat{n}_p$ est une constante du mouvement :

$$\hat{n}_s(t) + \hat{n}_i(t) + 2\hat{n}_p(t) = \hat{n}_s(0) + \hat{n}_i(0) + 2\hat{n}_p(0) \quad (1.7)$$

L'équation 1.7 reflète le fait qu'un photon de pompe donne naissance à deux photons : c'est *l'interaction paramétrique optique*.

Maintenant que les bases de l'interaction paramétrique sont posées, concentrons nous sur le cas particulier de la fluorescence paramétrique pour laquelle seul un faisceau de pompe continu est injecté dans le cristal. Ce faisceau de pompe est tellement intense que nous pouvons le traiter classiquement pour simplifier les calculs et remplacer $\hat{a}_p$ par un champ complexe d'amplitude : $a_p = v_p e^{-i\omega_p t}$. À ce moment, notre nouvel Hamiltonien s'écrit :

$$\hat{H}_{tot} = \sum_{j=s,i} \hbar \omega_j \left(\hat{a}_j^\dagger \hat{a}_j + \frac{1}{2} \right) + i\hbar g \left(\hat{a}_s^\dagger \hat{a}_i^\dagger v_p e^{-i\omega_p t} - \hat{a}_s \hat{a}_i v_p e^{i\omega_p t} \right) \quad (1.8)$$

Les photons sont toujours créés par paire

Dans le cas de la fluorescence paramétrique, on a la relation de commutation $[\hat{n}_s - \hat{n}_i, \hat{H}] = 0$ qui indique cette fois que c'est l'opérateur $\hat{n}_s - \hat{n}_i$ qui est une constante du mouvement. On peut donc logiquement écrire que :

$$\boxed{\hat{n}_s(t) - \hat{n}_i(t) = \hat{n}_s(0) - \hat{n}_i(0)} \quad (1.9)$$

Cette relation nous indique que les photons signal et idler sont *toujours* émis par paire.

Les paires de photons signal/idler sont initiées à partir du bruit

Nous allons chercher à montrer que le nombre de photons dans les modes *signal* et *idler* au bout d'un temps t n'est pas nul malgré des conditions initiales nulles. Écrivons les équations d'évolution des opérateurs $\hat{a}_s(t)$ et $\hat{a}_i(t)$:

$$\begin{cases} \frac{d\hat{a}_s(t)}{dt} = g\nu_p \hat{a}_i^\dagger(t) e^{-i\omega_p t} \\ \frac{d\hat{a}_i(t)}{dt} = g\nu_p \hat{a}_s^\dagger(t) e^{-i\omega_p t} \end{cases}$$

En suivant la résolution détaillée dans la référence [60], nous pouvons écrire les solutions au temps t sous la forme :

$$\begin{cases} \hat{a}_s(t) = \left[\hat{a}_s(0) \cosh(g\nu_p t) + \hat{a}_i^\dagger(0) \sinh(g\nu_p t) \right] e^{-i\omega_s t} \\ \hat{a}_i(t) = \left[\hat{a}_i(0) \cosh(g\nu_p t) + \hat{a}_s^\dagger(0) \sinh(g\nu_p t) \right] e^{-i\omega_i t} \end{cases}$$

Nous pouvons alors facilement calculer le nombre de photons moyens contenus dans les modes à ω_s et ω_i à l'instant t sachant que :

$$\langle N_j(t) \rangle_{j=s,i} = \langle \psi(0) | \hat{a}_j^\dagger(t) \hat{a}_j(t) | \psi(0) \rangle$$

À l'aide des relations de commutation :

$$\begin{aligned} \hat{a}_j^\dagger |N_j\rangle &= \sqrt{N_j + 1} |N_j + 1\rangle & \hat{a}_j |0\rangle &= 0 \\ \hat{a}_j |N_j\rangle &= \sqrt{N_j} |N_j - 1\rangle & \langle N_k | N_l \rangle &= \delta_{k,l} \end{aligned} \quad (1.10)$$

nous allons calculer $\langle N_j(t) \rangle$ en supposant que le champ injecté à $t=0$ ne contenait aucun photon à ω_j et que les nombres de photons à ω_s et ω_i sont à priori décorrélés. La fonction d'onde initiale s'écrit donc $|\psi(0)\rangle = |0, 0\rangle = |0\rangle |0\rangle$. Nous arrivons après de rapides calculs à :

$$\boxed{\langle N_s(t) \rangle = \langle N_i(t) \rangle = \sinh^2(g\nu_p t)} \quad (1.11)$$

Nous constatons que $\langle N_j(t) \rangle$ est non nul, donc que nous pouvons quand même obtenir des photons dans les champs signal et idler même si aucun photon n'était présent à l'origine pour initier la conversion à ces fréquences.

Cela traduit le fait que les fluctuations quantiques du vide fournissent un champ d'entrée effectif d'une intensité suffisante pour initier la conversion paramétrique de la pompe vers les modes signal et idler.

Statistique de la génération des paires de photons

Jusqu'à présent, nous avons considéré que le processus de conversion paramétrique donnait naissance à deux photons signal et idler parfaitement monochromatiques. Cette approximation était acceptable pour justifier la création des paires des photons, mais nous allons voir que son rôle sur la statistique des photons est important³. Nous montrerons que des paires de photons monochromatiques suivent une statistique thermique et que l'élargissement spectral des photons conduit à une distribution poissonnienne⁴. Dans cette partie, nous allons adopter la représentation de Schrödinger dans laquelle les fonctions d'ondes évoluent dans le temps suivant la relation :

$$|\psi(t)\rangle = e^{-\frac{i}{\hbar}\hat{H}t}|\psi(0)\rangle \quad (1.12)$$

Reprenons l'Hamiltonien d'interaction de l'équation 1.8 en supposant que la pompe est classique et que le système ne produit qu'une paire unique de modes signal et idler. L'état du système après un temps t s'écrit dans notre cas :

$$|\psi(t)\rangle = e^{g v_p (\hat{a}_s^\dagger \hat{a}_i^\dagger - \hat{a}_s \hat{a}_i) t} |0, 0\rangle \quad (1.13)$$

A l'aide

- du « disentangling theorem » [61] stipulant que :

$$e^{\theta(\hat{a}_1^\dagger \hat{a}_2^\dagger - \hat{a}_1 \hat{a}_2)} = e^{\Gamma \hat{a}_1^\dagger \hat{a}_2^\dagger} e^{-\Omega(\hat{a}_1^\dagger \hat{a}_2^\dagger - \hat{a}_1 \hat{a}_2)} e^{-\Gamma \hat{a}_1 \hat{a}_2} \quad (1.14)$$

où θ est une constante, $\Gamma = \tanh(\theta)$ et $\Omega = \ln(\cosh(\theta))$

- du développement en série de Taylor,

$$e^x = \sum_{i=0}^{\infty} \frac{x^i}{i!}$$

- et des relations 1.10,

³Le lecteur pourra se référer à la section 2.1.2 du chapitre 2 pour avoir une description de notre source et en apprendre plus sur la largeur spectrale des photons émis. En attendant nous supposons que les photons émis possèdent une largeur $\Delta\omega$.

⁴Calcul inspiré par le professeur C. Fabre de l'université de Paris VI.

nous aboutissons à

$$|\psi(t)\rangle = \frac{1}{\cosh(b)} \sum_{n=0}^{\infty} \tanh(b)^n |n, n\rangle$$

où

$$b = g v_p t$$

représente « l'efficacité totale » de l'interaction, avec g l'efficacité, v_p l'amplitude de l'onde de pompe et t le temps d'interaction dans le cristal.

Calculons à présent la matrice densité du signal (ou de l'idler) :

$$\rho_{s,i}(t) = \text{Tr} \{ |\psi(t)\rangle \langle \psi(t)| \} = \frac{1}{\cosh(b)^2} \sum_{n=0}^{\infty} \tanh(b)^{2n} |n\rangle \langle n| \quad (1.15)$$

La probabilité de créer n photons suit une loi géométrique :

$$P_n = \frac{\tanh(b)^{2n}}{\cosh(b)^2}$$

Ce comportement, $P_n = K \eta^n$ où K et η sont des constantes, est caractéristique des distributions de type « *Bose-Einstein* » [60]. Ainsi, une source de paires de photons ayant un temps de cohérence infini (correspondant à des photons parfaitement monochromatiques), présentera une statistique des paires de type *Bose-Einstein*.

Supposons maintenant que notre source ne soit plus monochromatique, mais que sa largeur spectrale soit en fait due à un grand nombre N de modes possibles pour le signal et l'idler, notés s_m et i_m . L'Hamiltonien s'écrit alors :

$$\hat{H}_{int} = i\hbar g v_p \sum_{m=1}^N \left(\hat{a}_{s_m}^\dagger \hat{a}_{i_m}^\dagger - \hat{a}_{s_m} \hat{a}_{i_m} \right)$$

Il faut noter que les opérateurs d'indices m différents commutent, $[\hat{a}_{s_m}, \hat{a}_{s_n}^\dagger] = [\hat{a}_{i_m}, \hat{a}_{i_n}^\dagger] = \delta_{m,n}$ ce qui nous permet d'écrire l'état du système comme le produit tensoriel de N états identiques à celui que l'on vient de considérer :

$$|\psi(t)\rangle = \prod_{m=1}^N \left[\frac{1}{\cosh(b)} \sum_{n_m=0}^{\infty} \tanh(b)^{n_m} |n_m, n_m\rangle \right]$$

Pour voir simplement ce qu'il se passe, restreignons nous aux états à 0, 1 et 2 photons répartis dans m modes. Il vient :

$$\begin{aligned}
 |\psi(t)\rangle &= \prod_{m=1}^N \left[\frac{1}{\cosh(b)} (|0_m, 0_m\rangle + \tanh(b)|1_m, 1_m\rangle + \tanh(b)^2|2_m, 2_m\rangle + \dots) \right] \\
 |\psi(t)\rangle &= \frac{1}{\cosh(b)^N} \left[|0_1, 0_1\rangle \dots |0_N, 0_N\rangle \right. \\
 &+ \sum_{m=1}^N \tanh(b) |0_1, 0_1\rangle \dots |0_{m-1}, 0_{m-1}\rangle |1_m, 1_m\rangle |0_{m+1}, 0_{m+1}\rangle \dots |0_N, 0_N\rangle \\
 &+ \sum_{m=1}^N \tanh(b)^2 |0_1, 0_1\rangle \dots |0_{m-1}, 0_{m-1}\rangle |2_m, 2_m\rangle |0_{m+1}, 0_{m+1}\rangle \dots |0_N, 0_N\rangle \\
 &\left. + \sum_{m=1}^N \sum_{j>m}^N \tanh(b)^2 |0_1, 0_1\rangle \dots |1_m, 1_m\rangle \dots |1_j, 1_j\rangle \dots |0_N, 0_N\rangle \right] \quad (1.16)
 \end{aligned}$$

À présent, prenons la largeur spectrale⁵ de nos photons, $\Delta\lambda = 20 \text{ nm}$. Il vient un temps de cohérence des photons émis $\tau_c = \frac{\lambda^2}{c\Delta\lambda} = 0,4 \text{ ps}$ que nous pouvons considérer comme une durée caractéristique durant laquelle les photons sont émis dans le même mode. Nous pouvons affirmer qu'en pratique la durée d'une mesure complète (quelques nanosecondes) est beaucoup plus grande que le temps de cohérence de la source et que le détecteur des photons ne discriminera pas si tous les photons mesurés proviennent du même mode ou de modes différents. Nous pouvons donc écrire pour un nombre de mode N grand et une efficacité totale d'interaction b petite (l'estimation de g dans le chapitre 2 confirmera cette hypothèse) :

$$P_1 = \frac{N \tanh(b)^2}{\cosh(b)^{2N}} \sim Nb^2$$

$$P_2 = \frac{N \tanh(b)^4 + \frac{N(N-1)}{2} \tanh(b)^4}{\cosh(b)^{2N}} \sim \frac{N^2 b^4}{2}$$

Il vient alors simplement la relation

$$P_2 = \frac{P_1^2}{2}$$

qui est caractéristique des sources ayant une statistique poissonnienne.

⁵Elle sera déterminée expérimentalement dans le chapitre 2.

Mesure du temps effectif qui sépare les deux photons d'une paire

Il est important de préciser que les photons signal et idler ne sont pas émis tout à fait simultanément, mais de façon « quasi simultanée ». En effet, bien que l'interaction soit globalement régie par l'équation de conservation de l'énergie, l'un des deux photons est forcément émis avant l'autre. Après l'émission du premier photon, l'état intermédiaire dans lequel est placé le système champ-matière n'existe pas⁶. C'est pourquoi le second photon doit suivre le premier afin de ne pas violer la loi de conservation de l'énergie. A ce titre, Hong, Ou et Mandel [62] furent les premiers à mesurer le temps séparant l'émission des deux photons. Ils utilisèrent pour cela une mesure de corrélation de photons interférant sur une lame semi-réfléchissante (voir figure 1.2).

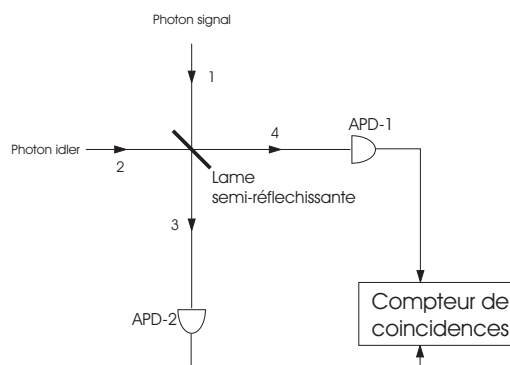


FIG. 1.2 – Montage interférométrique utilisé par Hong, Ou et Mandel pour mesurer l'intervalle de temps entre les instants de création des photons d'une paire [62].

Si deux photons identiques arrivent simultanément sur les entrées 1 et 2 de la lame, ceux-ci interfèrent destructivement, les obligeant à sortir ensemble par le bras 3 ou 4. Dans ce cas, le taux de coïncidences entre les détecteurs 1 et 2 est évidemment nul : nous sommes dans le « trou » de la courbe 1.3. Cependant si un des photons est retardé par rapport à l'autre, le recouvrement de leur paquet d'onde diminue jusqu'au point où le retard est tel que les deux photons « ne se voient plus » : ils prennent aléatoirement la sortie 3 ou 4 sans tenir compte du choix de l'autre, signifiant l'apparition de coïncidences. L'apparition progressive des coïncidences correspond aux flancs du « trou » et les photons « ne se voient plus » quand la courbe 1.3 atteint son maximum. Considérons que le photon signal a parcouru une distance $c\tau_1$ avant d'arriver

⁶On le qualifie d'état virtuel.

sur l'entrée 1 du séparateur de faisceau et que le photon idler a lui parcouru une distance $c\tau_2$ pour arriver sur l'entrée 2. Le taux de coïncidence entre les sorties 3 et 4 est donnée par :

$$R_{3,4} = \eta_1\eta_2|M|^2[1 - e^{(-\Delta\omega(\tau_2-\tau_1))^2}] \quad (1.17)$$

où $\eta_{1,2}$ représente l'efficacité des détecteurs et M est un coefficient exprimé en photons par seconde. La mesure par interférométrie permet d'obtenir une résolution de l'ordre de $1/\Delta\omega$ en s'affranchissant des limitations dues au temps de réponse⁷ des détecteurs utilisés. Au final, la confrontation de la théorie et de l'expérience a permis d'estimer que le temps qui sépare les deux photons d'une paire est inférieur à la centaine de femtosecondes.

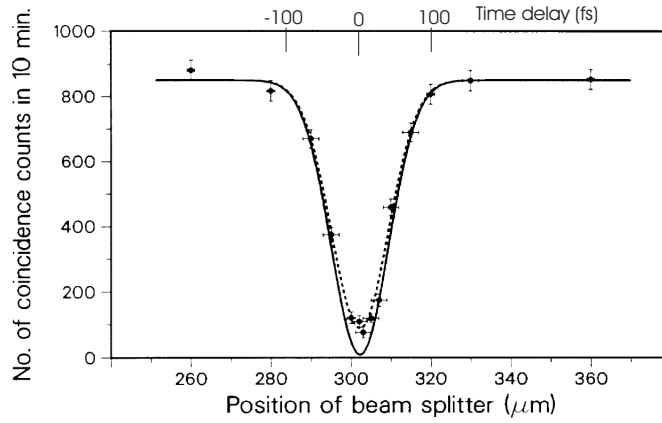


FIG. 1.3 – Résultat de l'expérience d'interférométrie à deux photons de Hong, Ou et Mandel [62]. La courbe en trait plein correspond à un ajustement selon la loi 1.17, tandis que celle en pointillé correspond à l'expérience avec ses barres d'erreurs.

En conclusion

Pour résumer, les interactions paramétriques optiques sont à l'origine de paires de photons signal/idler créés quasi-simultanément dont la statistique est poissonnienne. Fort de ces prédictions, nous n'avons plus qu'à utiliser un laser pulsé pour créer des impulsions contenant, au plus, une paire de photons par impulsions comme nous l'avons présenté dans le prologue. Pourtant, nous

⁷qui correspond à environ un million de fois plus. . .

avons choisi d'utiliser un laser continu qui implique alors une méconnaissance totale de l'instant de création des paires et du temps qui sépare deux paires successives. Nous verrons que l'utilisation d'un laser continu, n'empêche en rien l'utilisation des paires de photons pour réaliser une source de photons uniques annoncés. Au-delà de la séparation des photons signal/idler de la paire pour avoir un photon annoncé, nous allons expliquer dans la suite comment distinguer et isoler les paires les unes des autres pour obtenir un photon unique annoncé. On dit, dans ce cas, que la source est de type *asynchrone*.

1.2 Une source de photons asynchrone

Mode de Transfert Asynchrone, ATM : adj., [télécom] : *Mode de communication dans lequel deux ordinateurs s'échangent des informations sans être synchronisés : les informations sont transmises et traitées à intervalles variables selon les ressources disponibles, en utilisant des références temporelles différentes. Chaque caractère à transmettre est précédé d'un signal de départ qui prépare l'appareil récepteur à la réception et à l'enregistrement du caractère puis est suivi d'un signal d'arrêt qui met au repos l'appareil récepteur pour le préparer à recevoir le caractère suivant [63].*

Quelque soit le moyen, supposons pour l'instant que nous puissions séparer les deux photons de la paire. En pratique, une fois les deux photons séparés, le photon signal est utilisé pour prévenir de l'arrivée du photon idler qui a été préalablement retardé et « *préparer l'appareil récepteur à la réception et à l'enregistrement du caractère* » (*sic*). L'appareil récepteur en optique est une *photodiode à avalanche ou APD* qui peut « voir » l'arrivée d'un photon⁸. Munissons-nous alors de deux détecteurs de ce type : nous en gardons un constamment allumé pour détecter l'arrivée éventuelle d'un photon signal et pour avertir le second détecteur qui fonctionne, lui, en mode déclenché (« *gated* » en anglais). Ceci signifie qu'à la réception d'une impulsion électrique, appelée *trigger*, il est mis en marche pendant une fenêtre temporelle ΔT durant laquelle la détection du photon idler correspondant va être possible. Dans ce contexte, au lieu de s'imaginer des impulsions lumineuses contenant un (ou plusieurs) photons, il faut se représenter des fenêtres temporelles contenant un (ou plusieurs) photons.

Jusqu'à présent, nous n'avons pas encore parlé de l'unicité des photons annoncés, pourtant sur le schéma 1.4 nous comprenons aisément que si les paires sont suffisamment séparées temporellement (c'est-à-dire par un temps

⁸Voir l'annexe B pour plus d'informations sur les compteurs de photons utilisés.

supérieur à ΔT) les photons idler annoncés arrivant sur le détecteur déclenché seront *uniques*. À l'opposé, deux paires trop proches l'une de l'autre (inférieur à ΔT) conduiront à des événements où deux photons idler arriveront sur le détecteur pendant qu'il est allumé.

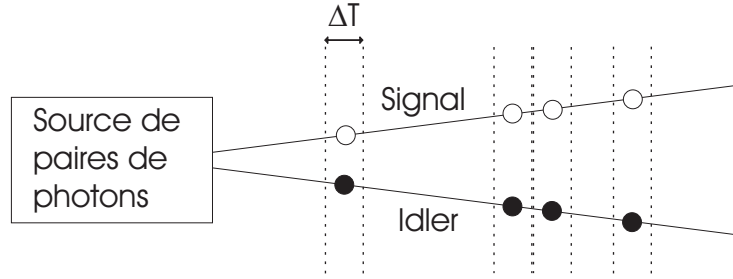


FIG. 1.4 – Schéma représentant des paires créées aléatoirement au sein du cristal.

Que faut-il comprendre par « suffisamment séparées temporellement ? » Sur le schéma 1.4, l'intervalle temporel entre deux paires de photons n'est pas constant. En effet, comme nous utilisons un laser de pompe continu, l'instant de création des paires n'est pas défini et la statistique générale des paires est poissonnienne. Évidemment, à puissance de pompe constante, le taux moyen de paires créées est constant, mais nous ne pouvons pas pour autant définir un intervalle de temps fixe entre les paires. Toutes les durées t de zéro à l'infini sont donc possibles avec une probabilité plus ou moins grande. Toutefois, nous pouvons facilement calculer la probabilité d'obtenir un intervalle d'une durée supérieure ou égale à ΔT connaissant le taux moyen μ_p de paires créées dans le cristal [61]. Il s'agit de la probabilité poissonnienne de n'avoir aucune paire sachant qu'on a un nombre moyen de paires égal à $\bar{n} = \mu_p \Delta T$ durant cet intervalle de temps :

$$P_{\bar{n}}(0) = e^{-\mu_p \Delta T} \quad (1.18)$$

Cette fonction est représentée graphiquement sur la figure 1.5 pour des taux moyens de paires $\mu_p = 1 \cdot 10^6 s^{-1}$ et $4 \cdot 10^6 s^{-1}$. Il faut comprendre à partir du schéma que la probabilité pour qu'un intervalle infiniment petit ne contienne aucun photon est proche de un. Aussi plus l'intervalle ΔT est grand, plus la probabilité qu'il soit vide est petite. De même, plus le taux moyen de paires est grand, plus la probabilité qu'un intervalle ΔT soit vide est petite. Nous proposerons dans l'annexe C, un protocole pour tracer expérimentalement cette courbe.

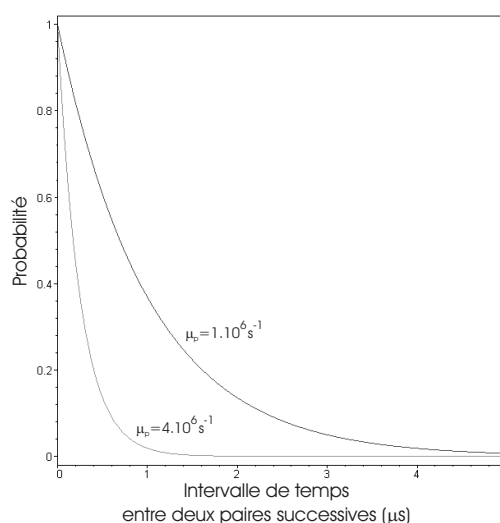


FIG. 1.5 – Probabilité d’avoir un intervalle d’une durée ΔT vide pour un taux moyen de paires $\mu_p = 1 \cdot 10^6 s^{-1}$ et $4 \cdot 10^6 s^{-1}$.

Pour l’instant, notons qu’à l’opposé du régime pulsé où la durée entre deux paires successives est connue (c’est le taux de répétition du laser impulsionnel de pompe⁹), elle est dans notre cas totalement inconnue. Mais l’instant de création des paires n’est pas important car la détection d’un photon annoncé revient à mesurer une coïncidence entre les deux photons d’une même paire et le temps relatif d’un photon par rapport à l’autre est bien plus important. À ce titre, rappelons que ce temps est quasi-nul (plus ou moins 100 fs). Une fois le photon unique annoncé, il sera séparé des suivants (ou précédents) grâce à la détection déclenchée. Nous pouvons imaginer que du point de vue du compteur de photon qui n’est allumé qu’à la réception d’un trigger, la source se comporte comme une source impulsionnelle dont le taux de répétition serait aléatoire, mais prévenue par un signal : *c’est le principe de la communication asynchrone*.

Prenons tout de même le temps de noter quelques différences par rapport au mode de fonctionnement impulsionnel :

En régime impulsionnel : la fréquence moyenne d’émission des paires par la source est définie par celle du laser de pompe. En changeant simplement la puissance du laser, nous influons uniquement sur la probabilité de créer une ou plusieurs paires par impulsion.

En régime continu : la fréquence moyenne de répétition de la source est

⁹ou du moins ses harmoniques.

fixée par le taux de création de paires μ_p lui même directement lié à la puissance du laser de pompe. Cependant la probabilité de créer une ou plusieurs paires par fenêtre temporelle est aussi liée à μ_p . Dans notre cas, la fréquence d'émission des photons et leur probabilité d'être uniques¹⁰ sont liées à la puissance du laser de pompe.

1.3 Une source de photons uniques !

Avant de s'aventurer dans le domaine des sources de photons uniques, nous souhaitons d'abord définir les termes usuels de cette communauté. Tout d'abord les probabilités d'avoir 0, 1 ou 2 photons seront couramment notées P_0 , P_1 et P_2 , puis il existe une fonction qui quantifie le caractère « unique » des photons émis : c'est la fonction d'autocorrélation d'ordre deux à $\tau = 0$ que nous noterons $g^{(2)}(0)$. Une fois les bases posées, nous étudierons l'influence des paramètres de fonctionnement de la source, comme par exemple la puissance du laser de pompe, sur les performances globales que nous pouvons espérer.

1.3.1 Définition des paramètres utiles

Pour commencer, il nous semble judicieux de définir tout d'abord quelques probabilités qui nous serviront de support pour développer plus facilement les formules correspondant aux probabilités P_0 , P_1 , P_2 et $g^{(2)}(0)$. Nous allons d'abord relier μ_p , le taux de paires créées au sein du cristal, à $\mu_{s,i}$, les taux de photons signal et idler présents en un point du montage 1.7, puis nous calculerons les probabilités P_a^1 et P_{na}^{1+} que nous définissons à l'aide de la figure 1.6 comme :

- La première probabilité, P_a^1 , correspond à la probabilité que le photon idler annoncé soit présent dans la fenêtre temporelle ΔT , qui a été ouverte suite à la détection du photon signal correspondant. L'indice 1 est associé au nombre de photons tandis que a correspond au caractère *annoncé* du photon.
- La seconde probabilité, P_{na}^{1+} , correspond à la probabilité qu'un ou plusieurs photons *non-annoncés* se glissent dans la fenêtre temporelle ΔT qui a été ouverte pour la détection d'un autre photon.

¹⁰qui correspond à la probabilité de créer deux paires plus proches que ΔT .

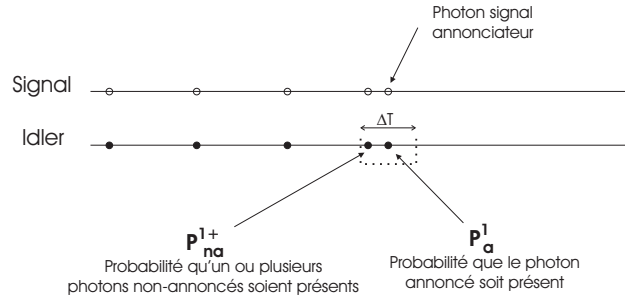


FIG. 1.6 – Schéma de principe sur la probabilité d'avoir un ou plusieurs photons dans une fenêtre temporelle ΔT .

Taux moyen de paires créées : μ_p

Dans le cristal, le taux moyen de photons signal et idler $\mu_{s,i}$ est égal au taux moyen de paires créées μ_p . Cette relation, valable au point de création de la paire (cf. équation 1.7), ne l'est plus dès lors que l'on sort du cristal et que l'on souhaite collecter les photons. Le taux moyen effectif de photons signal et idler collectés dépend toujours de μ_p , mais il faut maintenant rajouter un coefficient dû à une mauvaise collection des photons et aux pertes vues par les photons. Les facteurs $\gamma_{s,i}$ représentent justement les coefficients de collection des photons signal et idler dans lesquels les pertes sont incluses. En pratique, ces coefficients sont inférieurs à 1 et, de plus, il n'y a aucune raison que γ_s soit égal à γ_i :

$$\mu_s = \gamma_s \mu_p \quad (1.19)$$

$$\mu_i = \gamma_i \mu_p \quad (1.20)$$

Ces relations nous permettront par la suite d'exprimer les probabilités P_0 , P_1 et P_2 en fonction de μ_p , en écrivant le taux de détection des photons signal dans l'APD de la figure 1.7 :

$$S_{dét} = \mu_s \eta_{dét} = \gamma_s \mu_p \eta_{dét} \quad (1.21)$$

avec $\eta_{dét}$ correspondant à l'efficacité du détecteur.

Probabilité P_a^1 d'avoir le photon annoncé dans sa fenêtre ΔT

Pour estimer P_a^1 , il est astucieux d'écrire cette probabilité d'avoir le photon annoncé sous la forme opposée¹¹ :

$$P_a^1 = 1 - P_a^0$$

¹¹La probabilité de « ne pas avoir de photon annoncé » : P_a^0

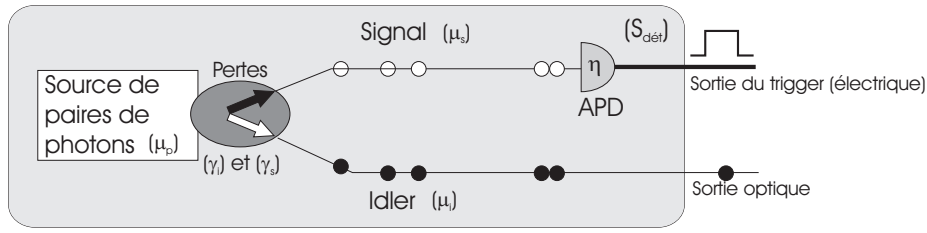


FIG. 1.7 – Schéma simplifié d'une source de photons uniques annoncés en régime de pompe continu. Les taux μ_p , μ_s , μ_i et $S_{dét}$ s'expriment en s^{-1} tandis que γ_s et γ_i correspondent aux pertes vues par les photons et η à l'efficacité du détecteur.

Comme les photons sont *toujours* créés par paires, l'origine des événements à zéro photon est de deux natures. Elle peut provenir des cas où le photon idler n'a pu être collecté (tandis que le photon signal l'a été) ou des imperfections du système de détection. À l'heure actuelle, toutes les photodiodes à avalanches présentent des « coups sombres », qui ne correspondent effectivement à aucune détection de photon mais qui sont uniquement dus au bruit thermique, et notre APD associée aux photons signal ne déroge pas à la règle. Cette dernière présente un taux de « coups sombres », notés D_{Ge}^c et expérimentalement, nous retrancherons sa valeur au taux brut de détection, pour avoir accès au taux de coups net S_{Ge}^{net} dans ce détecteur¹² :

$$S_{Ge}^{net} = S_{Ge}^{brut} - D_{Ge}^c$$

Mathématiquement, les probabilités de ne pas avoir de photons en sortie de la source correspondent aux cas où :

- le photon signal a été détecté, mais l'idler a été perdu, ce qui s'écrit en normalisant par S_{Ge}^{brut} :

$$\frac{(1 - \gamma_i) S_{Ge}^{net}}{S_{Ge}^{brut}}$$

- un coup sombre est survenu dans l'APD, il n'y aura jamais de photon idler correspondant. Nous obtenons simplement :

$$\frac{D_{Ge}^c}{S_{Ge}^{brut}}$$

¹²Nous avons pris le parti d'affecter un indice Ge , pour Germanium, au détecteur associé aux photons signal bien que nous soyons encore dans un chapitre qui traite de considérations d'ordre général. Nous espérons ainsi assurer une cohérence entre les chapitres et familiariser le lecteur avec la notation. En attendant, pour de plus amples informations sur les performances et l'origine de coups sombres de ce détecteur en germanium, le lecteur peut se reporter à l'annexe B.

Nous trouvons donc pour P_a^0 , qui correspond à la probabilité de ne pas avoir le photon annoncé dans sa fenêtre de détection :

$$P_a^0 = \frac{D_{Ge}^c + (1 - \gamma_i)S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c}$$

Nous pouvons maintenant écrire simplement la probabilité, P_a^1 , d'avoir le photon annoncé dans la fenêtre ΔT .

$$P_a^1 = \frac{\gamma_i S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c} \quad (1.22)$$

Remarquons que cette probabilité ne dépend pas de la durée de la fenêtre. Comme les deux photons de la paire sont simultanés à moins de 100 fs , le photon idler annoncé n'a aucune raison d'être en retard ou en avance par rapport au photon signal, au point d'être en dehors de la fenêtre temporelle ΔT , quelle que soit la durée¹³.

Probabilité P_{na}^{1+} d'avoir un ou plusieurs photons non-annoncés dans la fenêtre ΔT

Cherchons, à présent, la probabilité d'avoir un photon idler, contenu dans la fenêtre temporelle ΔT , mais qui ne soit pas annoncé par un photon signal. Le détecteur de photons idler fonctionne en mode « déclenché », ce qui signifie qu'à la réception d'un trigger électrique, il est mis en marche durant une fenêtre temporelle ΔT . À l'évidence, la probabilité d'avoir un ou plusieurs photons *non-annoncés* dans cette fenêtre est liée à la relation 1.18. Calculons donc la probabilité que cet intervalle contienne un ou plusieurs photons idler en définissant par $\bar{n}_i$ le nombre moyen de photons idler durant ΔT :

$$P_{na}^{1+} = 1 - P_{\bar{n}_i}(0) = 1 - e^{-\mu_i \Delta T}$$

Pour $\mu_i \Delta T$ petit (ce qui sera le cas en pratique), P_{na}^{1+} se simplifie à :

$$P_{na}^{1+} \approx \gamma_i \mu_p \Delta T \quad (1.23)$$

Nous pouvons à présent établir les formules associées à P_2 , P_1 et P_0 .

¹³Dans la limite où $\Delta T \gg 100 \text{ fs}$.

1.3.2 Probabilité P_2 d'avoir deux photons

Cette probabilité correspond simplement au cas où un photon *annoncé* et un photon *non-annoncé* se retrouvent ensemble dans la fenêtre de détection associée au premier photon. On a donc :

$$P_{2+} = P_a^1 \times P_{na}^{1+}$$

Il ne nous reste plus qu'à remplacer P_a^1 et P_{na}^{1+} par leur valeur respective à l'aide des formules 1.22 et 1.23, ce qui donne :

$$P_{2+} \approx \gamma_i^2 \mu_p \Delta T \left(\frac{S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c} \right) \quad (1.24)$$

$$\boxed{P_{2+} \approx \gamma_i^2 \mu_p \Delta T \left(\frac{\gamma_s \mu_p \eta_{Ge}}{\gamma_s \mu_p \eta_{Ge} + D_{ge}^c} \right)} \quad (1.25)$$

1.3.3 Probabilité P_1 d'avoir un photon unique

Cette probabilité s'écrit tout simplement comme la probabilité de détecter le photon idler annoncé ou bien de le rater, tout en détectant un autre photon *non-annoncé*¹⁴ :

$$P_1 = (1 - P_{na}^{1+}) \times P_a^1 + (1 - P_a^1) \times P_{na}^{1+}$$

Comme précédemment, en insérant les valeurs de P_a^1 et P_{na}^{1+} , nous trouvons facilement les expressions suivantes :

$$P_1 \approx \gamma_i \left\{ \left(\frac{S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c} \right) [1 - 2\gamma_i \mu_p \Delta T] + \mu_p \Delta T \right\} \quad (1.26)$$

$$\boxed{P_1 \approx \gamma_i \left\{ \left(\frac{\gamma_s \mu_p \eta_{Ge}}{\gamma_s \mu_p \eta_{Ge} + D_{ge}^c} \right) [1 - 2\gamma_i \mu_p \Delta T] + \mu_p \Delta T \right\}} \quad (1.27)$$

Au final, remarquons que μ_p pourra être aussi grand que possible (dans la limite $\mu_p \Delta T \ll 1$), la probabilité d'avoir un photon unique ne dépassera jamais la limite imposée par le couplage γ_i : « *rien ne sert de pomper, il faut collecter à point !* »

¹⁴Ce deuxième cas doit être pris en compte car expérimentalement nous n'avons aucun moyen de différencier un *vrai* photon annoncé d'un photon *non-annoncé* qui prendrait sa place si le premier venait à se perdre.

1.3.4 Probabilité P_0 de ne pas avoir de photon

Finalement, la probabilité de n'avoir aucun photon dans la fenêtre temporelle ΔT correspond aux événements où aucune détection n'a eu lieu (photon idler annoncé ou non). P_0 s'écrit alors simplement :

$$P_0 = (1 - P_a^1) \times (1 - P_{na}^{1+})$$

$$P_0 = \left(1 - \frac{\gamma_i S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c}\right) \times (1 - \gamma_i \mu_p \Delta T) \quad (1.28)$$

$$P_0 = \left(1 - \frac{\gamma_i \gamma_s \mu_p \eta_{Ge}}{\gamma_s \mu_p \eta_{Ge} + D_{ge}^c}\right) \times (1 - \gamma_i \mu_p \Delta T) \quad (1.29)$$

Il est rassurant de noter que la somme de P_0 , P_1 et P_2 vaut bien 1, ce qui confirme au passage la justesse des formules 1.25, 1.27 et 1.29.

1.3.5 La fonction d'autocorrélation d'ordre deux : $g^{(2)}(0)$

Nous venons de définir les probabilités P_i que i photons arrivent dans une fenêtre temporelle ΔT . Si nous souhaitons classer les sources de photons uniques en fonction de leurs performances, il est assez facile de placer en tête la source idéale de photons uniques où

$$P_0 = 0 \quad P_1 = 1 \quad P_2 = 0 \quad (1.30)$$

Mais ensuite d'un point de vue performance, *comment comparer les différentes solutions retenues ?* Faut-il :

- favoriser les sources qui présentent un P_1 très élevé ?
- plutôt favoriser celles qui offrent le plus faible P_2 ?

La solution serait de définir une source référence et de rapporter les performances de toutes les sources de photons uniques à cette référence. Nous allons montrer que la fonction $g^{(2)}(0)$ remplit à elle seule ce rôle et pour cela nous allons tout d'abord l'introduire dans sa forme la plus générale, c'est-à-dire $g^{(2)}(\tau)$.

$g^{(2)}(\tau)$ correspond à la fonction d'autocorrélation normalisée qui quantifie le taux de corrélation qui existe au sein d'un rayonnement $I(t)$ à un temps τ [64, 65, 60]. Cette fonction s'écrit :

$$g^{(2)}(\tau) = \frac{\langle I(t)I(t+\tau) \rangle}{\langle I(t) \rangle \langle I(t) \rangle} \quad (1.31)$$

En se plaçant à $\tau = 0$, la fonction d'autocorrélation normalisée devient

$$g^{(2)}(0) = \frac{\langle I(t)I(t) \rangle}{\langle I(t) \rangle \langle I(t) \rangle} \quad (1.32)$$

et peut s'interpréter comme le taux de *détections simultanées*, normalisé par le taux *moyen* de détections au carré.

À l'aide d'un montage de type *Hanbury-Brown & Twiss*, représenté sur la figure 1.8, les chercheurs du même nom ont démontré qu'il était possible d'observer la fonction d'autocorrélation $\langle I(t)I(t + \tau) \rangle$ [66]. Dans le montage

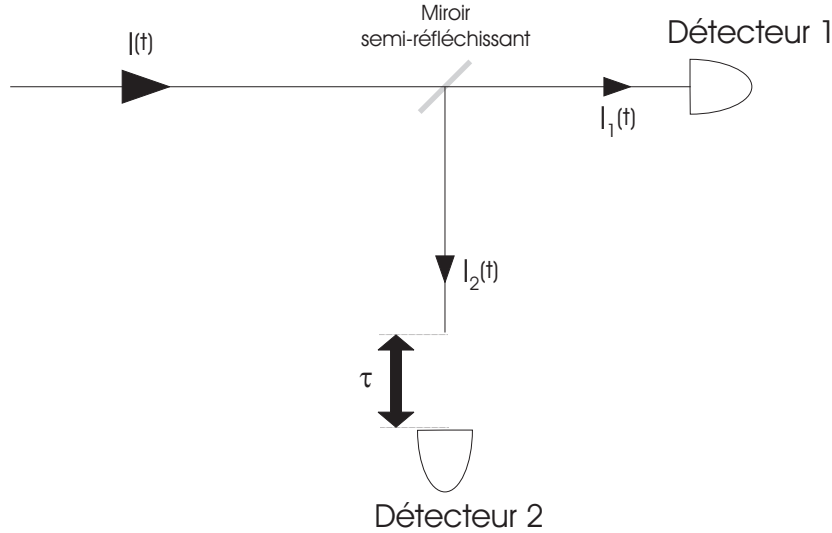


FIG. 1.8 – Schéma simplifié du montage de type *Hanbury-Brown & Twiss* permettant de mesurer la fonction d'autocorrélation d'ordre 2 $g^{(2)}(\tau)$. Dans le cas d'un miroir semi-réfléchissant, notons que $I_1(t) = I_2(t) = \frac{I(t)}{2}$.

expérimental, $\langle I_1(t)I_2(t) \rangle$ correspond au taux de détections simultanées dans les détecteurs 1 et 2, tandis que $\langle I_1(t) \rangle \langle I_2(t) \rangle$ correspond simplement au produit des taux de détections moyens associés aux détecteurs 1 et 2.

Nous allons réécrire $g^{(2)}(0)$ en fonction des probabilités P_1 et P_2 , en prenant soin de noter que le comportement classique du miroir semi-réfléchissant introduit des facteurs $1/2$ dans les formules suivantes :

$$\begin{aligned} \langle I_1(t)I_2(t) \rangle &\propto \frac{P_2}{2} \\ \langle I_1(t) \rangle \langle I_2(t) \rangle &\propto \frac{P_1}{2} \times \frac{P_1}{2} \end{aligned}$$

Nous obtenons alors une forme de $g^{(2)}(0)$ aisément calculable à partir des valeurs P_1 et P_2 . Nous estimerons à l'aide de la formule suivante le $g^{(2)}(0)$ associé à notre source :

$$\boxed{g^{(2)}(0) = \frac{2P_2}{(P_1)^2}} \quad (1.33)$$

Un moyen de comparaison ?

Nous pouvons souligner le caractère comparatif associé à la fonction $g^{(2)}(0)$. Dans l'introduction nous avons vu que la manière la plus simple de simuler une source de photons uniques est d'utiliser une source poissonnienne atténuée à 0,1 photon par impulsion. Cette faible valeur est imposée par le fait que la probabilité d'avoir deux photons est liée à celle d'en avoir un par :

$$P_2^{poisson} \approx \frac{(P_1^{poisson})^2}{2} \quad (1.34)$$

Pour les sources poissonniennes, il est donc extrêmement facile de calculer P_2 (ou P_1) connaissant uniquement P_1 (ou P_2 respectivement). Cela confère à ce type de source l'avantage d'être un point de référence universel. Imaginons une source poissonnienne équivalente à celle étudiée, c'est-à-dire possédant le même P_1 . Nous pouvons alors réécrire $g^{(2)}(0)$ comme :

$$g^{(2)}(0) = \frac{2P_2^{source}}{(P_1^{source})^2} = \frac{2P_2^{source}}{(P_1^{poisson})^2} \Rightarrow \boxed{g^{(2)}(0) \approx \frac{P_2^{source}}{P_2^{poisson}}}$$

$g^{(2)}(0)$ est donc simplement le rapport entre le P_2 de la source étudiée et celui d'une source poissonnienne équivalente. La fonction d'autocorrélation d'ordre deux $g^{(2)}(0)$ permet donc de comparer aisément les performances de n'importe quelle source étudiée par rapport à une source poissonnienne équivalente.

Nous voyons ainsi que, pour une source parfaite de photons uniques, le $g^{(2)}(0)$ vaut zéro tandis que pour une source poissonnienne celui-ci vaut un. Nous obtenons donc une borne à cette fonction ainsi qu'une échelle qui quantifie le caractère unique des photons émis. *Plus le $g^{(2)}(0)$ est proche de zéro, plus la source se comporte comme une source de photons uniques.* On dit d'une source possédant un $g^{(2)}(0)$ inférieur à un qu'elle présente une statistique *sub-poissonnienne*.

Pour notre source de photons annoncés,

$$g^{(2)}(0) \approx \frac{2\mu_p \Delta T \left(\frac{S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c} \right)}{\left(\mu_p \Delta T + \left[1 - 2\gamma_i \mu_p \Delta T \right] \left(\frac{S_{Ge}^{net}}{S_{Ge}^{net} + D_{Ge}^c} \right) \right)^2} \quad (1.35)$$

$$\boxed{g^{(2)}(0) \approx \frac{2\mu_p \Delta T \left(\frac{\gamma_s \mu_p \eta_{Ge}}{\gamma_s \mu_p \eta_{Ge} + D_{ge}^c} \right)}{\left(\mu_p \Delta T + \left[1 - 2\gamma_i \mu_p \Delta T \right] \left(\frac{\gamma_s \mu_p \eta_{Ge}}{\gamma_s \mu_p \eta_{Ge} + D_{ge}^c} \right) \right)^2}} \quad (1.36)$$

1.3.6 Étude de l'influence des paramètres de la source sur ses performances

Une source de photons uniques est dite performante comparée à un laser atténué si elle présente une probabilité P_1 supérieure à 0,1 et une valeur de $g^{(2)}(0)$ inférieure à 1. Dans l'étude qui suit, nous allons introduire et quantifier les paramètres de notre source (tels que la collection et les pertes sur l'idler et les coups sombres D_{Ge}^c dans le détecteur des photons signal) et chercher à savoir lesquels sont déterminants sur ses performances. Définissons tout d'abord les paramètres pertinents et leurs valeurs caractéristiques qui seront utilisées pour les graphiques :

- *la collection + les pertes γ_i et γ_s* : ces facteurs compris entre 0 et 1 quantifient la probabilité de collecter (et de ne pas perdre) les photons idler et signal. En pré-supposant la valeur du couplage mesurée dans le chapitre 2, nous prendrons $\gamma_i = 0,45$ et $\gamma_s = 0,29$ comme valeurs de références.
- *Le taux de paires créées μ_p* : cette valeur sera directement reliée au taux de répétition de la source. Plus nous créons de paires par seconde, plus la fréquence de détection des photons signal est élevée. Nous choisirons arbitrairement la valeur expérimentale du chapitre 3 : $\mu_p = 6,6 \cdot 10^6 \text{ s}^{-1}$
- *Le taux de coups sombres D_{Ge}^c* : ce facteur est associé au taux de détections qui ne correspondent pas à l'arrivée d'un photon signal. Leur taux est dans notre cas compris entre 10000 et 30000 s^{-1} .
- *Durée d'allumage des détecteurs ΔT* : une valeur représentative de cet intervalle pour un compteur de photons en mode déclenché est 3 ns.

Influence du taux de collection des paires γ_i et γ_s

Le taux de collection γ_i va caractériser la probabilité que nous avons de collecter le photon idler, quand le photon signal l'a déjà été. Nous insistons sur le fait que le photon signal a déjà été collecté, car dans le cas contraire il n'y a pas de trigger émis par la source, donc pas de détection de photon idler possible.

Mais il ne faut pas croire pour autant que γ_s n'a alors aucune influence. Au même titre que l'efficacité de détection η_{Ge} , la probabilité γ_s de collecter le photon signal, va influencer le taux de répétition de la source ou bien, si ce dernier est fixé, sur la probabilité d'avoir deux photons à l'intérieur d'une fenêtre de détection ΔT .

Notre configuration est un peu particulière, car les valeurs de γ_s et de γ_i sont liées entre elles. Il ne sera pas possible expérimentalement d'améliorer

l'un sans l'autre¹⁵.

Sur le graphique 1.9, il s'agit bien du coefficient γ_i qui est utilisé en abscisse, mais nous avons considéré γ_s comme une fonction de γ_i . Dans cette étude nous avons donc supposé en prévision des mesures expérimentales du chapitre 2 que le rapport $\frac{\gamma_s}{\gamma_i}$ vaut 0,65. Nous considérons donc que cette figure représente l'influence de la collection globale des paires sur les performances de la source.

Pourtant, elle ne permet pas d'identifier l'effet précis de la récolte des photons idler γ_i ou des photons signal γ_s ce que nous proposons de regarder dans les paragraphes suivants.

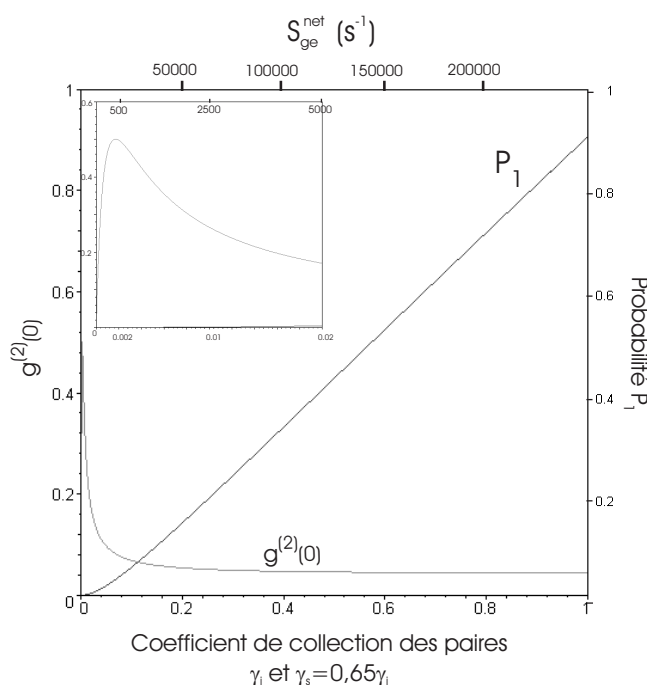


FIG. 1.9 – Évolutions de P_1 et $g^{(2)}(0)$ en fonction du coefficient de collection des paires. Il faut voir ici γ_i en abscisse et une dépendance implicite de γ_s en fonction de γ_i . Les courbes ont été tracées pour une fenêtre de 3 ns avec $\mu_p = 6,6 \cdot 10^6\text{ s}^{-1}$ et $D_{Ge}^c = 20000\text{ s}^{-1}$. À titre indicatif sont affichés sur l'axe supérieur le taux de comptage effectif dans l'APD germanium.

¹⁵Nous verrons au chapitre 2 que nous collectons les paires avec une seule et même fibre optique. Dans ce cas, bien que les coefficients de couplage soient différents, ils sont expérimentalement liés.

Influence du taux de collection des photons idler γ_i

Cependant, nous pouvons constater grâce à la formule 1.36 que $g^{(2)}(0)$ ne dépend pas directement de γ_i , tandis que P_1 en dépend directement. À ce titre, il semble assez intuitif que l'amélioration de la récolte des photons annoncés doit affecter tant la probabilité d'avoir un photon que celle d'en avoir deux. C'est pourquoi nous pouvons affirmer que l'amélioration de la récolte des photons idler augmentera la probabilité P_1 et n'aura aucune influence sur $g^{(2)}(0)$. Ce comportement est confirmé par la figure 1.10 ci-dessous qui, à γ_s constant, reproduit les performances de la source uniquement en fonction de γ_i . Nous retrouvons bien sur cette figure une probabilité P_1 proportionnelle à γ_i tandis que $g^{(2)}(0)$ reste quasiment constant.

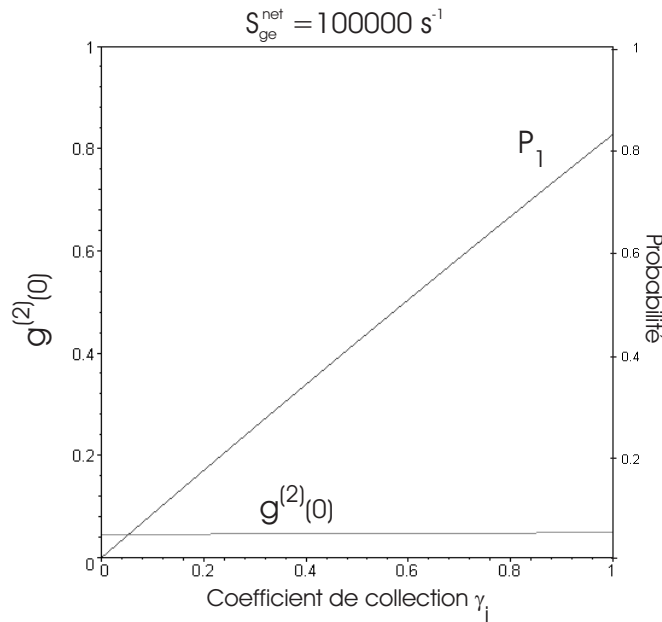


FIG. 1.10 – Évolutions de P_1 et $g^{(2)}(0)$ en gardant, cette fois, γ_s constant et en faisant varier uniquement γ_i .

Influence du taux de collection des photons signal γ_s

Plaçons-nous dans le cas idéal où la collection des photon signal et l'efficacité du détecteur *Germanium* sont égales à l'unité. Si nous sommes seulement capable de collecter les photons idler avec un coefficient γ_i , alors la probabilité P_1 d'avoir un photon unique ne peut excéder cette valeur limite. Maintenant, supposons $\gamma_s < 1$, certains photons signal peuvent être perdus, il n'y aura pas

de trigger émis par la source pour annoncer le photon idler correspondant, donc il n'y a donc aucune raison pour que la probabilité P_1 en souffre. À μ_p constant, il n'y a de même aucune raison que la probabilité d'avoir deux photons dans une fenêtre de détection dépende de γ_s .

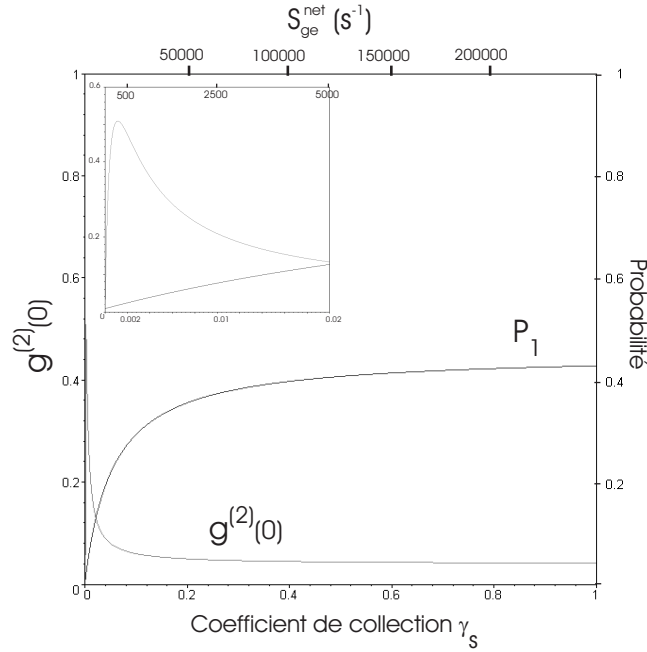


FIG. 1.11 – Évolutions de P_1 et $g^{(2)}(0)$ en gardant, cette fois, γ_i constant et en faisant varier uniquement γ_s .

La figure 1.11 représente les variations de P_1 et $g^{(2)}(0)$ vis à vis de γ_s pour μ_p constant. Nous retrouvons bien, pour les fortes valeurs de S_{Ge}^{net} , des courbes P_1 et $g^{(2)}(0)$ relativement plates donc indépendantes de la valeur de γ_s . Pourtant, le lecteur peut constater que, pour les très faibles taux de collection, P_1 tend vers 0 tandis que $g^{(2)}(0)$ passe par un point maximal avant de tendre finalement vers 0 à cause des coups sombres. En l'absence de ceux-ci, P_1 et $g^{(2)}(0)$ devraient être constants quelle que soit la valeur de γ_s . Il suffit de regarder simplement les équations 1.27 et 1.36 pour $D_{Ge}^c = 0$ pour s'en convaincre.

En pratique, il se trouve que nous ne gardons pas μ_p constant, mais plutôt S_{Ge}^{net} qui correspond au taux de comptage net dans le détecteur signal et qui proportionnel à $\gamma_s \times \mu_p$ (voir équation 1.21). De cette façon, quand nous

améliorons la récolte des photons signal, nous pouvons diminuer la puissance de pompe, donc la valeur de μ_p .

Penchons-nous donc sur l'influence de μ_p sur la statistique de la source.

Influence du taux moyen de création des paires μ_p

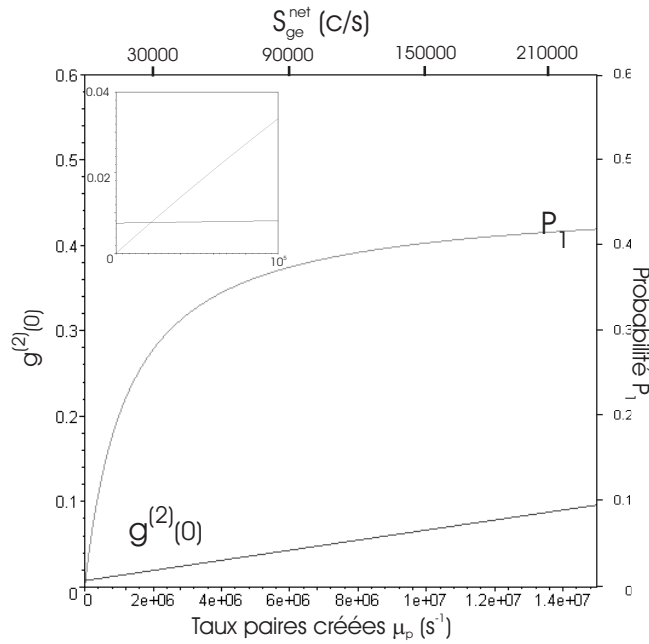


FIG. 1.12 – Évolutions de P_1 et $g^{(2)}(0)$ en fonction du taux moyen de création des paires de photons (μ_p) pour une gate de 3 ns , des coefficients de collection $[\gamma_i = 0, 45, \gamma_s = 0, 29]$ et $D_{Ge}^c = 20000\text{ s}^{-1}$.

D'après la formule 1.23, plus le taux moyen de création des paires est faible, plus la probabilité P_1^{na} d'avoir un photon supplémentaire non-annoncé dans la fenêtre ΔT est faible tandis que P_1 ne dépend que très peu de μ_p . évidemment, sur le graphique 1.12, pour les très faibles valeurs de μ_p , la probabilité d'avoir un photon unique chute vers zéro, mais ce comportement est encore une fois dû à la présence de coups sombres qui tendent à augmenter la part de P_0 dans la statistique de la source.

Nous comprenons alors que la plus faible valeur possible de μ_p est préférable pour la statistique globale de la source, mais que la plus petite valeur, que nous pouvons lui donner, est fortement limitée par la valeur de γ_s et D_{Ge}^c . Ainsi nous pouvons conclure qu'une valeur élevée de γ_s permet de garder le

rapport *signal/bruit* $\frac{S_{Ge}^{net}}{D_{Ge}^c}$ le plus élevé possible, tout en ayant le plus petit μ_p possible.

Influence du taux de coups sombres D_{Ge}^c

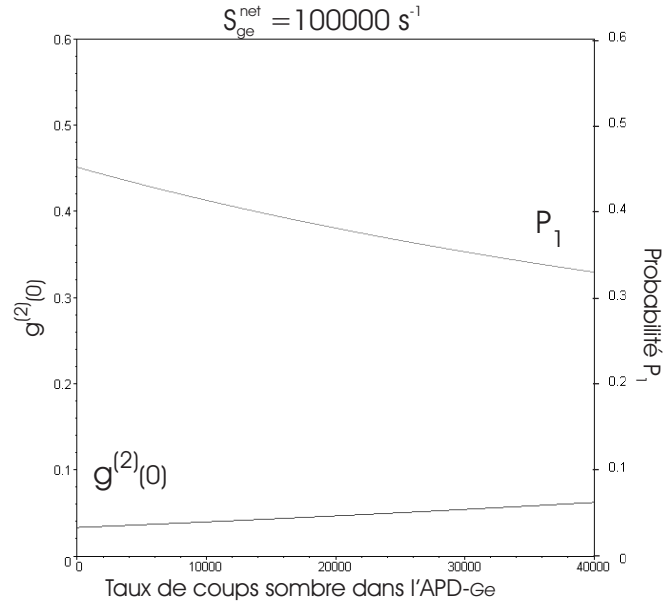


FIG. 1.13 – Évolutions de P_1 et $g^{(2)}(0)$ en fonction du taux de coups sombres (D_{Ge}^c) pour une gate de 3 ns , des coefficients de collection $[\gamma_i = 0,45, \gamma_s = 0,29]$ et $\mu_p = 6,6 \cdot 10^6\text{ s}^{-1}$.

Bien que les coups sombres ne soient pas intrinsèquement liés au processus de création des paires, ils jouent un rôle dans le sens où ils augmentent essentiellement P_0 et réduisent indirectement P_1 et P_2 comme le montre la figure 1.13 ($P_0 + P_1 + P_2 = 1$). Il se trouve que leur présence a constitué, dans le paragraphe précédent, une barrière à un choix arbitraire du paramètre μ_p qui nous aurait permis d'obtenir les performances les plus élevées.

La qualité du détecteur des photons signal a donc une importance capitale dans le comportement de la source. Expérimentalement, nous choisirons une valeur pour D_{Ge}^c la plus faible possible, en gardant tout de même une efficacité correcte. Les valeurs retenues seront discutés dans la chapitre 2.

Influence de la durée d'allumage du détecteur déclenché ΔT

Contrairement au régime impulsionnel, où les paires sont contenues à l'intérieur de l'impulsion lumineuse donc bien séparées temporellement les unes des autres, dans le cas d'une source comme la notre, fonctionnant en continu, la durée de la gate est un paramètre important sur les performances de la source. C'est en réglant la durée d'ouverture des détecteurs déclenchés que nous allons changer la probabilité, P_2 , qu'un photon supplémentaire vienne se glisser à l'intérieur de la fenêtre de détection. Au fur et à mesure que la durée de la gate diminue, la probabilité que le photon annoncé soit présent (P_1) ne doit pas changer ; par contre, la probabilité qu'un photon supplémentaire vienne s'insérer dans cette fenêtre (P_2) doit tout naturellement diminuer. Ce comportement est représenté sur le graphique 1.14 où nous observons bien une courbe associée à P_1 constante, tandis que $g^{(2)}(0)$ dépend proportionnellement de ΔT .

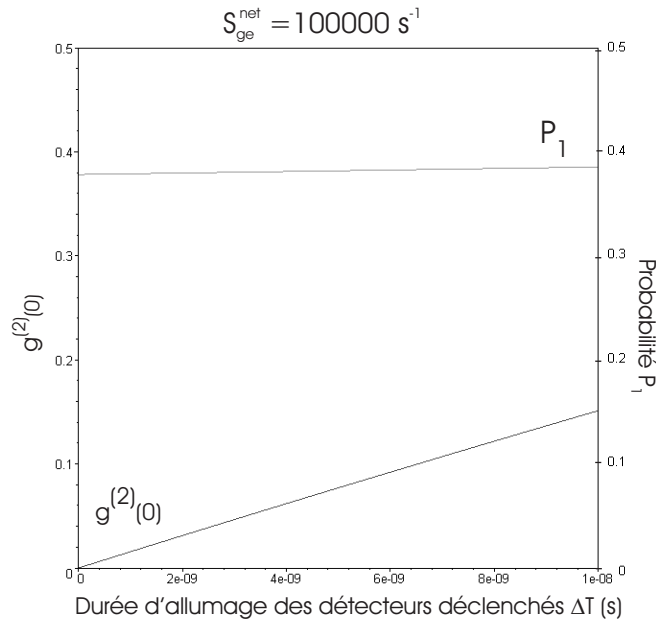


FIG. 1.14 – Évolutions de P_1 et $g^{(2)}(0)$ en fonction de la durée de la gate (ΔT) pour $D_{Ge}^c = 20000 \text{ s}^{-1}$, des coefficients de collection $[\gamma_i = 0, 45, \gamma_s = 0, 29]$ et $\mu_p = 6, 6 \cdot 10^6 \text{ s}^{-1}$.

1.3.7 Rappel de l'influence des paramètres μ_p , γ_i , γ_s et D_{Ge}^c

- μ_p est le paramètre le plus important, car de sa valeur dépend la probabilité qu'une fenêtre contienne deux photons. Dans le cas idéal, il faudrait travailler avec la plus faible valeur possible, mais, comme le taux de comptage S_{Ge}^{net} dépend aussi de μ_p , il nous est impossible de diminuer μ_p à volonté en raison des coups sombres dans le détecteur signal qui font alors tendre la probabilité d'avoir un photon vers zéro.
- γ_i représente le deuxième paramètre important de la source, tant il permet d'améliorer proportionnellement la probabilité d'avoir un photon unique P_1 .
- γ_s intervient aussi sur la valeur de $S_{Ge}^{net} = \gamma_s \mu_p \eta_{Ge}$. Son amélioration permet donc de diminuer μ_p tout en gardant un taux de comptage constant pour s'affranchir de l'effet néfaste des coups sombres.
- D_{Ge}^c correspond au taux de coups sombres dans le détecteur signal et constitue (avec les pertes) la principale origine des événements à zéro photon P_0 .
- ΔT constitue enfin le paramètre crucial concernant les performances de la source. En effet, la durée d'allumage du détecteur déclenché influe directement sur la probabilité P_2 qu'un photon supplémentaire vienne se glisser à l'intérieur de la fenêtre de détection. Les performances rencontrées durant ce manuscrit seront donc toutes données pour une durée ΔT donnée.

1.3.8 Choix des paramètres de fonctionnement

En fonction de l'utilisation qui va être faite de la source, nous allons favoriser la probabilité d'avoir un photon (P_1) ou celle de ne pas en avoir deux ($g^{(2)}(0)$). À titre d'exemple, intéressons nous à la distribution quantique de clé et utilisons la figure 1.15 issue de l'annexe A.

Nous voyons qu'une source présentant une forte probabilité P_1 permet d'échanger une clé avec un très haut débit, par contre pour échanger une clé sur une plus longue distance il faut impérativement diminuer la quantité d'impulsions contenant deux photons, c'est-à-dire $g^{(2)}(0)$.

Nous avons vu précédemment que, dans notre cas, à cause de l'existence de coups sombres dans notre détecteur signal, nous ne pouvons pas utiliser la source avec simultanément le P_1 maximum et le $g^{(2)}(0)$ minimum. Il nous faut trouver un compromis. Si nous souhaitons transmettre une clé sur une longue distance, il sera donc préférable de diminuer le taux moyen de paires créées par seconde μ_p à condition d'accepter une diminution de P_1 , donc du

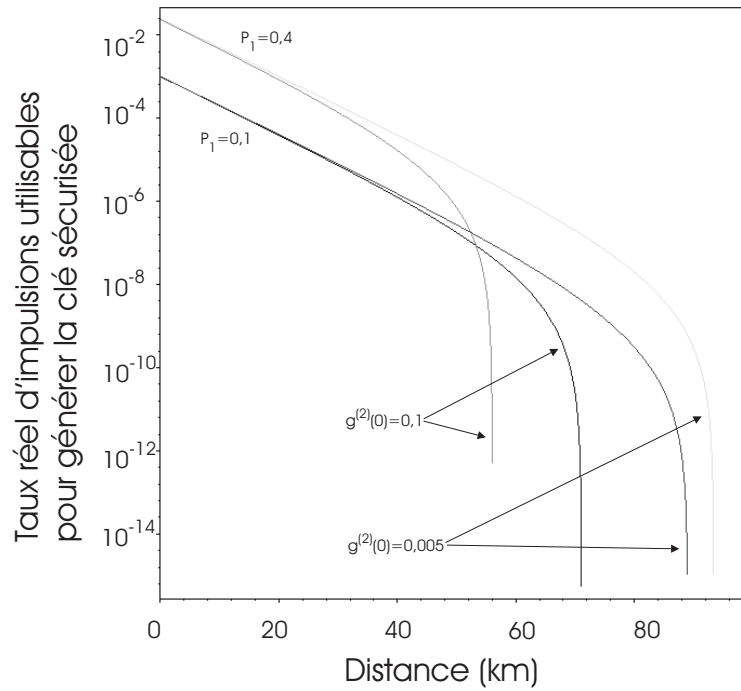


FIG. 1.15 – Taux de bits utilisables par impulsion en fonction de la distance pour des sources de photons uniques présentant un P_1 de 0,1 ou 0,4 et un $g^{(2)}(0)$ de 0,1 ou 0,005.

débit. Par contre, pour un échange de clé sur courte distance, où le débit sera primordial, nous pouvons nous autoriser un $g^{(2)}(0)$ qui ne soit pas à son minimum et chercher à obtenir un P_1 qui soit le plus élevé possible ($P_1 = \gamma_i$).

1.4 Conclusion du chapitre 1

Ce chapitre nous a permis de montrer que la conversion paramétrique était un excellent moyen de produire de paires de photons et que le principe d'utiliser un des photons de la paire pour prévenir de l'existence de l'autre permettait de réaliser une source de photons uniques annoncés. En effet, la détection du photon signal permet d'obtenir un signal électrique, dit *trigger*, pour déclencher le détecteur qui servira à mesurer ce photon ainsi *annoncé*. Contrairement au régime impulsionnel où les paires sont uniquement créées à des temps bien précis, nous avons choisi de travailler en régime continu dans lequel l'instant de création n'est pas défini. Nous avons qualifié ce régime d'asynchrone par identification au protocole actuel de télécommunica-

tion *ATM* pour *Asynchronous Transfert Mode*, en vigueur dans les réseaux actuels. Ce mode de fonctionnement pour la source est original à plusieurs points de vue en comparaison au régime pulsé :

- les paires sont séparées les unes des autres grâce à une détection déclenchée. C'est donc en partie la durée d'ouverture du détecteurs de utilisés à la sortie de la source qui va établir ses performances,
- la puissance de pompe du laser influe sur le taux de répétition de la source et simultanément sur la probabilité P_2 d'avoir deux photons dans une fenêtre de détection ΔT ,

Nous avons étudié les facteurs qui influencent les performances de la source, comme le coefficient de collection des paires créées au sein du cristal, le taux de coups sombres dans notre détecteur de photons signal ou encore le taux de paires créées dans le cristal. Cette étude théorique nous a permis de prévoir les performances que nous pouvons attendre de la source et de déterminer les facteurs auxquels il faudra prêter une grande attention afin d'optimiser la source.

Chapitre 2

La technologie de l'optique guidée

Dans le chapitre précédent, nous avons introduit les principes fondamentaux de la conversion paramétrique ainsi que les caractéristiques générales d'une source de photons annoncés, mais dans un but pédagogique les démonstrations reposaient sur le cas idéal d'une source de paires de photons par interaction paramétrique¹. Nous entendons par là que la condition d'accord de phase était toujours remplie et que les photons créés étaient monochromatiques. Il existe pourtant différents moyens de réaliser l'accord de phase. Nous verrons que le Quasi-Accord de Phase (QAP) introduit par Armstrong [57] offre une très large plage d'accordabilité en longueur d'onde et qu'il s'associe parfaitement à la configuration guidée. Dans un premier temps, nous décrirons la source de paires de photons réalisée à partir de guides d'ondes sur *Niobate de Lithium Périodiquement Polarisé* (PPLN) fabriqués au laboratoire et nous verrons les caractéristiques particulières de ces paires. Nous déterminerons la largeur spectrale et la localisation spatiales des paires au sein de ces microguides. En lui-même, le guide PPLN n'est qu'un générateur de paires de photons et dans la deuxième partie nous décrirons et caractériserons les composants optiques fibrés issus de la technologie des télécommunications que nous lui associerons pour en faire une source de photons uniques annoncés. Ce sera finalement à la fin de ce chapitre que nous quantifierons la compatibilité des deux technologies, à savoir le guide PPLN et les composants fibrés en mesurant les pertes totales que subissent les paires jusqu'à la sortie de la source.

¹De nombreuses considérations technologiques avaient d'ailleurs incité le lecteur à se reporter directement à ce chapitre, nous allons donc tâcher de répondre à ses interrogations.

2.1 Des guides d'ondes sur PPLN

Dans cette section, nous allons nous attacher à décrire la source de paires de photons que nous avons réalisée au LPMC. Cette source est le fruit de deux concepts :

- *L'optique intégrée* : cette technique permet de confiner et de guider la lumière à l'intérieur du matériau. Dans le cas de la génération paramétrique de paires de photons, cette méthode permet de confiner l'interaction sur la longueur totale de l'échantillon afin d'obtenir une plus grande efficacité et de faciliter la récolte des paires de photons.
- *Le Quasi-Accord de Phase (QAP)* : cette technique au cœur des interactions paramétriques permet d'obtenir une interaction efficace, grâce à un accord de phase artificiel. De cette manière, les longueurs d'ondes des photons de la paire peuvent être aisément choisies et l'utilisation du QAP où les ondes de pompe, signal et idler, sont co-linéaires s'associe parfaitement à la configuration guidée.

Volontairement nous n'aborderons que très peu la fabrication des guides², ce qui nous permettra plutôt de décrire les généralités de la propagation de la lumière dans les guides et le principe du QAP en configuration guidée. Nous finirons cette section par la caractérisation précise notre interaction paramétrique.

2.1.1 Les guides d'ondes SPE

Pour faciliter la récolte des paires de photons par une fibre optique mono-mode et aussi augmenter l'efficacité de l'interaction paramétrique, il est judicieux de confiner la propagation des ondes à l'intérieur d'un guide. Il s'agit alors de réaliser une structure guidante en forme de canal qui respectera le coefficient non-linéaire du $LiNbO_3$. Ce dernier est un cristal biréfringent et possède donc un indice extraordinaire et un indice ordinaire notés n_e et n_o correspondant à l'indice « vu » par l'onde incidente, si sa polarisation est parallèle ou perpendiculaire à l'axe optique Z du cristal. La méthode employée au laboratoire pour réaliser ce guide canal est *l'échange protonique* qui permet d'accroître l'indice selon l'axe extraordinaire et de le diminuer selon l'axe ordinaire. De ce fait, seule la lumière polarisée selon l'axe extraordinaire (mode TM) sera guidée. Les caractéristiques opto-géométriques auxquels nous nous attendons sont les suivantes :

²Une description détaillée du protocole de fabrication des guides SPE est fourni en annexe D. De même nous invitons le lecteur à se référer aussi à cette annexe pour une description détaillée du protocole expérimental de fabrication du PPLN.

- une largeur dont la valeur à $1/e$ est comprise entre 4 et 8 μm et le profil d'indice est de type gaussien,
- une profondeur dont la valeur à $1/e$ vaut typiquement 2,1 μm et dont le profil d'indice est exponentiel,
- enfin, un accroissement d'indice δn aux alentours de 0,024.

Mais attention, nos guides ne sont pas circulaires, ce sont en réalité des « demi-cylindres » délimités par la surface du cristal (voir figure 2.1). De ce fait, nous verrons lors de la caractérisation de l'interaction paramétrique en mode guidé que les modes sont asymétriques et que bien qu'ils soient tous des modes fondamentaux, leur distribution spatiale n'est pas de type gaussien à symétrie cylindrique comme dans une fibre optique.

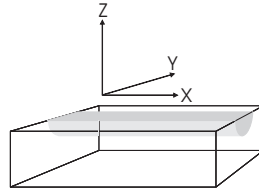


FIG. 2.1 – Guide d'onde intégré sur niobate de lithium en coupe Z .

Pour l'instant penchons nous sur le moyen d'obtenir l'accord de phase dans un guide d'onde. Dans le chapitre 1, nous avons admis que *la conservation de l'énergie et l'accord de phase* étaient satisfaits :

$$\hbar\omega_p = \hbar\omega_s + \hbar\omega_i \quad (2.1)$$

$$\vec{k}_p = \vec{k}_s + \vec{k}_i \quad (2.2)$$

La deuxième condition traduisait le fait que chacune des ondes se propageait à la même vitesse que la polarisation non-linéaire qui lui sert de source et ainsi évitait l'apparition d'interférences destructives. Ici nous allons nous pencher sur la manière d'obtenir l'accord de phase dans un guide d'onde avec un couple signal et idler satisfaisant la conservation de l'énergie. À partir de maintenant nous travaillerons dans une configuration co-linéaire pour les ondes pompe, signal et idler et nous projeterons l'équation de l'accord de phase sur l'axe de propagation. Réécrivons différemment l'équation 2.2 pour bien comprendre la façon d'obtenir l'accord de phase en configuration guidée :

$$n_p(\omega_p)\omega_p = n_s(\omega_s)\omega_s + n_i(\omega_i)\omega_i \quad (2.3)$$

qui peut alors se représenter schématiquement ainsi :

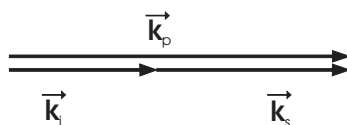


FIG. 2.2 – Accord de phase co-linéaire dans un guide d'onde.

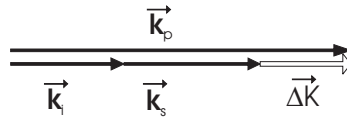
En pratique, l'indice de réfraction du matériau dépend toujours de la longueur d'onde. Alors, dans ces conditions, obtenir l'accord de phase n'est pas trivial. Il faut, dans ce cas, utiliser des astuces pour obtenir l'accord de phase, comme par exemple chercher des matériaux dont l'indice de réfraction aurait un comportement « particulier ». Il se trouve que les matériaux biréfringents ont un indice ordinaire légèrement différent de l'indice extraordinaire et nous pouvons nous en servir pour trouver des indices égaux à des longueurs d'ondes différentes : *c'est l'accord de phase par biréfringence*.

Malheureusement nos guides ne supportent que les modes TM et nous ne pourrions donc pas utiliser la biréfringence dans ce cas. Mais « à chaque chose, malheur est bon », nous aurons la possibilité de tirer parti du coefficient non-linéaire d_{33} qui couple les ondes de pompe, signal et idler polarisés selon Z et qui s'avère être le plus élevé du matériau.

2.1.2 Le Quasi Accord de Phase

Pour contourner l'utilisation de la biréfringence qui n'autorise qu'une plage restreinte de longueurs d'ondes de pompe pour lesquelles les couples signal et idler étaient parfaitement définis, Armstrong et ses collaborateurs ont suggéré l'idée de garder une configuration de « désaccord de phase », mais de régulièrement compenser le déphasage accumulé entre les ondes et la polarisation qui leur sert de source. Nous verrons que le Quasi accord de phase permet de compenser le déphasage entre n'importe quel couple signal et idler, vérifiant la conservation de l'énergie, en ne faisant varier qu'un seul paramètre : *le pas d'inversion du signe du coefficient non linéaire* du matériau. Pour bien comprendre ce qu'il se passe considérons que l'accord de phase n'est pas atteint. Notons ΔK ce désaccord, que nous représentons mathématiquement et schématiquement par :

$$\vec{k}_p = \vec{k}_s + \vec{k}_i + \overrightarrow{\Delta K} \quad (2.4)$$



Sans justification pour l'instant, admettons que, dans ce cas, il apparaît un terme $e^{j\Delta K \times L}$, où L est la longueur de l'échantillon qui témoigne du déphasage accumulé entre les ondes de pompe, signal et idler dans l'équation qui régit l'amplitude des ondes. Il peut être intéressant de définir $\ell_c = \frac{\pi}{\Delta K}$ la *longueur de cohérence* de sorte à faire ressortir la nature oscillante du facteur de phase. Regardons sur la figure 2.3 comment varie la puissance de l'onde signal en fonction de la longueur du cristal pour un désaccord de phase ΔK :

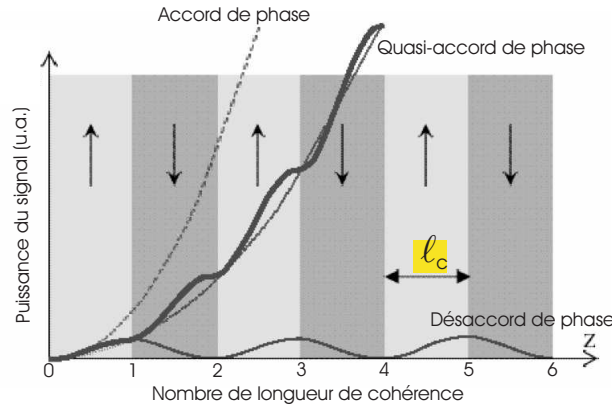


FIG. 2.3 – Puissance de l'onde signal dans le cristal en fonction de la longueur dans 3 cas : accord de phase – quasi accord de phase – désaccord de phase.

- À l'accord de phase, l'onde signal interfère constructivement sur toute la longueur du cristal. Alors, assez logiquement, la puissance du signal augmente quadratiquement.
- En cas de désaccord de phase, nous constatons que, sur une longueur ℓ_c , la puissance transférée de la pompe vers le signal augmente (interférences constructives), puis le déphasage accumulé devient tel que l'onde et la source commencent à interférer destructivement pour s'annuler complètement au bout d'une longueur $2\ell_c$.
- Le Quasi accord de phase consiste à ajouter au bout d'une longueur ℓ_c une phase de π à la polarisation pour la remettre en phase avec l'onde qu'elle a créée et ainsi continuer à interférer constructivement jusqu'à $2\ell_c$.

Sur le schéma 2.3, les flèches orientées périodiquement vers le haut et le bas représentent le coefficient non-linéaire $\chi^{(2)}$ du cristal, car effectivement renverser le signe du $\chi^{(2)}$ reviendra à adjoindre, dans les équations qui régissent l'amplitude des ondes pompe, signal et idler, un facteur π à ΔK puisque $e^{j\pi} = -1$. Plongeons nous dans quelques équations pour justifier ceci en détail.

Principe

Pour plus de clarté, nous allons utiliser le système de coordonnées défini par les axes cristallins du $LiNbO_3$, ci-dessous sur la figure 2.4. Dans les guides d'ondes, seuls les modes TM sont guidés et de ce fait nous prendrons uniquement en compte des champs électriques polarisés selon l'axe z .

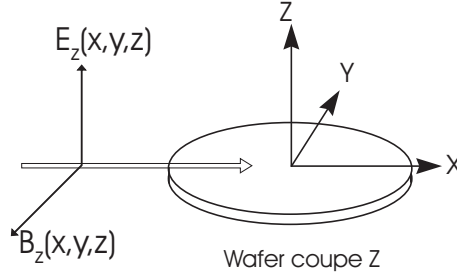


FIG. 2.4 – Définition des axes cristallins du Niobate de Lithium et de la polarisation du champ électrique pompe, signal et idler incidents.

Lorsque les champs optiques sont importants (de l'ordre du champ interatomique, soit 10^{11} V.m^{-1}), la polarisation $\vec{P}$ du milieu devient une fonction non-linéaire du champ électrique [57] :

$$\vec{P}(\vec{E}) = \epsilon_0[\chi^{(1)}]\vec{E} + \epsilon_0[\chi^{(2)}]\vec{E}\vec{E} + \dots = \epsilon_0\chi^{(1)}\vec{E} + \vec{P}^{NL}(\vec{E}) \quad (2.5)$$

où $[\chi^{(2)}]$ est le tenseur susceptibilité d'ordre deux.

Prenons comme point de départ la très classique équation de propagation relative au champ électrique que l'on appelle encore *équation de Helmholtz*³ :

$$\Delta\vec{E} - \vec{\text{grad}}\left(\text{div}(\vec{E})\right) - \frac{\epsilon_r}{c^2}\frac{\partial^2\vec{E}}{\partial t^2} = \mu_0\frac{\partial^2\vec{P}^{NL}}{\partial t^2} \quad (2.6)$$

³Nous ne justifierons pas son écriture ici et le lecteur intéressé pourra trouver une démonstration approfondie dans les références [67, 68], ainsi qu'un rappel complet des hypothèses considérées pour la suite.

Les indices p , i et s seront maintenant réservés pour référencer les ondes de *pompe*, *signal* et *idler* et nous allons chercher des solutions à l'équation de propagation 2.6 sous la forme :

$$E_z(\omega_n) = \sum_{\text{modes } m} A_{n,m}(x) \mathbb{E}_{n,m}(y, z) e^{j(\omega_n t - k_{n,m} x)} \quad (2.7)$$

Cette écriture générale correspond à un développement du champ sur la base propre des modes guidés notés m . Ici, nous allons choisir de nous concentrer sur le cas où seul le mode fondamental est présent $m = 0$ à l'intérieur du guide pour la pompe, le signal et l'idler (nous omettrons par la suite d'indiquer $m=0$ pour alléger l'écriture des équations). En émettant quelques hypothèses sur les champs considérés, comme par exemple celle de l'enveloppe lentement variable, nous aboutissons au système d'équations suivant :

$$\begin{aligned} \frac{dA_p(x)}{dx} &= -\alpha_p A_p(x) - j\chi_{zzz}\omega_p I_p A_s(x) A_i(x) e^{j\Delta K x} \\ \frac{dA_s(x)}{dx} &= -\alpha_s A_s(x) - j\chi_{zzz}\omega_s I_s A_p(x) A_i^*(x) e^{j\Delta K x} \\ \frac{dA_i(x)}{dx} &= -\alpha_i A_i(x) - j\chi_{zzz}\omega_i I_i A_p(x) A_s^*(x) e^{j\Delta K x} \end{aligned} \quad (2.8)$$

pour lequel :

- χ_{zzz} est la susceptibilité du matériau qui lie les composantes z des champs entre elles. χ_{zzz} et le coefficient d_{33} (le plus élevé du tenseur associé au $LiNbO_3$) sont liés par la relation $\chi_{zzz}=2d_{33}$.
- Les α_k représentent les coefficients d'atténuation pour l'amplitude $A_k(x)$.
- $I_k = \frac{1}{n_k^{guide} \epsilon_0 c} \frac{\int \int \mathbb{E}_p(y, z) \mathbb{E}_s(y, z) \mathbb{E}_i(y, z) dx dy}{\int \int \mathbb{E}_k^2(y, z) dx dy}$ représente l'intégrale de recouvrement entre les modes guidés.
- Et enfin, $\Delta K = 2\pi \left(n_p^{guide} \omega_p - n_s^{guide} \omega_s - n_i^{guide} \omega_i \right)$ témoigne du désaccord de phase entre les ondes dans le guide⁴.

Nous retrouvons dans notre système le terme $e^{j\Delta K \times x}$, introduit au début du chapitre, qui témoigne du déphasage entre les ondes pompe, signal et idler accumulé au bout d'une longueur d'interaction x .

Dans l'hypothèse du QAP, considérons par exemple le cas simple d'une inversion périodique du signe du coefficient non-linéaire où les domaines ont une forme rectangulaire. Une représentation en est donnée sur la figure 2.5.

⁴Soulignons que la notation n_k^{guide} indique que les ondes « voient » un indice effectif qui n'est plus celui du matériau brut mais un indice qui dépend dorénavant des paramètres opto-géométrique du guide, du mode dans lequel elles se trouvent et bien sûr de leur longueur d'onde. Par la suite nous omettrons l'indice *guide* pour alléger les notations.

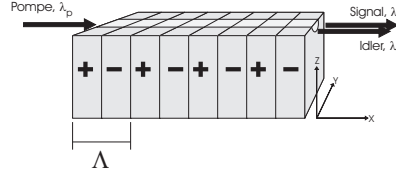


FIG. 2.5 – Guide d'onde intégré sur substrat PPLN. L'inversion du signe apparaît sur toute l'épaisseur du cristal et les domaines sont des parallélépipèdes rectangles de taille identique ($a = \frac{1}{2}$).

Le coefficient χ_{zzz} peut se voir alors comme une fonction périodique de pas Λ et de facteur de forme a , correspondant au rapport de la largeur des domaines retournés sur la période d'inversion. Alors $\chi_{zzz}(x)$ peut se décomposer en série de Fourier :

$$\chi_{zzz}(x) = \sum_{n \in \mathbb{Z}} \chi_n e^{-j2\pi \frac{n}{\Lambda} x} \quad (2.9)$$

avec

$$\begin{aligned} \chi_n &= \frac{2d_{33}}{n\pi} \sin(n\pi a) \\ \chi_0 &= 0 \end{aligned}$$

Dans notre cas, a vaut $1/2$ ce qui permet de maximiser la valeur du coefficient non-linéaire effectif que nous utiliserons. Le quasi accord de phase va consister à choisir correctement Λ pour compenser le désaccord de phase ΔK :

$$\boxed{\frac{2\pi n}{\Lambda} = \Delta K} \quad (2.10)$$

Dans ce cas, d'un point de vue mathématique, nous voyons que les termes $e^{j\Delta K x}$ et $e^{-j2\pi \frac{n}{\Lambda} x}$ vont se regrouper et s'annuler, conduisant à la disparition des termes oscillants du système 2.8. Dans la décomposition en série de Fourier, n va correspondre à l'ordre utilisé pour remettre en phase les ondes. Grâce à l'équation 2.10 nous voyons que compenser un très grand désaccord de phase va nécessiter un très petit pas d'inversion. D'un point de vue technologique, c'est ici que se trouve la limitation du QAP : la taille minimale du pas d'inversion du $\chi^{(2)}$ réalisable. Pour compenser cela, on peut utiliser les ordres supérieurs ... au détriment de l'efficacité car nous n'utiliserons alors qu'un coefficient non-linéaire effectif χ_n qui correspond au d_{33} diminué d'un facteur $\frac{2}{n\pi}$. Il existera donc un facteur $(\frac{2}{n\pi})^2$ de réduction sur la puissance

obtenue à l'ordre n . Notons que ce facteur de réduction provient du fait que les puissances des champs sont proportionnelles au carré de leurs modules. Pour notre source, nous n'irons pas jusqu'à l'ordre n , le premier nous suffira largement⁵, mais cette remarque nous permet de voir que l'évolution de notre puissance ne bénéficiera pas du même taux de variation que celui donné, sur la figure 2.3, pour l'accord de phase parfait à un facteur $(\frac{2}{\pi})^2$.

Expérimentalement, les domaines d'inversion étant beaucoup plus larges que les micro-guides, les coefficients χ_n ne dépendent pas de y et z , ce qui nous permet de réécrire simplement le système 2.8 :

$$\begin{aligned}\frac{dA_p(x)}{dx} &= -\alpha_p A_p(x) - j\chi_1 \omega_p I_p A_s(x) A_i(x) \\ \frac{dA_s(x)}{dx} &= -\alpha_s A_s(x) - j\chi_1 \omega_s I_s A_p(x) A_i^*(x) \\ \frac{dA_i(x)}{dx} &= -\alpha_i A_i(x) - j\chi_1 \omega_i I_i A_p(x) A_s^*(x)\end{aligned}\quad (2.11)$$

À partir de maintenant, les fréquences des ondes de pompe, signal et idler seront régies par la conservation de l'énergie et une équation de conservation de l'impulsion quelque peu modifiée par l'adjonction d'un vecteur de type réseau $\frac{2\pi c}{\Lambda} \cdot \vec{u}_x$:

$$\begin{aligned}\hbar\omega_p &= \hbar\omega_s + \hbar\omega_i \\ n_p^{guide} \cdot \omega_p &= n_s^{guide} \cdot \omega_s + n_i^{guide} \cdot \omega_i + \frac{2\pi c}{\Lambda}\end{aligned}\quad (2.12)$$

$$\vec{k}_i + \vec{k}_s = \vec{k}_p + \frac{2\pi c}{\Lambda} \cdot \vec{u}_x$$

Ainsi pour une température T de fonctionnement et un pas d'inversion Λ donnés, l'ensemble des triplets $\{\omega_p, \omega_s, \omega_i\}$ solutions des équations 2.12 forment ce que l'on appelle les courbes de QAP que nous tracerons expérimentalement et numériquement dans le prochain paragraphe. Ces courbes sont d'une importance capitale car elles permettent de trouver le point de fonctionnement expérimental de notre source.

⁵Il est intéressant de constater qu'à l'ordre 1, le pas d'inversion vaut exactement 2 fois la longueur de cohérence ℓ_c .

La largeur à mi-hauteur de la fluorescence

Comment est-il possible d'avoir une largeur spectrale différente de zéro sachant qu'il n'y a qu'un seul triplet de longueur d'onde qui remplisse les conditions 2.12 ?

Partons d'un triplet pompe, signal et idler dont le désaccord de phase ΔK_0 est compensé par une inversion périodique du coefficient non-linéaire $\frac{2\pi c}{\Lambda_0}$:

$$\begin{aligned}\omega_p &= \omega_s + \omega_i \\ n_p \cdot \omega_p &= n_s \cdot \omega_s + n_i \cdot \omega_i + \Delta K_0\end{aligned}\quad (2.13)$$

Imaginons, à présent, un nouveau triplet $\{\omega_p, \omega_s + \delta\omega, \omega_i - \delta\omega\}$ qui remplisse encore la conservation de l'énergie et imaginons que $\delta\omega$ soit petit devant ω_s que le désaccord de phase associé à ce triplet vaille maintenant $\Delta K = \Delta K_0 + \delta K$, avec δK petit devant ΔK_0 . Dans ce cas le système 2.13 devient :

$$\begin{aligned}\omega_p &= (\omega_s + \delta\omega) + (\omega_i - \delta\omega) \\ n_p \cdot \omega_p &\approx n_s \cdot \omega_s + n_i \cdot \omega_i + \Delta K_0 + \underbrace{\delta\omega \cdot (n_s - n_i)}_{\delta K}\end{aligned}\quad \begin{aligned}(2.14) \\ (2.15)\end{aligned}$$

Alors l'inversion périodique Λ_0 précédente compensera exactement ΔK_0 mais plus le terme δK supplémentaire. Nous nous retrouvons dans le cas du désaccord de phase rencontré au début de la section 2.1.2 et si nous cherchons à connaître la longueur de cohérence de ce terme ($\ell_c = \frac{\pi}{\delta K}$), nous trouverons des longueurs plus grandes que l'échantillon.

Que faut-il comprendre à ce résultat ?

Simplement que ce triplet, qui va entraîner un élargissement spectral proportionnel à $\delta\omega$, interférera constructivement tout le long de l'échantillon et sera ressorti du guide avant même de commencer à interférer destructivement. Continuons ce petit raisonnement pour sentir que plus l'échantillon sera long, plus le spectre sera étroit. En effet, dans un échantillon long, seuls les très faibles élargissements spectraux possèdent des longueurs de cohérence suffisamment grandes pour interférer tout le long. Ce comportement a été représenté sur le schéma 2.6, mais il ne faut y voir aucune étude quantitative. C'est pourquoi nous nous proposons maintenant de nous plonger dans quelques équations qui vont nous permettre de justifier plus proprement cet élargissement et de lui attribuer une valeur.

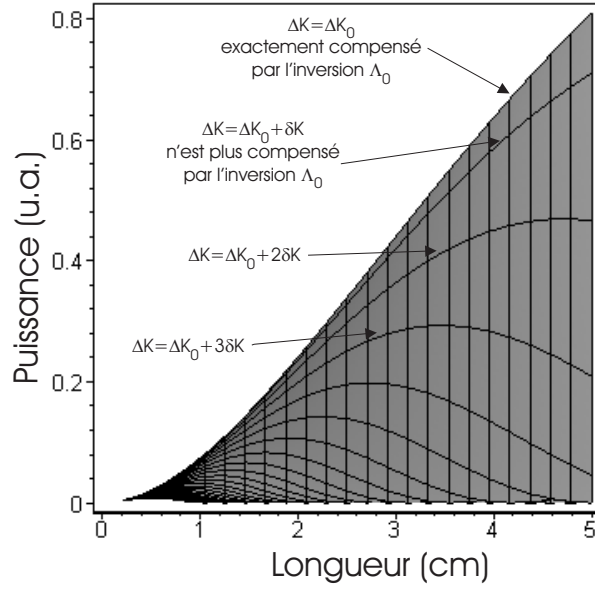


FIG. 2.6 – Évolution de la puissance transférée au couple signal/idler en fonction de la longueur de l'échantillon. La courbe supérieure correspond à l'accord de phase quand $\Delta K - \frac{2\pi n}{\Lambda_0} = 0$, tandis que les suivantes correspondent à un désaccord de phase $\Delta K - \frac{2\pi n}{\Lambda_0} = \delta K$ croissant par petits incréments δK .

Dans le cas de la fluorescence paramétrique, faisons l'hypothèse raisonnable que les amplitudes des ondes signal et idler sont petites devant celles de l'onde de pompe. De ce fait, il y a non-déplétion de la pompe par conversion paramétrique et son amplitude s'écrit donc simplement à partir de l'équation 2.11 :

$$A_p(x) = A_p(0)e^{-\alpha_p x} \quad (2.16)$$

En définissant⁶ alors

$$g^2 = \frac{2}{(\epsilon_0 c)^3} \cdot \omega_s \omega_i \cdot \frac{\chi_1^2 I_s I_i}{\int \int \mathbb{E}_p^2(y, z) dx dy} \cdot P_p(0) \quad (2.17)$$

$$\gamma^2 = \frac{\omega_i I_i A_p^*(0)}{\omega_s I_s A_p(0)} \quad (2.18)$$

nous pouvons réécrire le système 2.8 d'une manière simplifiée pour un désaccord de phase qui n'est plus totalement compensé par l'inversion périodique

⁶Il faut noter que la puissance est définie par $P_p = \frac{n_p^{guide} c}{2} \int \int |A_p(x)|^2 \mathbb{E}_p^2(y, z) dy dz$

Λ_0 du $\chi^{(2)} \Rightarrow \Delta K - \frac{2\pi c}{\Lambda_0} = \delta K$:

$$\begin{aligned} \frac{dA_s(x)}{dx} &= -\alpha_s A_s(x) - jg \frac{A_i^*(x)}{\gamma} e^{j\delta K x} \\ \frac{d\left(\frac{A_i^*(x)}{\gamma}\right)}{dx} &= -\alpha_i \frac{A_i^*(x)}{\gamma} - jg A_s^*(x) e^{j\delta K x} \end{aligned} \quad (2.19)$$

Assez simplement, viennent les solutions suivantes :

$$\begin{aligned} A_s(x) &= e^{-\alpha_s x} e^{-j\delta K x} \left[A_s(0) \left(\cosh(bx) + j \frac{\delta K}{b} \sinh(bx) \right) - jg \frac{A_i^*(0)}{b\gamma} \sinh(bx) \right] \\ A_i^*(x) &= e^{-\alpha_i x} e^{-j\delta K x} \left[A_i^*(0) \left(\cosh(bx) - j \frac{\delta K}{b} \sinh(bx) \right) + jg \frac{\gamma A_s(0)}{b} \sinh(bx) \right] \end{aligned}$$

où $b = \sqrt{g^2 - \left(\frac{\delta K}{2}\right)^2}$.

Rappelons-nous que, dans le cas de la fluorescence paramétrique, la puissance signal croît uniquement à partir des fluctuations de l'idler ; nous avons alors :

$$A_s(0) = 0$$

et

$$A_s(x) = -je^{-\alpha_s x} e^{-j\delta K x} \frac{g A_i^*(0)}{b\gamma} \sinh(bx)$$

Il en découle la densité de puissance suivante dans le cas où $g \ll \delta K^7$ pour une longueur d'interaction l :

$$P_s(l) = e^{-2\alpha_s l} g^2 x^2 \frac{\omega_s}{\omega_i} dP_i(0) \text{sinc}^2 \left(\sqrt{\left(\frac{\delta K}{2}\right)^2 - g^2} l \right) \quad (2.20)$$

Cette expression montre que les pics de fluorescence attendus auront la forme du *carré d'un sinus cardinal* dont la largeur à mi-hauteur sera donnée par :

$$\sqrt{\left(\frac{\delta K}{2}\right)^2 - g^2} l = \pm \frac{\pi}{2} \quad \Longleftrightarrow \quad \delta K = 4\sqrt{g^2 + \frac{\pi^2}{4l^2}} \quad (2.21)$$

Nous savons donc à ce stade que l'interaction paramétrique accepte un désaccord de phase δK déterminé par l'équation 2.21. Essayons de le relier à la

⁷ce qui est expérimentalement le cas, étant donnée la faible valeur expérimentale de $g \approx 10^{-2} - 10^{-1} \text{ cm}^{-1}$.

largeur de ligne $\delta\omega$ et pour cela, développons ΔK au voisinage du quasi accord de phase, que nous identifierons à $\Delta K_0 + \delta K$. En considérant la pompe monochromatique⁸, nous pouvons écrire :

$$\Delta K = k_{p_0} - \left(k_{s_0} + \frac{\partial k_s}{\partial \omega_s} \Big|_{\omega_{s_0}} d\omega_s \right) - \left(k_{i_0} + \frac{\partial k_i}{\partial \omega_i} \Big|_{\omega_{i_0}} d\omega_i \right) = \Delta K_0 + \delta K$$

Par identification des termes, nous obtenons la relation :

$$\delta K = 4\sqrt{g^2 + \frac{\pi^2}{4l^2}} = \frac{\partial k_s}{\partial \omega} \Big|_{\omega_{s_0}} d\omega_s - \frac{\partial k_i}{\partial \omega} \Big|_{\omega_{i_0}} d\omega_i \quad (2.22)$$

En dérivant maintenant l'équation de conservation de l'énergie, nous pouvons noter que $|d\omega_i| = |d\omega_s|$ ce qui nous permet de calculer :

$$|d\omega_i| = |d\omega_s| = \frac{4c\sqrt{g^2 + \frac{\pi^2}{4l^2}}}{\left| n^{guide}(\omega_{s_0}) - n^{guide}(\omega_{i_0}) - \frac{\partial n^{guide}}{\partial \omega} \Big|_{\omega_{s_0}} \omega_{s_0} + \frac{\partial n^{guide}}{\partial \omega} \Big|_{\omega_{i_0}} \omega_{i_0} \right|} \quad (2.23)$$

La formule peut se simplifier pour les faibles gains tels que $g \ll \frac{\pi}{2l}$. Dans notre cas, l'approximation est justifiée par les valeurs typiques de $g \approx 10^{-2} - 10^{-1} \text{ cm}^{-1}$ et $\frac{\pi}{2l} \approx 0, 1-1 \text{ cm}^{-1}$. La largeur à mi-hauteur s'écrit finalement :

$$\boxed{\delta\omega_s = \delta\omega_i = \frac{2\pi c}{\mathcal{N}l}} \quad (2.24)$$

avec $\mathcal{N} = \left| n^{guide}(\omega_{s_0}) - n^{guide}(\omega_{i_0}) - \frac{\partial n^{guide}}{\partial \omega} \Big|_{\omega_{s_0}} \omega_{s_0} + \frac{\partial n^{guide}}{\partial \omega} \Big|_{\omega_{i_0}} \omega_{i_0} \right|$

La forme de l'équation 2.24 nous conforte dans l'explication que nous avons essayé de donner *avec les mains* en début de paragraphe. Nous constatons bien une dépendance de la largeur du spectre en fonction de l'inverse de la longueur de l'échantillon l . Nous allons donc voir maintenant comment caractériser expérimentalement la fluorescence paramétrique et confronter les résultats avec les prévisions théoriques que nous venons de démontrer.

2.1.3 Caractérisation expérimentale de la fluorescence paramétrique

Dans l'introduction, nous avons résolument choisi de tourner notre source vers les longueurs d'onde des télécommunications optiques. En effet, le but

⁸Notre laser de pompe possède une largeur de raie inférieure à $0,05 \text{ nm}$, dont l'effet est négligeable dans l'équation 2.22.

est d'envoyer le photon unique via des fibres monomodes standards présentant de faibles pertes à la propagation et une faible dispersion dans l'infrarouge. L'intérêt d'utiliser ce type de fibres associées à des photons à 1550 nm est de pouvoir très aisément insérer la source sur les réseaux télécoms déjà existants, si, dans un futur proche, nous désirons réaliser un protocole de communication quantique longue distance en dehors du laboratoire. L'autre raison qui nous a poussé à choisir la *seconde fenêtre des télécommunications* est l'existence de photodiodes à avalanche déclenchées⁹, entièrement fibrées, compactes et « correctement efficaces » pour compter les photons. Ces détecteurs sont fabriqués à partir de semi-conducteurs en *InGaAs* et sont vendus par la société *Id-Quantique*.

⇒ La longueur d'onde de nos photons idler qui deviendront par la suite les photons uniques annoncés, sera donc centrée autour de 1550 nm .

L'idée était de pouvoir simplement récolter les paires de photons provenant du guide d'onde en utilisant la même fibre optique, puis d'utiliser des composants fibrés pour filtrer les photons de pompe résiduels et ensuite séparer les paires. Plusieurs raisons pratiques (et techniques...) nous ont poussé à choisir un photon signal dans la première fenêtre des télécommunications (1310 nm). La principale réside dans l'existence de coupleurs directionnels en longueur d'onde (couramment appelés démultiplexeurs d'où l'acronyme *WDM*) qui présentent d'excellentes performances pour la séparation $1310/1550\text{ nm}$, tout en restant des produits largement répandus, donc peu onéreux.

Notons aussi, nous disposions au laboratoire de photodiodes à avalanches en germanium qui peuvent être passivement déclenchées et sont capables de détecter des photons dont la longueur d'onde est centrée sur 1310 nm .

⇒ La longueur d'onde des photons signal qui serviront à annoncer les photons uniques sera centrée sur 1310 nm .

Caractérisation spectrale

La conversion de longueur d'onde utilisée sera donc du type 710 nm à la pompe $\Rightarrow 1310\text{ nm}$ au signal et 1550 nm à l'idler. Il y a encore une contrainte imposée par le Niobate de Lithium : pour éviter les phénomènes photo-réfractifs à l'intérieur du cristal, nous devons élever sa température aux alentours de 340 K [69, 67].

⁹Le concept de détecteur déclenché a été introduit dans le chapitre 1 pour nous permettre d'isoler les photons uniques des suivants.

Fort de ces paramètres, nous procédons généralement à quelques simulations numériques afin de définir une plage de période d'inversion confortable et susceptible de nous garantir les interactions recherchées et nous disposons, par ailleurs, de cinq largeurs de guide nous permettant d'augmenter les chances de succès. Nous invitons le lecteur à lire l'annexe D pour savoir comment le guide d'onde sur PPLN a été préparé car, dans la présente section, nous partirons du stade où l'échantillon a été fabriqué et se présente comme un produit fini avec comme principales caractéristiques :

- longueur totale de l'échantillon 11 mm,
- des guides d'onde monomodes dans l'infrarouge de 4, 5, 6, 7 et 8 μm de large,
- une inversion périodique du $\chi^{(2)}$ comprise entre 14,2 et 12,8 μm ,
- les faces d'entrée et de sortie du guide sont polies parallèlement et l'angle des faces avec la surface fait exactement 90°.

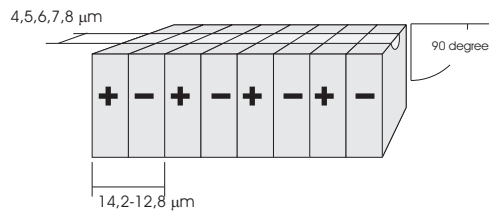


FIG. 2.7 – Échantillon fini avec un guide d'onde monomode parfaitement perpendiculaire aux domaines et aux faces d'entrée et de sortie

DISPOSITIF EXPÉRIMENTAL

Il est nécessaire de caractériser expérimentalement la fluorescence paramétrique pour déterminer les courbes de QAP du guide PPLN et mesurer la largeur de son spectre.

Pour cela, nous avons à notre disposition au laboratoire le protocole expérimental représenté sur la figure 2.8 dont les principaux éléments constitutifs sont :

Une diode laser de pompe *TOPTICA Photonics-DL 100* à cavité externe, monomode fréquentiel et longitudinal, accordable sur la plage 707-713 nm grâce à un réseau fermant la cavité. Ce laser est capable de délivrer quelques milliwatts.

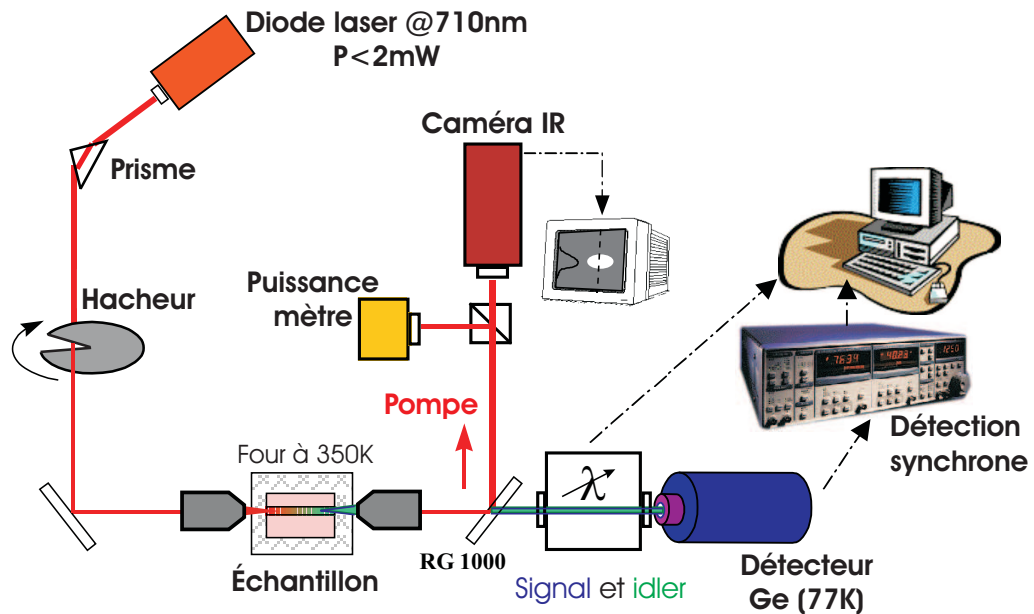


FIG. 2.8 – Montage expérimental utile à la caractérisation de la fluorescence paramétrique

Un **prisme de dispersion** qui permet d'éliminer l'infrarouge résiduel du faisceau de pompe qui pourrait perturber les mesures de fluorescence que nous attendons justement à ces longueurs d'onde.

Un **système de chauffage complet** comprenant un four en dural recouvert de bakélite et de mousse polyuréthane, taillé aux dimensions de l'échantillon. La température y est contrôlée à l'aide d'un *TEC-2000 – Thermoelectric temperature controller* de *Thorlabs* associé à une sonde de type *Analog-Device-590*.

Une **caméra de contrôle** du mode de pompe injecté dans le guide. Grâce à l'*Hamamatsu-C1000*, nous pouvons nous assurer que nous injectons bien la pompe dans le mode fondamental.

Un **système complet d'analyse** du signal de fluorescence constitué d'un monochromateur (*SPEC-270M*) dont le réseau est monté sur un moteur pas à pas piloté par ordinateur, d'un détecteur Germanium¹⁰ refroidi à l'azote liquide (*North-Coast-EO-817L*) et enfin d'un dispositif

¹⁰ Attention, bien que son principe de fonctionnement soit très proche, ce détecteur n'est pas un compteur de photon, mais il a cependant l'aptitude à mesurer d'extrêmement faibles signaux dans l'infrarouge.

(*Stanford-SR 830*) permettant de synchroniser la détection avec celle du hacheur.

Un filtre *RG-1000* utilisé pour stopper le reste du faisceau de pompe qui ressort du guide. C'est aussi par son intermédiaire que nous observons les modes de pompes injectés grâce à l'onde qui s'y réfléchit.

Deux objectifs de couplage, l'un traité pour le visible et utile à l'injection de la pompe et le second, traité pour l'infrarouge, utile à la récolte du signal de fluorescence.

Un banc de couplage *Elliot-Martock* qui représente le summum en déplacement de précision dans les 3 directions de l'espace.

Un spectromètre *Anritsu MS9701B* non représenté sur la figure nous indique la valeur de la longueur d'onde de pompe utilisée.

Le déroulement de l'expérience est relativement simple. Nous fixons la température à 340 K , la longueur d'onde du laser à 710 nm et nous nous lançons à la recherche d'un signal de fluorescence en « jouant » sur la largeur des guides et sur la période d'inversion du $\chi^{(2)}$.

Avec un peu de chance et beaucoup de patience, le détecteur confirme ce que les simulations avaient prédites... des puissances signal et idler sont enfin détectées, il ne nous reste plus qu'à faire des acquisitions de spectres pour différentes longueurs d'onde de pompe. Ceci nous permettra de tracer par la suite les courbes de QAP.

Nous venons de trouver l'interaction tant désirée dans un guide large de $8\text{ }\mu\text{m}$ et pour une période d'inversion de $13,6\text{ }\mu\text{m}$. Nous allons vous présenter les courbes de fluorescence et les fameuses courbes de QAP qui en sont extraites.

La confrontation des courbes de QAP expérimentales avec celles simulées numériquement nous permet de trouver les paramètres opto-géométriques exacts du guide. Nous entendons par là que, bien que le protocole de fabrication des guides *SPE* soit assez fiable au niveau de la reproductibilité, à partir du moment où deux échantillons n'ont pas été fabriqués ensemble, au vu des nombreuses étapes de fabrication, des différences minimales peuvent apparaître entre eux. Si la profondeur et la largeur du guide sont des paramètres très bien connus (une étude approfondie de la reproductibilité de ces paramètres est d'ailleurs disponible dans la référence [69]), l'accroissement d'indice δn est le paramètre le plus incertain. Une fois ce paramètre en main, nous calculerons la largeur spectrale attendue et nous mesurerons la largeur à mi-hauteur du spectre de fluorescence. La confrontation des deux résultats nous fournira l'ultime recoupement validant la théorie développée précédemment.

DESCRIPTION GÉNÉRALE DES COURBES

La figure 2.9 ci-dessous représente les spectres expérimentaux obtenus pour des longueurs d'onde de pompe différentes.

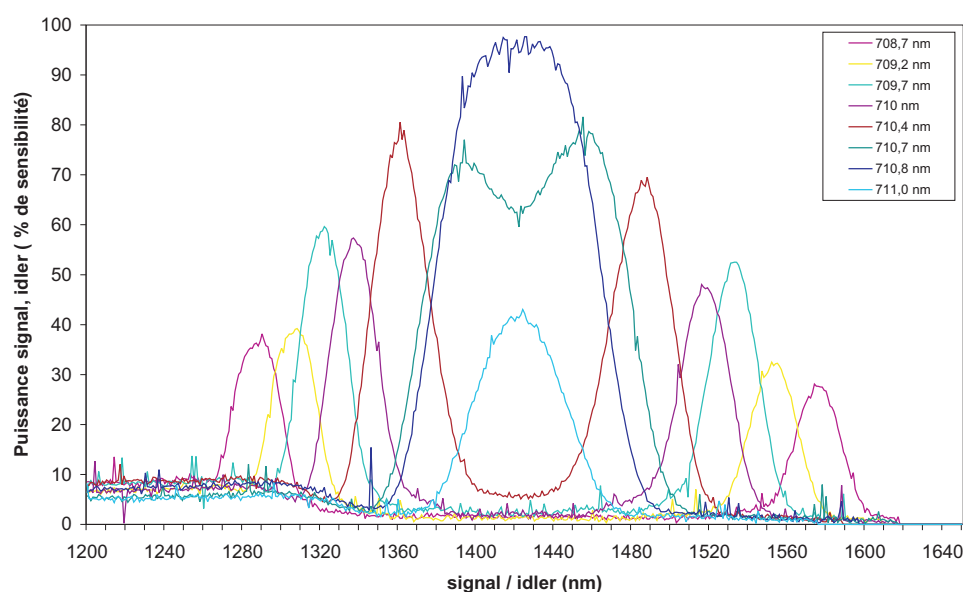


FIG. 2.9 – Spectres de fluorescence expérimental relatif au guide de $8\ \mu\text{m}$ avec un pas d'inversion de $13,6\ \mu\text{m}$. Les courbes correspondent à des longueurs d'ondes de pompe différentes.

Notons que les puissances de signal et d'idler sont exprimées en % de la sensibilité s correspondant au calibre choisi sur le système de détection synchrone. Nous observons donc, à λ_p donnée, deux pics de fluorescence qui correspondent à une configuration de quasi-accord de phase la plus efficace où les trois ondes en interaction sont guidées. Le recouvrement intermodal est alors important et associé à un transfert d'énergie du mode de la pompe vers les modes signal et idler. D'une courbe à l'autre, la longueur d'onde de pompe augmente et les pics correspondants au signal (longueur d'onde la plus faible par convention) et à l'idler se rapprochent. Aussi, mises à part quelques courbes qui dérogent à la règle à cause d'une puissance de pompe injectée que nous n'avons pas parfaitement contrôlée, nous pouvons constater que l'interaction est de plus en plus efficace au fur et à mesure que les pics signal et idler se rapprochent du point de dégénérescence pour lequel $2\lambda_p = \lambda_s = \lambda_i$, où ils « fusionnent ». Ce comportement s'explique par le meilleur recouvrement des modes du signal et de l'idler à l'intérieur du guide au fur et à

mesure que leur longueur d'onde se rapproche. Nous verrons dans la partie « caractérisation spatiale » qui suivra, cette notion de recouvrement *imparfait* pour des modes à des longueurs d'ondes différentes. Il faut cependant noter que le plus petit spectre, pour $\lambda_p=711\text{ nm}$, n'obéit pas à cette loi. Sa faible puissance est justifiée par le fait qu'une fois le point de dégénérescence dépassé, si nous continuons à augmenter la longueur d'onde de pompe, l'efficacité de l'interaction s'effondre¹¹. La raison pour laquelle l'efficacité ne chute pas brutalement une fois le point de dégénérescence passé est la même que pour l'élargissement spectral des pics de fluorescence : l'interaction reste quand même efficace pour des désaccords de phase suffisamment petits. De plus, nous remarquons que la valeur trouvée pour le point de dégénérescence $\lambda_s = \lambda_i$ ne correspond pas exactement au double de la valeur de la pompe $\lambda_p=710,8\text{ nm}$, le décalage provient du fait que nous utilisons deux appareils différents pour analyser les longueurs d'ondes de pompe et de fluorescence.

Cette mesure nous permet ensuite de tracer sur la figure 2.10, la courbe de QAP associée à notre guide de $8\text{ }\mu\text{m}$ de large avec un pas d'inversion $\Lambda = 13,6\text{ }\mu\text{m}$.

Nous disposons au laboratoire d'un programme informatique développé dans le cadre d'études antérieures (se référer à la thèse [67]) traitant les interactions paramétriques à partir des équations d'évolution dans le cas d'un guide homogène et dans des conditions de température homogène. Il permet ainsi de modéliser l'interaction entre le cristal et les champs électromagnétiques de pompe, signal et idler via les équations de Sellmeier [71] et la méthode des indices effectifs [72] dans le cas de guides aux profils d'indice quelconques. Ce programme a plusieurs entrées :

- l'accroissement d'indice du guide δn ,
- la largeur et la profondeur du guide,
- le profil du guide (en largeur et en profondeur),
- la température de l'échantillon,

et nous permet d'identifier les paramètres opto-géométriques exacts du guide, en ajustant la courbe de QAP simulée à celle tracée à partir des mesures expérimentales.

¹¹Nous invitons le lecteur à regarder l'animation multimédia associée à la référence [70] à l'adresse <http://www.opticsexpress.org/abstract.cfm?URI=OPEX-10-19-990>

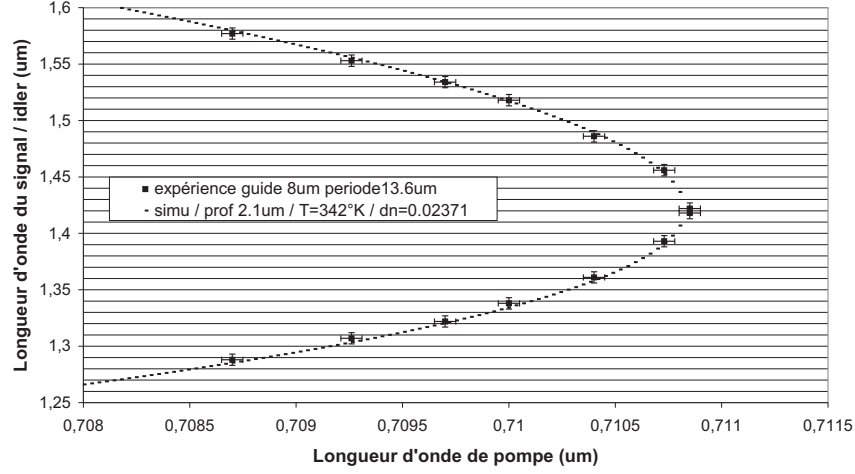


FIG. 2.10 – Courbe de quasi-accord de phase pour le guide de $8\ \mu m$ de large avec un pas d'inversion de $13,6\ \mu m$, réunissant tous les triplets $\{\lambda_p, \lambda_s, \lambda_i\}$ qui peuvent entrer en interaction dans le guide.

Nous visualisons sur le schéma 2.10 un parfait accord entre la courbe et les points expérimentaux, nous autorisant à dire que notre guide possède les caractéristiques suivantes :

Profil en largeur : gaussien avec ouverture à $1/e$ de $8\ \mu m$

Profil en profondeur : exponentiel avec une profondeur caractéristique à $1/e$ de $2,1\ \mu m$

Accroissement d'indice : $\delta n = 0,02371$

Température de fonctionnement : $342\ K$

Période du pas d'inversion : $\Lambda = 13,6\ \mu m$

LARGEUR SPECTRALE DE LA FLUORESCENCE

Maintenant que nous connaissons parfaitement les paramètres opto-géométriques de notre guide, nous pouvons les introduire dans la formule 2.24, pour calculer la largeur spectrale attendue et confronter le résultat aux mesures expérimentales. Pour cela, il est intéressant de plutôt utiliser la formule 2.24 sous la forme

$$\delta\lambda_{s,i} = \frac{\lambda_{s,i}^2}{Nl} \quad (2.25)$$

avec cette fois $\mathcal{N} = \left| n^{guide}(\lambda_{s_0}) - n^{guide}(\lambda_{i_0}) - \frac{\partial n^{guide}}{\partial \lambda} \Big|_{\lambda_{s_0}} \lambda_{s_0} + \frac{\partial n^{guide}}{\partial \lambda} \Big|_{\lambda_{i_0}} \lambda_{i_0} \right|$

La principale incertitude provient de la longueur de l'échantillon qui vaut $11 \pm 0,2 \text{ mm}$ et nous négligeons celle due au calcul des indices du guide, sachant qu'ils sont définis avec 8 chiffres après la virgule. À l'aide de la formule 2.25, nous estimons donc la largeur spectrale des ondes signal (1310 nm) et idler (1550 nm), générées à l'intérieur du guide, égales à :

$$\begin{aligned} \delta \lambda_s^{th} &= 14,6 \text{ nm} \pm 0,3 \text{ nm} \\ \delta \lambda_i^{th} &= 20,5 \text{ nm} \pm 0,4 \text{ nm} \end{aligned}$$

Il est finalement intéressant de mesurer expérimentalement cette largeur pour savoir si les paramètres qui ont servi à l'estimation de $\delta \lambda_{s,i}$ sont corrects.

QUELQUES CONSIDÉRATIONS SUR LA MESURE DES SPECTRES

En réalité, le spectre du signal mesuré $M(\lambda)$ correspond à la convolution de la résolution du système $R(\lambda)$ avec le signal réel $S(\lambda)$

$$M(\lambda) = S(\lambda) * R(\lambda) \quad (2.26)$$

Alors si on considère que le signal de fluorescence et la résolution du système peuvent être assimilés à des gaussiennes :

$$S(\lambda) = S_0 e^{-\frac{\lambda^2}{\Delta \lambda_s^2}} \quad (2.27)$$

$$R(\lambda) = e^{-\frac{\lambda^2}{\Delta \lambda_r^2}} \quad (2.28)$$

le spectre mesuré s'écrit comme suit :

$$M(\lambda) = \left(S_0 \frac{\sqrt{\pi} \Delta \lambda_r \Delta \lambda_s}{\sqrt{\Delta \lambda_r^2 + \Delta \lambda_s^2}} \right) e^{-\frac{\lambda^2}{\Delta \lambda_r^2 + \Delta \lambda_s^2}} \quad (2.29)$$

Ainsi, pour toutes mesures faites sur le spectre, la largeur apparente sera $\Delta \lambda_m = \sqrt{\Delta \lambda_r^2 + \Delta \lambda_s^2}$

Pour bien comprendre l'influence de la résolution du monochromateur, nous avons détaillé sur la figure 2.11 le spectre obtenu à une longueur d'onde de pompe donnée mais pour deux résolutions différentes. En effet, le monochromateur possède une résolution $\Delta \lambda_r$ ajustable et il apparaît de façon claire que, pour deux spectres identiques à l'origine, la figure enregistrée est

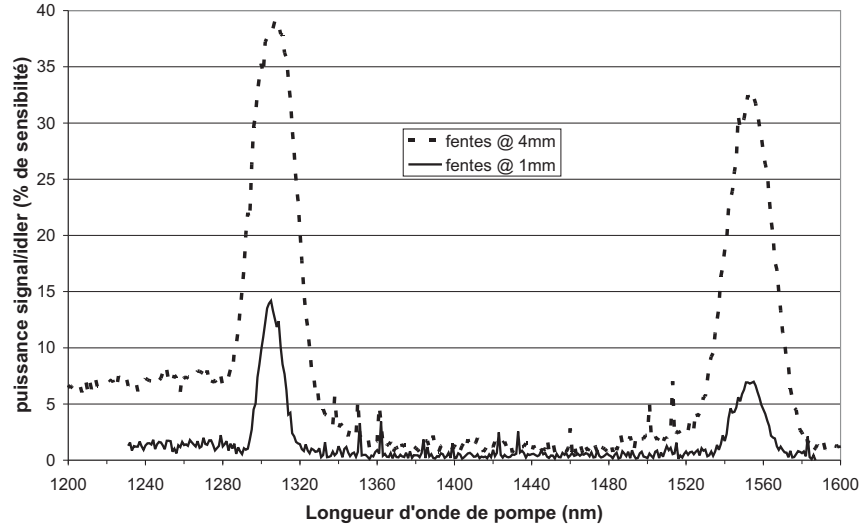


FIG. 2.11 – Spectres de fluorescence obtenus pour une interaction à longueur d'onde de pompe donnée, mais pour des résolutions différentes. Notons que la résolution est proportionnelle à l'ouverture des fentes du monochromateur d'analyse.

dissemblable aussi bien en intensité, qu'en largeur. Finalement, à partir de la figure 2.11 ci-dessus, nous mesurons dans les deux cas une largeur spectrale réelle de :

$$\begin{aligned}\delta\lambda_s^{exp} &= 14,8 \text{ nm} \pm 0,5 \text{ nm} \\ \delta\lambda_i^{exp} &= 20,1 \text{ nm} \pm 0,5 \text{ nm}\end{aligned}$$

Cette dernière mesure est tout à fait en rapport avec les estimations faites précédemment et nous permet de conclure sur la justesse des paramètres opto-géométriques que nous avons déduits de la courbe de QAP (figure 2.10).

Depuis le début de ce chapitre, nous avons occulté les termes de *photons de pompe* et de *paires de photons signal et idler* pour décrire les interactions paramétriques guidées comme un phénomène mettant en jeu des ondes. Ce choix est délibérément pédagogique et nous informons le lecteur qu'il est très facile de passer d'un point de vue à l'autre en associant¹² « *puissance* » de l'onde avec « *probabilité* » d'avoir un photon. À titre d'exemple, nous n'émettons pas des photons qui possèdent une largeur spectrale $\delta\lambda$, mais nous

¹²Attention cependant à la fragilité de cette association qui reste extrêmement naïve.

avons une probabilité donnée d'émettre un photon dont la longueur d'onde est comprise dans une plage $[\lambda - \frac{\delta\lambda}{2}, \lambda + \frac{\delta\lambda}{2}]$. Pour la suite, nous proposons au lecteur d'associer « *densité de puissance* » avec assez logiquement « *densité de probabilité* » d'avoir un photon s'il préfère penser *corpuscule* plutôt qu'*onde*.

À plusieurs reprises, nous avons laissé entendre que nos guides présentaient une asymétrie spatiale (en forme de demi-cylindre) et qu'il existait, en plus, une forte asymétrie au niveau des indices de réfraction puisque les guides sont échangés en surface du substrat. En effet, la différence d'indice au « dessus » du guide ($n_{air} - n_{guide}$) est beaucoup plus importante que celle en « dessous » ($n_{guide} - n_{substrat}$). Cette asymétrie nous a fait dire dans le paragraphe précédent que le recouvrement spatial entre les modes guidés n'était pas optimal pour des longueurs d'ondes différentes. Nous nous proposons d'expliquer ce phénomène un peu plus en détail et montrerons que récolter 100% des paires de photons émises ne sera malheureusement pas possible.

2.1.4 Caractérisation spatiale transverse des modes

Pour cela nous disposons au laboratoire d'un autre programme qui permet de simuler la distribution spatiale de la densité de puissance de n'importe quel mode, dans n'importe quel profil de microguide et ce à n'importe quelle longueur d'onde...

En utilisant les paramètres opto-géométriques obtenus précédemment et pour le mode fondamental à la pompe, au signal et enfin à l'idler, nous avons tracé simultanément sur la figure 2.13 les densités de puissances des trois modes à l'intérieur du microguide.

Nous avons réuni toutes les courbes en couleur de cette section sur les pages 91 et 92. *Attention les couleurs des modes n'ont aucune vraisemblance et n'ont été choisies que pour la clarté des schémas.* Le code des couleurs suivi est le suivant :

Rouge correspond au faisceau de pompe dont la longueur d'onde est 710 nm

Bleu correspond quant à lui au signal à 1310 nm

Vert est associé à l'idler à 1550 nm

Notons que sur les figures 2.13, 2.15 et 2.14, il s'agit des densités de puissance qui sont représentées ($\approx |E_Z|^2$) et c'est donc leurs intégrales sur une section transverse du guide qui sont égales.

Sur la figure 2.13, nous constatons aisément que les trois modes ont des formes très proches, mais ne se recouvrent effectivement pas totalement et possèdent tous des densités de puissance différentes.

Vue de coté $\Rightarrow$ figure 2.14

Sur cette figure, nous pouvons nous concentrer sur la distribution spatiale en fonction de la profondeur. Il se trouve que cette composante est asymétrique, c'est donc dans cette configuration que nous risquons d'observer les effets les plus gênants pour récolter simultanément les photons signal et idler.

Pour tenter d'expliquer ce phénomène, nous allons considérer un guide avec un profil en saut d'indice¹³. Cette approximation permet de simplifier les formules que nous allons utiliser, sans pour autant changer le phénomène qui nous intéresse ici : *le guidage de la lumière dans un guide asymétrique*. Aussi, il sera suffisant de considérer un milieu isotrope, non-conducteur et non-magnétique pour s'affranchir des phénomènes non-linéaires qui ne nous intéressent pas ici.

Nous ne détaillerons pas ce calcul qui part des équations de Maxwell pour ce type de milieu (développé dans la section 7 de la référence [73]), mais donnerons directement l'équation de propagation dont la forme finale est la suivante :

$$\Delta \vec{\mathcal{E}} - \nabla \left(\frac{1}{n^2} \nabla n^2 \cdot \vec{\mathcal{E}} \right) - \epsilon_0 \mu_0 n^2 \frac{\partial^2 \vec{\mathcal{E}}}{\partial t^2} = 0 \quad (2.30)$$

$$\Delta \vec{\mathcal{H}} + \frac{1}{n^2} \nabla n^2 \times (\nabla \times \vec{\mathcal{H}}) - \epsilon_0 \mu_0 n^2 \frac{\partial^2 \vec{\mathcal{H}}}{\partial t^2} = 0 \quad (2.31)$$

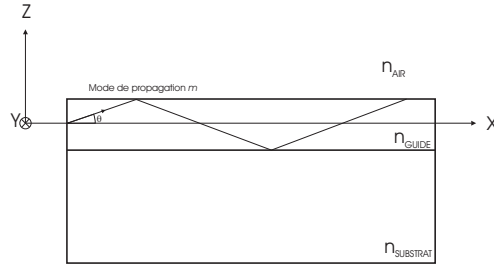


FIG. 2.12 – Schéma microguide asymétrique dans le plan vertical.

Plaçons-nous dans le cas d'une onde se propageant selon l'axe X comme représenté sur la figure 2.12. Si l'indice optique du milieu est une fonction qui ne dépend que des coordonnées transverses alors :

$$n^2 = n^2(y, z)$$

¹³Rappelons que notre guide présente, quant à lui, un profil en gradient d'indice qui peut être représenté comme un empilement de guides à saut d'indice et d'épaisseur dz .

et les solutions des l'équation 2.30 et 2.31 seront de la forme générale

$$\vec{\mathcal{E}} = \vec{E}(y, z)e^{i(\omega t - \beta x)} \quad \text{et} \quad \vec{\mathcal{H}} = \vec{H}(y, z)e^{i(\omega t - \beta x)} \quad (2.32)$$

avec $\beta = \frac{n_{mode}\omega}{c}$ la constante de propagation du mode guidé. n_{mode} représente ici l'indice optique que « voit » le champ guidé par la structure et qui est généralement appelé « indice effectif » et qui est différent¹⁴ de l'indice optique du guide n_{guide} . En introduisant ces solutions générales dans l'équation 2.31, nous trouvons de nouvelles équations pour le champ H_Y associées aux modes TM de chaque région {Air-Guide-Substrat}. Notons aussi que pour que le mode soit guidé par la structure, il faut que $n_{guide} > n_{mode} > n_{sub} > n_{air}$.

$$\frac{d^2 H_Y}{dz^2} + \left(\frac{\omega}{c}\right)^2 [n_{mode}^2 - n_{air}^2] H_Y = 0 \quad (2.33)$$

$$\frac{d^2 H_Y}{dz^2} + \left(\frac{\omega}{c}\right)^2 [n_{guide}^2 - n_{mode}^2] H_Y = 0 \quad (2.34)$$

$$\frac{d^2 H_Y}{dz^2} + \left(\frac{\omega}{c}\right)^2 [n_{mode}^2 - n_{sub}^2] H_Y = 0 \quad (2.35)$$

À l'aide des conditions limites aux deux interfaces du guide (continuité de H_Y et de $\frac{1}{n^2} \frac{dH_Y}{dz}$), nous trouvons un système de solutions pour chaque domaine Air-Guide-Substrat, ainsi que des valeurs discrètes de β qui correspondent aux m modes guidés par la structure. On a donc :

$$H_y(z) = \begin{cases} Ae^{-\frac{\omega}{c}\sqrt{n_{mode}^2 - n_{air}^2}z} & z \text{ dans l'air} \\ B \cos\left(\frac{\omega}{c}\sqrt{n_{guide}^2 - n_{mode}^2}z\right) & z \text{ dans le guide} \\ Ce^{-\frac{\omega}{c}\sqrt{n_{mode}^2 - n_{sub}^2}z} & z \text{ dans le substrat} \end{cases} \quad (2.36)$$

Évidemment, ces équations sont valables pour une longueur d'onde donnée ; il existe donc trois systèmes comme celui-ci correspondant au guidage de l'onde de pompe, signal et idler.

À chaque longueur d'onde sont associées des¹⁵ valeurs de n_{mode} qui lui sont propres et qui diminuent (donc se rapprochent de la valeur de n_{sub})

¹⁴En effet, un mode guidé se propageant dans le guide peut être considéré comme une superposition d'ondes planes se propageant dans le guide avec un angle $\pm\theta$. De par ses multiples réflexions sur les interfaces du guide, l'onde plane « voit » un *indice effectif* associé à son mode de propagation qui s'écrit : $n_{mode} = n_{guide} \cos(\theta)$.

¹⁵Une seule dans le cas de nos guides qui sont « monomodes » dans l'infrarouge.

au fur et à mesure que la longueur d'onde est grande : *c'est la dispersion géométrique du guide*. Ainsi l'équation correspondant à l'intérieur du guide, nous apprend que les ondes qui possèdent les plus grandes longueurs d'ondes, vont avoir tendance à s'élargir et occuper le plus de place possible jusqu'à ce qu'elles ne soient plus guidées ($n_{mode} = n_{sub}$) : c'est la longueur d'onde de coupure.

La forme des champs H_Y à l'extérieur des guides sont dits « *évanescents* » et nous constatons aisément que plus la longueur d'onde sera grande plus l'onde va s'étendre en dehors du guide. Ce phénomène est d'autant plus visible en profondeur que la différence $n_{mode} - n_{substrat}$ est petite et très peu visible en surface tant la différence $n_{mode} - n_{air}$ est grande en comparaison. Ainsi, l'onde idler, sera l'onde dont la puissance va le plus « s'étaler » dans le substrat.

Nous pouvons aussi prévoir qu'avec l'augmentation de la différence de longueur d'onde entre les trois ondes de pompe, signal et idler, les différences spatiales entre les modes soient d'autant plus prononcées.

Pour l'explication qui va suivre, il est intéressant d'associer « *densité de puissance* » avec « *densité de probabilité d'avoir un photon* ». En effet, nous pouvons interpréter la figure 2.14 comme suit : « lors de l'émission d'une paire de photons à l'intérieur du guide, la localisation spatiale des deux photons suit la probabilité décrite par les courbes bleue et verte ». Dans l'optique où nous cherchons à collecter les paires de photons grâce à une unique fibre optique, nous constatons qu'optimiser la récolte des photons idler à 1550 nm, en se plaçant au sommet de la courbe verte, ne concordera pas avec le maximum de probabilité de collecter le photon signal à 1310 nm. En résumé, l'optimisation de la récolte de l'un des deux photons se fera toujours au détriment de l'autre. Voyons par ailleurs dans cette discussion la justification théorique des hypothèses posées dans la section 1.3 du chapitre 1 : « *les valeurs des coefficients de récolte des photons signal et idler (γ_s et γ_i) sont liées entre elles car nous verrons au chapitre 2 que nous collectons les paires avec une seule et même fibre optique. Pour l'instant, nous admettrons en prévision des mesures expérimentales du chapitre 2 que le rapport $\frac{\gamma_s}{\gamma_i}$ vaut 0,65* ».

Vue de dessus $\Rightarrow$ figure 2.15

Sur cette figure, seule la composante en largeur nous est accessible. Le profil d'indice gaussien est le même à droite comme à gauche du guide et spatialement le guide est bordé des deux cotés par du Niobate. De par sa construction, *le guide est donc totalement symétrique*, nous en attendons de même des modes guidés.

La justification est sensiblement la même que précédemment à la diffé-

rence près que le guide est maintenant symétrique dans ce plan. Il suffit alors de de supprimer l'interface *guide-air* et de symétriser le raisonnement fait juste avant. Selon toute attente, la distribution des modes est symétrique et nous observons un comportement classique, c'est-à-dire dicté par la dispersion géométrique du guide : le mode à l'idler (vert-1550 nm) est celui qui occupe la plus grande largeur et présente la plus faible densité de puissance devant le signal (bleu-1310 nm) et enfin le mode de pompe (rouge-710 nm) qui possède quant à lui la plus grande densité de puissance.

RAPPELS DES FIGURES 2.13, 2.14 ET 2.15 EN COULEURS

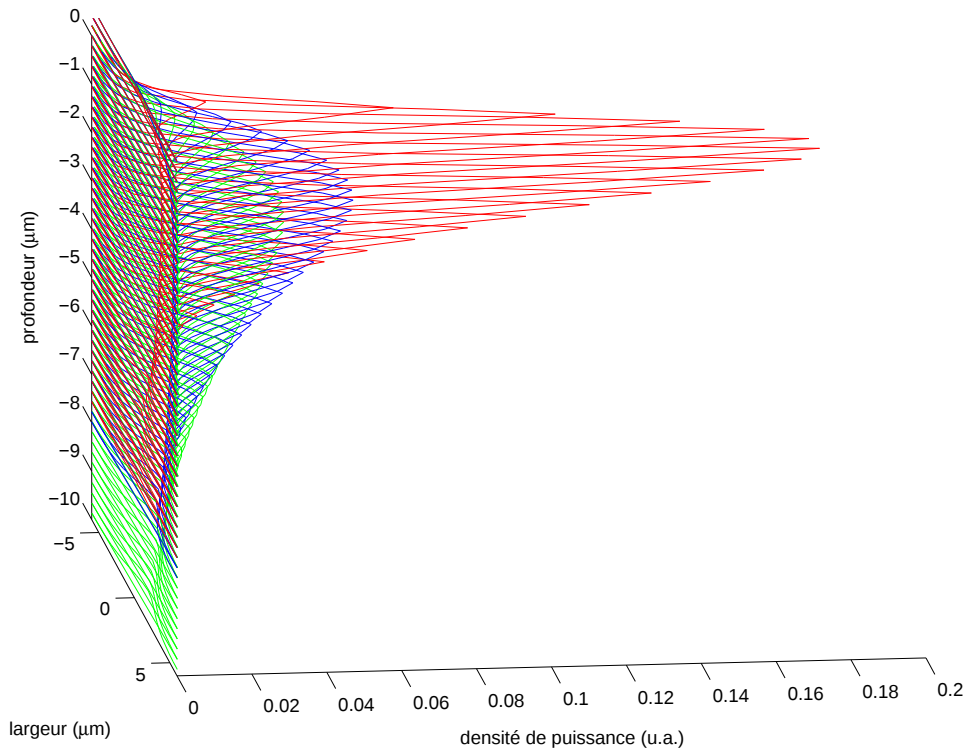


FIG. 2.13 – Représentation de la densité de puissance $3D$ associés aux 3 modes pompe (rouge), signal (bleu) et idler (vert) à l'intérieur de notre microguide *SPE* de $8\ \mu\text{m}$ de large et $2,1\ \mu\text{m}$ de profondeur à $1/e$.

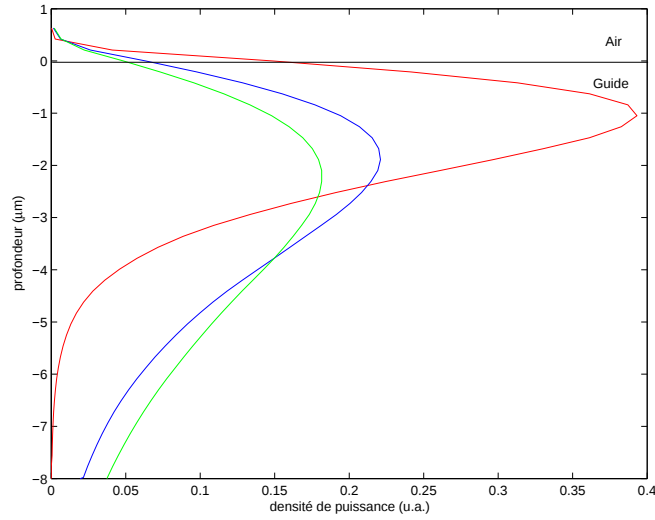


FIG. 2.14 – Représentation de la densité de puissance $2D$ associés aux 3 modes pompe (rouge), signal (bleu) et idler (vert) à l'intérieur de notre microguide *SPE* présentant une profondeur à $1/e$ de $2,1 \mu m$ de large en fonction de la profondeur.

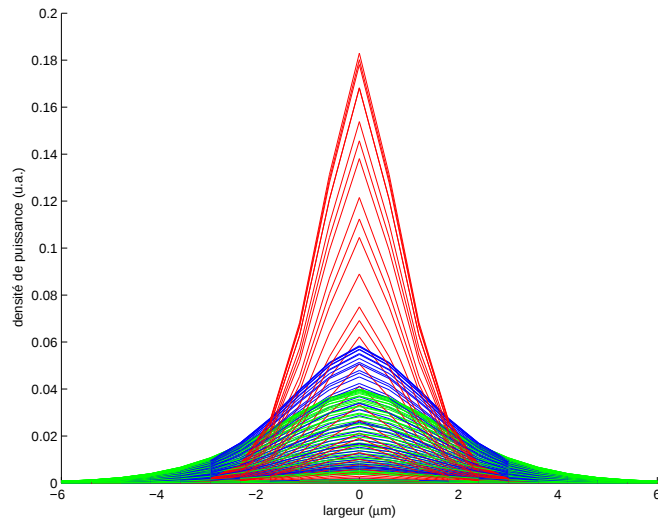


FIG. 2.15 – Représentation de la densité de puissance $2D$ associés aux 3 modes pompe, signal et idler à l'intérieur de notre microguide de surface en fonction de la largeur.

Jusqu'à maintenant, nous avons seulement parlé du guide PPLN qui est le résultat d'un travail technologique développé de *A* à *Z* au LPMC, dans lequel de nombreux paramètres interviennent et peuvent changer totalement le comportement de la source. Ainsi, dans un soucis de reproductibilité et de comparaison avec d'autres guides existants ou à venir, il nous a semblé nécessaire de s'attarder longuement sur la description de la source, de la largeur spectrale des photons émis et de leur localisation spatiale à la sortie du guide.

Cependant, il ne faut pas perdre de vue que le guide d'onde PPLN ne constitue en lui-même que la moitié du projet final. Il est donc temps de commencer la description des composants optiques qui vont intervenir pour récolter, filtrer et séparer les paires de photons, sans oublier l'électronique pour transformer et synchroniser le photon signal en impulsion électrique.

2.2 Les composants fibrés de la source : description et caractérisation

Nous avons représenté sur la figure 2.16 l'intérieur complet de la source de photons uniques annoncés.

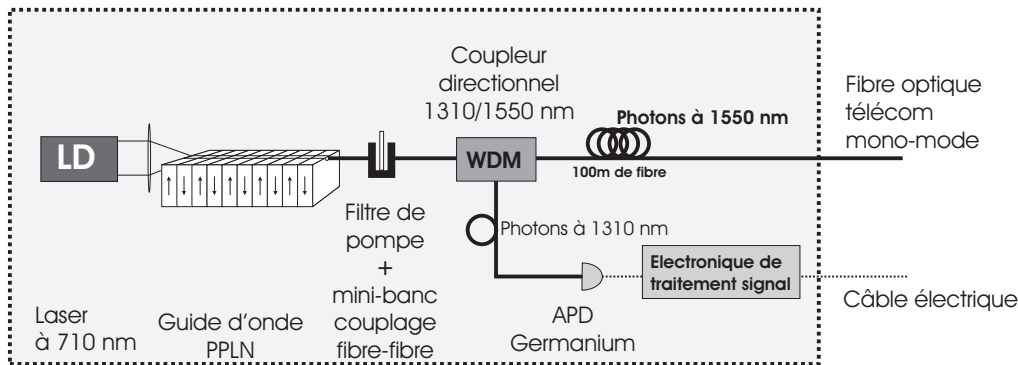


FIG. 2.16 – Schéma de l'intérieur de la source de photon uniques annoncés.

Les paires de photons, créées au sein du guide PPLN, sont récoltées à l'aide d'une fibre optique monomode télécom, accolée au guide, qui récolte par la même occasion une grande part des photons de pompe, qui n'ont pas été convertis en paires. Nous filtrons donc les photons pour ne garder que les paires de photons. Un coupleur directionnel 1310/1550 nm sépare les photons signal et idler pour envoyer les premiers vers le détecteur *Germanium* tandis que les seconds sont retardés dans 100 m de fibre optique. Le signal électrique

issu d'une détection servira à annoncer l'arrivée d'un photon idler à la sortie de la fibre optique. Notre source dispose donc de deux sorties : une électrique constituée par le détecteur de photons et une sortie optique correspondant à la fibre optique monomode.

Nous avons introduit dans le chapitre 1 le paramètre γ_i comme le coefficient de *couplage+pertes* des photons idler et nous avons montré qu'il établissait la limite supérieure de la probabilité P_1 . En pratique, c'est l'addition des pertes des composants fibrés et du coefficient réel de collection des photons du guide vers la fibre qui vont définir la valeur du γ_i expérimental. Il va donc être très important de choisir des composants présentant le moins de pertes possibles. C'est pour cela que nous essayerons aussi à chaque fois de quantifier les pertes des composants retenus.

Dans la section qui va suivre, nous allons entreprendre de présenter et décrire les éléments qui constituent notre source.

2.2.1 La fibre optique monomode

Le choix de la fibre optique constituant notre source a été fort simple car uniquement imposé par le standard des télécommunications dans les fibres optiques. Nous cherchons en effet une fibre monomode à 1310 nm et 1550 nm, présentant les pertes minimales à ces longueurs d'onde. Le modèle que nous utilisons est la fibre *SMF-28* pour *Single Mode Fiber*. Ses principales caractéristiques opto-géométriques sont les suivantes :

- un diamètre de cœur de 8,2 μm dans lequel la lumière sera guidée. Cette valeur n'est pas anodine car rappelons-nous que la taille caractéristique de notre guide est de 8 μm et en choisissant un guide d'une taille équivalente à la fibre, nous espérons optimiser la probabilité de récolter les photons du guide vers la fibre,
- le diamètre total de la fibre est de 100 μm ,
- l'indice de réfraction optique de la fibre vaut $\approx 1,461$,

Dans le tableau suivant, nous avons synthétisé les performances associées à cette fibre :

λ (nm)	Pertes maximales (dB/km)	Dispersion $\left(\frac{ps}{nm \cdot km}\right)$	Diamètre du mode (μm)
1310	$0,35 \pm 0,02$	0	$9,2 \pm 0,4$
1550	$0,20 \pm 0,02$	≤ 18	$10,4 \pm 0,5$

TAB. 2.1 – Caractéristiques générales de la fibre optique SMF-28.

Cette fibre servira aussi bien à connecter entre eux les différents éléments qui constituent la source, qu'à récolter les paires de photons à la sortie du guide. À ce titre, soulignons que nous récoltons les paires de photons directement à l'aide de la fibre optique qui viendra se placer en sortie du guide. L'avantage de cette méthode est de pouvoir accoler la fibre au guide. Dans notre montage expérimental, la fibre est simplement appuyée contre le guide, mais généralement la fibre est « *pigtailed* », c'est-à-dire collée au guide. Cette technique permet d'augmenter la stabilité du montage, mais surtout d'améliorer la collection des photons. En effet, grâce à l'utilisation d'une colle possédant un indice optique équivalent à celui de la fibre, nous allons permettre à cette dernière de littéralement épouser la face de sortie du guide et ainsi éliminer les éventuels « gaps » d'air entre le guide et la fibre.

Essayons d'estimer le gain du « pigtailing » en terme de récolte des photons.

À chaque fois qu'un faisceau lumineux passe d'un milieu d'indice optique n_1 vers un autre milieu d'indice n_2 , il subit des pertes à l'interface. Plus exactement, une partie du faisceau est transmise tandis que le reste est réfléchi à l'interface. La proportion de l'onde réfléchie est déterminée par le *coefficient de Fresnel* :

$$R = \frac{(n_1 - n_2)^2}{(n_1 + n_2)^2}$$

Pour les photons, on dit simplement qu'ils subissent des pertes R au changement d'interface ou plus généralement qu'ils ont une probabilité R d'être réfléchis.

Comparons les pertes aux interfaces dans les deux cas représentés sur la figure 2.17. Premièrement, la fibre est simplement appuyée contre le guide avec entre les deux un « gap » d'air et deuxièmement, la fibre est « *pigtailed* » au guide. Pour cela, nous prendrons les valeurs d'indices optiques suivantes :

$$n_{\text{guide}} = 2,2 \qquad n_{\text{fibre}} = 1,5 \qquad n_{\text{air}} = 1$$

Avec le tableau 2.2, nous constatons que l'utilisation d'une fibre collée au guide permet de réduire d'environ 15% les pertes associées à la récolte des photons.

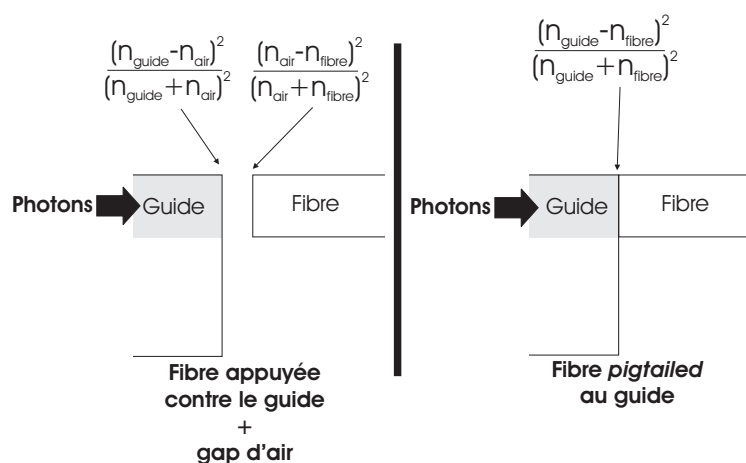


FIG. 2.17 – Synoptique des pertes lors du passage *guide–air–fibre* comparé au passage *guide–fibre*.

	pertes (dB)
Guide–air–fibre	0,84
Guide–fibre	0,16

TAB. 2.2 – Comparaison des pertes entre une fibre simplement appuyée contre le guide et une fibre littéralement collée sur le guide.

2.2.2 Les connecteurs optiques

S'il est un composant qui n'apparaît pas sur le schéma 2.16 tellement son utilisation est courante mais qui joue pourtant un rôle primordial au sein de la source : c'est le connecteur optique (voir figure 2.18). Les faibles pertes que rajoutent ces connecteurs sont compensées par la souplesse d'utilisation et la modulabilité qu'ils apportent en contrepartie. Ce dernier est précisément poli afin que les deux cœurs des fibres entrent en contact à l'intérieur du « *I optique* » garantissant ainsi des pertes inférieures à 0,3 dB. Tous les composants internes à la source devront donc être adaptés à la norme FC-PC imposée par les connecteurs.



FIG. 2.18 – Connecteur optique standard (FC-PC), couramment appelé *I-optique*.

2.2.3 Le U-bench + filtre passe-bas

À la sortie du guide et après avoir été collectées dans la fibre optique, les paires de photons sont encore accompagnées par un grand nombre de photons de pompe résiduels. Comme leur longueur d'onde est très différente de celle des photons signal et idler, nous pouvons assez simplement les filtrer à l'aide de filtres *passe-haut* en longueur d'onde qui ont à la fois des pertes très élevées dans le visible et très faibles dans l'infrarouge. Nous avons utilisé un filtre *DJ1160a* de chez *MTO* dont la courbe de transmission est représentée sur la figure 2.20.



FIG. 2.19 – Le U-bench OFR connectorisé.

Notons qu'il s'agit ici du seul point du montage où les photons voyageront à l'air libre pendant 3 cm !!! En effet, la société *OFR* commercialise des coupleurs fibre à fibre dont les pertes sont garanties inférieures à $0,5 \pm 0,2\text{ dB}$, qui permettent d'insérer des composants optiques massifs (tel le filtre *MTO*) sur un réseau fibré. Le coupleur réglable démontre une stabilité sur plusieurs mois, ce qui est d'autant plus impressionnant que le banc de couplage fibre à fibre est connectorisé à la norme FC-PC et peut accepter durant ce temps une quantité infinie de vissage-dévisage.

Le composant est optimisé pour garantir des pertes minimales à une longueur d'onde précise. Nous avons donc choisi d'optimiser le passage des photons idler à 1550 nm . Voici ce que nous avons obtenu :

$$\gamma_{\text{coupleur}}^{1550} = 0,47 \pm 0,01\text{ dB}$$

$$\gamma_{\text{coupleur}}^{1310} = 0,78 \pm 0,02\text{ dB}$$

$$\gamma_{\text{coupleur}}^{710} = 1,36 \pm 0,03\text{ dB}$$

Nous avons par ailleurs mesuré les pertes de l'ensemble *coupleur + filtre passe-bas* :

$$\gamma_{\text{coupleur+filtre}}^{1550} = 0,58 \pm 0,01\text{ dB}$$

$$\gamma_{\text{coupleur+filtre}}^{1310} = 1,66 \pm 0,02\text{ dB}$$

Enfin, les pertes de l'ensemble à 710 nm ont été estimées comme valant :

$$\gamma_{\text{coupleur+filtre}}^{710} \approx 50 \pm 2\text{ dB}$$

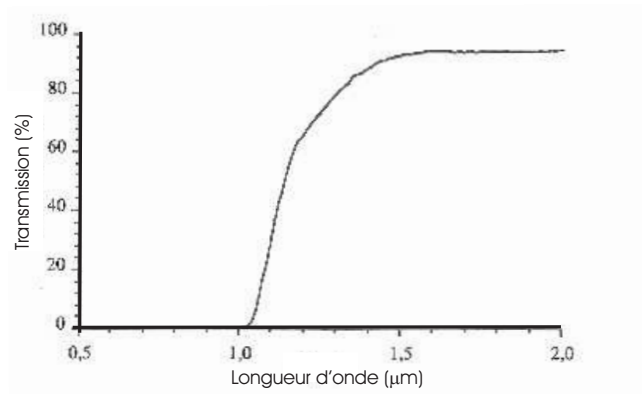


FIG. 2.20 – Caractéristique de la transmission du filtre *DJ1160a* de chez *MTO*.

2.2.4 Le WDM (wavelength division multiplexer) 1310/1550 nm

Ce qu'il faut retenir de ce produit, c'est :

- la grande maturité de la technologie utilisée,
- leur compacité et leur stabilité,
- leur faible coût de revient,
- les excellentes performances aux longueurs d'ondes télécom.

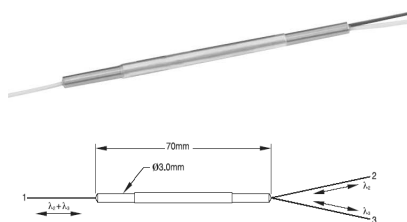


FIG. 2.21 – WDM 1310/1550 nm.

fenêtre pour λ	Pertes mesurées	Isolation
(nm)	(dB)	dB
1260-1350	$0,34 \pm 0,01$	≥ 17
1530-1600	$0,34 \pm 0,01$	≥ 17

2.2.5 Les composants d'optique guidée au complet

Au final, nous avons mesuré expérimentalement les pertes totales associées au montage suivant qui correspond exactement à ce qui sera ajouté au guide PPLN pour en faire une source de photons uniques annoncés :

- coupleur fibre à fibre avec le filtre de pompe inséré,
- coupleur directionnel 1310/1550 nm,
- 100 mètres de fibre optique monomode.

$$\gamma_{pertes\ totales}^{1550} = 1,12 \pm 0,02\ dB$$

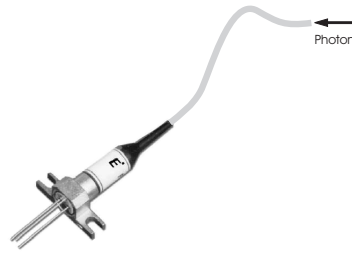
Cette valeur correspond bien à l'addition des pertes du banc de couplage fibre à fibre, du coupleur directionnel 1310/1550 nm, de 100 m de fibre et de quelques connecteurs optiques...

Pour leur part, les pertes totales à 1310 nm valent seulement la somme des pertes du coupleur et du coupleur fibre à fibre avec le filtre de pompe inséré, les photons signal ne passant pas dans la fibre de 100 m.

$$\gamma_{pertes\ totales}^{1310} = 2,0 \pm 0,04\ dB$$

2.2.6 La photodiode à avalanche en Germanium

Depuis quelques années, certains fabricants (comme *EG&G* ou *Id-quantique*) commercialisent des modules de détection « prêts à l'emploi », c'est à dire montés sur un circuit de polarisation adéquat et la plupart du temps refroidis par effet Peltier. C'est le cas par exemple des photodiodes en *silicium* ou en *InGaAs* qui ont fait l'objet de larges études grâce notamment

FIG. 2.22 – Photodiode à avalanche fibrée en *germanium*.

à leur utilisation au sein des réseaux télécom. Cependant, ces détecteurs disponibles sur le marché ne sont pas prévus pour réunir les deux conditions de fonctionnement qui nous intéressent, à savoir une détection ciblée sur la longueur d'onde de 1310 nm et l'utilisation d'une extinction passive¹⁶ (c'est-à-dire non déclenchée). En effet, si les APD-*silicium* acceptent le mode passif et sont d'excellentes photodiodes dans le visible (plus de 60% d'efficacité), elles deviennent complètement aveugles à partir de 1000 nm . Pour leur part, les APD-*InGaAs* sont capables de voir les photons à 1310 nm mais ne peuvent fonctionner en passif qu'au prix d'un taux de coups sombres très élevé. Ainsi, les seuls détecteurs susceptibles de remplir le cahier des charges sont faits de *germanium*. Nos détecteurs (figure 2.22) proviennent de chez *Fujitsu* et présentent l'avantage d'être entièrement fibrés et connectés.

Pour fonctionner en mode passif, ces APD-*Ge* doivent par contre être impérativement refroidies par un bain d'azote liquide à 77 K . De plus, elles présentent des efficacités plutôt faibles $\eta_{ge} \approx 10\%$ et un taux de coups sombres assez élevé D_{Ge}^c aux environs de 30 kHz pour une tension de biais donnée. À cette température, ces dispositifs présentent une longueur d'onde de coupure à 1450 nm rendant impossible leur utilisation pour la détection de photons à 1550 nm . Dans le cadre de notre source, nous avons préféré légèrement diminuer l'efficacité de l'APD dans le but de diminuer les coups sombres qui étaient trop élevés d'après les considérations du chapitre 1. Ainsi, nous avons fixé aux alentours de 20 kHz le taux de coups sombres, ce qui impose une efficacité de détection aux alentours de 6%.

Aussi, le lecteur intéressé pourra se référer à une publication de Cova et ses collaborateurs [74] qui traitent, en un seul papier très détaillé, l'ensemble des utilisations possibles des photodiodes à avalanches disponibles dans le commerce, ainsi que les divers régimes de fonctionnement associés.

¹⁶Cette expression vient du fait que leur extinction est gérée passivement par la charge d'un circuit RC , en opposition au mode *gated* où l'allumage comme l'extinction sont forcés par le trigger.

2.2.7 L'électronique de traitement du signal

Après détection, les signaux délivrés par l'APD germanium sont d'abord amplifiés par des dispositifs rapides : des amplificateurs *VT120C* de chez *EG&G-Ortec*. Les signaux traités ici sont analogiques et non homogènes les uns par rapport aux autres. Ensuite, il faut donc prévoir un dispositif de mise en forme des signaux afin d'obtenir des formes identifiables à la sortie de la source de photon unique : c'est maintenant le rôle des discriminateurs. Ici, nous utilisons un discriminateur *EG&G-Ortec* 4 entrées *Quad-935* qui fonctionne suivant un critère de seuil donnant lieu à un 1 logique si le signal d'entrée franchit le seuil ou à un 0 logique dans le cas contraire et ce jusqu'à 200 MHz.

Le standard utilisé pour le comptage des photons est désigné par le sigle *NIM* venant de l'anglais "Nuclear Instrumentation Module". Le *NIM* est un standard de logique négatif dont le 0 vaut 0 V et le 1 $-0,8$ V. L'écart entre les niveaux de tension étant faible, il est possible de basculer les niveaux logiques très facilement et donc très rapidement, le tout sous haute impédance.

Il ne faut pas oublier que durant le temps où le photon signal a été « transformé » en impulsion électrique, le photon idler (annoncé) a dû être retardé par 100 m de fibre optique. À la sortie du discriminateur, le photon idler est synchronisé avec le signal électrique, non pas en jouant sur la longueur de la fibre optique, mais appliquant un délai à l'impulsion électrique. Pour cela, nous avons choisi d'utiliser le *Digital Delay Generator/Pulse Generator DG-535* de chez *Stanford Research Systems*. Nous avons fait ce choix pour plusieurs raisons :

- Cet appareil permet de jouer beaucoup plus aisément sur le délai et avec beaucoup plus de reproductibilité que nous n'aurions pu le faire en jouant sur la longueur de la fibre optique.
- De plus, ce n'est pas une simple boîte à délais, mais un générateur de délais. La différence réside dans le simple fait que l'impulsion électrique n'est pas simplement retardée, mais mesurée puis détruite pour qu'ensuite une nouvelle impulsion électrique soit émise, après un délai T choisi par l'opérateur. Ce mode de fonctionnement apporte, en outre, la possibilité de choisir le standard électronique (*NIM*, *TTL*, *ECL* ...) de l'impulsion que recevra le détecteur déclenché.

Il est important de noter pour la suite que ce composant possède un temps mort de $1,3 \mu s$ après chaque conversion durant lequel il ne pourra traiter une éventuelle impulsion électrique en provenance du discriminateur. Au final, cela signifie que l'APD déclenchée ne recevra jamais deux triggers plus proches que $1,3 \mu s$. Nous verrons que cette valeur est très importante pour ce

type d'APD qui va être utilisée en mode déclenchée.

2.3 Détermination du coefficient de couplage guide-fibre

À ce stade du manuscrit, nous possédons deux technologies complémentaires : un guide PPLN et des composants fibrés monomodes au standard télécom mais le résultat de leur association est encore incertain. Nous venons de déterminer les pertes inhérentes aux composants fibrés que nous utilisons, mais le point crucial garant de la réussite du projet va être la connexion du guide PPLN à la fibre monomode SMF-28 et la valeur du coefficient de transfert des photons idler de l'un vers l'autre.

Notons que nous avons pris soin jusque-là de nommer γ_i le *coefficient de couplage+pertes*. Ce n'est pas anodin car, au final, c'est ce coefficient global qui va définir la probabilité maximale de récupérer un photon unique. En effet, que le photon ne soit pas couplé dans la fibre ou qu'il ait été perdu une fois dans la fibre, font qu'au final, l'utilisateur ne dispose pas de ce photon. Nous avons donc pris le parti de regrouper sous le terme γ_i tous les types de pertes qui diminuent les performances de la source. Toutefois ne les confondons pas, les pertes des composants étant constantes et mesurées auparavant, nous calculerons à la fin de cette partie le « vrai » coefficient de couplage *guide-fibre* noté Γ_i .

2.3.1 Le montage et les paramètres utilisés

Commençons par nous pencher sur le montage expérimental représenté sur la figure 2.23, grâce auquel nous allons mesurer γ_i .

Sur ce schéma expérimental, nous reconnaissons rapidement la source de photons uniques annoncés, pourtant il apparaît de nouveaux composants que nous n'avons pas encore présentés :

- la *photodiode à avalanche en InGaAs déclenchée (APD-InGaAs)*, fournie par le *GAP-Optique* et *Id-Quantique*. Ces APDs ont la possibilité de détecter des photons uniques à 1550 nm , au même titre que l'APD en *Germanium* détecte les photons à 1310 nm dans la source. Toutefois elles nécessitent d'être déclenchées pour réduire le nombre de « coups sombres » très élevés dans ce type de détecteur. En pratique, ces détecteurs seront donc la plupart du temps éteints et leur allumage, pendant une durée ΔT , sera conditionné par l'arrivée d'une impulsion électrique provenant de la détection d'un photon signal. Il est un point auquel il faut prêter une attention toute particulière : les *after-pulses*

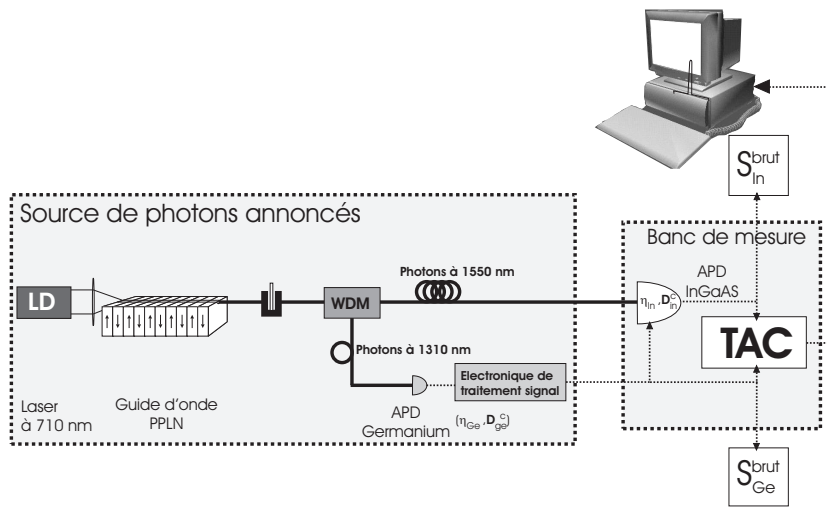


FIG. 2.23 – Expérience de comptage de paires de photons permettant d’optimiser puis de déterminer le coefficient de couplage+pertes γ_i . L’optimisation est réalisée à l’aide d’un banc de couplage (non représenté sur la figure) qui permet d’ajuster précisément la position de la fibre par rapport au guide PPLN.

ou *échos d’avalanches*. Juste après la détection d’un photon, la probabilité d’obtenir une seconde avalanche accidentelle lors de l’allumage suivant décroît exponentiellement en fonction du temps. Aussi, en cas d’allumage trop fréquent d’une APD-*InGaAs*, son nombre de détection sera biaisé par un taux anormalement élevé de coups qui ne sera pas quantifiable. Dans notre cas, nous avons vérifié que le temps mort du générateur de délai ($1,3\ \mu\text{s}$) était suffisant pour éviter ce problème.

- un *convertisseur temps-amplitude* de chez *EG&G-Ortec*, plus connu sous l’appellation anglaise de *Time to Amplitude Converter (TAC)*, que nous utilisons pour visualiser les coïncidences. Ce dernier dispositif joue en fait le rôle d’une porte *ET* capable d’identifier deux événements (bits de start et de stop) lui arrivant simultanément. Mieux qu’une simple porte *ET*, le *TAC* est capable de donner, via le couplage avec une carte d’acquisition et un *PC*, l’histogramme temporel des événements qui lui arrivent en coïncidence sur ses deux entrées.
- Nous tenons à souligner qu’il manque sur le dessin le banc de couplage *Elliot-Martock* permettant d’ajuster avec précision la position de la fibre par rapport au guide PPLN.

Grâce au *TAC* couplé à une carte d’acquisition, nous pouvons visualiser en

temps réel l'histogramme des coïncidences entre les impulsions électriques et les photons annoncés, représenté sur la figure 2.24. Nous pouvons distinguer un pic qui correspond aux photons qui arrivent tous précisément $T = 26 \text{ ns}$ après qu'il aient été annoncés. La largeur du pic fait 600 ps et est due au *jitter* dans les APDs qui correspond au fait que les avalanches au sein de la photodiode sont ponctuelles à plus ou moins 600 ps après l'arrivée d'un photon. Notons aussi que ce pic de détection des photons est entouré par une base d'une largeur de 10 ns qui correspond en réalité aux coups sombres pouvant (contrairement aux photons) survenir n'importe quand durant la durée d'allumage ΔT de l'APD-*InGaAs*.

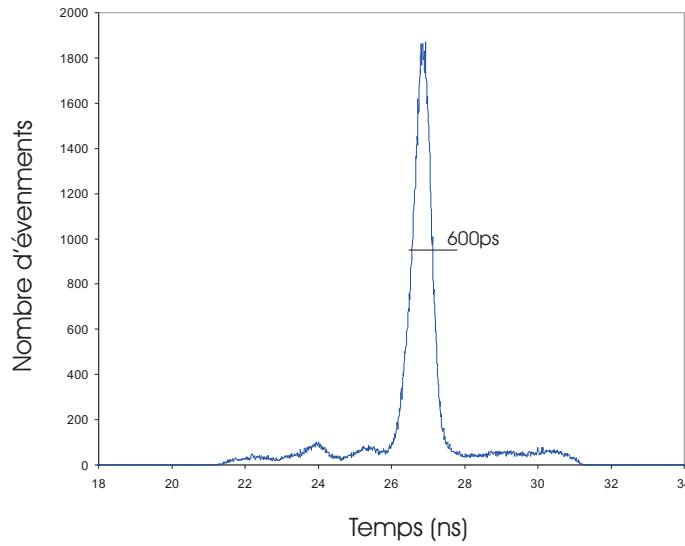


FIG. 2.24 – Diagramme de coïncidence entre les impulsions électriques et les photons annoncés pour une durée d'ouverture de l'APD-*InGaAs* de 10 ns . L'acquisition a duré environ 20 secondes.

Maintenant, définissons les notations et les termes que nous allons utiliser pour la suite :

- $S_{In}^{brut} = S_{In}^{net} + D_{In}^c$ représente les coups bruts dans l'APD-*InGaAs* qui sont l'addition des coups nets S_{In}^{net} et des coups sombres D_{In}^c ,
- $S_{Ge}^{brut} = S_{Ge}^{net} + D_{Ge}^c$ est leur équivalent pour l'APD-*Germanium*,
- η_{In} et η_{Ge} représentent respectivement les efficacités de détection des APD-*InGaAs* et APD-*Ge*.

Un coup net correspond à la « vraie » détection d'un photon qui déclenche l'avalanche au sein de l'APD, tandis qu'un coup sombre correspond justement au cas contraire où aucun photon n'est à l'origine de l'avalanche. La

probabilité d'avoir un coup sombre dans une APD passivement déclenchée est facilement mesurable car elle ne dépend que de sa tension de biais (donc de son efficacité) qui est constante. Alors, la simple mesure du nombre de détections, pendant une seconde, tandis que l'on empêche tous les photons signal d'arriver sur l'APD-Ge, définit D_{Ge}^c en s^{-1} . Pour une APD déclenchée, la durée d'allumage ΔT et le nombre d'ouvertures par seconde jouent un rôle important dans la valeur des coups sombres¹⁷. À fréquence et à durée moyenne d'allumage constantes, nous mesurons expérimentalement D_{In}^c en bloquant cette fois l'arrivée des photons idler.

2.3.2 Les résultats

Commençons tout d'abord par donner l'expression du taux de comptage net sur l'APD-InGaAs en fonction du taux de photons signal détectés. Pour détecter un photon idler, il faut tout d'abord avoir détecté le photon signal associé (S_{Ge}^{net}), puis que le photon idler ait été collecté (γ_i) et enfin qu'il soit lui-même détecté (η_{In}), ce qui correspond mathématiquement à :

$$S_{In}^{net} = \gamma_i \times \eta_{In} \times S_{Ge}^{net} \quad (2.37)$$

Alors très rapidement si nous connaissons l'efficacité de notre APD, nous pouvons calculer :

$$\gamma_i = \frac{S_{In}^{brut} - D_{In}^c}{\eta_{In} \times (S_{Ge}^{brut} - D_{Ge}^c)} \quad (2.38)$$

Lors du comptage de N événements aléatoires, nous prévoyons toujours une incertitude poissonnienne en $\sqrt{N}$. Aussi, les incertitudes indiquées dans le tableau 2.3 ont été calculées, sauf celle sur l'efficacité de l'APD-InGaAs qui a été tirée des spécifications données sur la documentation. Calculons donc l'erreur type que nous faisons sur la détermination de γ_i :

$$\Delta\gamma_i = \gamma_i \left[\frac{(\Delta S_{In}^{brut} + \Delta D_{In}^c)}{S_{In}^{net}} + \frac{\Delta\eta_{In}}{\eta_{In}} + \frac{(\Delta S_{Ge}^{brut} + \Delta D_{Ge}^c)}{S_{Ge}^{net}} \right] \quad (2.39)$$

En relevant les valeurs S_{In}^{brut} , S_{Ge}^{brut} ainsi que les taux de coups sombres et les efficacités des APDs associées, nous obtenons le tableau suivant :
À partir du tableau 2.3, nous calculons

$$\boxed{\gamma_i = 0,45 \pm 0,03}$$

¹⁷La probabilité d'avoir un coup sombres est donnée par ns d'ouverture de la fenêtre de détection et pour une seule fenêtre. Pour estimer D_{In}^c l'utilisateur n'a qu'à multiplier cette probabilité par la fréquence d'allumage de son détecteur et la durée d'ouverture.

	APD- <i>InGaAs</i>	APD- <i>Ge</i>
$S^{brut} (s^{-1})$	5750 ± 76	135000 ± 367
$D^c (s^{-1})$	320 ± 18	19000 ± 138
$Eff. APDs$	$0,103 \pm 0,005$	$0,055 \pm 0,001$

TAB. 2.3 – Données expérimentales associées à la mesure du coefficient de couplage γ_i .

Nous voulons insister encore une fois sur l'interprétation qui doit être faite de γ_i . Ce coefficient correspond à la *probabilité de récolter et de ne pas le perdre* le photon idler dans les composants fibrés. Si nous souhaitons identifier la valeur du couplage réel « *guide-fibre* », il suffit de retirer les pertes des composants, mesurées page 99, à la valeur de γ_i . On a alors :

$$\Gamma_{guide/fibre} = \frac{0,45}{10^{-0,112}} = 0,58 \pm 0,04$$

Il aurait été pratique de pouvoir suivre un protocole identique pour mesurer le couplage Γ_s « *guide-fibre* » des photons signal à 1310 nm . Malheureusement, cette mesure n'est possible que pour les photons idler car leur détection est *conditionnée* par la détection préalable d'un photon signal et faire l'inverse, c'est-à-dire conditionner la détection d'un photon signal à celle d'un idler, n'est pas possible à cause de la plage de travail de l'APD-*Ge* qui est quasiment aveugle vis à vis des photons à 1550 nm .

Le protocole expérimental que nous allons suivre nous permettra uniquement de déterminer « grossièrement » le rapport entre les coefficients Γ_i et Γ_s . La méthode consiste à injecter dans le guide par la face de sortie, à l'aide de la fibre, la lumière provenant d'un laser à 1310 nm puis à 1550 nm . À l'aide d'une lentille caractérisée par une grande ouverture numérique¹⁸ et par un traitement anti-reflet adapté à l'infrarouge, nous mesurons la puissance qui ressort du guide grâce à un puissance-mètre. Connaissant la puissance injectée, nous pouvons remonter au taux de couplage à 1310 nm et 1550 nm . Pourtant ces valeurs ne correspondent pas directement aux Γ_i et Γ_s attendus, car il y a trop de paramètres extérieurs qui rajoutent des pertes non contrôlées, comme la lentille qui sert à récolter la lumière issue du guide par exemple mais aussi les pertes dans le guide lui-même à ces longueurs d'ondes. Dans ce cas, nous pouvons seulement supposer que leur action est relativement identique pour les deux longueurs d'onde et calculer le rapport $\frac{\Gamma_i}{\Gamma_s}$. Il faut donc garder à l'esprit la fragilité du raisonnement et considérer

¹⁸supérieure à celle du guide.

2.3. DÉTERMINATION DU COEFFICIENT DE COUPLAGE GUIDE-FIBRE 107

ce rapport comme une simple indication. Toutefois, pour pouvoir utiliser la valeur de γ_s dans les formules théoriques du chapitre 1, nous allons associer à cette mesure une grande incertitude.

	1550 nm	1310 nm
Puissance injectée	1,2 mW	240 μ W
Puissance récoltée	550 μ W	98 μ W

$$\frac{\Gamma_s}{\Gamma_i} \approx 0,8 \pm 0,1 \quad (2.40)$$

À ce sujet, notons qu'une estimation théorique de Γ_i et Γ_s par le calcul du recouvrement des modes signal et idler du guide avec ceux d'une fibre *SMF-28* a donné¹⁹ :

$$\Gamma_i = 0,71 \quad \text{et} \quad \Gamma_s = 0,61$$

ce qui nous laisse penser que notre mesure expérimentale, bien qu'à prendre avec prudence, reflète une certaine réalité physique.

Maintenant, nous pouvons donc évaluer la valeur γ_s qui correspond cette fois au *coefficient de couplage + pertes* associé aux photons signal. Pour cela il faut faire le rapport des pertes vues par le photon idler sur celles vues par le photon signal.

$$\gamma_s = \gamma_i \times \frac{10^{-0,2}}{10^{-0,112}} \times \frac{\Gamma_s}{\Gamma_i} \approx 0,29 \pm 0,05 \quad (2.41)$$

Voici la justification expérimentale du rapport $\frac{\gamma_s}{\gamma_i}$ utilisée dans le chapitre 1 pour tracer les évolutions de P_1 et $g^{(2)}(0)$ en fonction du coefficient de collection des paires (voir page 56). Dans le chapitre suivant, nous utiliserons donc $\gamma_s = 0,29$ pour calculer P_0 , P_1 et P_2 à partir des formules théoriques 1.29, 1.27 et 1.25. Nous vérifierons ainsi si les données expérimentales suivent les prédictions théoriques du premier chapitre.

¹⁹La position de la fibre était optimisée pour la récolte des photons idler comme dans notre configuration expérimentale et nous avons tenu compte de la présence d'un éventuel gap d'air entre la fibre et le guide.

2.4 Conclusion du chapitre 2

Ce chapitre orienté « technologie », nous a permis de décrire précisément les différents aspects de la source de photons uniques annoncés suivante.

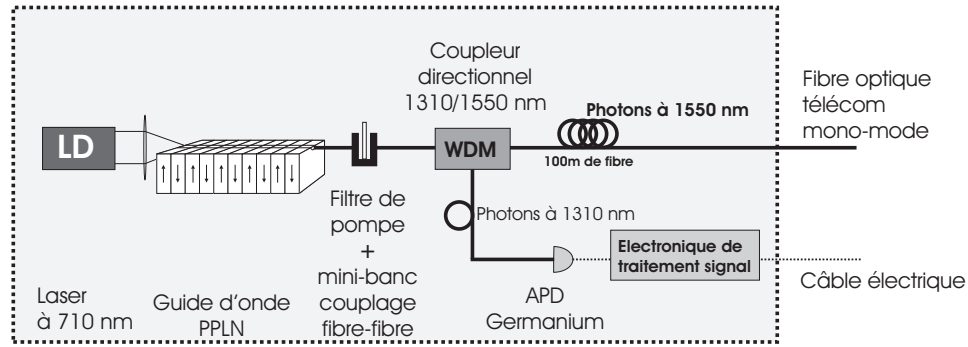


FIG. 2.25 – Source de photon uniques annoncés « Niçoise ».

Dans un premier temps, nous avons étudié les caractéristiques des paires de photons émises par le guide PPLN. Il s'est avéré que la longueur d'onde des photons idler était centrée sur 1550 nm et présentait une largeur spectrale de 20 nm . Nous nous sommes aussi intéressé au guidage de la lumière et nous avons montré que la distribution spatiale des paires de photons à l'intérieur du guide ne nous permettrait pas de collecter 100% des paires émises à l'aide d'une fibre conventionnelle. Une mesure expérimentale a permis d'établir que le coefficient de couplage « *guide-fibre* » des photons idler était égal à 0,58 ce qui est un très bon résultat.

La seconde partie du chapitre s'est focalisée sur le choix des composants associés à la source. Nous avons décidé d'utiliser des composants conventionnels car selon nous leur disponibilité, leur faible coût et les faibles pertes que présentent finalement ces composants sont autant d'atouts pour assurer la reproductibilité de la réalisation de cette source. Dans cette optique, nous avons aussi fait le pari d'utiliser des composants entièrement connectés au détriment des pertes, ce qui se justifie aisément par la construction modulaire de la source qui accroît sa souplesse d'utilisation et de diagnostic. Expérimentalement, nous avons estimé les pertes totales à 1550 nm associées aux composants égales à $1,12\text{ dB}$.

Par la suite, l'association du guide PPLN et des composants fibrés nous a permis de déterminer lors d'une première mesure en mode comptage de photon, la valeur du coefficient de *couplage+perles* γ_i qui détermine la probabilité maximale qu'un photon annoncé soit présent au bout de la fibre optique de sortie. Il se trouve que la valeur de $\gamma_i = 0,45$ correspond tout à fait à la

probabilité de collecter le photon à la sortie du guide multipliée par celle de ne pas le perdre dans les composants optiques. Pour introduire le chapitre suivant, nous pouvons alors espérer au mieux une probabilité P_1 d'obtenir un photon unique de 0,45, ce qui correspondrait à un des meilleurs résultats à l'heure actuelle pour une source de photons uniques annoncés entièrement fibrée aux longueurs d'ondes télécoms.

Chapitre 3

Mesures expérimentales des performances de la source

Le premier chapitre nous a permis de mettre en place la théorie associée à la *source de photons uniques annoncés* en régime continu et d'établir, en fonction des pertes vues par les photons ou de l'efficacité de l'APD servant à annoncer l'arrivée d'un photon, les figures de mérites P_1 et $g^{(2)}(0)$ associées à notre source. Pour sa part, le second chapitre a été l'occasion de quantifier expérimentalement les pertes γ_i , que subissent les photons annoncés, qui correspond à la limite supérieure de la probabilité d'obtenir un photon unique.

Toutefois, il n'y a pas que la probabilité d'avoir un photon unique qui compte mais surtout celle de ne pas en avoir deux...

Pour cela nous aurons à déterminer P_2 , la probabilité d'avoir deux photons, qui devra être la plus faible possible. Nous montrerons donc, dans la première partie de ce troisième chapitre, que moyennant un montage de type « Hanbury-Brown & Twiss » nous pouvons remonter, via l'utilisation d'un modèle théorique, à la statistique des photons (P_0^{exp} , P_1^{exp} et P_2^{exp}) qui sont à l'origine des taux de détection sur le montage.

Il faut voir cette méthode expérimentale comme une approche inverse à la théorie développée dans le chapitre 1 où, en partant de la distribution poissonnienne des paires de photons au sein du guide PPLN, nous estimions les probabilités P_i^{th} d'avoir i photon annoncés en fonction des pertes γ_i ou encore de l'efficacité η_{Ge} de l'APD germanium.

La confrontation des probabilités P_i expérimentales et théoriques constituera donc le point fort de ce chapitre où les deux approches se compléteront.

Par ailleurs, dans la communauté des photons uniques, l'observation de la fonction d'autocorrélation d'ordre deux $g^{(2)}(\tau)$ des photons est souvent considérée comme une référence pour déterminer leur « caractère unique ». Dans la seconde partie de ce chapitre, nous montrerons d'abord qu'il n'est pas possible de mesurer « classiquement » $g^{(2)}(\tau)$ pour une source asynchrone, mais que grâce à un nouveau protocole expérimental, développé en collaboration avec le GAP-Optique de Genève, il est possible de contourner le problème et de tracer l'équivalent *asynchrone* de la courbe d'autocorrélation d'ordre deux.

À partir de maintenant et jusqu'à la fin du manuscrit, la source de photons uniques annoncés sera considérée comme une boîte noire avec deux sorties, représentées sur la figure 3.1 :

- Une sortie *électrique* pour l'impulsion servant à annoncer l'arrivée d'un photon.
- Une sortie *optique* constituée d'une fibre monomode délivrant le « photon unique annoncé ». Cette fibre monomode nous assure une excellente localisation spatiale du photon et une intégration facile dans toute expérience de communication quantique.

Pour cette « boîte » les probabilités d'avoir 0, 1 ou 2 photons pour un trigger électrique sont données par P_0 , P_1 ou P_2 .

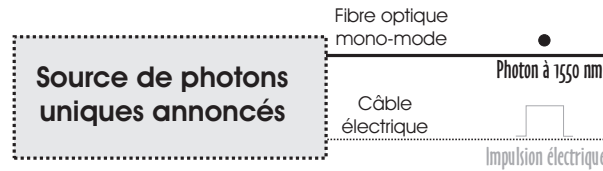


FIG. 3.1 – Représentation globale de la source de photon uniques annoncés.

3.1 Les mesures de P_0 , P_1 et P_2

Pour déterminer P_0 , P_1 et P_2 , nous ne pouvons pas utiliser le montage 2.23 utilisé au chapitre précédent. En effet, l'APD-*InGaAs* ne peut résoudre le nombre de photons incidents à l'intérieur de sa fenêtre de détection ΔT et nous n'avons alors aucun moyen d'identifier les événements où deux paires créées successivement se suivent dans un intervalle de temps inférieur à ΔT .

Alors, le seul moyen de déterminer la statistique de la source est de faire une mesure directe des événements à 2 photons en utilisant deux détecteurs.

Pour cela nous allons faire évoluer le montage 2.23 en considérant une lame séparatrice 50/50 qui va distribuer les photons incidents aléatoirement selon deux routes possibles, aboutissant chacune à une APD. Les événements à deux photons peuvent être identifiés par les détections simultanées dans les deux APDs, et les cas où les photons sont uniques, sont identifiés par les détections simples dans une seule des deux APDs. Ce montage est connu sous le nom de « Hanbury-Brown & Twiss » (HBT) [66].

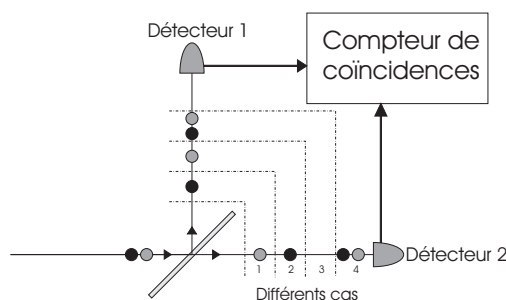


FIG. 3.2 – Principe du montage de type *HBT*. Les photons sont séparés par un miroir semi-réfléchissant aussi appelé lame séparatrice 50/50. La séparation de deux photons incidents est aléatoire et nous avons représenté les différents cas possibles ayant tous une probabilité $1/4$.

Le miroir semi-réfléchissant transmet ou réfléchit le photon incident de façon aléatoire. Si deux photons sont incidents simultanément, alors dans 50% des cas ils seront séparés correctement (c'est-à-dire un photon dans chaque bras) comme représenté pour les cas 1 et 2 du schéma 3.2.

3.1.1 Le montage

Bien évidemment, il existe un équivalent fibré à ce miroir semi-réfléchissant : c'est le *coupleur directionnel 50/50* qui possède exactement le même comportement. Nous avons donc réalisé le montage 3.3, que nous allons détailler :

- Pour ce montage, nous utilisons un coupleur directionnel dont les sorties sont reliées à deux photodiodes à avalanches notées par la suite APD-1 et APD-2.
- Ces deux APDs fonctionnent en mode déclenché. Cela signifie qu'à l'arrivée d'un signal électrique, elles s'allument durant une durée ΔT réglable par l'utilisateur.
- Ces APDs sont caractérisées par une efficacité respective η_1 et η_2 et un taux de coups sombres D_1^c et D_2^c exprimé en s^{-1} .

- S_1^{brut} et S_2^{brut} , leur taux brut respectif de détections, exprimé en s^{-1} , est accessible grâce à deux compteurs.
- Les éventuelles coïncidences entre les deux APDs sont analysées par un *TAC/SCA* et leur taux accessible via un troisième compteur. Nous avons déjà rencontré le *TAC* à la fin du chapitre précédent, pour visualiser les coïncidences, mais en lui adjoignant un analyseur de voie unique ou *Single Channel Analyser (SCA)*, il devient possible de faire l'inventaire, en temps réel, du nombre d'événements simultanés reçus par le *TAC* grâce à une fenêtre temporelle ajustable en position et en largeur. Le taux d'événements, R_c^{brut} , est exprimé en s^{-1} .
- Notons enfin, un quatrième compteur sur le schéma 3.3 servant à enregistrer le taux d'impulsions électriques (trigger) envoyées par la source noté N_T , exprimé aussi en s^{-1} .

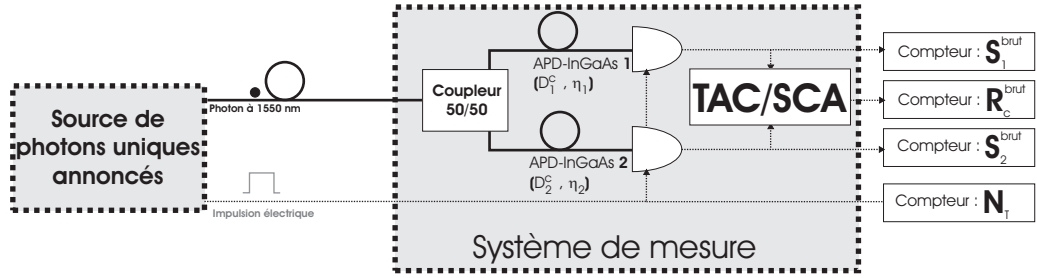


FIG. 3.3 – Expérience de comptage de photon pour déterminer la statistique de la source. Ce montage fibré est dit de type « Hanbury-Brown & Twiss ».

Ainsi à première vue, la probabilité :

- P_2 est accessible en calculant $\frac{R_c^{brut}}{N_T}$,
- P_1 est liée à $\frac{S_1^{brut} + S_2^{brut} - R_c^{brut}}{N_T}$,
- P_0 est finalement égale à $1 - P_1 - P_2$.

Mais ce ne sera pas aussi simple, rappelons-nous le comportement aléatoire du miroir semi-réfléchissant qui ne sépare correctement deux photons que dans 50% des cas, la faible efficacité et les coups sombres dans les APD-*InGaAs*₁ et APD-*InGaAs*₂ qui ne correspondent à aucun « vrai » photon...

Ces APDs et le miroir semi-réfléchissant faisant partie du système de mesure, il faut soustraire les erreurs qui en découlent. Pour cela, il nous faut écrire proprement les probabilités d'obtenir une coïncidence, un coup simple et aucun coup, en fonction des probabilités P_0 , P_1 , P_2 d'avoir 0, 1 ou 2 photons à la sortie de la source et des taux de coups sombres D_1^c et D_2^c dans les deux APDs du montage HBT.

3.1.2 Le Modèle d'analyse des données

Étant donnée, une impulsion électrique envoyé par la source, nous allons noter :

- $\mathbb{P}_2$ la probabilité d'obtenir une coïncidence entre les deux APDs ;
- $\mathbb{P}_1$ celle de n'avoir qu'une seule détection sur une des deux APDs ;
- puis assez logiquement $\mathbb{P}_0$ sera la probabilité de ne pas avoir de détections sur les deux APDs.

Ces valeurs $\mathbb{P}_2$, $\mathbb{P}_1$ et $\mathbb{P}_0$ sont quand à elles respectivement égales aux mesures expérimentales :

$$\begin{aligned} & - \frac{R_c^{brut}}{N_T} \\ & - \frac{S_1^{brut} + S_2^{brut}}{N_T} \\ & - \frac{(N_T - S_1^{brut}) + (N_T - S_2^{brut})}{N_T} \end{aligned}$$

et ne peuvent pas être normalisées à l'unité, dû à la présence de deux détecteurs.

N'ayant pas trouvé de notation plus appropriée, il faudra faire attention de ne pas les confondre avec P_0 , P_1 , P_2 , qui rappelons-le sont les probabilités d'avoir 0, 1 ou 2 photons annoncés à la sortie de la source.

Nous chercherons donc à exprimer les probabilités $\mathbb{P}_0$, $\mathbb{P}_1$ et $\mathbb{P}_2$ en fonction des probabilités P_0 , P_1 , P_2 , des efficacités de détection η_1 et η_2 et des probabilités d'avoir un coup sombre $p_{APD-1}^{dc} = \frac{D_1^c}{N_T}$ et $p_{APD-2}^{dc} = \frac{D_2^c}{N_T}$ dans les APDs 1 et 2.

Calcul d'introduction

Nous allons tout d'abord chercher à écrire la probabilité totale, p_{APD-1}^{phot} , d'avoir une « vraie » détection provenant de j photons¹ sortant du bras 1 du coupleur 50/50. En définissant $p(j)$ la probabilité d'avoir j photons incidents sur l'APD-1, nous pouvons écrire :

$$\boxed{p_{APD-1}^{phot} = \sum_{j=0}^{\infty} p(j) \left(1 - (1 - \eta_1)^j\right)} \quad (3.1)$$

Nous allons à présent exprimer $p(j)$ en fonction de P_i la probabilité d'avoir

¹Il est astucieux d'écrire la probabilité d'avoir une détection sous la forme $p_{dét}^{(j)} = 1 - p_{pasdét}^{(j)}$ avec $p_{pasdét}^{(j)} = (1 - \eta_1)^j$. Il vient alors :

$$p_{dét}^{(j)} = 1 - (1 - \eta_1)^j$$

i photons émis par la source en amont du coupleur 50/50. La probabilité d'avoir j photons incidents dans le bras 1 du coupleur s'écrit comme une loi binômiale :

$$p(j) = \sum_{i=j}^{\infty} P_i \mathcal{C}_j^i \left(\frac{1}{2}\right)^j \left(\frac{1}{2}\right)^{i-j} \quad (3.2)$$

$$\boxed{p(j) = \sum_{i=j}^{\infty} P_i \mathcal{C}_j^i \left(\frac{1}{2}\right)^i} \quad (3.3)$$

Ce qui permet d'écrire la formule 3.1 sous la forme

$$p_{APD-1}^{phot} = \sum_{j=0}^{\infty} \left[\sum_{i=j}^{\infty} \left\{ P_i \mathcal{C}_j^i \left(\frac{1}{2}\right)^i \right\} \left(1 - (1 - \eta_1)^j\right) \right] \quad (3.4)$$

Nous pouvons inverser les sommes sur les indices en faisant attention à leurs bornes

$$p_{APD-1}^{phot} = \sum_{i=0}^{\infty} \left\{ P_i \left(\frac{1}{2}\right)^i \times \sum_{j=0}^i \left[\mathcal{C}_j^i \left(1 - (1 - \eta_1)^j\right) \right] \right\} \quad (3.5)$$

puis faire disparaître la somme sur j , grâce aux identités remarquables²

$$p_{APD-1}^{phot} = \sum_{i=0}^{\infty} \left\{ P_i \left(\frac{1}{2}\right)^i \times \left[2^i - 2^i \left(1 - \frac{\eta_1}{2}\right)^i \right] \right\} \quad (3.6)$$

$$\boxed{p_{APD-1}^{phot} = \sum_{i=0}^{\infty} \left\{ P_i \times \left[1 - \left(1 - \frac{\eta_1}{2}\right)^i \right] \right\}} \quad (3.7)$$

Pour finir, nous allons exprimer la probabilité totale, $\mathbb{P}_{APD-1}$, d'avoir un « clic » dans l'APD-1 que ce soit un « vrai » photon (p_{APD-1}^{phot}) ou un coup sombre (p_{APD-1}^{dc}) :

$$\begin{aligned} \mathbb{P}_{APD-1} &= (1 - p_{APD-1}^{dc}) p_{APD-1}^{phot} + p_{APD-1}^{dc} (1 - p_{APD-1}^{phot}) \\ &= p_{APD-1}^{phot} (1 - 2p_{APD-1}^{dc}) + p_{APD-1}^{dc} \end{aligned} \quad (3.8)$$

2

$$\sum_{j=0}^i \mathcal{C}_j^i x^j y^{i-j} = (x + y)^i$$

À l'aide de l'équation 3.7 nous pouvons écrire :

$$\mathbb{P}_{APD-1} = \sum_{i=0}^{\infty} \left\{ P_i \times \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i \right] \right\} (1 - 2p_{APD-1}^{dc}) + p_{APD-1}^{dc} \quad (3.9)$$

ou encore en notant que

$$\sum_{i=0}^{\infty} P_i = 1$$

$$\boxed{\mathbb{P}_{APD-1} = \sum_{i=0}^{\infty} P_i \times \left\{ \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i \right] (1 - 2p_{APD-1}^{dc}) + p_{APD-1}^{dc} \right\}} \quad (3.10)$$

Bien entendu, un calcul similaire nous permet de déterminer la même formule pour $\mathbb{P}_{APD-2}$.

Notons tout de même une petite astuce à propos de l'efficacité de détection des APDs 1 et 2. En effet, à aucun moment dans le développement théorique nous n'avons tenu compte des pertes du coupleur directionnel 50/50. Étant donné que la probabilité finale que possède un photon d'être détecté correspond à la probabilité d'être collecté *et* détecté, les valeurs de η_1 et η_2 seront simplement affectées d'un coefficient correspondant aux pertes du coupleur pour rétablir l'exactitude du modèle. Pour se convaincre de la justesse de cette astuce, il suffirait de refaire le calcul 3.1 en tenant compte des pertes du coupleur, et de constater qu'un terme supplémentaire $\gamma_{50/50}$ apparaîtrait³ uniquement devant les termes $\eta_{1,2}$.

Maintenant, il ne nous reste plus qu'à exprimer $\mathbb{P}_0$, $\mathbb{P}_1$ et $\mathbb{P}_2$ en fonction de $\mathbb{P}_{APD-1}$ et $\mathbb{P}_{APD-2}$.

Contributions à la probabilité de n'avoir aucun coup $\mathbb{P}_0$

Finalement, il suffit de se rappeler que les deux compteurs sont décorrélés, ce qui nous permet d'écrire simplement :

$$\boxed{\mathbb{P}_0 = (1 - \mathbb{P}_{APD-1}) + (1 - \mathbb{P}_{APD-2})}$$

³ Nous avons fait le choix de ne pas tenir compte de ce terme dans nos calculs pour favoriser la lisibilité, au détriment de l'exactitude.

Nous pouvons enfin finir par écrire la dernière équation, associée à $\mathbb{P}_0$:

$$\mathbb{P}_0 = \sum_{i=0}^{\infty} P_i \times \left\{ (1 - p_{APD-1}^{dc}) - \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i \right] (1 - 2p_{APD-1}^{dc}) \right. \\ \left. + (1 - p_{APD-1}^{dc}) - \left[1 - \left(1 - \frac{\eta_2}{2} \right)^i \right] (1 - 2p_{APD-2}^{dc}) \right\} \quad (3.11)$$

Contributions à la probabilité d'obtenir des coups simples $\mathbb{P}_1$

Il faut noter dans la définition expérimentale de $\mathbb{P}_1$ que les coïncidences sont implicitement incluses car les compteurs S_1^{brut} et S_2^{brut} sont décorrélés et indiquent *toutes* les détections faites, sans tenir compte d'une possible coïncidence avec son voisin. Ainsi, la probabilité $\mathbb{P}_1$ d'enregistrer une détection sur les APDs 1 et 2 s'écrit simplement :

$$\boxed{\mathbb{P}_1 = \mathbb{P}_{APD-1} + \mathbb{P}_{APD-2}}$$

Ainsi, il vient :

$$\mathbb{P}_1 = \sum_{i=0}^{\infty} P_i \times \left\{ \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i \right] (1 - 2p_{APD-1}^{dc}) + p_{APD-1}^{dc} \right. \\ \left. + \left[1 - \left(1 - \frac{\eta_2}{2} \right)^i \right] (1 - 2p_{APD-2}^{dc}) + p_{APD-2}^{dc} \right\} \quad (3.12)$$

Contributions à la probabilité d'obtenir une coïncidence $\mathbb{P}_2$

La probabilité d'obtenir une coïncidence s'écrit :

$$\boxed{\mathbb{P}_2 = \mathbb{P}_{APD-1} \times \mathbb{P}_{APD-2}}$$

On a donc :

$$\mathbb{P}_2 = (1 - 2p_{APD-1}^{dc}) (1 - 2p_{APD-2}^{dc}) p_{APD-1}^{phot} p_{APD-2}^{phot} \\ + p_{APD-1}^{dc} (1 - 2p_{APD-2}^{dc}) p_{APD-2}^{phot} \\ + p_{APD-2}^{dc} (1 - 2p_{APD-1}^{dc}) p_{APD-1}^{phot} \\ + p_{APD-1}^{dc} \cdot p_{APD-2}^{dc} \quad (3.13)$$

Ce qui conduit, grâce à la formule 3.7, à :

$$\begin{aligned} \mathbb{P}_2 = \sum_{i=0}^{\infty} P_i \times \left\{ \begin{aligned} & \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i - \left(1 - \frac{\eta_2}{2} \right)^i \right. \\ & \left. + \left(1 - \frac{\eta_1 + \eta_2}{2} \right)^i \right] (1 - 2p_{APD-1}^{dc}) (1 - 2p_{APD-2}^{dc}) \\ & + \left[1 - \left(1 - \frac{\eta_1}{2} \right)^i \right] (1 - 2p_{APD-1}^{dc}) p_{APD-2}^{dc} \\ & + \left[1 - \left(1 - \frac{\eta_2}{2} \right)^i \right] (1 - 2p_{APD-2}^{dc}) p_{APD-1}^{dc} \\ & \left. + p_{APD-1}^{dc} \cdot p_{APD-2}^{dc} \right\} \end{aligned} \quad (3.14) \end{aligned}$$

Nous allons maintenant supposer que $P_2 \ll P_1$, donc que toutes les probabilités P_i , pour i supérieur à 2, sont complètement négligeables⁴. Ceci nous permet de réécrire sur la page suivante les formules 3.12, 3.14, 3.11 simplifiées qui dépendront alors uniquement de P_0 , P_1 et P_2 . Notons toutefois que les probabilités $\mathbb{P}_i$ et p_{APD}^{dc} ont été remplacées par leurs valeurs expérimentales pour faire apparaître un système, à trois inconnues et 3 équations, assez simple à résoudre. En effet, seules les probabilités P_i sont inconnues, le reste étant des paramètres expérimentaux que nous mesurerons par la suite. Nous avons validé le système à l'aide de valeurs numériques simples et nous ne présenterons pas le système résolu, car *MATLAB* le fait bien plus efficacement que nous et sans risque d'erreur. Pour notre travail, nous avons donc programmé une routine sous *MATLAB* qui nécessite simplement 6 *entrées* qui sont :

- N_T le nombre total de *trigger* fournis par la source,
- S_1^{brut} le taux de détections sur l'APD-1,
- D_1^c le taux de coups sombres sur l'APD-1,
- S_2^{brut} le taux de détections sur l'APD-2,
- D_2^c le taux de coups sombres sur l'APD-2,
- R_c^{brut} le taux de coïncidences entre les APDs 1 et 2.

⁴Nous verrons que l'approximation $P_2 \ll P_1$ sera justifiée par les résultats expérimentaux.

Le modèle final

Les formules finales sont donc :

$$\begin{aligned}
\frac{R_c^{brut}}{N_T} &= P_2 \times \left[\left(1 - \frac{D_1^c}{N_T}\right)\left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{2}\eta_1\eta_2\right] \right. \\
&\quad + \frac{D_1^c}{N_T}\left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{4}2\eta_2 + \frac{1}{2}(1 - \eta_1)\eta_2\right] \\
&\quad + \left(1 - \frac{D_1^c}{N_T}\right)\frac{D_2^c}{N_T} \cdot \left[\frac{1}{4}2\eta_1 + \frac{1}{2}(1 - \eta_2)\eta_1\right] \\
&\quad \left. + \frac{D_1^c D_2^c}{N_T^2} \cdot \left[\frac{1}{4}(1 - \eta_1)^2 + \frac{1}{4}(1 - \eta_2)^2 + \frac{1}{2}(1 - \eta_1)(1 - \eta_2)\right] \right] \\
&+ P_1 \times \left[\frac{D_1^c}{N_T}\left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{2}\eta_2\right] \right. \\
&\quad + \left(1 - \frac{D_1^c}{N_T}\right)\frac{D_2^c}{N_T} \cdot \left[\frac{1}{2}\eta_1\right] \\
&\quad \left. + \frac{D_1^c D_2^c}{N_T^2} \cdot \left[\frac{1}{2}(1 - \eta_1) + \frac{1}{2}(1 - \eta_2)\right] \right] \\
&+ P_0 \times \frac{D_1^c D_2^c}{N_T^2} \\
\frac{S_1^{brut} + S_2^{brut}}{N_T} &= P_2 \times \left[\left(1 - \frac{D_1^c}{N_T}\right) \cdot \left[\frac{1}{4}2\eta_1 + \frac{1}{2}\eta_1\right] \right. \\
&\quad + \frac{D_1^c}{N_T} \cdot \left[\frac{1}{4}(1 - \eta_1)^2 + \frac{1}{2}(1 - \eta_1) + \frac{1}{4}\right] \\
&\quad + \left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{4}2\eta_2 + \frac{1}{2}\eta_2\right] \\
&\quad \left. + \frac{D_2^c}{N_T} \cdot \left[\frac{1}{4}(1 - \eta_2)^2 + \frac{1}{2}(1 - \eta_2) + \frac{1}{4}\right] \right] \\
&+ P_1 \times \left[\left(1 - \frac{D_1^c}{N_T}\right) \cdot \left[\frac{1}{2}\eta_1\right] \right. \\
&\quad + \frac{D_1^c}{N_T} \cdot \left[\frac{1}{2}(1 - \eta_1) + \frac{1}{2}\right] \\
&\quad + \left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{2}\eta_2\right] \\
&\quad \left. + \frac{D_2^c}{N_T} \cdot \left[\frac{1}{2}(1 - \eta_2) + \frac{1}{2}\right] \right] \\
&+ P_0 \times \left[\frac{D_1^c}{N_T} \right. \\
&\quad \left. + \frac{D_2^c}{N_T} \right] \\
\frac{2 - S_1^{brut} - S_2^{brut}}{N_T} &= P_2 \times \left[\left(1 - \frac{D_1^c}{N_T}\right) \cdot \left[\frac{1}{2}(1 - \eta_1) + \frac{1}{4}(1 - \eta_1)^2 + \frac{1}{4}\right] \right. \\
&\quad \left. + \left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{2}(1 - \eta_2) + \frac{1}{4}(1 - \eta_2)^2 + \frac{1}{4}\right] \right] \\
&+ P_1 \times \left[\left(1 - \frac{D_1^c}{N_T}\right) \cdot \left[\frac{1}{2}(1 - \eta_1) + \frac{1}{2}\right] \right. \\
&\quad \left. + \left(1 - \frac{D_2^c}{N_T}\right) \cdot \left[\frac{1}{2}(1 - \eta_2) + \frac{1}{2}\right] \right] \\
&+ P_0 \times \left[\left(1 - \frac{D_1^c}{N_T}\right) \right. \\
&\quad \left. + \left(1 - \frac{D_2^c}{N_T}\right) \right]
\end{aligned} \tag{3.15}$$

3.1.3 Les résultats expérimentaux

Voyons tout d'abord comment se déroule une mesure expérimentale des 6 entrées nécessaires aux formules 3.15 : nous disposons de 4 compteurs sur la figure 3.3, qui nous permettent de réaliser une mesure complète en deux étapes. Nous commençons par mesurer :

- N_T le taux de *trigger* fournis par la source,
- S_1^{brut} le taux de détections sur l'APD-1,
- S_2^{brut} le taux de détections sur l'APD-2,
- R_c^{brut} le taux de coïncidences entre les APDs 1 et 2,

puis en débranchant la fibre reliant la source de photons uniques au coupleur directionnel, nous mesurons les taux de coups sombres et relevons alors :

- D_1^c le taux de coups sombres sur l'APD-1,
- D_2^c le taux de coups sombres sur l'APD-2.

Le protocole d'acquisition, que nous suivons pour chaque mesure, est dicté par la volonté de diminuer au maximum les incertitudes et consiste à faire des acquisitions pendant 40 secondes⁵.

Pourquoi 40 s ? En mode comptage de photon, l'incertitude liée à la prise de mesure est égale à la racine carrée du nombre total de détections $\sqrt{N_{détections}}$ (cf. loi de poisson). Ainsi, plus nous enregistrons de points pour la mesure, plus l'incertitude relative est faible. Voyons cela, dans notre cas :

$$N_{détections}^{40s} \sim 40 \times N_{détections}^{1s}$$

et l'incertitude vaut alors :

$$\Delta N_{détections}^{40s} = \sqrt{N_{détections}^{40s}} \sim \sqrt{40} \times \Delta N_{détections}^{1s}$$

Nous voyons ainsi facilement que l'incertitude relative vaut seulement

$$\frac{\Delta N_{détections}^{40s}}{N_{détections}^{40s}} = \frac{1}{\sqrt{40}} \frac{\Delta N_{détections}^{1s}}{N_{détections}^{1s}}$$

soit une incertitude réduite d'un facteur 6 comparée à une simple mesure durant une seconde.

Nous pouvons finalement compléter le tableau expérimental suivant, dans lequel nous avons inséré la valeur mesurée de D_{ge}^c qui est facilement accessible

⁵La durée a été imposée par le système de mesure qui autorise des intégrations sur une durée maximale de 10 s que nous avons répétées 4 fois... Faire les mesures plus de fois n'aurait rien apporté car, dans ce cas, les incertitudes sur les autres paramètres tels que les pertes ou l'efficacité des APDs limiteraient de toutes façons l'incertitude totale.

en coupant le laser de pompe à 710 nm . Cette mesure n'est pour l'instant justifiée que pour calculer le couplage γ_i^{exp} qui est calculée à partir de la formule 2.38, à la différence près que, cette fois, le taux de détections totales de photons annoncés S_{In}^{net} , utilisé au chapitre précédent, est ici mesuré en additionnant les taux de coups nets dans les APDs 1 et 2. Cette mesure permettra de vérifier la cohérence de cette mesure avec celle faite lors du chapitre précédent. Nous verrons par ailleurs qu'elle nous sera également utile pour le rapprochement des résultats expérimentaux et des prédictions théoriques du chapitre 1. Bien qu'elle ne soit pas nécessaire non plus au modèle 3.15, la *durée d'allumage* des APDs 1 et 2 est très importante, car les performances affichées ici sont données pour un ΔT donné. Contrairement à une source de photon impulsienne, les performances de notre source dépend ici directement de la durée d'allumage des détecteurs.

	Mesure moyenne	Incertitude
Puissance avant le guide (μW)	80	± 4
N_T (s^{-1})	124300	± 56
D_{Ge}^c (s^{-1})	19660	± 22
S_1^{brut} (s^{-1})	2351	± 8
D_1^c (s^{-1})	67	± 1
S_2^{brut} (s^{-1})	2603	± 8
D_2^c (s^{-1})	185	± 2
R_c^{brut} (s^{-1})	8,0	$\pm 0,4$
Efficacité APD-1 η_1	0,104	$\pm 0,005$
Efficacité APD-2 η_2	0,103	$\pm 0,005$
Correction pertes coupleur 50/50	0,950	$\pm 0,005$
Durée allumage APD-1 et 2 (ns)	3,0	$\pm 0,1$

TAB. 3.1 – Tableau récapitulatif des mesures expérimentales. Les efficacités des APDs sont extraites de l'annexe B.

À partir de ces valeurs, nous obtenons des valeurs de P_0^{exp} , P_1^{exp} et P_2^{exp} à l'aide du système d'équation 3.15. Aussi, nous avons ajouté, parmi les résultats, le calcul de la valeur de $g^{(2)}(0) \approx \frac{2P_2}{P_1^2}$ qui quantifie le caractère unique des photons émis par rapport à une source poissonnienne atténuée⁶. Ce coefficient servira à comparer les performances des différentes techniques de génération de photons uniques que nous rappellerons dans le tableau 3.4 situé à la fin de cette section.

⁶Rappelons à ce sujet que plus le $g^{(2)}(0)^{exp}$ sera proche de zéro, plus la source se comportera comme une source de photons uniques idéale.

	Valeur	Incertitude
P_0^{exp}	0,62	$\pm 0,02$
P_1^{exp}	0,37	$\pm 0,02$
P_2^{exp}	0,005	$\pm 0,001$
$g^{(2)}(0)^{exp}$	0,08	$\pm 0,02$
γ_i^{exp}	0,46	$\pm 0,02$

TAB. 3.2 – Tableau récapitulatif des performances expérimentales.

Avant de commenter les valeurs des résultats du tableau 3.2 et d'en tirer des conclusions, nous voulons d'abord les comparer avec les prédictions théoriques du chapitre 1. En effet, grâce à des valeurs expérimentales telles que N_T , D_{Ge}^c , γ_i et η_{Ge} nous pourrions remonter aux valeurs attendues de P_0^{th} , P_1^{th} et P_2^{th} . Nous attribuons à cette approche le mérite de valider la démarche expérimentale, car c'est à partir de paramètres en partie différents de ceux utilisés dans le modèle précédent que nous espérons obtenir l'accord entre les valeurs P_0^{th} , P_1^{th} et P_2^{th} et respectivement P_0^{exp} , P_1^{exp} et P_2^{exp} .

3.1.4 Comparaisons avec les prévisions théoriques du chapitre 1

Commençons tout d'abord par rappeler les formules du chapitre 1 associées à P_0^{th} , P_1^{th} et P_2^{th} , ainsi que l'incertitude liée à leur calcul :

$$P_0^{th} = \left(1 - \frac{\gamma_i S_{Ge}^{net}}{N_T} \right) \times (1 - \gamma_i \mu_p \Delta T) \quad (3.16)$$

$$\begin{aligned} \Delta P_0^{th} = & \left(1 - \gamma_i \mu_p \Delta T \right) \left(\frac{\gamma_i S_{Ge}^{net}}{N_T} \right) \left(\frac{\Delta \gamma_i}{\gamma_i} + \frac{\Delta S_{Ge}^{net}}{S_{Ge}^{net}} + \frac{\Delta N_T}{N_T} \right) \\ & + \left(1 - \frac{\gamma_i S_{Ge}^{net}}{N_T} \right) \left(\gamma_i \mu_p \Delta T \right) \left(\frac{\Delta \gamma_i}{\gamma_i} + \frac{\Delta \mu_p}{\mu_p} + \frac{\Delta(\Delta T)}{\Delta T} \right) \end{aligned} \quad (3.17)$$

$$P_1^{th} \approx \gamma_i \left\{ \left(\frac{S_{Ge}^{net}}{N_T} \right) \left[1 - 2\gamma_i \mu_p \Delta T \right] + \mu_p \Delta T \right\} \quad (3.18)$$

$$\begin{aligned} \Delta P_1^{th} &= \left(1 - 2\gamma_i \mu_p \Delta T \right) \left[\frac{S_{Ge}^{net} \Delta \gamma_i}{N_T} + \frac{\Delta S_{Ge}^{net} \gamma_i}{N_T} + \frac{S_{Ge}^{net} \gamma_i (\Delta N_T)}{N_T^2} \right] \\ &+ \frac{\gamma_i^2 \mu_p \Delta T S_{Ge}^{net}}{N_T} \left[\frac{\Delta \gamma_i}{\gamma_i} + \frac{\Delta \mu_p}{\mu_p} + \frac{\Delta \Delta T}{\Delta T} \right] \\ &+ \Delta T \Delta \mu_p + \mu_p \Delta(\Delta T) \end{aligned} \quad (3.19)$$

$$P_2^{th} \approx \gamma_i^2 \mu_p \Delta T \left(\frac{S_{Ge}^{net}}{N_T} \right) \quad (3.20)$$

$$\Delta P_2^{th} \approx P_2^{th} \left[2 \frac{\Delta \gamma_i}{\gamma_i} + \frac{\Delta \mu_p}{\mu_p} + \frac{\Delta(\Delta T)}{\Delta T} + \frac{\Delta S_{Ge}^{net}}{S_{Ge}^{net}} + \frac{\Delta N_T}{N_T} \right] \quad (3.21)$$

À partir des valeurs de D_{Ge}^c et N_T issues du tableau 3.1, il est aisé de déduire le taux de coups nets dans l'APD-Ge à partir de la simple relation :

$$S_{Ge}^{net} = N_T - D_{Ge}^c = 104640 \pm 78 \text{ s}^{-1}$$

Ensuite, la valeur de μ_p va s'avérer la plus délicate à estimer. Pour cela nous aurons besoin de la valeur de η_{Ge} et surtout de γ_s afin de pouvoir utiliser la formule $\mu_p = \frac{S_{Ge}^{net}}{\gamma_s \eta_{Ge}}$. Si l'efficacité du détecteur germanium est parfaitement connue,

$$\eta_{Ge} = 0,055 \pm 0,001$$

il n'en est pas de même pour la valeur de γ_s qui a été estimée à la fin du chapitre 2, où nous avons signalé son caractère approximatif. Nous avons donc associé à cette valeur une grande incertitude, qui sera la principale source d'erreur⁷ pour le calcul de μ_p .

$$\mu_p = \frac{S_{Ge}^{net}}{\gamma_s \eta_{Ge}} = \frac{104640}{0,29 \times 0,055} \approx 6,6 \cdot 10^6 \pm 1 \cdot 10^6 \text{ s}^{-1}$$

Finissons par rappeler la durée d'ouverture des APDs 1 et 2 :

$$\Delta T = 3 \pm 0,05 \text{ ns}$$

⁷ $\Delta \mu_p = \mu_p \left[\frac{\Delta S_{Ge}^{net}}{S_{Ge}^{net}} + \frac{\Delta \gamma_s}{\gamma_s} + \frac{\Delta \eta_{Ge}}{\eta_{Ge}} \right]$.

Nous pouvons alors remplir le tableau suivant qui regroupe les probabilités expérimentales et théoriques P_i , d'avoir i photons émis par la source :

	Théorie	Expérience
P_0	$0,62 \pm 0,02$	$0,62 \pm 0,02$
P_1	$0,38 \pm 0,02$	$0,37 \pm 0,02$
P_2	$0,003 \pm 0,001$	$0,005 \pm 0,001$
$g^{(2)}(0)$	$0,05 \pm 0,02$	$0,08 \pm 0,02$

TAB. 3.3 – Tableau comparatif des performances théoriques et expérimentales.

Nous constatons alors que les valeurs théoriques sont en bon accord avec les valeurs expérimentales obtenues, ce qui renforce notre confiance quant à la validité des performances obtenues, qu'elles soient expérimentales ou théoriques.

Fort de ces résultats, il serait intéressant jouer maintenant avec les paramètres de fonctionnement de la source et voir si les statistiques des photons émis par la source suivent bien les évolutions prédites au chapitre 1. Pour cela, il y a un paramètre facilement accessible, tout en gardant les autres constants : c'est la puissance de pompe donc directement le taux moyen de paires créées μ_p . Nous allons donc reprendre les tracés de P_1 et $g^{(2)}(0)$ de la figure 1.12 et ajouter conjointement les points expérimentaux pour différentes valeurs de μ_p sur la page suivante.

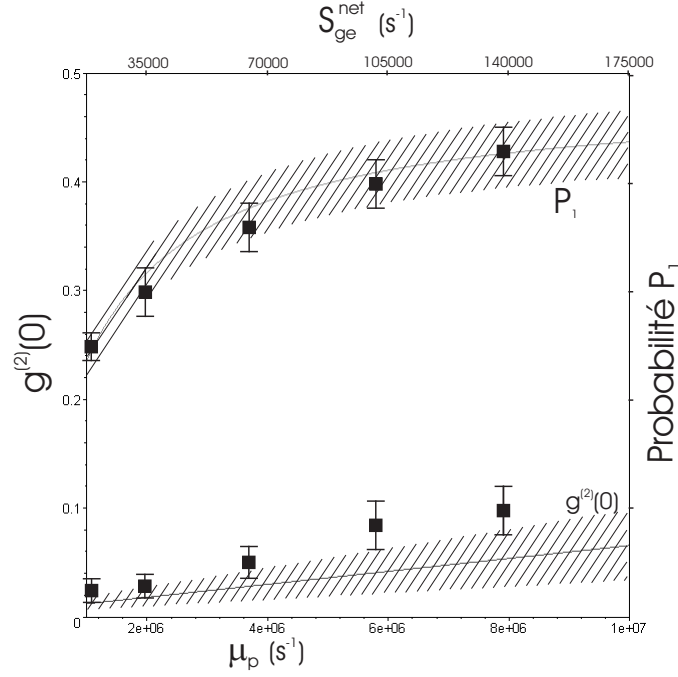


FIG. 3.4 – Évolution de la probabilité P_1 et de la valeur de $g^{(2)}(0)$ en fonction de la puissance. Les surfaces hachurées correspondent à l'incertitude liée au calcul, tandis que les rectangles noirs sont associés aux mesures expérimentales.

Nous constatons effectivement une assez bonne cohérence entre les résultats expérimentaux et les prédictions théoriques qui confirme une fois pour toute la justesse des prédictions théoriques 1.27 et 1.36 et du modèle d'analyse 3.15.

Nous avons réuni dans le tableau 3.4, les performances des principales réalisations expérimentales de sources de photons uniques. Il s'agit en fait des réalisations rencontrées dans le prologue de ce manuscrit et ce tableau sera l'occasion de situer nos performances par rapport à celles-ci. Ce sera aussi l'occasion de résumer les caractéristiques des photons émis par ces sources.

Avant de présenter les résultats du tableau, nous tenons tout d'abord à souligner deux points importants :

- les valeurs P_1^* associées aux expériences de Baltimore et de Stanford ont été recalculées pour afficher ici la probabilité brute d'obtenir un photon unique à la sortie de la fibre. En effet, ces équipes ont choisi de présenter des résultats où les pertes des composants servant à la récolte des photons à l'intérieur de la source sont corrigées et ne peuvent donc

pas être directement comparées aux autres expériences.

- Les résultats associés aux atomes froids et aux atomes en cavités ne représentent certainement pas ce qui peut se faire de mieux. Ce domaine expérimental est encore tout récent et les auteurs ne prétendent pas réaliser une source de photons uniques efficace mais simplement démontrer d'autres avantages que présentent les atomes. À ce titre nous invitons le lecteur à lire les références [54, 55, 56].

Note au lecteur

Nous voulons toutefois souligner un point qui, à nos yeux, justifie l'écart de nos performances avec ceux de la source réalisée par l'équipe d'Oxford [49]. Cette source repose sur la génération de paires de photons dans des guides, mais présente une configuration expérimentale bien étrange. Le protocole utilisé à Oxford consiste à récolter les paires de photons issues d'un microguide asymétrique à l'aide d'une fibre multimode. La valeur donnée de $P_1 = 0,85$ correspond alors à la probabilité d'obtenir un photon quel que soit le mode injecté dans la fibre et s'explique par un meilleur recouvrement spatial entre les nombreux modes de la fibre et le mode asymétrique issu du microguide. Pourtant, à l'heure actuelle, le contrôle du mode spatial du photon unique est primordial pour mener à bien une expérience de communication quantique [9, 20, 75, 76]. Contrairement aux équipes qui utilisent une fibre monomode, les photons émis ici ne pourront pas être insérés simplement dans un montage expérimental et, selon nous, la valeur $P_1 = 0,85$ obtenue ne reflète certainement pas la « probabilité expérimentale d'utiliser le photon » dans une expérience de communication quantique.

Il faut d'autant plus relativiser ce résultat qu'aucune mesure de $g^{(2)}(0)$ n'est publiée concernant cette source. Selon nous, l'utilisation d'un guide d'onde rend la génération de paires de photons très efficace et il nous semble plus qu'important d'estimer alors la probabilité de créer deux paires par impulsions laser qui peuvent contribuer à augmenter substantiellement la valeur de P_1 .

TAB. 3.4 – Comparaison des performances pour différents types de sources de photons uniques.

Type de source	λ (nm) $\Delta\lambda$ (nm)	T_{fonct}° (K)	Avantages	Inconvénients	P_1	$g^{(2)}(0)$
Paires de photons						
Nice	1550 20	342	Optique guidée Fibre monomode		0,37	0,08
Genève [51]	1550 7	Ambiante	Fibre monomode	Optique massive	0,61	0,0022
Baltimore [50]	780 10	Ambiante	Fibre monomode	Optique massive	0,50*	0,033
Oxford [49]	820 17	NC	Optique guidée	Fibre multimode	0,85	NC
Boîtes quantiques						
Stanford [38]	855 0,7	5	Largeur spectrale	Température de fonctionnement	0,083*	0,02
Centres NV						
Orsay [33]	637 75	Ambiante	Facilité de mise en œuvre [77]		0,022	0,07
Molécules						
Cachan [31]	570 50	Ambiante		« Durée de vie » de la source	0,047	0,046
Atomes froids						
Pasadena [55]	894 NC	Ambiante		« Taille » du montage	0,03	0,24
Atomes en cavité						
Garching [54]	Visible NC	Ambiante		« Taille » du montage	0,1	0,25

Commentaires sur le tableau 3.4

Il faut noter l'écart des performances entre les *sources de photons uniques annoncés* réalisées à partir de sources de paires de photons et les *sources de photons uniques à la demande* qui sont représentées par les autres techniques. Cette différence de performance se justifie par la très bonne localisation spatiale des paires qui sont émises par la conversion paramétrique, tandis que pour les sources basées sur les centres colorés ou les molécules, les photons sont émis dans un angle solide de 4π stéradians ce qui rend impossible une collection efficace des photons. Toutefois, les boîtes quantiques constituent les premières sources de photons uniques à la demande en configuration guidée. De cette façon, il est possible de favoriser l'émission des photons selon une direction privilégiée et l'équipe de Stanford et de Paris [78, 37] ont démontré une probabilité d'émettre le photon dans le mode du micropilier de 0,80. Cette valeur est tout à fait comparable avec celles obtenues par l'utilisation de paires de photons, mais, à l'heure actuelle, ces boîtes quantiques nécessitent d'être refroidies aux alentours de 5 K , ce qui implique un appareillage lourd et des pertes élevées sur le parcours du photon entre le micropilier et l'utilisateur. Des travaux sont actuellement en cours pour ramener la température de fonctionnement de ces boîtes quantiques à température ambiante [79, 80].

Enfin, notons les longueurs d'ondes respectives des photons émis par les sources : la plupart des sources dans le visible se destinent à des communications « *courtes distances* » à l'air libre, tandis que celles dans l'infrarouge sont plus adaptées à des communications longues distances par fibre optique. De ce point de vue, nous partageons avec l'équipe Genevoise la première réalisation d'une source de photons uniques aux longueurs d'ondes télécom. La principale différence entre les deux sources est une configuration guidée pour la source Niçoise, tandis que la seconde repose sur l'utilisation d'un cristal de KNbO_3 massif associé à des composants optiques massif pour coupler finalement les photons dans une fibre monomode.

3.1.5 Conclusion

Nous avons mesuré les performances (P_0 , P_1 , P_2 et $g^{(2)}(0)$) d'une source de photons uniques présentant les caractéristiques suivantes :

- Les photons annoncés sont à 1550 nm , ce qui correspond à la longueur d'onde utilisée pour les télécommunications optiques dans les fibres.
- Les photons sont récoltés directement à la sortie du guide par une fibre optique monomode qui constitue le support de base des futurs réseaux de communications quantiques longues distances.
- Le taux de répétition moyen de la source est de 100 kHz ce qui corres-

pond à la fréquence maximale de fonctionnement des APD-*InGaAs*.

- La source est entièrement fibrée et n'utilise que des composants optiques standards connectés apportant stabilité, compacité et une souplesse d'utilisation accrue.

Nous avons obtenu une probabilité d'avoir un photon unique P_1 de 0,37 tandis que la probabilité d'avoir deux photons P_2 est diminuée d'un facteur 10 comparée à une source poissonnienne équivalente.

Le travail décrit dans la première partie de ce chapitre a permis d'identifier, à partir de deux méthodes complémentaires, les performances de notre source de photons uniques :

- En partant du processus fondamental de création des paires auquel nous ajoutons les pertes probables que ces paires vont expérimenter puis l'efficacité de détection de l'APD-*Ge*, nous avons élaboré un calcul théorique grâce auquel nous estimons les probabilités P_i^{th} d'avoir i photon annoncés.
- Dans l'autre sens, un montage de type « Hanbury-Brown & Twiss » nous permet d'avoir, expérimentalement, accès à une distribution $\mathbb{P}_i$ correspondant à la probabilité d'avoir i détections dans le système de mesure. Ici, nous avons développé un modèle d'analyse qui, partant de ces probabilités et corrigeant les défauts du système de mesure, remonte à la probabilité P_i^{exp} d'avoir i photons annoncés.

Enfin, les deux méthodes complémentaires et originales permettent se s'assurer de la validité des résultats et peuvent être utilisées pour n'importe quelle source de photon uniques annoncés en régime continu⁸.

Nous avons vu que les *sources de photons uniques annoncés* suivent un fonctionnement *asynchrone* qui ne remet pas en cause leur statut de « sources de photons uniques » à condition, bien sûr, d'utiliser un protocole de communication adapté. Si ce mode de communication existe déjà sur les réseaux actuels, il n'existe toutefois aucune méthode expérimentale permettant de tracer la fonction d'autocorrélation d'ordre deux $g^{(2)}(\tau)$ pour les sources asynchrone. Cette fonction est largement utilisée par la communauté des « *sources de photons uniques à la demande* » pour vérifier expérimentalement le comportement « unique » des photons émis. Grâce à la collaboration

⁸Si le modèle d'analyse des données 3.15 fonctionne pour n'importe quel type de source asynchrone (continu ou impulsionnelle), les prédictions théoriques du chapitre 1 nécessitent un petit changement dans la forme de P_2 pour une source impulsionnelle.

avec le *GAP-Optique* de l'Université de Genève, le développement et la réalisation expérimentale d'un nouveau protocole ont été possibles. Nous nous proposons de découvrir dans la prochaine section le principe de la mesure ainsi que les principaux résultats expérimentaux.

3.2 La mesure directe du $g^{(2)}(0)$ en mode asynchrone

Nous allons commencer par rappeler le principe de la mesure de $g^{(2)}(\tau)$ pour les sources de photons uniques à la demande. Le montage, de type Hanbury-Brown & Twiss [66], utilisé pour cette mesure est identique à celui précédemment décrit, mais le caractère *asynchrone* de la source « brouille », en quelque sorte, l'accès aux informations. Nous montrerons alors qu'une modification du montage classique permet d'avoir finalement accès à ces informations en mode asynchrone.

3.2.1 Rappel sur la mesure expérimentale de $g^{(2)}(\tau)$

La caractérisation des sources de photons uniques à la demande requiert un montage de type *HBT*, identique au montage 3.3 utilisé dans la section précédente. Ce dernier nous donnait accès aux probabilités $\mathbb{P}_i$ d'avoir i détections et à partir de cette statistique nous remontions, via un modèle théorique, aux probabilités P_i d'avoir i photons annoncés. L'estimation de $g^{(2)}(0) \approx \frac{2P_2}{P_1^2}$ s'ensuivait alors naturellement.

Pourtant, l'utilisation d'un *TAC* permet bien plus que le simple enregistrement des coïncidences entre deux APDs. Il faut revenir à la section 2.3 du chapitre 2 pour se rappeler que « *mieux qu'une simple porte ET, le TAC est capable de donner, via le couplage avec une carte d'acquisition et un PC, l'histogramme temporel des événements qui arrivent en coïncidence sur ses deux entrées* ». Placé entre les APD 1 et 2, il devrait donc permettre de tracer la *fonction d'autocorrélation d'ordre deux*, qui normalisée, aboutit à la valeur de $g^{(2)}(\tau)$:

$$g^{(2)}(\tau) = \frac{\langle I_1(t)I_2(t+\tau) \rangle}{\langle I_1(t) \rangle \langle I_2(t) \rangle} \quad (3.22)$$

Essayons de comprendre l'origine et la forme de la courbe d'autocorrélation d'ordre deux obtenue dans le cas des sources de photons uniques à la demande.

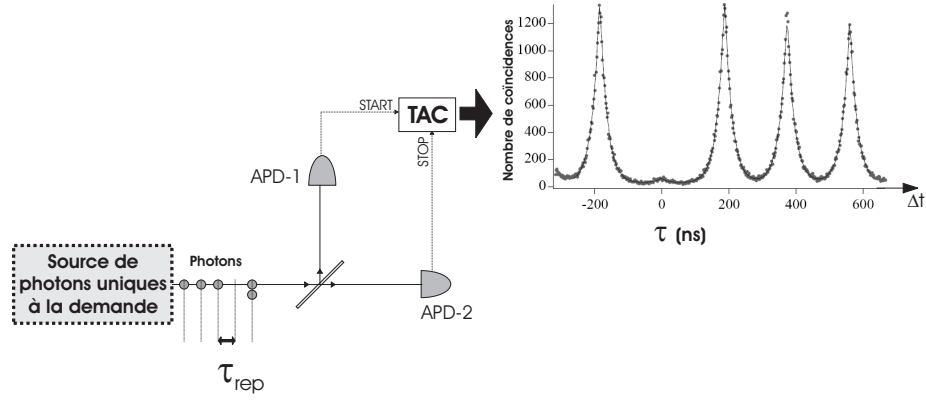


FIG. 3.5 – Montage de type « Hanbury-Brown & Twiss » permettant de tracer l'histogramme temporel des coïncidences, proportionnel à la *fonction d'autocorrélation d'ordre deux* $g^{(2)}(\tau)$. Nous avons noté τ_{rep} le taux de répétition de la source de photons unique à la demande. L'histogramme présenté ici est extrait de la référence [33].

Dans le cas du montage 3.5, le *TAC* accumule dans un histogramme les intervalles de temps qui séparent deux événements successifs arrivant entre ses entrées *start* et *stop*. Ainsi, deux événements simultanés seront rangés dans la case $\Delta t = 0$, tandis que deux événements séparés par τ seront rangés dans la case $\Delta t = \tau$ et ainsi de suite...

Nous sommes en présence d'une source qui ne peut émettre qu'à intervalles réguliers (τ_{rep}) 0, 1 ou 2 photons. Ainsi, la période des pics sur la courbe 3.5 correspond à la période de répétition de la source⁹.

Si la source émet deux photons simultanés (P_2), le *TAC* va comptabiliser un événement à $\Delta t = 0$. Le pic central va se remplir progressivement...

$\Rightarrow$ La hauteur du pic central, $\langle I_1(t)I_2(t) \rangle$, est proportionnelle à P_2 .

À présent, si elle émet deux photons uniques successifs ($P_1 \times P_1$), le *TAC* va comptabiliser un événement à $\Delta t = \tau_{rep}$. Notons qu'il y a un pic à droite comme à gauche du pic central dû à la probabilité 1/2 que le premier photon unique aille d'abord vers l'entrée *start* ou *stop* du *TAC*, que ce dernier interprète alors comme un temps positif ou négatif...

⁹L'élargissement des pics, quant à lui, correspond à la durée de vie du niveau excité qui est à l'origine de l'émission du photon unique [33].

$\Rightarrow$ La hauteur du pic adjacent de droite, $(\langle I_1(t)I_2(t + \tau_{rep}) \rangle)$, est proportionnelle à $\frac{P_1^2}{2}$.

Enfin, les autres pics à $\Delta t = n \times \tau_{rep}$ correspondent au cas où un photon unique a été enregistré sur le *start*, mais aucun sur le *stop* les $n - 1$ coups suivants.

$\Rightarrow$ Le caractère unique des photons, émis par une source de photon uniques, sera donc mis en évidence par l'absence du pic central ($\Delta t = 0$).

Notons une astuce qui nous permet de lire directement sur l'histogramme, la valeur normalisée de $g^{(2)}(0)$. En supposant qu'il n'y ait aucune corrélation entre deux photons successivement émis par la source et que la probabilité d'avoir un photon unique est relativement faible, nous avons

$$g^{(2)}(\tau_{rep}) \approx 1 \quad \Leftrightarrow \quad \langle I_1(t)I_2(t + \tau_{rep}) \rangle \approx \langle I_1(t) \rangle \langle I_2(t) \rangle$$

et le principal intérêt de l'histogramme est de pouvoir alors observer graphiquement la valeur de la fonction d'autocorrélation d'ordre deux $g^{(2)}(0)$ en faisant simplement au rapport des aires du pic central et de l'un des deux pic adjacents. On a donc :

$$g^{(2)}(0) = \frac{\langle I_1(t)I_2(t) \rangle}{\langle I_1(t) \rangle \langle I_2(t) \rangle} \approx \frac{\langle I_1(t)I_2(t) \rangle}{\langle I_1(t)I_2(t + \tau_{rep}) \rangle}$$

Notons que cette méthode pour calculer $g^{(2)}(0)$ n'est valable qu'à la condition où $g^{(2)}(\tau_{rep}) \approx 1$ et que, pour certains types de sources [36], normaliser la courbe expérimentale reste le seul moyen pour connaître la valeur exacte de $g^{(2)}(0)$.

Depuis le début du manuscrit, nous associons aux sources poissonniennes un $g^{(2)}(0)=1$. Il nous semble intéressant de voir ici à quoi ressemble un histogramme expérimental associé à une telle source. Nous avons donc extrait de la référence [33] un histogramme expérimental typique d'une source cohérente atténuée...

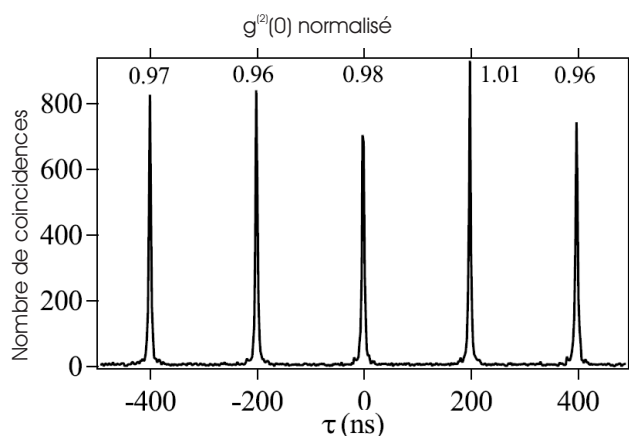


FIG. 3.6 – Histogramme expérimentale donnant accès à $g^{(2)}(\tau)$ pour un laser atténué, source typiquement poissonnienne. Notons que cette courbe est extraite de la référence [33].

À la vue des explications précédentes, il paraît très important de tracer cette courbe dans notre configuration.

S'il ne fait aucun doute qu'une source de paires de photons utilisée en régime continu constitue une source asynchrone, nous insistons sur le fait que l'utilisation d'un laser impulsif ne la rend pas synchrone pour autant. Même si nous possédons une information supplémentaire sur les temps de création des paires de photons (grâce au trigger issu du laser de pompe), nous ne connaissons toujours pas l'intervalle de temps exact qu'il y a entre deux paires successives. Nous sommes donc toujours dans une configuration de communication asynchrone et montrerons que dans un cas comme dans l'autre il ne sera pas possible d'observer « classiquement » la fonction $g^{(2)}(0)$.

Étude la forme des histogrammes obtenus dans le cas de sources asynchrones

Nous allons, dans le premier temps, nous pencher sur la configuration rencontrée depuis le début de ce manuscrit :

★ UNE SOURCE DE PAIRES DE PHOTONS EN RÉGIME DE POMPAGE CONTINU

Dans notre cas précis, la source a un taux de répétition moyen constant (cf. section 1.2), mais l'intervalle de temps entre deux photons successifs est aléatoire. De ce fait, l'utilisation d'un *TAC* conventionnel ne nous permet pas d'identifier l'équivalent d'un pic central, $\langle I_1(t)I_2(t) \rangle$, ni d'un pic adjacent

$\langle I_1(t)I_2(t + \tau_{rep}) \rangle$. Nous avons représenté sur la figure 3.7 la courbe obtenue expérimentalement : le pic central est identifiable grâce au temps mort du générateur de délai qui ne peut fournir aux APD-*InGaAs* deux *triggers* successifs dans un temps inférieur à $1,3 \mu s$. Par contre, le temps séparant deux événements successifs étant aléatoire, il est impossible d'identifier une valeur de $\langle I_1(t)I_2(t + \tau_{rep}) \rangle$ à partir de la courbe située à droite du pic.

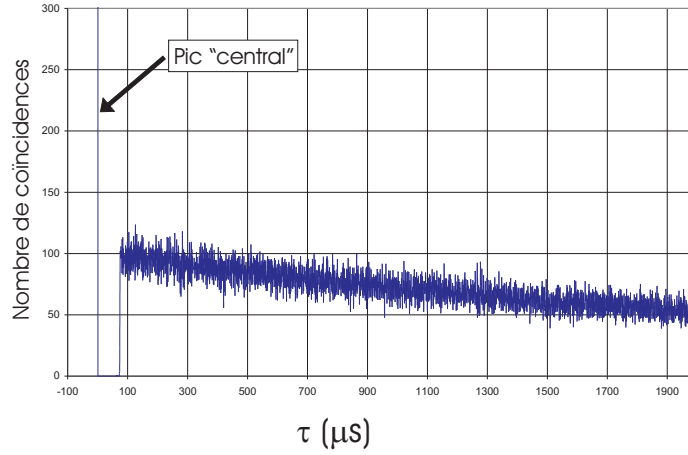


FIG. 3.7 – Représentation de l'histogramme obtenu sur un TAC avec une source de photons unique asynchrone fonctionnant en régime continu.

Pour résoudre le problème précédent, nous sommes tenté de dire qu'il suffit simplement d'imposer une base de temps fixe au système en utilisant un laser impulsionnel...

★ UNE SOURCE DE PAIRES DE PHOTONS EN RÉGIME DE POMPAGE IMPULSIONNEL

Intuitivement et à la vue de l'explication précédente, nous devrions pouvoir identifier un pic central et un pic adjacent. Ce sera effectivement le cas, pourtant la mesure de $g^{(2)}(0)$ ne reflétera pas pour autant les performances de la source. En effet, le véritable obstacle à la mesure est bel et bien le régime asynchrone quelque soit le mode de pompage du cristal non-linéaire. Pour nous en rendre compte, reprenons la description des sources de photons annoncés faite dans le prologue ainsi que le schéma qui l'accompagnait :

« En partant d'une source poissonnienne atténuée impulsionnelle, nous n'allons pas réduire le nombre d'impulsions contenant deux photons mais chercher à obtenir une information sur celles qui contenaient au moins un

photon puis de ne garder que celles là. Bien évidemment la probabilité d'avoir un photon par impulsion est toujours égale à 0,1 en moyenne, mais la probabilité d'avoir au moins un photon par impulsion annoncée est artificiellement égale à un. »

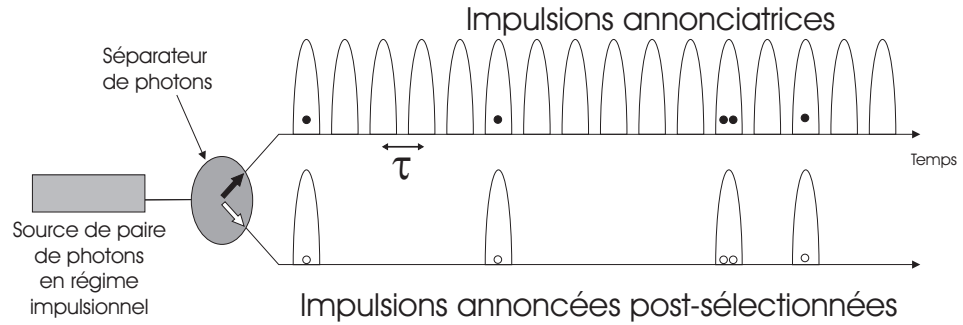


FIG. 3.8 – Principe de la post-sélection des impulsions contenant au moins un photon. Les impulsions originales ont un temps de répétition τ et une probabilité $P_1^{poisson} = 0,1$, tandis que les impulsions post-sélectionnées présentent un temps de répétition aléatoire (qui reste cependant un multiple de τ) avec une probabilité $P_1^{annoncé} \approx 1$.

Comme l'explique la définition, la probabilité que la première impulsion, qui suit l'arrivée d'un photon unique, contiennent elle-même un photon est purement poissonnienne et la technique d'annoncer l'arrivée d'un photon consiste juste à ne pas allumer le détecteur pour rien. Alors d'un point de vue statistique, la probabilité d'avoir un événement dans le pic central est donc bien égale à $P_2^{poisson}$, tandis que celle d'avoir un événement dans le premier pic adjacent à $\Delta t = \tau$ est directement proportionnel à la probabilité que deux impulsions successives contiennent un photon, c'est-à-dire $\frac{1}{2} \times P_1^{poisson} \times P_1^{poisson}$. Selon toute logique, nous observerions alors une courbe identique à celle d'une source poissonnienne représentée sur la figure 3.9.

Une telle source de source est effectivement une source de photon uniques mais le caractère unique des photons émis est masqué par un taux de répétition que nous pourrions qualifier de « poissonnien ». Le problème est simplement que nous n'observons pas cette courbe dans la bonne « base de temps » et nous allons proposer un moyen de visualiser la fonction d'autocorrélation en régime asynchrone. L'idée est venue du groupe du GAP-Optique à Genève où nous avons réalisé l'expérience [51].

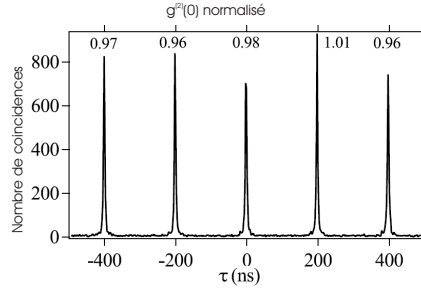


FIG. 3.9 – Histogramme expérimental pour une source typiquement poissonnienne.

3.2.2 Le principe de la mesure asynchrone de $g^{(2)}(0)$

Quelque soit le mode de pompage utilisé, le problème rencontré avec un *TAC* conventionnel, est la base de temps qui empêche de classer correctement les événements. En effet, pour les sources asynchrones, la durée entre deux impulsions successives est aléatoire, mais n'a aucune importance du point de vue du détecteur déclenché, car de son point de vue, il n'y a que des allumages qu'il référence $a_0, a_1, \dots, a_n \dots$ sans information sur le temps où l'événement s'est produit. Il faudrait donc pouvoir utiliser un *TAC* « spécial » qui raisonnerait de manière « asynchrone », c'est-à-dire sans attacher d'importance au temps effectif qui sépare deux photons uniques successifs.

Aussi, nous pourrions tracer la fonction d'autocorrélation $g^{(2)}(\tau)$, à condition de s'affranchir de notre référence temporelle habituelle. Pour cela, il va falloir enregistrer en temps réel tous les événements arrivant sur les APDs 1 et 2 dans un fichier et, une fois sauvegardé, nous le relirons et retranscrivons les événements qui se sont déroulés, en supprimant toute notion de temps entre deux détections successives. Maintenant les événements n'ont plus de label temporel, mais sont simplement étiquetés $(n+1)$ ou $(n-1)$ par rapport au $n^{\text{ème}}$ événement contenu dans le fichier.

3.2.3 Le montage expérimental

La source que nous utiliserons pour la mesure aurait pu être celle développée à Nice, mais pour des raisons pratiques, nous utiliserons la source de Genève qui est identique à la notre d'un point de vue fonctionnement. Le montage permettant de tracer l'histogramme $g^{(2)}(\tau)$, quant à lui, nécessitera quelques aménagements comme le remplacement du *TAC* par une carte d'acquisition spéciale, couplée à un algorithme « maison » de traitement des données capable de tracer la fonction d'autocorrélation.

La source de photons uniques Genevoise

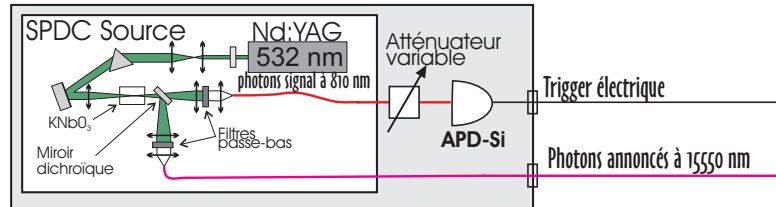


FIG. 3.10 – Source de photons uniques annoncés utilisée à Genève.

Nous disposons à Genève d’une source de photons uniques annoncés qui fonctionne sur le même principe que celle développée à Nice, avec toutefois quelques différences. Nous nous proposons de les souligner succinctement à l’aide de la figure 3.10 et nous précisons que le lecteur pourra trouver une description détaillée de la source dans la référence [81] :

- Un laser *Nd : YAG* à 532 nm , avec une puissance de sortie continue de 60 mW , est utilisé pour créer les paires de photons à 1550 nm et 810 nm au travers d’un cristal massif de KNbO_3 .
- Les paires de photons sont séparées à l’aide d’un miroir dichroïque, qui les envoie respectivement vers deux systèmes de couplages indépendants utilisés pour récolter séparément les photons dans des fibres optiques monomodes. Notons l’utilisation de deux filtres passe-bas, devant chaque lentille de couplage, pour supprimer les photons de pompe résiduels.
- La longueur d’onde du photon signal, centrée sur 810 nm , permet d’utiliser une APD en *silicium* qui présentent une très bonne efficacité de détection ($\eta_{\text{Si}} \approx 0,6$), ainsi qu’un taux de coups sombres très faibles (de l’ordre d’une centaine de s^{-1}).
- La présence d’un atténuateur variable permettra de s’assurer que la photodiode en silicium ne sature pas, nous autorisant ainsi à explorer sur une plus grande plage de puissance les performances de la source.

Reposant sur l’optique massive, cette source présente des possibilités de réglage différentes de celle de Nice, mais d’un point de vue extérieur elle lui est complètement semblable se résumant, elle aussi, à une boîte noire disposant de deux sorties : une première, électrique, pour l’émission d’un trigger annonçant la présence d’un photon unique sur la seconde, constituée d’une fibre optique télécom monomode.

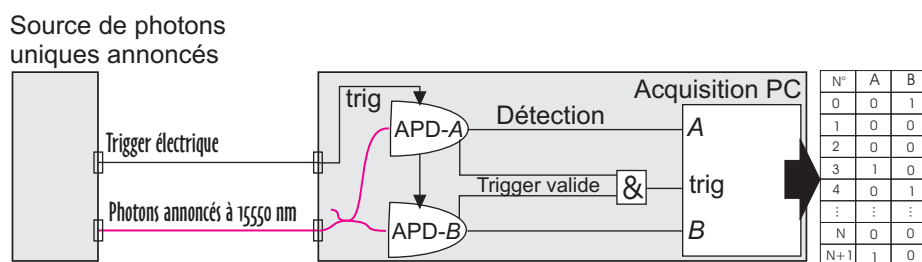


FIG. 3.11 – Montage de type Hanbury-Brown & Twiss, couplé au système d'acquisition « spécial » utilisé pour remplacer le *TAC*.

Le système de mesure

La principale différence entre le nouveau montage de la figure 3.11 et le précédent de la figure 3.3 réside dans la disparition du *TAC* au profit d’une carte d’acquisition *National Instruments* couplée à un ordinateur. Les APDs *A* et *B* sont les mêmes que celles utilisées au LPMC à Nice, à une différence près : à Genève, les APD-*InGaAs* fournies par *Id-Quantique* possèdent chacune leur propre générateur de délai interne qui applique seulement un temps mort de $10\,\mu\text{s}$ après une détection¹⁰ (due à un photon ou à un coup sombre). Celles de Nice en étaient dépourvues, l’unique générateur de délai était alors inclus dans la source et appliquait un temps mort, entre deux allumages, aux APDs 1 et 2 qu’il y ait eu détection ou non. Donc à Genève, il se pourrait maintenant qu’une des deux APDs ne soit pas déclenchée malgré la réception d’un trigger issu de la source, si ce dernier intervient à moins de $10\,\mu\text{s}$ après une détection. Pour s’assurer que nous faisons bien de « vraies » mesures simultanées et qu’une des deux APDs n’a pas manqué un photon à cause du temps mort, nous disposons d’une sortie, appelée « *trigger valide* », qui fournit une valeur logique 1 si le trigger a été accepté par l’APD ou un 0 s’il est arrivé durant l’application d’un temps mort. Il faudra donc au final que les APDs *A* et *B* aient validé toutes les deux leur trigger pour que la mesure soit prise en compte par la carte d’acquisition.

La carte d'acquisition asynchrone

La carte est une « High speed digital I/O device » capable d'acquérir 32 *bits* d'information simultanément à une fréquence maximum de 20 *MHz*. Inutile de dire que nous n'irons pas jusque là, car les APD-*InGaAs* sont données pour une fréquence de trigger maximale de 100 *kHz* (pour un temps

¹⁰Pour s'assurer de l'élimination des *after-pulses* dans les APDs qui faussent les mesures.

mort de $10\ \mu s$). La carte possède une horloge interne réglable qui définit le départ d'une acquisition, mais possède aussi une entrée pour un déclenchement externe de l'acquisition. C'est mode de fonctionnement qui va plus particulièrement nous intéresser car c'est ici que viendra se connecter la sortie de la porte *ET* qui témoigne de la validité simultanée des triggers issus des APDs. De cette façon, nous nous assurons que lors de l'acquisition d'une mesure, les deux APD-*InGaAs* sont disponibles pour détecter un ou plusieurs éventuels photons incidents.

Le diagramme de la figure 3.12 permet de saisir le principe de la conversion *analogique/numerique* lors d'une acquisition. La carte a un fonctionnement *asynchrone* qui lui permet de s'affranchir de la base de temps pour faire ses mesures. Elle attend l'arrivée d'un signal déclencheur (TRIGGER) pour commencer l'acquisition durant laquelle elle va enregistrer la valeur logique des entrées *A* et *B* jusqu'à la disparition du trigger. Pour qu'une acquisition soit validée par la carte, il faut que le signal mesuré soit constant durant toute la durée de l'acquisition. Pour cela il est important que les signaux électriques envoyés par les APDs *A* et *B* soient beaucoup plus longs que ceux du trigger afin d'éviter des erreurs d'acquisition. Nous choisirons des impulsions de $20\ ns$ de large pour le trigger et de $100\ ns$ pour les APDs. Il est intéressant de noter que ce mode de fonctionnement est en tout points identique à la définition des communications asynchrones faite dans la section 1.2.

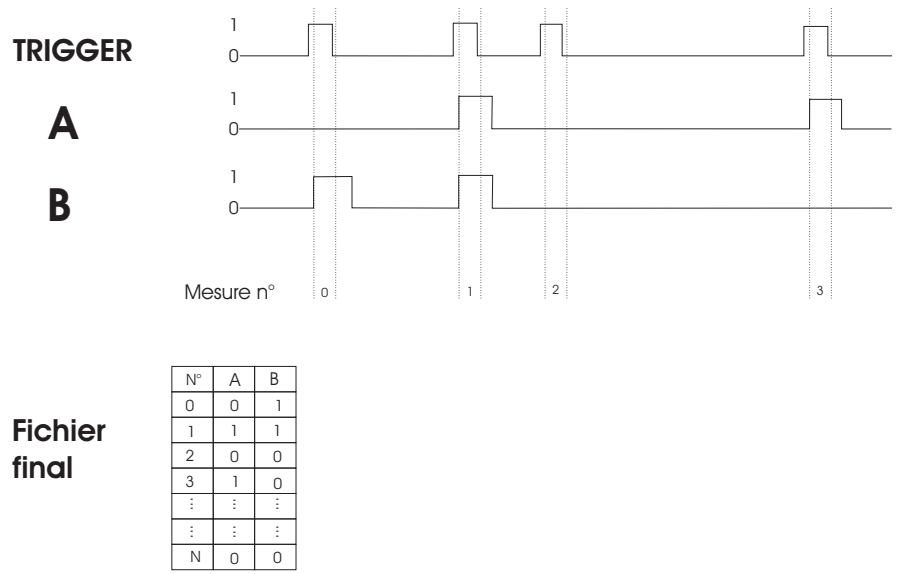


FIG. 3.12 – Diagramme de fonctionnement de la carte d'acquisition « High speed digital I/O device ».

3.2.4 Le programme d'acquisition et de traitement des données

À ce stade, la mesure nous fournit un fichier qu'il nous faut ensuite traiter « à la façon d'un TAC ». Le but du programme va être de simuler exactement le fonctionnement d'un TAC lors de la lecture des données. Commençons d'abord par rappeler le fonctionnement d'un TAC, ce qui nous permettra ensuite d'expliquer le fonctionnement du programme.

LE TAC (TIME TO AMPLITUDE CONVERTER)

Il dispose de deux entrées *start* et *stop* que nous nommerons respectivement *A* et *B*.

- La finalité du TAC est de mesurer le temps qui sépare deux événements (ou bits), l'un arrivant sur l'entrée *A* et l'autre sur l'entrée *B*.
- Pourtant si un événement se présente en premier sur l'entrée *B* et ensuite sur l'entrée *A*, alors le TAC devrait considérer un temps négatif. En réalité, il ne peut pas compter un temps négatif, mais en retardant d'un délai τ l'entrée *stop*, tous les intervalles de temps compris entre $-\tau$ et $+\infty$ seront accessibles. C'est pour cette raison que ses entrées portent les noms *start* et *stop* et ne peuvent inverser leur rôle.
- Après la réception d'un bit *start* le TAC est aveugle à tous les *start* suivant tant qu'un bit *stop* n'a pas été enregistré. Ces *start* refusés constituent les « *starts non-valides* » qui doivent rester faibles pour qu'une mesure reflète la réalité physique.
- Pour justement éviter d'accumuler trop d'« *starts non-valides* », le TAC dispose d'une consigne temporelle. Cette dernière définit le temps au-delà duquel s'il n'y a pas eu de signal *stop*, le TAC se réinitialise et recommence une nouvelle acquisition, sans tenir compte de la précédente.

L'ALGORITHME DE SIMULATION DU TAC DÉVELOPPÉ À GENÈVE

Les deux entrées *A* et *B* du TAC sont l'équivalent des colonnes *A* et *B* du fichier enregistré pour le programme.

- Le programme devra entreprendre la lecture des valeurs logiques des colonnes *A* et *B* pour les n acquisitions que contient le fichier. En lisant le fichier de l'acquisition 0 jusqu'à n , il va chercher le premier 1 qu'il va rencontrer dans la colonne *A* ou *B*. Imaginons qu'il le trouve d'abord dans la colonne *A*, il va alors chercher la valeur 1 suivante dans la colonne *B*. Le nombre d'acquisitions entre ces deux événements

permettra de déterminer à quel pic appartiendra cette mesure. Il va s'en dire que le programme devra, durant ce temps, être aveugle aux événements qui pourraient avoir lieu sur la voie A . Selon qu'un 1 est détecté en premier sur la voie A ou B , le programme devra associer au résultat respectivement un signe $(+)$ ou $(-)$.

- Au cas où le programme rencontre deux valeurs 1 simultanément, il devra ranger cet événement dans le pic central.
- Nous pouvons attribuer au programme une consigne (en nombre d'acquisitions maximum) au-delà de laquelle il doit arrêter la recherche du bit de *stop* et se ré-initialiser.
- Enfin, le programme fournit à l'utilisateur les valeurs respectives de $\mathbb{P}_A$ et $\mathbb{P}_B$, qui correspondent aux probabilités d'avoir une détection dans les voies A et B respectivement. Elles sont calculées simplement en sommant le nombre de 1 dans chaque colonne normalisé par le nombre total de trigger.

Au final, le programme trace un histogramme, représentant le nombre de fois où il a rencontré des 1 pour A et B simultanés, ou séparés par i acquisitions. La figure 3.13 représente justement l'histogramme que nous obtenons typiquement, accompagné des probabilités $\mathbb{P}_A$ et $\mathbb{P}_B$:

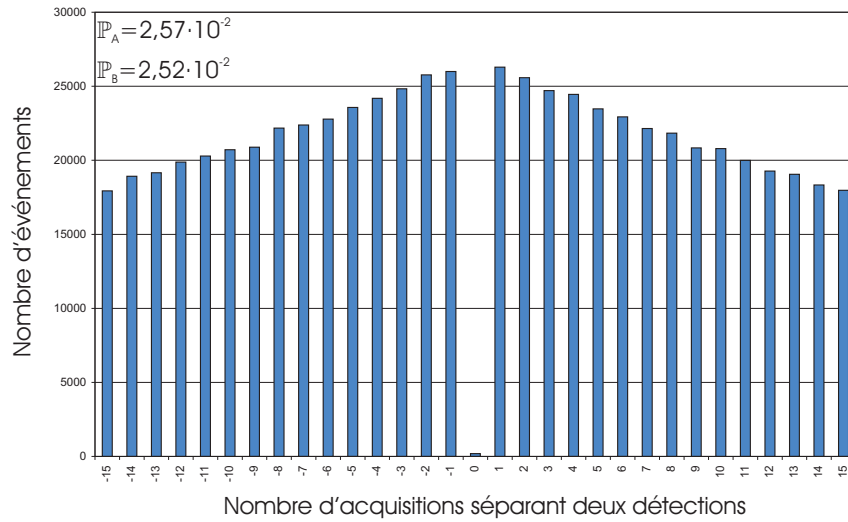


FIG. 3.13 – Histogramme du nombre d'événements séparés par 0, 1 ..., 15 acquisitions. $\mathbb{P}_A$ et $\mathbb{P}_B$ représentent les probabilités d'avoir une détection dans la voie A et B .

À cet instant, rien ne nous autorise à simplement faire le rapport entre le pic central et le pic adjacent pour calculer $g^{(2)}(0)$. En effet, nous n'avons

aucun moyen de savoir si $g^{(2)}(1)$ est bien égal à la valeur 1 qui, rappelons-le, est la condition pour pouvoir utiliser ce raccourci. Nous allons donc devoir normaliser la courbe obtenue et le paragraphe suivant sera donc l'occasion de déterminer la bonne méthode à suivre pour y parvenir.

3.2.5 La théorie

Rappelons la définition du $g^{(2)}(0)$ que nous avons adopté : *la fonction d'autocorrélation d'ordre deux permet de quantifier les performances de la source étudiée par rapport à une source poissonnienne équivalente. Elle vaut donc simplement le rapport entre le P_2 de la source étudiée et celui d'une source poissonnienne possédant le même P_1 .*

Pour pouvoir calculer $g^{(2)}(0)$, il nous faut déterminer quelle serait la forme de l'histogramme 3.13 (et surtout la valeur du point central) pour une source poissonnienne équivalente à la notre. Cette source équivalente présenterait donc des probabilités $\mathbb{P}_A$ et $\mathbb{P}_B$ identiques à celles de notre source et l'histogramme $M_{poisson}(n)$ qui lui est associée nous permettra de normaliser correctement l'histogramme expérimental $M_{source}(n)$ associé à notre source de photons.

En définissant par $M_{source}(0)$ la valeur de la barre centrale de l'histogramme qui correspond à $\langle I_A(t)I_B(t) \rangle$, nous devons définir $M_{poisson}(0)$ qui sera l'équivalent de $\langle I_A(t) \rangle \langle I_B(t) \rangle$. Les valeurs de $M_{poisson}(n)$ serviront à normaliser $M_{source}(n)$ et nous pourrons calculer $g^{(2)}(0)$ comme suit :

$$g^{(2)}(0) = \frac{\langle I_A(t)I_B(t) \rangle}{\langle I_A(t) \rangle \langle I_B(t) \rangle} = \frac{M_{source}(0)}{M_{poisson}(0)}$$

Écrivons que $M_{poisson}(n)$ correspond à une détection sur la voie A puis seulement à une détection sur la voie B , n acquisitions plus tard. Cela se traduit mathématiquement par l'expression :

$$M_{poisson}(n) = C_a \mathbb{P}_A (1 - \mathbb{P}_B)^n \mathbb{P}_B \quad (3.23)$$

Pour les valeurs négatives de n la formule devient :

$$M_{poisson}(n) = C_a \mathbb{P}_B (1 - \mathbb{P}_A)^{|n|} \mathbb{P}_A \quad (3.24)$$

avec C_a une constante de normalisation qui dépend principalement de la durée de l'acquisition. Nous voyons que les deux courbes se rejoignent pour $n = 0$ où $M_{poisson}(0)$ vaut simplement la probabilité d'obtenir deux détections simultanées :

$$M_{poisson}(0) = C_a \mathbb{P}_A \mathbb{P}_B \quad (3.25)$$

La constante C_a est déterminée en ajustant l'histogramme 3.13 avec les formules 3.23 et 3.24. Nous calculons ensuite $g^{(2)}(0)$ comme suit :

$$g^{(2)}(0) = \frac{M_{source}(0)}{M_{poisson}(0)} = \frac{M_{source}(0)}{C_a \mathbb{P}_A \mathbb{P}_B} \quad (3.26)$$

Le lecteur pourra trouver sur la figure 3.14 une représentation de l'histogramme expérimental $M_{source}(n)$ accompagnée de la courbe $M_{poisson}(n)$.

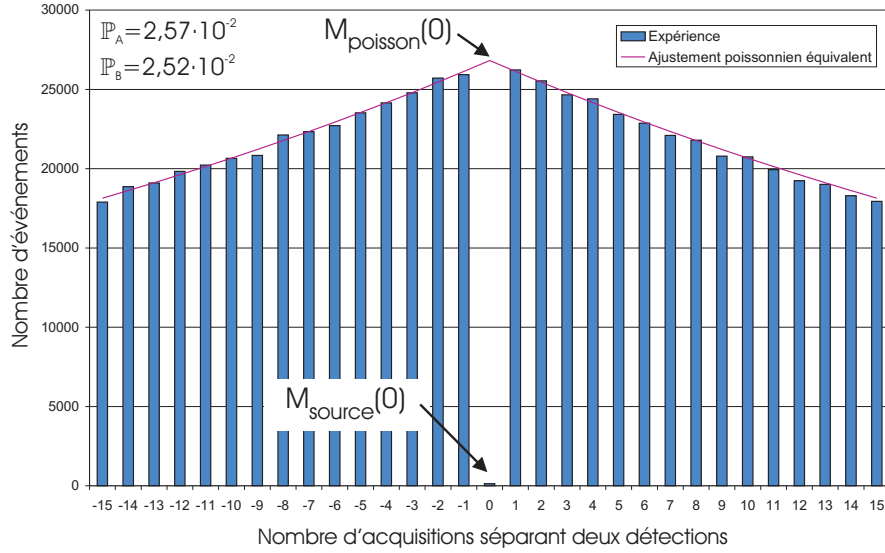


FIG. 3.14 – Histogramme du nombre d'événements séparés par 0, 1 ..., 15 acquisitions et son ajustement poissonnien. $\mathbb{P}_A$ et $\mathbb{P}_B$ représentent les probabilités d'avoir une détection dans la voie A et B .

L'intérêt de tracer cette courbe est de pouvoir calculer $g^{(2)}(0)$ en n'ayant besoin de connaître aucun paramètre expérimental. Contrairement à la mesure faite dans la section précédente, où les efficacités de détection $\eta_{1,2}$ des APDs 1 et 2 intervenaient dans le calcul de P_1 et P_2 , donc de $g^{(2)}(0)$, ici celles-ci n'interviennent pas. En effet, nous travaillons avec les probabilités d'avoir une détection (et non d'avoir un photon) qui contiennent implicitement les efficacités de détections des APDs A et B . Cette non-dépendance de $g^{(2)}(0)$ vis à vis de $\eta_{1,2}$ est facilement compréhensible à partir de sa formule simplifiée $\frac{2P_2}{P_1^2}$. En effet, la probabilité P_2 fait intervenir l'efficacité de détection η

au carré, qui se simplifie avec η^2 contenu dans P_1^2 . Nous possédons donc un moyen très simple de faire une mesure absolue de $g^{(2)}(0)$, qui ne dépend pas des conditions expérimentales.

3.2.6 La mesure expérimentale et l'évaluation du bruit

L'approche théorique a été faite en considérant que les mesures étaient toutes dues à la détection de photons. Pourtant cette expérience n'est pas différentes des autres, les coups sombres apportent aussi leur contribution. Les valeurs $M_{source}(n)$ expérimentales ne sont donc pas représentatives des performances de la source et incluent les erreurs du système de mesure. Il faut comprendre que les erreurs du système s'ajoutent aux « valeurs mesurées » :

$$M_{source}^{brut}(n) = M_{source}^{net}(n) + M_{bruit}(n) \quad (3.27)$$

Selon toute logique, $\mathbb{P}_A^{brut}$ et $\mathbb{P}_B^{brut}$ ne sont alors plus représentatifs des performances de la source et leur utilisation pour ajuster la courbe par une source poissonnienne n'est plus correcte. Il faut alors décomposer l'ajustement poissonnien comme suit :

$$M_{poisson}^{brut}(n) = M_{poisson}^{net}(n) + M_{bruit}(n) \quad (3.28)$$

En prenant soin de définir auparavant $\mathbb{P}_A^{dc}$ et $\mathbb{P}_B^{dc}$ les probabilités d'avoir un coup sombre dans les APDs A et B , nous pouvons alors introduire les probabilités nettes de bruit $\mathbb{P}_A^{net}$ et $\mathbb{P}_B^{net}$ qui valent simplement :

$$\mathbb{P}_A^{net} = \mathbb{P}_A^{brut} - \mathbb{P}_A^{dc} \quad \text{et} \quad \mathbb{P}_B^{net} = \mathbb{P}_B^{brut} - \mathbb{P}_B^{dc}$$

Maintenant, la valeur corrigée de la fonction d'autocorrélation d'ordre deux normalisée se calcule ainsi :

$$g^{(2)}(0)^{net} = \frac{M_{source}^{net}(0)}{M_{poisson}^{net}(0)} = \frac{M_{source}^{brut}(0) - M_{bruit}(0)}{C_a \mathbb{P}_A^{net} \mathbb{P}_B^{net}} \quad (3.29)$$

Toute l'art de réaliser une mesure de $g^{(2)}(0)^{net}$, va consister à mesurer correctement $\mathbb{P}_A^{dc}$ et $\mathbb{P}_B^{dc}$ et les utiliser judicieusement afin d'estimer $M_{bruit}(0)$.

Calcul de la valeur de $M_{bruit}(0)$

Le bruit dans les APDs A et B est un processus aléatoire régi par une distribution poissonnienne, ce qui nous autorise à utiliser les formules 3.23 et 3.24 pour réaliser un ajustement du bruit à partir des probabilités $\mathbb{P}_A^{dc}$

et $\mathbb{P}_B^{dc}$ qui correspondent respectivement aux probabilités d'avoir un coup sombre dans les APDs A et B . Expérimentalement, en empêchant les photons d'accéder aux APDs A et B , nous pouvons mesurer les valeurs de $\mathbb{P}_A^{dc}$ et $\mathbb{P}_B^{dc}$ et vérifier par la même occasion le comportement poissonnien du bruit dans les APDs à l'aide d'un ajustement représenté sur la figure 3.15.

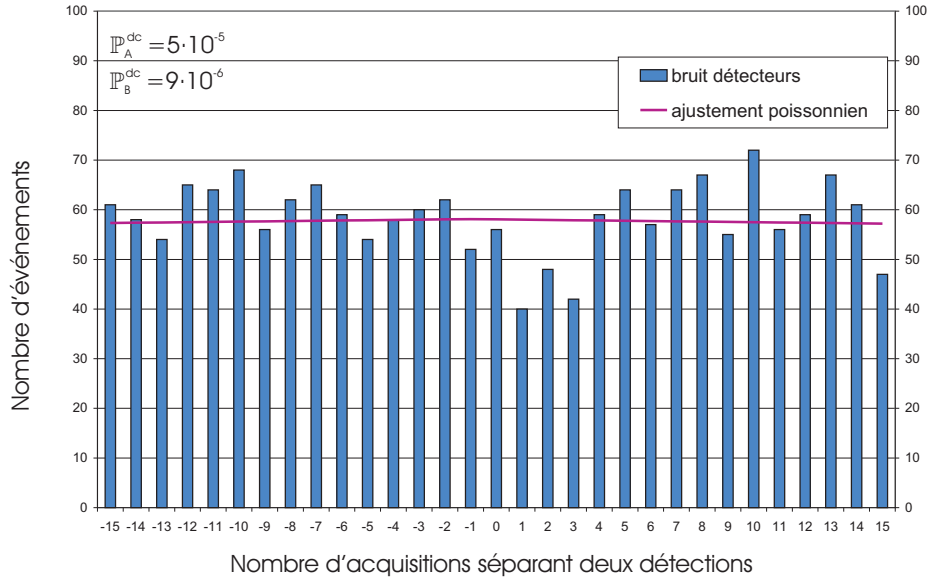


FIG. 3.15 – Histogramme d'anti-coïncidences pour le bruit dans les APDs A et B et son ajustement poissonnien. $\mathbb{P}_A^{dc}$ et $\mathbb{P}_B^{dc}$ représentent la probabilité d'avoir un coup de bruit dans les voies A et B .

Au premier abord, nous pourrions être tenté de dire simplement que la part du bruit, dans l'histogramme 3.14, s'élève à $M_{bruit}(0) = C_a \mathbb{P}_A^{dc} \times \mathbb{P}_B^{dc}$.

Malheureusement ce n'est aussi simple !

Aidons nous de la figure 3.16 pour comprendre exactement à quelle quantité correspond $M_{bruit}(0)$. Les deux ensembles supérieurs représentent les probabilités brutes $\mathbb{P}^{brut}$ avec à l'intérieur les parties qui proviennent du bruit $\mathbb{P}^{dc}$. Sur la partie inférieure de la figure, *toutes* les parties grisées, qui correspondent à l'intersection des deux ensembles, représentent exactement les coïncidences brutes : $\mathbb{P}_A^{brut} \times \mathbb{P}_B^{brut}$. Cependant, les coïncidences nettes (provenant de « vrais » photons) sont seulement désignées par les ensembles gris clairs. Nous souhaitons donc retrancher non seulement la partie en noir

qui correspond à $\mathbb{P}_A^{dc}\mathbb{P}_B^{dc}$ mais aussi celles en gris foncé qui correspondent à $\mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} - 2\mathbb{P}_A^{dc}\mathbb{P}_B^{dc}$.

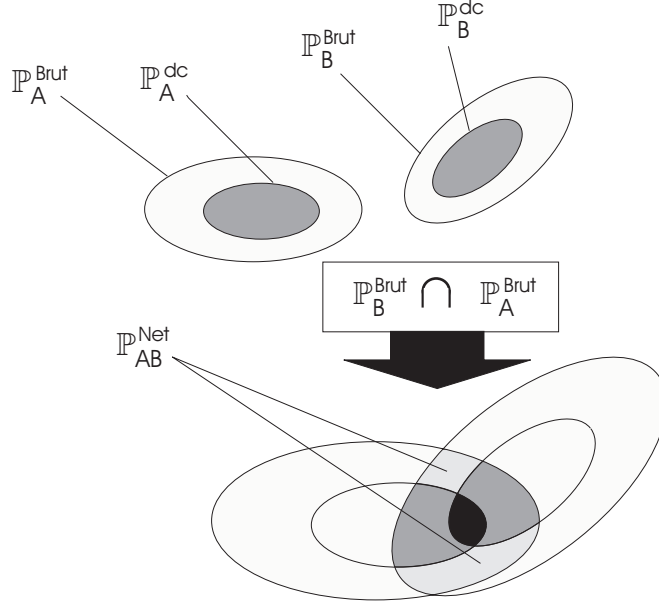


FIG. 3.16 – Représentation de la probabilité d'obtenir une vraie coïncidence à partir de deux ensembles qui contiennent chacun une part de bruit.

Finalement nous pouvons écrire

$$M_{bruit}(0) = C_a \left(\mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} - \mathbb{P}_A^{dc}\mathbb{P}_B^{dc} \right) \quad (3.30)$$

qui permettra de calculer la valeur de $g^{(2)}(0)$ nette de bruit par la formule,

$$g^{(2)}(0)^{net} = \frac{M_{source}(0) - C_a \left(\mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} - \mathbb{P}_A^{dc}\mathbb{P}_B^{dc} \right)}{C_a \left(\mathbb{P}_A^{brut} - \mathbb{P}_A^{dc} \right) \left(\mathbb{P}_B^{brut} - \mathbb{P}_B^{dc} \right)} \quad (3.31)$$

tandis que l'incertitude sera estimée à partir de la relation :

$$\begin{aligned} \frac{\Delta g^{(2)}(0)^{net}}{g^{(2)}(0)^{net}} = & \left[\frac{\Delta M_{source}(0)}{M_{source}(0) - C_a \left(\mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} - \mathbb{P}_A^{dc}\mathbb{P}_B^{dc} \right)} \right. \\ & + \frac{C_a \left(\Delta \mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{dc}\Delta \mathbb{P}_B^{brut} + \Delta \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} + \mathbb{P}_A^{brut}\Delta \mathbb{P}_B^{dc} + \Delta \mathbb{P}_A^{dc}\mathbb{P}_B^{dc} + \mathbb{P}_A^{dc}\Delta \mathbb{P}_B^{dc} \right)}{M_{source}(0) - C_a \left(\mathbb{P}_A^{dc}\mathbb{P}_B^{brut} + \mathbb{P}_A^{brut}\mathbb{P}_B^{dc} - \mathbb{P}_A^{dc}\mathbb{P}_B^{dc} \right)} \\ & \left. + \frac{\Delta \mathbb{P}_A^{brut} + \Delta \mathbb{P}_A^{dc}}{\mathbb{P}_A^{brut} - \mathbb{P}_A^{dc}} + \frac{\Delta \mathbb{P}_B^{brut} + \Delta \mathbb{P}_B^{dc}}{\mathbb{P}_B^{brut} - \mathbb{P}_B^{dc}} \right] \quad (3.32) \end{aligned}$$

Dans l'introduction de cette section, nous avons introduit la courbe d'autocorrélation d'ordre deux en expliquant que $g^{(2)}(0)$ pouvait se calculer en faisant le rapport du pic central sur le pic adjacent, l'un étant proportionnel à P_2 et l'autre à $\frac{P_1^2}{2}$. Nous avons souligné que cette méthode approximative n'était justifiée que dans le cadre de sources dont l'émission de deux photons successifs n'est pas corrélée ($g^{(2)}(1) = 1$). Dans notre cas, cette relation requiert une simplification qui permet uniquement de l'utiliser pour estimer approximativement $g^{(2)}(0)$. Pour cela reprenons la formule 3.26 qui permet de calculer exactement $g^{(2)}(0)$:

$$g^{(2)}(0) = \frac{M_{source}(0)}{C_a \mathbb{P}_A \mathbb{P}_B}$$

En écrivant la valeur de l'ajustement poissonnien pour la première barre de l'histogramme :

$$M_{source}(1) = M_{poisson}(1) = C_a \mathbb{P}_A (1 - \mathbb{P}_B) \mathbb{P}_B$$

Nous pouvons alors écrire :

$$g^{(2)}(0) = \frac{M_{source}(0)}{M_{source}(1)} (1 - \mathbb{P}_B)$$

Finalement, comme $\mathbb{P}_B \ll 1$, l'équation précédente se simplifie en

$$\boxed{g^{(2)}(0) \approx \frac{M_{source}(0)}{M_{source}(1)}} \quad (3.33)$$

qui correspond bien au rapport du pic central par le pic adjacent de droite au coefficient $(1 - \mathbb{P}_B)$ près.

3.2.7 Les résultats

La courbe 3.14 correspond à une mesure de la fonction d'autocorrélation des photons émis par la source et la courbe 3.15 est associée à la fonction d'autocorrélation du bruit dans les détecteurs. Cette dernière servira uniquement à corriger l'erreur due aux défauts du système de mesure dans l'estimation de $g_{net}^{(2)}(0)$.

Nous souhaitons ajouter le tableau 3.5 qui regroupe les données numériques relatives aux acquisitions des courbes 3.14 et 3.15. Comme le montage, utilisé ici est identique en tout point à celui utilisé à Nice nous pourrions calculer les probabilités P_0 , P_1 et P_2 pour ensuite remonter à $g^{(2)}(0)$ en utilisant la méthode développée dans la section précédente.

Temps d'intégration $\approx 1200\text{ s}$	Mesure moyenne	Incertitude
Puissance de pompe (mW)	6,0	$\pm 0,2$
Nombre total d'événements	116070700	± 10774
Coups sombres dans l'APD- <i>Silicium</i>	60101	± 245
$Total_A$	2986104	± 1728
$Total_B$	2927585	± 1711
$Totalcoincidence$	593	± 24
$M_{source}(0)$	191	± 14
$\mathbb{P}_A^{brut}$	0,02573	$\pm 0,00002$
$\mathbb{P}_B^{brut}$	0,02522	$\pm 0,00002$
$\mathbb{P}_A^{dc}$	0,00005	$\pm 7 \cdot 10^{-7}$
$\mathbb{P}_B^{dc}$	0,00001	$\pm 1 \cdot 10^{-6}$
C_a	40369897	—
Efficacité APD-1 η_A	0,101	$\pm 0,005$
Efficacité APD-2 η_B	0,101	$\pm 0,005$
Correction pertes coupleur 50/50	0,93 (<i>A</i>) et 0,85 (<i>B</i>)	$\pm 0,005$
Durée d'ouverture des APDs (ns)	5	$\pm 0,5$

TAB. 3.5 – Tableau récapitulatif des mesures expérimentales faites à Genève.

Nous insistons sur le fait que les résultats exprimés ici concernent la source Genevoise fonctionnant à une fréquence de triggers de 100 kHz (ce qui représente le maximum supporté par les APD-*InGaAs*) et que de plus, les performances sont données pour une durée d'allumage des APDs de 5 ns .

Estimation de $g^{(2)}(0)$ à partir des histogrammes 3.14 et 3.15

La courbe 3.14 présente clairement une disparition du pic central comparativement aux pics latéraux. Ce comportement est la signature de photons présentant un comportement « uniques ». Nous pouvons estimer graphiquement la valeur de $g_{net}^{(2)}(0)$ en mesurant le rapport entre le pic central et son plus proche voisin :

$$g_{brut}^{(2)}(0) = 0,0070 \pm 0,0005$$

Pour sa part, la courbe 3.15 présente une courbe d'autocorrélation quasiment plate, typique des phénomènes poissonniens. L'utilisation de ces deux courbes permet de remonter à la valeur de $g_{net}^{(2)}(0)$ en s'appuyant sur la formule 3.31. L'incertitude quant à elle a été estimée à l'aide de la formule 3.32

développée au paragraphe précédent. Nous avons donc ici :

$$g_{net}^{(2)}(0) = 0,0049 \pm 0,0005 \quad (3.34)$$

Cet excellent résultat peut être interprété comme une diminution du nombre de doubles photons par fenêtre de détection par un facteur 200 comparé à un laser atténué.

Il est intéressant de chercher maintenant à connaître la valeur de P_1^{exp} à l'aide de la routine Niçoise et surtout de comparer la valeur de $g^{(2)}(0)$ calculée à l'aide du modèle d'analyse.

Estimation de $g^{(2)}(0)$ à partir du modèle d'analyse Niçois

Le programme nous fournit les valeurs suivantes :

	Valeur	Incertitude
P_0^{exp}	0,43	$\pm 0,03$
P_1^{exp}	0,57	$\pm 0,03$
P_2^{exp}	0,0009	$\pm 0,0001$
$g^{(2)}(0)^{exp}$	0,006	$\pm 0,001$
γ_i^{exp}	0,57	$\pm 0,03$

TAB. 3.6 – Tableau récapitulatif des performances de la source de Genève à partir du programme développé à Nice.

Nous constatons que la valeur calculée par le programme, en utilisant une forme de $g^{(2)}(0)$ simplifiée pour les sources (sub-)poissonniennes $\frac{2P_2}{P_1^2}$ est relativement proche de la valeur déduite de l'histogramme qui repose sur la définition originale de $g^{(2)}(0)$. Pourtant au niveau des incertitudes, le protocole Genevois améliore d'un ordre de grandeur la précision de $g^{(2)}(0)$. À ce titre, il faut accorder plus de valeur à la mesure de $g^{(2)}(0)$ à partir de l'histogramme 3.14 qui ne dépend pas de l'efficacité et des pertes associées au banc de mesure.

Récapitulatif des résultats associés à la source Genevoise

Ces résultats concernent la source de photons uniques annoncés du GAP-Optique de Genève, qui émet des paires de photons dont la longueur d'onde du signal est dans le visible (810 nm) tandis que la longueur d'onde des photons annoncés est à 1550 nm. Cette source tire donc profit d'une détection

très efficace des photons signal à l'aide d'une APD-*Silicium* qui lui permet d'afficher une probabilité de 0,57 d'obtenir un photon unique tandis que la statistique générale de la source est amélioré d'un facteur 200 par rapport à une source poissonnienne classique. Nous avons récemment amélioré les résultats en travaillant sur la récolte des photons signal et idler et les performances maximales de la source [51] sont aujourd'hui :

$$P_1 = 0,61 \quad \text{et} \quad g^{(2)}(0) = 2,2 \cdot 10^{-3}$$

La source présente une probabilité d'obtenir un photon unique très élevée de 0,61 ce qui est un des meilleurs résultats à l'heure actuel [50, 44, 45] pour une source de photons annoncés dans une fibre monomode et une diminution du nombre de doubles photons par un facteur 500 qui est tout simplement la meilleur valeur jamais rencontrée à notre connaissance.

3.3 Conclusion du chapitre 3

Nous pouvons décomposer ce chapitre expérimental en deux parties distinctes.

La première aura été l'occasion de mesurer expérimentalement les performances de la source de photon uniques que nous avons développé à Nice. Le protocole suivi à cet effet est inspiré de l'article de Mandel [19], où nous remontons à l'aide d'un calcul à la statistique d'émission des photons dans un montage de type *Hanbury-Brown & Twiss*. Nous avons obtenus par cette méthode d'excellents résultats :

$$P_1 = 0,37 \quad \text{et} \quad g^{(2)}(0) = 0,08$$

Cette méthode permet de déduire la valeur de $g^{(2)}(0)$ à partir des probabilités P_1 et P_2 mais ne permet pas de constater directement le comportement « unique » des photons émis par la source. Pour cela, il est plutôt habituel d'utiliser le montage de type *Hanbury-Brown & Twiss* pour tracer la fonction d'autocorrélation d'ordre deux $g^{(2)}(\tau)$, ce qu'il n'est pas possible de faire pour les sources asynchrones comme la notre.

Nous avons donc développé dans la seconde partie une méthode originale qui permet de tracer la fonction d'autocorrélation d'ordre deux $g^{(2)}(\tau)$ pour les sources de photons uniques annoncés asynchrone. Le montage expérimental est toujours de type *Hanbury-Brown & Twiss* mais le *convertisseur*

temps-amplitude (TAC) habituel est remplacé ici, par une carte d'acquisition particulière, couplée à un programme qui simule le comportement du TAC en s'affranchissant de toute base de temps. Nous avons réalisé ce montage au GAP-Optique et l'avons utilisé pour estimer les performances de la source de photons uniques « Genevoise ». Cette nouvelle mesure a été par ailleurs l'occasion de vérifier la justesse du modèle développé à Nice pour calculer P_0 , P_1 et P_2 puis remonter à $g^{(2)}(0)$. Nous avons trouvé pour cette source :

$$P_1 = 0,61 \quad \text{et} \quad g^{(2)}(0) = 2,2 \cdot 10^{-3}$$

qui sont, à notre connaissance, les meilleurs résultats toutes sources confondues.

La comparaison des performances de cette source avec la notre sont l'occasion de mettre le doigt sur les points faibles de la configuration niçoise. En prenant en compte les pertes vues par les photons à 1550 nm entre le cristal et la lentille de couplage, le couplage Γ_i de la source genevoise s'élève à environ 0,65 ce qui est tout à fait comparable avec la valeur obtenue à Nice ($\Gamma_i = 0,58$). Cela signifie que nous récoltons quasiment aussi bien les photons en configuration entièrement guidée à l'aide d'une simple fibre qu'en configuration massive avec des lentilles

L'écart de performance le plus significatif se situe au niveau de la valeur de $g^{(2)}(0)$ qui quantifie la réduction d'événements à deux photons par rapport à une source poissonnienne. Alors que notre source réduit cet quantité d'un facteur 10, la source genevoise la réduit d'un facteur 500... Comme le coefficient de récolte des paires de photons est sensiblement le même, cette écart de performances se justifie uniquement par l'utilisation de photons signal à 810 nm qui permettent l'utilisation d'une APD *silicium* beaucoup plus performante que l'APD germanium. Le chapitre suivant s'intitulera logiquement « perspectives » et montrera le gain qu'il sera possible d'obtenir en changeant quelques éléments de la source.

Chapitre 4

Perspectives

Nous avons montré dans le chapitre 3 que le taux de récolte des photons idler du guide vers la fibre avoisine les 60%, ce qui correspond à une très bonne valeur à l'heure actuelle, toutes sources de paires de photons confondues [49, 50, 51]. Nous savons donc bien fabriquer ces guides et nous arrivons à un bon confinement sans trop de pertes, pourtant la probabilité finale d'obtenir un photon annoncé à la sortie de la fibre est seulement de 45% justifiée par les pertes qu'introduisent les composants optiques placés en aval du guide PPLN. Par ailleurs, l'utilisation d'un détecteur Germanium handicape irrémédiablement les performances de la source, en raison de son taux de coups sombres élevé qui réduit la probabilité P_1 à 37% et de sa faible efficacité qui ne permet pas d'obtenir des valeurs de $g^{(2)}(0)$ minimales.

Nous allons étudier dans ce chapitre le gain qu'il est possible d'obtenir en améliorant la détection des photons signal. Tout d'abord nous envisagerons de réduire le bruit intrinsèque de notre APD germanium, puis nous envisagerons de diminuer la longueur d'onde du photon signal, de 1310 *nm* vers 810 *nm*, pour pouvoir tirer parti des excellentes performances des APDs en Silicium.

Dans les expériences de communications quantiques, la largeur spectrale des photons utilisés est un paramètre crucial. Dans les fibres optiques, la dispersion chromatique représente un frein aux communications longues distances qu'une faible largeur spectrale permet de contourner. Nous étudierons alors les largeurs spectrales de nos photons en fonction de la longueur du guide PPLN.

4.1 Le pompage en régime impulsif : un moyen de réduire les coups accidentels dans l'APD-Ge ?

Dans l'approche théorique du chapitre 1, nous avons souligné que P_0 la probabilité de ne pas avoir de photon était directement alimentée par le taux de coups sombres dans l'APD et nous avons insisté sur l'importance de travailler avec un rapport signal/bruit élevé. Pour assurer la détection des photons signal à 1310 nm , nous utilisons dans notre montage expérimental une APD en germanium qui fonctionne en mode passif et présente un taux de coups sombres assez élevé par rapport au taux maximal de détections qu'il est possible d'atteindre.

En quelques chiffres, notre APD-Ge présente un taux de coups sombres de 20 kHz et sature aux alentours de 150 kHz ce qui limite fortement la valeur du rapport $\text{signal/bruit} = \frac{S_{ge}^{net}}{D_{Ge}^c}$ à une valeur inférieure à 10. À ce sujet, nous estimons élevé un rapport signal/bruit supérieur à 10^2 et, à titre d'exemple dans le cas d'une APD silicium passivement déclenchée, il se situe généralement supérieur à 10^4 .

Un moyen d'améliorer notre APD serait de s'inspirer du mode de fonctionnement des APD-InGaAs qui, pour réduire leur taux coups sombres, sont utilisées en mode déclenché. Ainsi, il suffirait d'utiliser un laser de pompe impulsif et d'allumer l'APD pendant un temps très court seulement lorsque la probabilité d'avoir un photon signal est grande. Dans cette configuration, la probabilité de détecter un photon est proportionnelle au nombre de photons dans l'impulsion et à l'efficacité de détection de l'APD ($\eta_{Ge} \approx 10\%$), tandis le taux de coups sombres diminue car l'APD n'est effectivement plus allumée pendant une seconde. Vérifions, à l'aide d'un petit calcul, le gain qu'il est possible d'obtenir : supposons que nous disposions d'un laser capable d'émettre des impulsions lumineuses (de 1 ps par exemple) avec un taux de répétition de 100 kHz^1 , en considérant que la quantité de paires par impulsion suit une distribution poissonnienne avec $\bar{n} = 0,1$ paires par pulse, le taux net de détection peut alors atteindre :

$$S_{Ge|déclenchée}^{net} \approx \bar{n} \times \gamma_s \times F_{rep} \times \eta_{Ge} = 300\text{ s}^{-1}$$

Contrairement à l'arrivée des photons, les coups sombres dans l'APD sont aléatoires avec cependant un taux moyen $D_{Ge}^c = 20000\text{ s}^{-1}$. Donc si nous allumons seulement l'APD-Ge pendant 1 ns pour détecter le photon signal

¹Ce qui correspond à la fréquence de répétition maximum d'ouverture d'une APD germanium.

contenu dans l'impulsion lumineuse, le taux de coups sombres effectifs par seconde chute à raison de :

$$D_{Ge|déclenchée}^c = D_{Ge}^c \times F_{rep} \times T_{allumage} = 2 \text{ s}^{-1}$$

Ce petit raisonnement a le mérite de montrer que nous pouvons augmenter rapport signal/bruit jusqu'à 150 simplement en utilisant différemment notre détecteur.

Alors avec un laser impulsionnel, nous espérons obtenir une source de photons uniques annoncés présentant les figures de mérite suivantes :

$$P_1 = 0,45 \quad \text{et} \quad g^{(2)}(0) = 0,04$$

Nous n'observons qu'un léger gain en performance sur la probabilité P_1 , qui s'explique par le fait qu'elle est, de toutes façons, limitée par γ_i le coefficient de récolte des photons annoncés. Aussi, le gain quasi nul sur $g^{(2)}(0)$ s'explique par le fait que cette méthode n'améliore pas l'efficacité de détection, mais réduit simplement le taux de coups sombres. Il n'y a donc pas de raisons d'observer une diminution de la probabilité d'avoir deux photons dans une fenêtre ΔT .

Seule une amélioration du couplage des photons signal ou de l'efficacité de leur détection pourrait influencer sur $g^{(2)}(0)$. Nous allons maintenant nous pencher sur une perspective qui devrait nous permettre d'obtenir un gain sur les deux tableaux : à savoir une augmentation de l'efficacité de détection et une diminution importante du bruit.

4.2 L'utilisation d'une APD–*Silicium* comme trigger

Il existe sur le marché des détecteurs fibrés identiques aux nôtres mais basés sur des semi-conducteurs en *silicium*. Typiquement, ces détecteurs fonctionnent en mode passif à des températures raisonnables comprises entre -40°C et $+25^\circ \text{C}$ et présentent des efficacités de détection et des taux de coups sombres qui valent :

$$\eta_{Si} \approx 60\% \quad \text{et} \quad D_{Si}^c \approx 100 \text{ s}^{-1}$$

Cette excellente efficacité est associée à la détection de photons dont la longueur d'onde est située dans le visible ($400 \text{ nm} < \lambda < 800 \text{ nm}$) et s'explique par le fait que :

- L'énergie portée par un photon à 800 nm est plus grande qu'à 1310 nm . Il est donc plus facile à détecter.
- La fabrication des semi-conducteurs en silicium est très bien maîtrisée et permet d'obtenir des composants avec une très faible densité de défauts qui sont la principale source de bruit dans les APDs.

Dans notre configuration actuelle, nous ne pouvons pas simplement remplacer notre APD *germanium* par un modèle *silicium* qui est complètement aveugle à 1310 nm sans réduire la longueur d'onde des photons signal. Dans notre configuration, il est assez simple de changer la longueur d'onde des photons de la paire (voir section 2.1.2 du chapitre 2) puisqu'il suffit de changer le pas du réseau de *Quasi-accord de phase*.

Une étude nous a permis de déterminer le pas d'inversion Λ qui convertit efficacement des photons de pompe à 532 nm en paires non-dégénérées à 810 nm et 1550 nm . En configuration guidée, Λ est compris entre 6 et $7\text{ }\mu\text{m}$ ce qui constitue une gamme de valeurs accessibles grâce à la technologie de polarisation périodique du LiNbO_3 par champ électrique [82, 83]. Dans ce cas, à l'aide de l'équation 2.24 utilisée sous la forme

$$\delta\lambda_{s,i} = \frac{\lambda_{s,i}^2}{\mathcal{N}l} \quad (4.1)$$

où $\mathcal{N} = \left| n^{\text{guide}}(\lambda_{s_0}) - n^{\text{guide}}(\lambda_{i_0}) - \frac{\partial n^{\text{guide}}}{\partial \lambda} \Big|_{\lambda_{s_0}} \lambda_{s_0} + \frac{\partial n^{\text{guide}}}{\partial \lambda} \Big|_{\lambda_{i_0}} \lambda_{i_0} \right|$, nous prévoyons, pour un guide identique au notre ($1,1\text{ cm}$ de long), d'obtenir des photons de largeurs spectrales :

$$\begin{aligned} \delta\lambda_{810}^{th} &\approx 0,7\text{ nm} \\ \delta\lambda_{1550}^{th} &\approx 2,6\text{ nm} \end{aligned}$$

Nous allons maintenant calculer les performances théoriques qu'il est possible d'obtenir avec une source de photons annoncés utilisant une APD silicium pour détecter les photons signal. À l'aides des formules 1.25, 1.27 et 1.29 du premier chapitre, nous pouvons facilement calculer P_2 , P_1 et P_0 en considérant les valeurs numériques suivantes :

- une efficacité de détection de $\eta_{Si} = 60\%$,
- un taux de coups sombres de $D_{Si}^c = 100\text{ s}^{-1}$,
- une taux de répétition moyen de la source de 100 kHz , qui rappelons-le est la fréquence maximale de fonctionnement des APD-*InGaAs*,
- un coefficient de collection des photons idler identique à celui obtenu actuellement, à savoir $\Gamma_i = 0,58$,

- un coefficient de collection² des photons signal $\Gamma_s = 0,33$.

En supposant que les pertes dans les composants restent approximativement les mêmes (1,1 dB à 1550 nm et 2 dB à 810 nm), les performances attendues pour cette source sont donc :

$$P_1 = 0,45 \quad \text{et} \quad g^{(2)}(0) = 0,005$$

Dans cette configuration, la source présente d'excellentes caractéristiques :

- une probabilité $P_1 = 0,45$ d'avoir un photon unique,
- une réduction des événements à deux photons par un facteur 200 comparé à un laser atténué,
- et enfin, les photons émis à 1550 nm présentent une largeur spectrale de l'ordre de 2,6 nm.

La principale difficulté technologique à la réalisation de cette source réside dans l'inversion des domaines du $LiNbO_3$ sur de courtes périodes. Cette source n'est plus une lointaine perspective, car la réalisation de sources de photons uniques suivant ce schéma est financée par le projet européen³ *SECOQC* et nos premiers essais de poling à des périodes comprises entre 6 et 7 μm se sont montrés très concluants.

Nous venons d'envisager l'évolution de notre source en améliorant la détection des photons signal et de montrer que le gain obtenu promettait la réalisation à court terme d'une source entièrement fibrée extrêmement performante. Par contre, il était plutôt non-intuitif que les largeurs spectrales de photons émis soient autant réduites simplement en changeant la longueur d'onde des photons signal. Ce comportement peut se comprendre à l'aide de l'équation 4.1 en notant la présence au dénominateur de la différence des indices optique des modes guidés au signal et à l'idler. Alors, plus les longueurs d'ondes des photons de la paire diffèrent, plus la différence des indices est grande donc au plus leur largeur spectrale respective diminue.

Nous invitons le lecteur à lire la référence [84] qui aborde de manière complète les inconvénients⁴ liés la dispersion chromatique dans les fibres pour les

²À partir du recouvrement des modes, à 810 et 1550 nm, entre le microguide et la fibre SMF-28, nous avons calculé que le rapport $\frac{\Gamma_{810}}{\Gamma_{1550}}$ vaut 0,58.

³Development of a Global Network for Secure Communication based on Quantum Cryptography $\Rightarrow$ [http ://www.secoqc.net/](http://www.secoqc.net/)

⁴Et des remèdes...

communications quantiques. Cependant, nous essayerons d'introduire simplement le problème que soulève cet article et nous verrons, à partir de la formule 4.1, que la manière la plus simple de réduire la largeur spectrale de nos photons (sans réduire notre efficacité de création des paires) est de jouer avec la longueur du guide PPLN.

4.3 Réduction de la largeur spectrale $\Delta\lambda$ des photons

Dans les fibres optiques télécom, la dispersion chromatique implique des vitesses de propagation des photons différentes en fonction de leurs longueurs d'ondes. Couplée à la largeur spectrale des photons, elle induit donc une incertitude sur l'instant d'arrivée de ces photons pour un parcours d'une longueur L . La dispersion chromatique D s'exprime en ps par km de fibre pour une largeur spectrale des photons en nm et dans les fibres télécom vaut typiquement $17 ps \cdot nm^{-1} \cdot km^{-1}$. Nous pouvons alors calculer l'incertitude sur le temps d'arrivée d'un photon grâce à la formule :

$$\Delta\tau_{disp} = D \times \Delta\lambda \times L \quad (4.2)$$

Plaçons-nous dans un cas concret pour mieux sentir l'importance de ce phénomène. Reprenons nos photons uniques annoncés qui présentent une largeur spectrale de $20 nm$ et sont attendus dans une fenêtre de $3 ns$ pour être détectés. À l'aide de formule 4.2, nous pouvons estimer qu'à partir de $9 km$ de fibre optique, le photon a une probabilité non nulle de ne plus être à l'intérieur de sa fenêtre de détection. Il peut être intéressant de voir à quelle distance correspond $3 ns$ d'incertitude : *au bout de 9 km de fibre optique la position du photon est connue à plus ou moins 50 cm.*

À l'instar de l'expérience de mandel pour mesurer l'intervalle de temps séparant l'instant de création des deux photons d'une paire (cf. chapitre 1), certaines communications quantiques reposent sur la mesure jointe de l'état deux photons indiscernables qui consiste à réunir sur une lame semi-réfléchissante deux photons pour qu'ils interfèrent et « perdent » leur caractère propre (voir annexe C de la référence [68]). Expérimentalement, tout l'art de réaliser une mesure jointe de deux photons uniques réside dans l'ajustement de leurs chemins optiques respectifs pour qu'ils arrivent simultanément sur la lame semi-réfléchissante. Il est assez facile d'imaginer la difficulté qu'il y a, à ajuster le recouvrement spatialement de deux photons dont la position est connue à $50 cm$ près...

L'objectif de l'article de S. Fasel est d'étudier les performances de deux méthodes pour s'affranchir de la dispersion chromatique appliquées à la dis-

tribution quantique de clé à l'aide de paires de photons⁵. La première méthode consiste simplement à filtrer spectralement les photons issus du cristal non-linéaire et la seconde à utiliser une méthode de compensation de dispersion⁶. Pourtant, quelque soient les solutions envisagées, toutes deux impliquent une augmentation des pertes et un taux de comptage plus faible où les coups sombres risquent de prendre une place prépondérante.

Dans la section 2.1.2 du chapitre 2, nous avons vu que l'élargissement spectral était inversement proportionnel à la longueur d'interaction dans le guide PPLN. À l'heure actuelle, nous pouvons fabriquer des guides PPLN pouvant mesurer jusqu'à 4,5 cm de long. Il est donc intéressant de discuter de la largeur spectrale des photons annoncés émis par de tels guides en gardant à l'esprit que notre probabilité de créer une paire de photons restera la même ($\eta_{conv} \approx 10^{-7}$) quelle que soit leur largeur spectrale.

Sur le graphique 4.1 de la page suivante, nous avons tracé l'évolution des largeurs spectrales des photons annoncés en fonction de la longueur du guide PPLN. La première courbe correspond à notre configuration actuelle, c'est-à-dire en utilisant des paires de photons à 1310/1550 nm, tandis que la seconde est associée à la source développée dans le cadre du projet *SECOQC* à partir de paires à 810/1550 nm.

⁵Il s'agit de la source de paires de photons Genevoise rencontrée dans le chapitre 3.

⁶En pratique, il s'agit d'une bobine de fibre à dispersion négative.

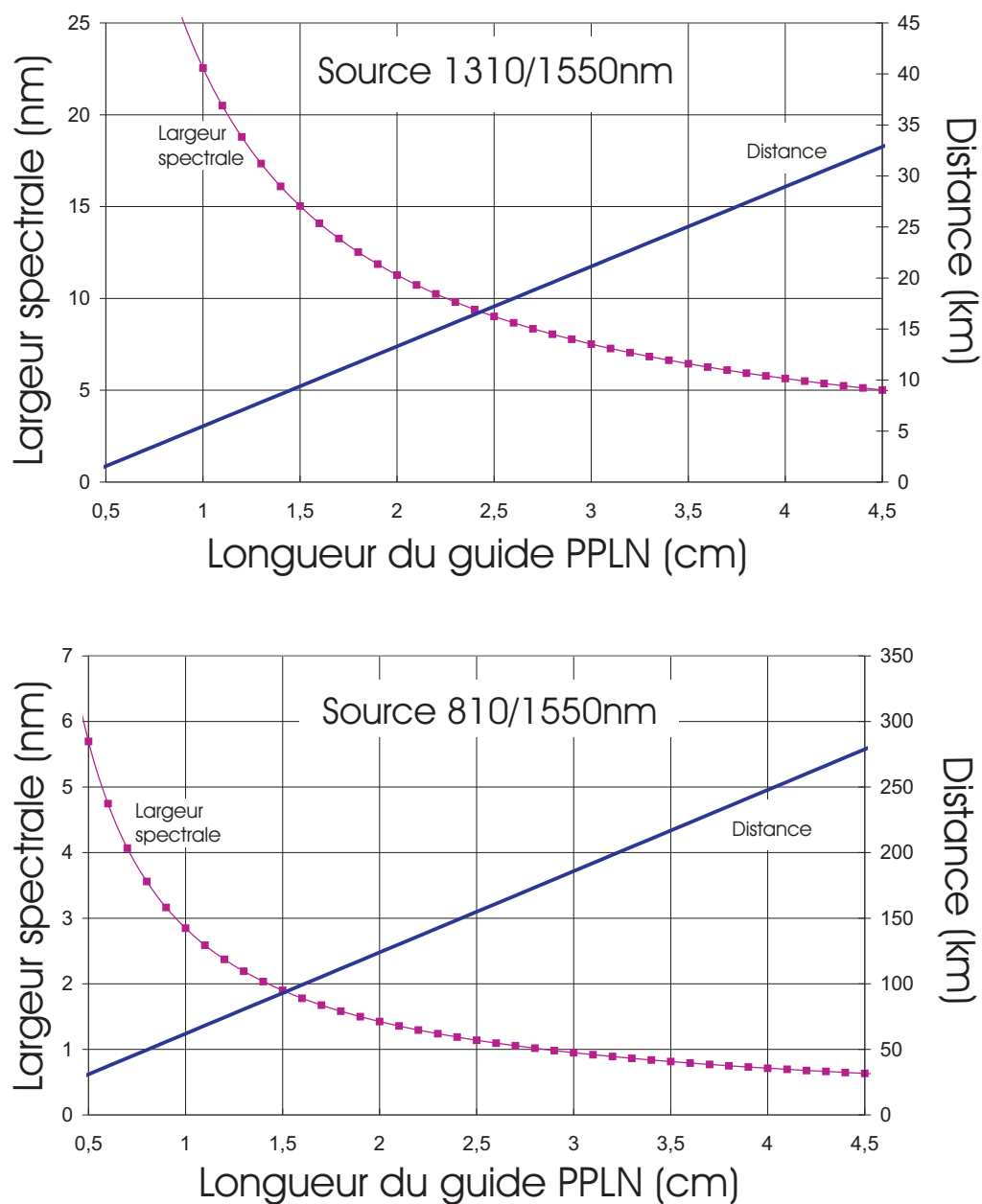


FIG. 4.1 – Évolutions des largeurs spectrales des photons annoncés en fonction de la longueur du guide pour des paires à 1310/1550 nm et 810/1550 nm. Nous avons rajouté à titre indicatif la distance au bout de laquelle nos photons annoncés ne sont plus contenus dans les fenêtres de détections de 3 ns qui leurs sont associées.

À la vue de ce graphique nous pourrions être tenté d'utiliser le guide le plus long possible. Pourtant, nous devons rester vigilant et faire attention aux *pertes à la propagation dans le guide*. Nous n'avions pas parlé de ce type de pertes dans le chapitre 1, pourtant elles étaient implicitement contenues dans la valeur de Γ_i mesurée au chapitre 2. À l'heure actuelle, les pertes à la propagation dans un guide *SPE* ont été estimées aux environs de $\alpha = 0,3 \text{ dB/cm} = 0,069 \text{ cm}^{-1}$ et comme elles participent directement à la valeur de Γ_i , il ne faudra peut être pas choisir un guide PPLN trop long afin de garder une bonne probabilité de collecter le photon idler.

Nous allons donc calculer la probabilité d'obtenir le photon annoncé à la sortie du guide, sachant que le photon signal a déjà été détecté⁷. En considérant que la paire a été créée à l'abscisse x à l'intérieur du guide d'une longueur L , la probabilité que le photon idler sorte du guide s'écrit :

$$P_{\text{sortie}} = \frac{1}{L} \int_0^L e^{-\alpha(L-x)} dx = \frac{1 - e^{-\alpha L}}{\alpha L}$$

Nous pouvons donc tracer la probabilité d'obtenir le photon idler en fonction de la longueur du guide sur la courbe 4.2.

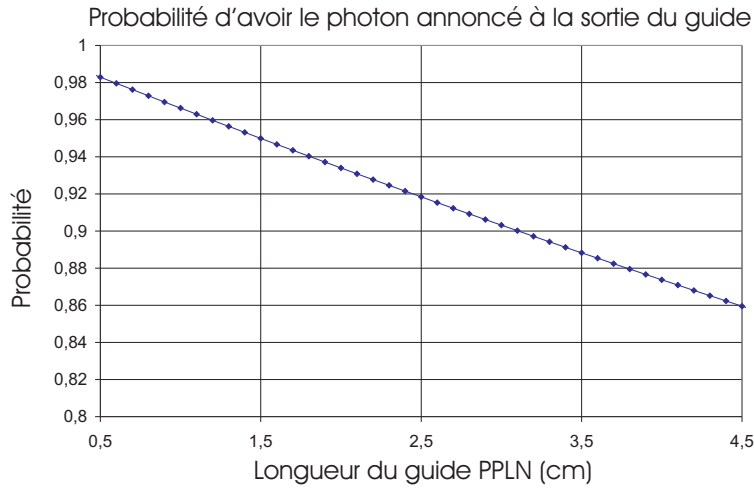


FIG. 4.2 – Probabilité d'obtenir le photon idler à la sortie du guide en fonction de la longueur de ce dernier.

⁷La formulation est importante, nous ne cherchons pas à connaître la probabilité d'obtenir les deux photons de la paire, mais uniquement l'idler sachant que le signal a été détecté. À ce sujet, la probabilité de collecter la paire est simplement déduite en élevant au carré la probabilité de collecter un photon.

Pour fixer quelques ordres de grandeur, restreignons les graphiques précédents à quelques cas concrets, c'est à dire à des paires de photons issues de guides d'une longueur de 1 *cm*, 3 *cm* et 4,5 *cm* :

Notre guide PPLN 1310/1550 *nm* actuel présente un coefficient de couplage *guide-fibre* de 0,58 et les photons idler annoncés présentent une largeur spectrale de 20 *nm*.

Un guide de 3 *cm* garderait un coefficient de couplage raisonnable de 0,54 et les photons annoncés auraient une largeur spectrale de 7,5 *nm*.

Un guide de 4,5 *cm* verrait son coefficient de couplage réduit à 0,52 mais, par la même occasion, la largeur de ses photons réduite à 5 *nm*.

Il est intéressant de voir les possibilités qu'offriront l'utilisation de nouveaux guides dont les paires de photons auraient une longueur d'onde centrée sur 810/1550 *nm*. Il va sans dire que les coefficients de couplage *guide-fibre* des photons idler sont identiques à ceux attendus pour un guide 1310/1550 *nm*.

Un guide PPLN 810/1550 *nm* présenterait un coefficient de couplage *guide-fibre* identique au notre ($\Gamma_i=0,58$) et les photons idler émis présenteraient alors une largeur spectrale de 2,6 *nm*.

Dans un guide de 3 *cm* les photons auraient une largeur spectrale de 1 *nm*.

Dans un guide de 4,5 *cm* la largeur des photons serait réduite à 0,6 *nm*.

Cette dernière configuration est plus qu'intéressante d'un point de vue expérimental, car elle permettrait de s'affranchir du filtrage spectral qui réduit l'efficacité globale du cristal en tant que générateur de paires de photons

4.4 Conclusion du chapitre 4

Le bruit dans notre détecteur germanium étant relativement élevé, nous nous sommes penchés sur un moyen de le réduire. Assez naturellement, l'utilisation de ce détecteur en mode déclenché est apparue être une évidence en ayant recours à l'utilisation conjointe d'un laser de pompe impulsionnel.

À partir de ce constat, nous avons donc envisagé de changer le type d'APD utilisée pour annoncer l'arrivée d'un photon et de tirer profit des hautes performances des APD en silicium. Pour ce faire, nous prévoyons de développer de nouveaux guides PPLN permettant d'obtenir l'interaction 532 *nm* $\Rightarrow$ 810/1550 *nm*. Ainsi, à l'aide d'un guide PPLN de 4,5 *cm* de long,

nous pouvons espérer annoncer l'arrivée de photons uniques avec une probabilité de 0,45 tout en réduisant la quantité d'événements à deux photons par un facteur 200 par rapport à une source cohérente atténuée.

Une dernière étude nous a permis d'envisager la réduction de la largeur spectrale des photons qu'il est possible d'obtenir en allongeant la longueur de nos guides. Il faut noter à ce sujet que cette technique n'est pas comparable au filtrage spectral des photons. En effet, nous ne venons pas sélectionner simplement une fenêtre spectrale à la sortie du guide, mais bien redistribuer la probabilité de créer une paires à l'intérieur d'une fenêtre spectrale réduite : *nos guides d'ondes gardent la même efficacité absolue* ($\eta_{conv} \approx 10^{-7}$)

Conclusion générale

En moins de cinquante ans, la physique quantique est passée de controverses philosophiques déclinées sous la forme « d'expériences de pensée » à des expériences mettant en évidence les propriétés quantiques des photons. Aujourd'hui, à l'heure de l'information et des télécommunications, cette discipline s'est très vite rendu compte de la véritable révolution qu'elle apporte à la théorie l'information. Depuis, un nombre impressionnant d'idées théoriques regroupées sous l'appellation « communications quantiques » ont vu le jour. Toutefois, mis à part la distribution quantique de clés ou la téléportation d'information, les réalisations expérimentales sont encore loin d'accomplir les prouesses escomptées par la théorie. Le développement expérimental des communications quantiques reposera sur l'existence d'outils capables de fournir efficacement des paires de photons intriqués et des photons uniques, qui sont les deux ressources de bases pour ce type d'expériences.

Expérimentalement, deux voies se sont dessinées :

- La première concerne les communications à l'air libre, utilisant des photons dont la longueur d'onde est située dans le visible. Ces communications sont limitées à quelques kilomètres, ce qui est amplement suffisant pour communiquer d'un immeuble à l'autre ou, plus ambitieux, pour établir une communication terre-satellite⁸.
- La deuxième a pour objectif de tirer parti des avantages qu'offrent les fibres optiques pour transmettre l'information sur plusieurs centaines de kilomètres. Les photons doivent avoir une longueur d'onde située dans l'infrarouge et surtout doivent être facilement insérable dans un réseau de fibre optique télécom.

Si la première catégorie dispose déjà de quelques sources performantes, la seconde est pour l'instant pas ou peu dotée de moyens à la mesure de ses ambitions.

Le but de ce travail fût de démontrer que *l'optique intégrée* est un atout

⁸Bien qu'aujourd'hui l'utilisation de l'infrarouge lointain semble plus prometteur...

technologique au service des communications quantiques en configuration guidée. En utilisant des paires de photons [19], nous avons pour objectif de montrer la faisabilité d'une source de photons uniques annoncés compacte et performante à partir d'un guide d'ondes intégré sur un substrat de niobate de lithium polarisé périodiquement. Grâce à la technique du Quasi-accord de phase, nous avons obtenu des paires de photons dont les longueurs d'ondes sont centrées sur 1310 nm et 1550 nm et la détection passive du photon à 1310 nm permet d'annoncer l'émission d'un photon unique dont la longueur d'onde est centrée sur la seconde fenêtre des télécommunications. La construction modulaire et l'utilisation de composants optiques au standard des télécommunications apportent une réelle souplesse d'utilisation et la source peut s'insérer de façon simple, stable et efficace dans les futurs réseaux de communications quantiques. D'un point de vue technique, le travail s'est articulé autour du couplage des photons depuis le guide PPLN vers une fibre optique monomode appuyée contre lui. Nous avons à ce titre un coefficient de couplage guide-fibre proche de 60%, comparable aux meilleurs résultats actuels.

Pour créer les paires de photons, nous avons choisi d'utiliser un laser de pompe continu. Dès lors, l'instant de création des paires est inconnu et le temps qui sépare deux paires successives est aléatoire. Dans cette configuration, les paires sont isolées les unes des autres grâce à une détection déclenchée. Il ne faut plus voir des impulsions, mais des fenêtres de détection contenant 0, 1 ou 2 photons. Nous avons dit que la source présentait un caractère *asynchrone* par identification au protocole de communication réseau : *ATM*, pour *asynchronous transfert mode* qui traite l'information sans référence temporelle simplement lorsque celle-ci est disponible. En prenant en compte ce schéma de fonctionnement et le formalisme associé à la production de paires de photons en régime continu, nous avons établi les formules théoriques qui permettent de prévoir les performances de la source (P_0 , P_1 et P_2) en fonction des pertes au sein de la source et des défauts du détecteur servant à annoncer les photons uniques.

Si le mode de fonctionnement asynchrone est parfaitement adapté aux communications quantiques, il ne l'est pas du tout pour les méthodes d'investigations « classiques » utilisées pour déterminer expérimentalement P_0 , P_1 et P_2 . Toute l'originalité de ce travail réside dans le développement d'un modèle théorique et de deux nouvelles méthodes expérimentales de caractérisation complémentaires et adaptées à n'importe quelle source asynchrone. Premièrement, nous avons établi théoriquement les probabilités P_1 et P_2 d'avoir un ou deux photons émis par la source en régime continu. Ces calculs prennent en compte, par exemple, les pertes vues par les paires de photons ou la faible

efficacité du détecteur servant à annoncer l'arrivée des photons uniques. Ensuite, nous avons développé deux nouvelles méthodes de caractérisation compatibles avec le schéma asynchrone de la source :

- La première permet de remonter à l'aide d'un algorithme à la statistique des photons qui sont à l'origine des détections sur un montage de type « Hanbury-Brown & Twiss ».
- La deuxième nous permet, à partir d'un montage légèrement différent, de tracer l'équivalent asynchrone de la fonction d'autocorrélation d'ordre deux $g^{(2)}(\tau)$, ce qui n'était pas possible avec le montage original. Notons que cette méthode permet de déterminer $g^{(2)}(0)$ en s'affranchissant des pertes, ce qui n'est pas le cas du premier montage.

Finalement, l'expérience confirme les figures de mérites entrevues grâce au modèle théorique et aboutit aux résultats suivants pour la source Niçoise :

$$P_1 = 0,37 \quad \text{et} \quad g^{(2)}(0) = 0,08$$

qui correspond à l'obtention d'un photon unique dans près de 40% des cas et à une réduction des événements à deux photons par un facteur 10 comparativement à une source cohérente atténuée.

Nous avons aussi cerné les points faibles du montage, ainsi que les éventuelles solutions réalisables à très court terme afin d'obtenir une diminution des événements à deux photons par un facteur 200 comme c'est le cas pour la source Genevoise. L'idée consiste à utiliser une photodiode silicium pour annoncer l'arrivée du photon unique à 1550 nm , ce qui nécessite que l'autre photon de la paire ait une longueur d'onde située dans le visible. Ce concept nécessite alors la création de nouveaux guides PPLN pour produire des paires de photons à 810 nm et 1550 nm à partir de photons de pompe à 532 nm . Les travaux de développement de ces échantillons sont actuellement supporté par le projet européen *SECOQC*.

L'ensemble des résultats que nous avons montrés dans ce travail constitue véritablement la première réalisation expérimentale d'une source de photons uniques intégrée [85]. Conjointement avec le travail de thèse précédent [68], nous pensons avoir démontré l'intérêt de l'optique intégrée et plus particulièrement des guides PPLN pour fournir efficacement les ressources nécessaires au développement expérimental des *communications quantiques* longues distances.

La dernière partie de ce travail de thèse, a permis de pointer du doigt la qualité de la collection des paires de photons dans les guides PPLN et

d'apercevoir une réduction importante de la largeur spectrale des photons émis ($\sim 1\text{ nm}$), tout en gardant une efficacité de conversion *photons de pompe* $\rightsquigarrow$ *paires de photons* très élevée⁹ (10^{-7}). Pour notre source de photon annoncé, cette efficacité nous a permis d'utiliser seulement quelques microwatts pour la faire fonctionner, mais notre futur projet serait d'utiliser cette efficacité afin d'obtenir efficacement plus d'une paire par impulsion. Grâce à la technique du quasi accord de phase, ces paires de photons seraient disponibles aux longueurs d'onde désirées et présenteraient une largeur spectrale réduite. Fort de ces guides PPLN, nous souhaitons obtenir de multiples paires cohérentes entre elles afin de créer efficacement des états intriqués à plusieurs photons qui permettent d'étendre plus encore les schémas expérimentaux des communications quantiques [86, 87, 88, 89].

⁹Cette dernière reste, à l'heure actuelle, la plus élevée au monde.

ANNEXES

Annexe A

Sécurité de la distribution quantique de clé face à l'espionnage

Cette annexe est très librement inspirée de l'article de N. Lütkenhaus [90], à partir duquel nous allons déterminer la formule permettant de connaître le taux de bits utilisables par Alice et Bob pour obtenir une clé secrète. Ce taux sera calculé pour une communication sur une distance D au travers d'une ligne réelle présentant des pertes α .

Supposons pour la suite qu'Alice utilise une source de photons uniques dont la statistique est la suivante : P_1 la probabilité que l'impulsion contienne un photon et $P_2 = g^{(2)}(0)\frac{P_1^2}{2}$ la probabilité qu'elle en contienne deux.

A.1 Qu'est ce le taux de bits utilisables ?

Il s'agit de la quantité de bit de la clé réellement utilisable par Alice et Bob pour crypter leur message. Il se trouve en pratique, une partie de la quantité de bit qu'ils ont échangés est souvent utilisée pour corriger les erreurs ou amplifier la confidentialité. Nous nous proposons justement d'introduire ces concepts expérimentaux.

A.1.1 L'amplification de confidentialité

Pour espionner l'échange de clé entre Alice et Bob, Ève va s'attaquer aux impulsions contenant à deux photons. En effet, les impulsions contenant deux photons vont être un moyen pour Ève d'avoir accès à l'information que partage Alice et Bob sans introduire d'erreur, donc sans révéler sa présence.

172 ANNEXE A. SÉCURITÉ DE LA DISTRIBUTION QUANTIQUE DE CLÉ

- Ève se place juste à la sortie des locaux d'Alice.
- Pour chaque impulsion contenant deux photons, elle les sépare pour n'en garder qu'un seul, qu'elle place dans une mémoire quantique, tandis que l'autre peut continuer son chemin vers Bob.
- Une fois qu'Alice et Bob auront révélé leurs bases, Ève va mesurer dans la bonne base le photon qu'elle a gardé dans la mémoire.

Alice connaissant la quantité d'impulsions contenant deux photons qu'a émise sa source durant la communication, peut facilement calculer le pourcentage d'information qu'Ève a potentiellement en sa possession. L'*amplification de confidentialité* consiste à recombinaison certains bits de la clé entre eux, pour réduire le pourcentage d'information qu'a obtenu Ève. Aussi, plus ce pourcentage initial est grand, plus le degré de recombinaison à effectuer va être élevé pour restaurer la confidentialité entre Alice et Bob, et plus la taille de la clé utilisable va être réduite par rapport à celle échangée initialement.

A.1.2 La correction d'erreurs

Notons aussi que dans la configuration réelle où les détecteurs de Bob peuvent avoir des coups sombres avec une probabilité p^{dc} , il apparaît des erreurs dans la clé¹. Toutes ces erreurs dans la clé doivent être corrigées au détriment de la taille de la clé finalement utilisable.

Le théorème de Shannon [91] donne le nombre de bits $N_{Shannon}$ que doivent sacrifier Alice et Bob afin corriger une quantité d'erreur e pour une clé de taille n :

$$\frac{N_{Shannon}}{n} = -e \log_2 e - (1 - e) \log_2 (1 - e) \quad (\text{A.1})$$

Estimons la valeur de e en considérant seulement les coups sombres dans les détecteurs de Bob comme source d'erreur. Supposons que la probabilité brute d'avoir une détection chez Bob s'écrive

$$p^{brut} = p^{net} + p^{dc} - p^{net} \times p^{dc}$$

- où p^{net} correspond à la probabilité que Bob détecte un vrai photon.
- et p^{dc} la probabilité d'avoir un coup accidentel qui sera interprété aléatoirement comme un des deux résultats possibles (bit 0 ou bit 1) pour Bob, ce qui introduit 50% d'erreur.

¹ Au même titre qu'Ève peut en introduire si elle s'attaque à une partie des impulsions contenant un seul photon.

Ainsi le taux d'erreur dans la clé est finalement modélisé par :

$$e \approx \frac{\frac{1}{2}p^{dc}}{p^{brut}}$$

A.2 La distance maximale de l'échange quantique de clé

Par ailleurs, il est intéressant de constater que la quantité d'impulsions contenant deux photons va limiter la distance maximale², au delà de laquelle Ève aurait accès à toute l'information que partageraient Alice et Bob sans introduire d'erreur, donc sans révéler sa présence. Nous allons tâcher d'expliquer comment : supposons que, lors de l'échange, les photons voient des pertes à la propagation α (en dB/km), telles que la probabilité d'avoir un photon chez Bob vaille :

$$10^{-\frac{\alpha D}{10}} P_1$$

La distance maximale est atteinte quand $10^{-\frac{\alpha D_{max}}{10}} P_1$ est égal à P_2 , car dans cette configuration :

- Bob s'attend à recevoir des photons avec un taux $10^{-\frac{\alpha D_{max}}{10}} P_1$.
- Ève se place juste à la sortie des locaux d'Alice.
- Pour chaque impulsion, elle effectue une mesure non-destructive du nombre de photons. S'il n'y a qu'un seul photon dans l'impulsion, alors elle la bloque. Sinon elle sépare les photons de l'impulsion pour n'en garder qu'un seul, qu'elle place dans des mémoires quantiques.
- Par le biais d'une fibre sans pertes³, ou bien d'un *téléporteur quantique*⁴, elle transfère le deuxième photon chez Bob qui recevra alors le nombre d'impulsions attendues (P_2) et ne se rendra compte de rien.
- Une fois qu'Alice et Bob auront révélé leur base, Ève mesurera dans la bonne base le photon qu'elle a gardé dans les mémoires quantiques.

Cette distance maximale est définie par les pertes maximales que la transmission sécurisée peut endurer αD_{max} . Nous trouvons la relation :

$$P_2 = 10^{-\frac{\alpha D_{max}}{10}} P_1 \quad (\text{A.2})$$

À partir de la relation A.2, nous pouvons aisément justifier l'assertion de l'introduction en constatant que la probabilité P_2 va limiter la distance D_{max} de l'échange.

²en réalité des pertes maximales

³Dans le domaine de la cryptographie, on considère toujours que l'espion dispose de moyens technologiques illimités . . .

⁴. . . vraiment illimités

En réécrivant la relation précédente en fonction de P_1 et $g^{(2)}(0)$ qui sont les figures de mérites associées aux sources de photons uniques, nous trouvons :

$$P_2 = 10^{-\frac{\alpha D_{max}}{10}} P_1 = g^{(2)}(0) \frac{P_1^2}{2}$$

$$\boxed{P_1 \times \frac{g^{(2)}(0)}{2 \cdot 10^{-\frac{\alpha D_{max}}{10}}} = 1} \quad (\text{A.3})$$

Nous pouvons voir maintenant à partir de la relation A.3, qu'une façon plus élégante d'accroître la distance de l'échange, serait de disposer d'une « vraie » source de photons uniques présentant le plus petit $g^{(2)}(0)$ possible, c'est-à-dire la plus petite probabilité P_2 pour une probabilité P_1 donnée.

Il se trouve que le bruit dans les détecteurs est le principal facteur limitant la distance d'échange et que la probabilité d'avoir deux photons vient simplement raccourcir cette distance maximale. Le principe du calcul élaboré par Lütkenhaus est de tenir compte de toutes les pertes possibles et de la fuite d'information vers Ève pour calculer le taux de bits échangés vraiment utilisable par Alice et Bob. Cette valeur est calculée après l'amplification de confidentialité et la correction d'erreurs et en fonction de la distance.

A.3 Évaluation de p^{net}

Sachant qu'Alice utilise une source qui envoie des impulsions, contenant un photon unique, avec une probabilité P_1 , la probabilité d'obtenir un « clic » de l'autre côté chez Bob s'écrit :

$$\boxed{p^{net} \approx P_1 \times 10^{-\frac{\alpha D}{10}} \times \eta_d}$$

avec α les pertes en dB/km sur la ligne et η_d l'efficacité des détecteurs de Bob.

A.4 Le taux de bits utilisables par impulsion

Nous avons extrait de la référence [90], le taux de Bits utilisables par impulsion vaut finalement :

$$G = \frac{1}{2}p^{brut} \times \left\{ \frac{p^{brut} - P_2}{p^{brut}} \left[1 - \log_2 \left(1 + 4e \frac{p^{brut}}{p^{brut} - P_2} - 4 \left(e \frac{p^{brut}}{p^{brut} - P_2} \right)^2 \right) \right] + e \log_2 e + (1 - e) \log_2 (1 - e) \right\} \quad (\text{A.4})$$

A.5 Calcul du taux de bits utilisables à partir de données expérimentales

Nous nous sommes évidemment placé dans le cas d'un échange de clé dans l'infrarouge à l'aide de fibre optique télécom. Les spécifications d'une telle communication sont :

- Les photons uniques ont une longueur d'onde de 1550 nm .
- À cette longueur d'onde les pertes dans la fibre optique valent généralement $0,3 \text{ dB/km}$.
- En contrepartie, l'efficacité η_d du détecteur de Bob vaut rarement plus de 0,1 et présente une probabilité p^{dc} d'avoir un coup accidentel égale à $5 \cdot 10^{-5}$

Nous avons représenté sur la figure A.1 le taux de bits utilisables par impulsion pour les sources de photons uniques annoncés de Nice et Genève, ainsi que pour un laser atténué. Il va sans dire que pour le laser atténué les valeurs numériques sont identiques à celles utilisées pour les sources de photons uniques, à la différence près que P_1 et P_2 valent maintenant⁵ :

$$P_1 = \mu e^{-\mu} \quad \text{et} \quad P_2 = \frac{\mu^2}{2} e^{-\mu}$$

avec μ le nombre moyen de photon par impulsion du laser atténué.

Remarquons à propos de cette figure que pour connaître le débit auquel une clé utilisable est échangée, il faut multiplier l'axe des ordonnées par la fréquence de répétition de la source de photons. Par exemple, pour les deux sources rencontrées dans le manuscrit il faudrait multiplier par 100 kHz ce qui correspondrait un débit efficace d'environ 500 Hz à 10 km et seulement 1 Hz à 50 km . Ces débits peuvent paraître faibles, mais il faut les comparer

⁵L'utilisation de ces données n'est valable que pour des valeurs de μ strictement inférieures à 0,1. Dans le cas contraire, les valeurs de P_3 , P_4 ... ne sont plus négligeables et doivent entrer en compte dans le calcul A.4.

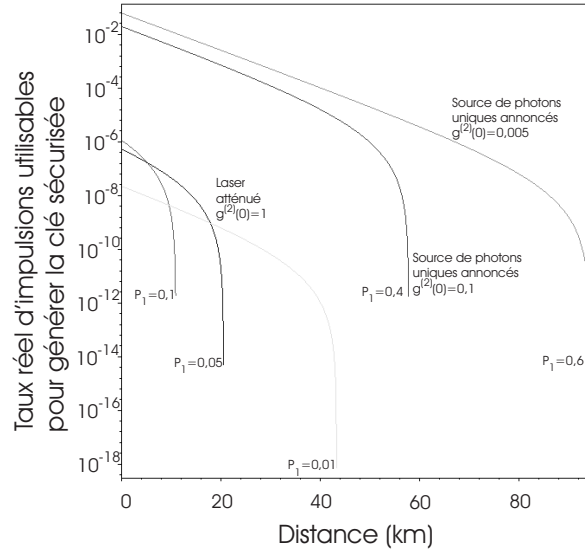


FIG. A.1 – Taux de bits utilisables par impulsion en fonction de la distance pour un laser atténué et pour les sources de photons uniques présentées dans ce manuscrit. Nous pouvons nous interroger sur la validité du calcul en constatant que les distances présentées ici ont toutes été dépassées expérimentalement [22, 23, 24]. Il faut garder à l'esprit que le calcul a été conduit ici avec des grandeurs expérimentales typique qui ne représentent pas forcément le meilleur de la technologie actuelle.

avec ceux d'un laser atténué. Dans ce cas, pour le laser atténué à 0,05 photon par impulsion, le débit efficace à 50 km est nul, tandis qu'à 10 km il faudrait utiliser un laser pulsé à 500 GHz!!! pour obtenir un débit équivalent à celui obtenu en utilisant notre source.

À la vue de ces rapides estimations, il est parfaitement clair qu'une source de photons uniques présentant un $g^{(2)}(0) < 1$ permet d'atteindre de plus grandes distances, tandis qu'une forte probabilité P_1 permet d'augmenter de manière considérable le débit efficace lors de l'échange de la clé.

Annexe B

Caractéristiques des APDs utilisées

Cette annexe n'a pas pour objectif d'expliquer précisément le fonctionnement d'une photodiode à avalanche. Nous souhaitons simplement introduire ici les connaissances « de base » qui permettent de comprendre et de distinguer les différents régimes de fonctionnement des APDs que nous avons utilisées durant ce travail. Cette annexe est inspirée des annexes D et E de la référence [68] et nous invitons le lecteur à s'y plonger pour obtenir des informations complètes et précises à ce sujet

B.1 Les photodiodes à avalanches

La détection de photons uniques se fait à l'aide de photodiodes à avalanches (APD) dont le type dépend de la longueur d'onde des photons que l'on souhaite détecter. Ce sont communément des jonctions semi-conductrices de type *PN* susceptibles de supporter des tensions de polarisation inverses. La caractéristique tension–courant inverse d'une jonction présente un front si raide (tension d'avalanche ou de claquage) que cette polarisation confère un véritable gain au système pour la détection (cf. figure B.1). L'idée est alors de donner à l'APD une tension V_{pol} juste inférieure à sa tension d'avalanche, l'énergie manquante pour déclencher l'avalanche étant fournie par le photon lui-même.

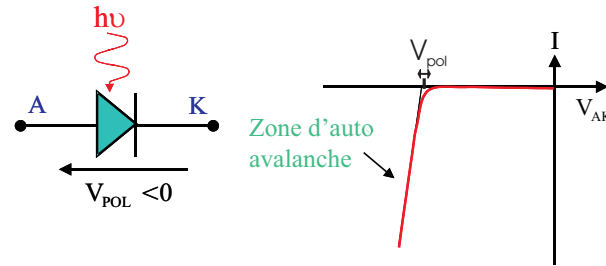


FIG. B.1 – Caractéristique schématique $I = f(V)$ où I est le courant inverse dans la jonction et V_{pol} la tension de polarisation inverse. De la valeur de V_{pol} dépend la facilité avec laquelle l'avalanche peut se déclencher suite à l'énergie qu'apporte un photon : c'est l'efficacité de détection η_{APD} . Attention aux conventions : A est l'anode de la diode et est reliée à la zone dopée P et K est la cathode reliée à la zone dopée N.

Les paramètres importants auxquels il faut s'intéresser sont :

L'efficacité quantique de détection qui dépend de la tension de polarisation V_{pol} appliquée et de la longueur d'onde des photons incidents. Cette efficacité correspond en quelque sorte au gain cité plus haut.

Le taux de coups sombres à tension de polarisation constante. Un coup sombre correspond à une avalanche déclenchée indépendamment d'un photon. Cela correspond à la rupture d'une liaison de valence par apport d'énergie thermique (vibration du réseau cristallin). Il s'ensuit la libération d'une paire électron-trou susceptible d'engendrer l'avalanche.

Comment stopper l'avalanche ? Pour cela il faut savoir quel est la densité de défauts au sein du réseau cristallin de la jonction considérée. En effet, si très peu de défauts sont présents comme pour le Silicium (*Si*) ou le Germanium (*Ge*), il est possible de polariser la jonction par le biais d'un simple pont de résistances fortement déséquilibré qui servira aussi à stopper l'avalanche comme le montre la figure B.2 ci-dessous. Une fois ce travail effectué, ce même circuit re-polarise automatiquement l'APD qui est de nouveau prête à déclencher. Et ainsi de suite... Ce principe est généralement désigné par le nom de « extinction passive », dont la figure B.2 montre le principe de fonctionnement.

Cependant certains matériaux comme les semi-conducteurs (*InGaAs* par exemple) possèdent beaucoup de défauts ce qui provoque le piégeage des porteurs lors du passage de l'avalanche. C'est pourquoi un tel dispositif de re-polarisation automatique ne saurait être adapté puisqu'il favoriserait, dès que le gain serait à nouveau suffisant, ce que l'on

appelle communément les « post-avalanches » (vient de l'anglais *after-pulses*). Ces effets correspondent en fait à de « fausses avalanches » uniquement dues à des paires e^-/h^+ piégées et ré-accélérées par le retour de la polarisation du mode passif. Pour ces dispositifs, la solution est un fonctionnement en mode déclenché (de l'anglais *gated*), qui privilégie une polarisation de la diode au moment même où le photon lui arrive. Cela nécessite évidemment la connaissance préalable du temps d'émission du photon...

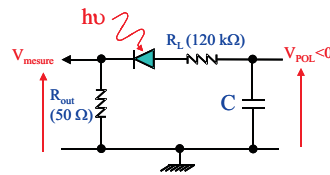


FIG. B.2 – Circuit de polarisation de type *extinction passive*. La tension inverse V_{pol} est appliquée de façon permanente via le pont de résistance. La capacité C joue un rôle de filtrage. Lorsqu'une avalanche se produit, l'APD se comporte comme un fil de jonction et on dit alors qu'elle se « vide » dans la résistance de $50\ \Omega$ aux bornes de laquelle se fait la mesure. La polarisation se trouve donc maintenant sur la résistance de $120\ k\Omega$ ce qui laisse à l'APD le temps de se décharger. Le temps de décharge est donnée par la constante de temps du circuit RC composé par la capacité interne de la photodiode et de la résistance de sortie et vaut typique quelques centaines de ns .

Pour finir sur ces détecteurs, reste à savoir quel matériau choisir. Il dépend bien entendu de l'application à laquelle on le destine et on peut dresser ainsi la classification suivante :

- Les APDs en Silicium (Si) seront utilisées dans le cas de photons émis dans le visible et le très proche infrarouge ($< 1\ \mu m$). Elles possèdent de très bonnes efficacités quantiques de détection ($> 60\%$), de très faibles taux de coups sombres et fonctionnent très bien en mode passif et à température ambiante. La fabrication de ces photodiodes a bénéficié de l'énorme effort qui fût fait pour le développement des semi-conducteurs en Silicium suite à la découverte du Transistor dans les années 60. Ce développement se fit bien entendu au détriment du Germanium.
- Les APDs en Germanium (Ge) sont très utiles pour la détection de photons appartenant à la première fenêtre télécom, c'est à dire autour de $1310\ nm$. Bien qu'elles possèdent l'énorme avantage de fonctionner en mode passif, elles doivent par contre être impérativement refroidie par un bain d'azote liquide ($77\ K$). Malgré cela elles présentent

des efficacités plutôt faibles ($\sim 10\%$ pour les meilleures) et de forts taux de coups sombres (typiquement environ 30 kHz pour 10% d'efficacité). A 77 K , ces dispositifs présentent une longueur d'onde de coupure à 1450 nm rendant impossible leur utilisation pour la détection de photons à 1550 nm . Il faut préciser ici que ces APDs sont actuellement très difficiles à trouver sur le marché puisque les grandes marques telles *NEC* et *FUJITSU* ont stoppé leur production faute de demande (seulement quelques laboratoires de recherche les utilisent dont le GAP et le LPMC).

- Les APDs en Arseniure de Gallium/Indium sur substrat de Phosphure d'Indium (*InGaAs*). Matériau en plein développement, l'*InGaAs* est très utilisé pour la détection de photons dans la seconde fenêtre télécom, c'est à dire autour de 1550 nm . Plus commodes à utiliser que les APDs en Germanium en termes de températures de fonctionnement (meilleur rendement obtenu typiquement entre -50°C et -60°C [92, 93]), elles possèdent cependant une densité de défauts excessivement importante interdisant toute extinction passive des avalanches. Le mode pulsé est donc requis ce qui oblige du coup les expérimentateurs à connaître les instants d'arrivée des photons et ce qui compromet, par conséquent, toute expérience continue.

Par ailleurs, ces détecteurs présentent typiquement 20% d'efficacité à 1550 nm et sont aussi deux fois meilleures que les Germanium à 1310 nm qui restent malheureusement les seules à pouvoir fonctionner en extinction passive. Le mode pulsé a cependant l'avantage de limiter les taux de coups sombres puisque les photodiodes *InGaAs* ne sont actives uniquement lorsqu'elles sont tenues de l'être.

B.2 L'APD *Germanium*

B.2.1 Mesure de l'efficacité de détection en mode passif

Pour déterminer l'efficacité de la photodiode, on l'éclaire avec une lumière très faible, c'est à dire dont l'intensité peut être exprimée en nombre de photons par seconde N_p . Aussi, il faut que cette quantité soit bien plus faible que le taux de comptage pour lequel le détecteur entre en saturation. À l'aide d'un compteur, nous accédons aux taux brut N_d d'avalanches par secondes sous éclairage ainsi que le taux N_s de coups sombres obtenu lorsque l'éclairage est coupé. À la sortie des photodiodes, les signaux traités sont analogiques et non homogènes les uns aux autres. Ils faut donc prévoir un dispositif de mise en forme des signaux afin d'obtenir des formes identifiables par les

systèmes de comptage : c'est le rôle des discriminateurs qui fonctionnent suivant un critère de seuil donnant lieu à un 1 logique si le signal franchit le seuil ou à un 0 logique dans le cas contraire. Ainsi, le discriminateur joue un rôle important sur l'efficacité globale de l'APD. Il est donc incorrect de parler d'*efficacité de détection* et de *taux de coups sombre* d'une APD sans inclure le système de mise en forme des signaux qui lui est associé.

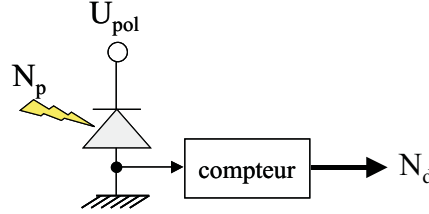


FIG. B.3 – Principe de caractérisation d'une photodiode en mode passif. U_{pol} représente la tension de polarisation inverse appliquée à la diode, N_p le nombre de photons incidents et N_d le nombre de coups détectés par seconde.

L'efficacité est alors donnée par la relation simple :

$$\eta = \frac{N_d - N_s}{N_p} \quad (\text{B.1})$$

Par la suite, le couple *APD-discriminateur* allant devenir un système de mesure à part entière, il faut déterminer l'incertitude liée à son efficacité :

$$\Delta\eta = \frac{1}{N_p} \cdot \left[\Delta N_d + \Delta N_s + \frac{N_d - N_s}{N_p} \cdot \Delta N_p \right] \quad (\text{B.2})$$

Sachant que les mesures des taux de comptage sont soumises à des distributions de type Poisson, l'erreur associée est de l'ordre de leur racine carrée. Il vient donc dans ce cas :

$$\Delta\eta = \frac{1}{N_p} \cdot \left[\sqrt{\frac{N_d}{t_d}} + \sqrt{\frac{N_s}{t_s}} + \frac{(N_d - N_s)}{N_p} \sqrt{\frac{N_p}{t_d}} \right] \quad (\text{B.3})$$

où cette fois

- N_p est le nombre de photons incidents¹ pendant le temps de détection t_d ,
- N_d est le nombre de coups détectés pendant le temps t_d ,
- et N_s le nombre de coups sombres comptés pendant le temps t_s .

¹Évidemment cette quantité n'est pas absolue est dépend de la précision de l'appareil utilisé pour l'estimer.

B.2.2 Courbe d'étalonnage de l'APD-Ge

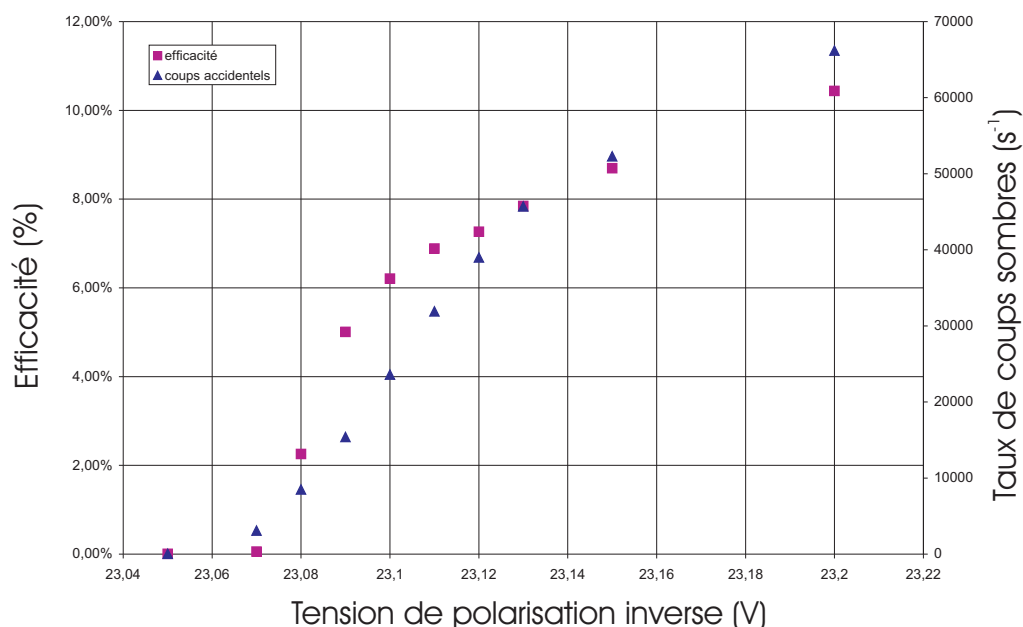


FIG. B.4 – Courbe d'étalonnage de l'efficacité et du taux de bruit de l'APD-Ge en mode passif associée au discriminateur *ORTEC*. L'erreur caractéristique sur l'efficacité est de l'ordre de 0,1%.

B.3 L'APD *InGaAs*

B.3.1 Mesure de l'efficacité de détection en mode déclenché

Par rapport à la méthode précédente, la façon de procéder est de pulser la lumière incidente et de profiter dans le même temps de ce déclenchement pour activer la détection. La photodiode est alors polarisée grâce à la superposition de deux tensions (voir figure ci-dessous) : l'une dite de "biais" correspondant à une composante continue permanente (U_{biais}), et l'autre qui consiste en un créneau capable d'amener ce qu'il manque de tension au moment où il faut (U_{pulse}). Le biais est juste là pour ne pas avoir à jouer avec de trop gros fronts de tension ; il est en effet plus facile de moduler une tension entre 0 et 5 V qu'entre 0 et 20 V.

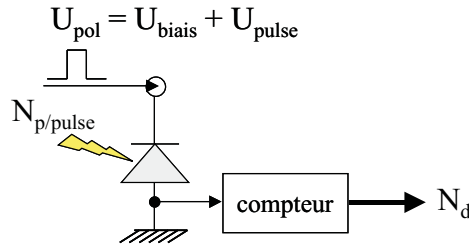


FIG. B.5 – Principe de caractérisation d'une photodiode en mode déclenché. U_{pol} représente la tension de polarisation inverse totale appliquée à la diode, $N_{p/pulse}$ le nombre de photons compris dans un pulse incident et N_d le nombre de photons détectés.

Que la lumière à analyser soit issue du rayonnement d'un laser ou que l'on ait affaire à un signal de fluorescence paramétrique, le nombre $N_{p/pulse}$ de photons contenus dans une impulsion est régit par une loi de distribution de Poisson de type :

$$p_{poisson}(n, m) = m^n \cdot \frac{e^{-m}}{n!} \quad (B.4)$$

où $p_{poisson}(n, m)$ représente la probabilité d'avoir n photons dans un pulse sachant qu'en moyenne il y en a m . De là, si l'on note η l'efficacité du détecteur, et b la probabilité d'obtenir un coup sombre dans la fenêtre d'activation, alors la probabilité p de "compter un coup" est donnée par :

$$p = e^{-m}b + \sum_{n=1}^{\infty} \frac{[1 - (1 - \eta)^n] \cdot (1 - b) \cdot m^n e^{-m}}{n!} \quad (B.5)$$

où l'on dénote, dans un ordre utile à la compréhension :

- (i) $(1 - \eta)^n$ la probabilité *de ne pas détecter* un photon lorsque n sont incidents,
- (ii) $1 - (1 - \eta)^n$ la probabilité *de détecter* un photon lorsque n sont incidents,
- (iii) $(1 - b)$ la probabilité de ne pas détecter un coup sombre,
- (iv) et enfin le produit $(1 - (1 - \eta)^n) \cdot (1 - b)$ qui représente donc la probabilité compter un coup dû à un photon et non au bruit.

Notons que la somme se fait alors sur toutes les valeurs possibles du nombre de photons n effectifs et que le terme $(1 - (1 - \eta)^n) \cdot (1 - b)$ est logiquement pondérée à chaque incrément par la valeur de la distribution de Poisson correspondante. Par ailleurs, la probabilité d'obtenir un coup sombre dans la

fenêtre active doit diminuer en fonction du nombre moyen de photons dans le pulse incident : c'est le rôle de la pondération e^{-m} .

En résolvant η selon n , il vient de façon non triviale :

$$\boxed{\eta = \frac{\ln\left(\frac{b-1}{p-1}\right)}{m}} \quad (\text{B.6})$$

L'intérêt de la méthode réside alors dans le fait que les probabilités p et b sont très facilement accessibles expérimentalement et s'obtiendront par :

$$p = \frac{\text{nombre de coups mesurés}}{\text{nombre de coups totaux}} = \frac{N_m}{N_t}, \text{ laser allumé} \quad (\text{B.7})$$

$$b = \frac{\text{nombre de coups sombres}}{\text{nombre de coups totaux}} = \frac{N_b}{N_t}, \text{ laser éteint} \quad (\text{B.8})$$

Aussi le nombre moyen de photons m contenus dans une impulsion peut être simplement contrôlé en jouant sur la transmission variable d'un atténuateur. On a alors :

$$m = N_{p/pulses} \cdot T_{\text{atténuateur}} \quad (\text{B.9})$$

L'expression de l'efficacité de détection devient alors :

$$\eta = \frac{1}{N_{p/pulse} \cdot T_{\text{atténuateur}}} \cdot \ln\left(\frac{N_b - N_t}{N_m - N_t}\right) \quad (\text{B.10})$$

De là, en définissant les erreurs sur p , b et m par

$$\begin{cases} \Delta p = \frac{\sqrt{N_m}}{N_t} \\ \Delta b = \frac{\sqrt{N_b}}{N_t} \end{cases} \quad (\text{B.11})$$

et

$$\Delta m = N_{p/pulse \text{ initial}} \cdot \Delta T_{\text{atténuateur}} \quad (\text{B.12})$$

on obtient, toujours à l'aide de la différentielle logarithmique, l'erreur sur η :

$$\Delta\eta = \frac{1}{m(b-1)} \cdot \Delta b + \frac{1}{m(p-1)} \cdot \Delta p + \frac{\ln\left(\frac{b-1}{p-1}\right)}{m^2} \cdot \Delta m \quad (\text{B.13})$$

$$\Rightarrow \Delta\eta = \frac{1}{m} \cdot \left[\frac{\sqrt{N_b}}{N_b - N_t} + \frac{\sqrt{N_m}}{N_m - N_t} + \ln\left(\frac{N_b - N_t}{N_m - N_t}\right) \cdot \frac{\Delta T_{\text{atténuateur}}}{T_{\text{atténuateur}}} \right] \quad (\text{B.14})$$

Rappelons si besoin est que toutes les quantités représentées dans les formules B.10 et B.14 sont accessibles expérimentalement.

B.3.2 Courbe d'étalonnage des APD-*InGaAs*

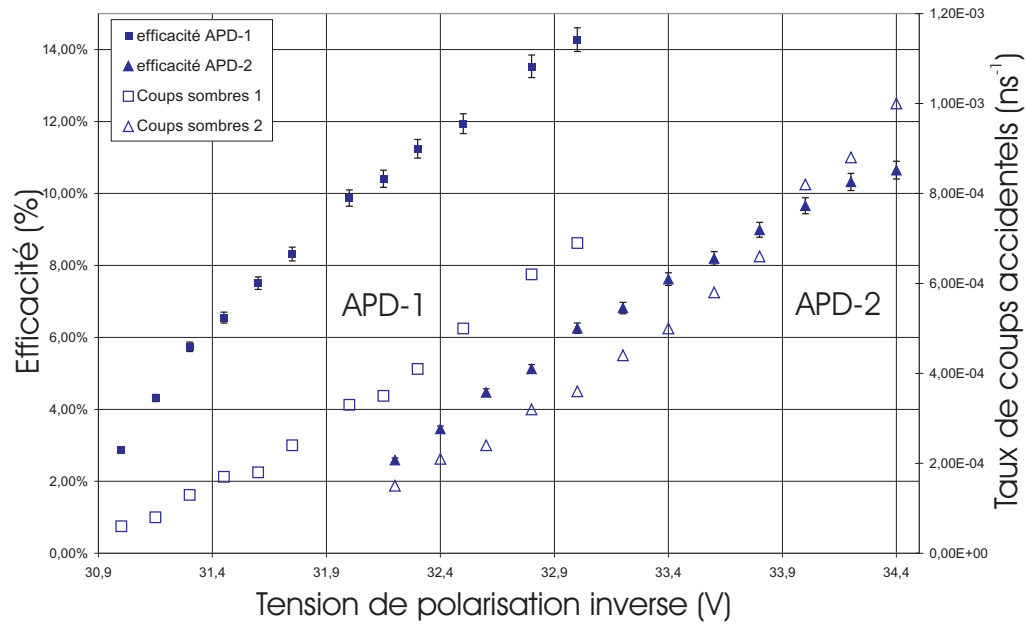


FIG. B.6 – Courbe d'étalonnage de l'efficacité et du taux de bruit des APD-*InGaAs* en mode déclenchée associées à leur carte de discrimination.

Annexe C

Mesure de l'intervalle de temps séparant deux paires

Cette annexe a pour objectif de présenter la méthode expérimentale que nous avons utilisée pour mesurer la statistique des intervalles de temps qui séparent l'instant de création de deux paires successives.

C.1 Rappel de la théorie dans le cas d'une distribution poissonnienne

Prenons des événements qui peuvent arriver n'importe quand (la création de paires de photons en régime continu par exemple... ou la détection de l'un d'eux). Ce type d'événement présente une statistique poissonnienne caractérisée par un taux moyen d'occurrences μ . Les paires peuvent donc être créées n'importe quand mais sur une très longue période de temps, nous attendons en moyenne $\bar{n} = \mu T$ paires. Il est alors assez facile de répondre à la question : « *Quelle est la probabilité d'avoir alors précisément N événements, alors que nous attendons un nombre moyen $\bar{n}$?* »

$$P_{\bar{n}}(N) = \frac{\bar{n}^N}{N!} e^{-\bar{n}} \quad (\text{C.1})$$

Pourtant, il est beaucoup moins évident de répondre à la question : « *Quelle est la probabilité qu'un intervalle de temps ΔT sépare alors deux événements successifs ?* »

Sous cette forme, nous ne pouvons pas répondre à la question, toutefois, il est possible d'estimer la probabilité qu'un intervalle de temps supérieur ou égal à ΔT sépare deux événements successifs. Il s'agit de la probabilité

(poissonnienne) d'avoir zéro événements sachant qu'on a un nombre moyen d'événements égal à $\bar{n} = \mu \Delta T$ durant cet intervalle de temps :

$$P(\Delta T) = P_{\bar{n}}(0) = e^{-\mu \Delta T} \quad (C.2)$$

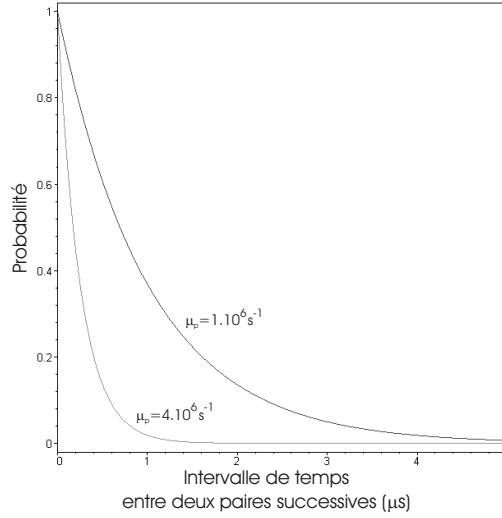


FIG. C.1 – Probabilité qu'un intervalle de temps supérieur ou égal à ΔT sépare deux événements successifs pour $\mu = 1 \cdot 10^6 s^{-1}$ et $4 \cdot 10^6 s^{-1}$.

Cette fonction est représentée graphiquement sur la figure C.1 pour des taux moyens de paires $\mu_p = 1 \cdot 10^6 s^{-1}$ et $4 \cdot 10^6 s^{-1}$. Il faut comprendre à partir du schéma que la probabilité pour qu'un intervalle infiniment petit ne contienne aucun photon est proche de un. Aussi plus l'intervalle ΔT est grand, plus la probabilité qu'il soit vide est petite. De même, plus le taux moyen de paires est grand, plus la probabilité qu'un intervalle ΔT soit vide est petite.

Fort de cela, nous pouvons tout de même estimer la probabilité qu'un intervalle ΔT précis sépare deux événements successifs. Dans ce cas nous devons considérer une tranche infinitésimale dt et nous poser la question : *Quelle est la probabilité (infinitésimale) que l'intervalle de temps tombe dans l'intervalle dt centré sur ΔT ?* La réponse s'écrit :

$$P(t)dt = \mu e^{-\mu t} dt$$

C.2 La mesure expérimentale de l'intervalle de temps séparant deux paires successives

Comment mesurer expérimentalement les intervalles de temps qui séparent deux détections successives ?

Supposons que notre photodiode détecte des photons qui donnent naissance à des impulsions électriques avec un taux moyen μ . Sur le schéma C.2, nous avons, pour l'exemple, numéroté cinq d'entre elles et nous cherchons donc le moyen de mesurer la durée ΔT qui sépare l'impulsion 1, de la 2 puis la 2 de la 3 et ainsi de suite... Afin de réaliser cette mesure, nous allons tirer parti des limitations du *TAC*¹ entrevues dans la section 3.2 du chapitre 3.

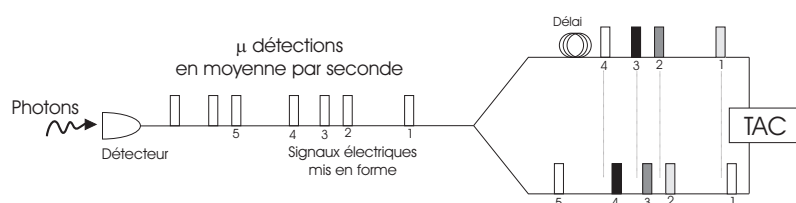


FIG. C.2 – Schéma expérimental pour mesurer l'intervalle de temps entre deux détections successives.

Le *TAC* permet de mesurer le temps qui sépare l'arrivée d'un signal sur son entrée *start* et d'un second signal sur son entrée *stop*. Nous allons pour cela commencer par dédoubler les impulsions électriques afin qu'elles soient dirigées vers le *start* et aussi vers le *stop* du *TAC*. Toute la technique réside dans le fait que le *TAC* attend aveuglement l'arrivée d'un *start* pour commencer la mesure, puis l'arrivée d'un *stop* pour enregistrer une durée.

Supposons que nous retardions faiblement (quelques nanosecondes suffisent) les signaux qui prennent la direction de l'entrée *start* du *TAC*. Dans cette configuration, l'impulsion 1 qui arrive d'abord sur l'entrée *stop* n'est pas prise en compte par le *TAC*, tandis que son homologue qui arrive quelques nanosecondes plus tard sur l'entrée *start* va déclencher le *TAC*. Maintenant, il se trouve que l'impulsion suivante, qui va arriver sur l'entrée *stop*, n'est autre que la numéro 2. Voici comment le *TAC* vient de mesurer la durée qui sépare l'impulsion 1 de l'impulsion 2. Il est assez facile d'imaginer que ce processus se répète autant de fois que nécessaire et permet de construire expérimentalement l'histogramme de la figure C.3. Attention toutefois, la valeur des intervalles de temps que nous mesurons alors n'est pas absolue

¹Time to Amplitude Converter

190 ANNEXE C. INTERVALLE DE TEMPS SÉPARANT DEUX PAIRES

et contient implicitement la valeur du retard que nous avons imposé sur la voie *stop* du *TAC*. Cependant, à la vue de la base de temps de l'ordre de la dizaine de microsecondes utilisée pour l'expérience, ce retard de quelques nanosecondes est négligeable.

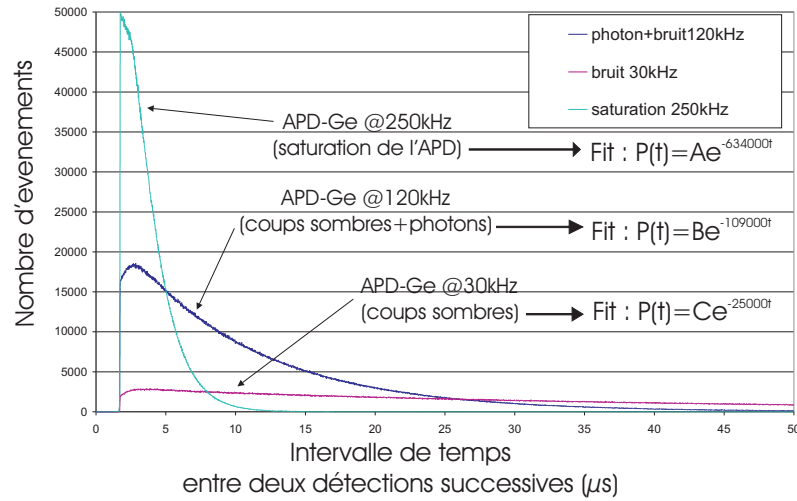


FIG. C.3 – Distributions des intervalles de temps pour différents taux de comptage dans le détecteur germanium.

Nous avons reporté sur le graphique C.3 trois mesures expérimentales, accompagnées par la formule du fit qui s'ajuste à la courbe. Il faut noter que les courbes ne continuent pas jusqu'aux très faibles valeurs de ΔT car le *TAC* ne peut mesurer deux impulsions incidentes plus proches que $2\mu s$ dans cette base de temps. Cependant, il aurait été inutile d'avoir une meilleure résolution, car nous pouvons noter un phénomène de saturation des photodiodes au alentours de $250kHz$. En effet, l'affaissement des courbes, pour des durées inférieures à $4\mu s$, n'est pas dû à un comportement particulier des photons, mais à l'impossibilité qu'ont les photodiodes de détecter successivement deux photons trop proches de l'un de l'autre. À ce titre, les ajustements mathématiques ont été fait à l'aide de la formule C.2 sur une plage temporelle qui s'étend de $5\mu s$ à $50\mu s$.

Une première mesure a été réalisée à partir du bruit dans l'APD-Ge. Le taux comptage dans l'APD était alors d'environ $30kHz$. Nous constatons que l'ajustement mathématique propose une valeur de $\mu \approx 25kHz$ qui correspond assez bien à la valeur indiquée par le détecteur plus ou moins l'incertitude. Le bruit dans le détecteur suit bien une statistique poissonnienne.

Pour la seconde mesure, nous avons réalisé l'acquisition alors que la pho-

todiode détectait des photons à un taux moyen de 120 kHz . Selon la formule C.2, cette forme de distribution correspond bien à des événements poissonniens avec un taux moyen $\mu \approx 109\text{ kHz}$ qui est tout à fait en adéquation avec le taux de comptage affiché par le détecteur.

À la vue de ces deux résultats, nous pouvons nous demander si la distribution temporelle des impulsions électriques traduit bien le comportement des photons incidents ou si nous observons depuis le début le comportement aléatoire de la conversion *photon-avalanche d'électrons*.

Une dernière mesure, nous donc a permis de vérifier que la distribution des intervalles de temps dépendait bien de la statistique des photons incidents. Nous avons lancé l'acquisition en plaçant la photodiode dans une configuration « saturée », c'est-à-dire que cette dernière affiche un taux moyen de détections aux alentours de 250 kHz , alors que le taux réel de photons incidents devrait être bien plus élevé. Dans cette configuration, l'ajustement mathématique nous apprend que la statistique, à l'origine de ce type de courbe, présente un taux moyen d'événements $\mu \approx 630\text{ kHz}$, ce qui correspond effectivement à la caractéristique des photons incidents et non à celle des conversions photon-avalanche d'électron qui sont, de toutes façons, limitées à 250 kHz .

192 *ANNEXE C. INTERVALLE DE TEMPS SÉPARANT DEUX PAIRES*

Annexe D

La fabrication des guides PPLN

Cette annexe va nous permettre de révéler tout le travail technologique dissimulé derrière un guide d'onde sur niobate de lithium périodiquement polarisé (PPLN). La construction de cette annexe suivra chronologiquement le protocole de fabrication d'un échantillon. Nous expliquerons par exemple comment inverser périodiquement le coefficient non-linéaire du LiNbO_3 et comment contrôler la périodicité ou l'art d'inscrire ensuite un guide d'onde sur le matériau. Pour une description approfondie de la technique de retournement périodique des domaines et de fabrication de guide SPE, le lecteur pourra se tourner vers les références [69, 94, 95]

D.1 Le substrat de niobate de lithium

Les guides d'ondes que nous avons réalisés pour notre source de paires de photons ont tous été intégrés sur des substrats de niobate de lithium (LiNbO_3) qui est aujourd'hui l'un des matériaux optiques les plus couramment utilisés. En effet, avec plus de 70 tonnes produites par an dans le monde, on retrouve ce matériau dans de plusieurs domaines de pointes ainsi que dans nombre de travaux de recherche en optique intégrée et en optique non-linéaire, au même titre que le *KTP* ou le *GaAs*. Il constitue l'un des substrats privilégiés pour la réalisation de modulateurs d'intensité utiles aux systèmes de télécommunications à hauts débits.

Rappelons ici quelques unes des propriétés du niobate de lithium :

- Tout d'abord, il est, comme le *KTP*, *ferro-électrique* ce qui autorise, comme nous le verrons par la suite, l'inversion périodique du signe de son coefficient non-linéaire le plus fort.
- C'est par ailleurs un matériau *biréfringent*. On peut donc, lorsque l'on dispose d'un cristal massif, utiliser l'accord de phase par biréfringence.

Notons que cette propriété est aussi propre au KTP.

- Son *coefficient non-linéaire le plus élevé* (d_{33}), utilisable avec le QAP, est proche des 30 pm/V . Il est donc environ 5 fois meilleur que le coefficient d_{31} utilisé avec une accordabilité par biréfringence.
- Il est aussi *électro-optique* ce qui peut permettre une quasi-accordabilité rapide par modulation de la dispersion.
- Il possède un *seuil de dommage* optique proche de 200 MW/cm^2 ce qui reste une valeur assez confortable.
- Il possède une *bande de transparence* comprise entre $0,4$ et $5 \mu\text{m}$ qui nous permet d'obtenir pratiquement toutes les interactions désirées.
- Par contre, c'est un matériau *photo-réfractif* dont les effets s'observent dès 10 kW de puissance optique par cm^2 à la température de 20°C . Bien que la photo-réfractivité puisse constituer un problème majeur pour les structures guidantes, il faut savoir qu'il est possible de s'en affranchir simplement en chauffant l'échantillon considéré à $\sim 70^\circ \text{C}$.

C'est naturellement pour l'ensemble de ces propriétés que le niobate de lithium constitue un très bon candidat pour réaliser des composants non-linéaires intégrés.

Aussi, il faut savoir qu'il existe aujourd'hui de nombreuses méthodes pour obtenir la configuration PPLN dont nous avons besoin pour utiliser la technique du quasi-accord de phase. Par exemple, on notera la polarisation par champ électrique ou par diffusion titane. Par ailleurs, l'intégration des guides peut aussi se faire via diverses méthodes, comme notamment la diffusion titane, l'implantation ionique, ou encore l'une des méthodes de type échange protonique dont le SPE fait partie.

Cependant, toutes les techniques de poling et d'intégration des guides ne sont pas compatibles et les composants non-linéaires intégrés ne sont pas si simples à réaliser. En effet, il convient au préalable de s'assurer que l'intégration des guides ne soit pas susceptible d'effacer l'inversion périodique et qu'elle respecte la valeur du coefficient non-linéaire.

Détaillons en quoi le caractère ferro-électrique du LiNbO_3 nous intéresse et examinons au passage quelques autres particularités du cristal non citées ci-dessus.

Le niobate de lithium appartient au groupe rhomboédrique. La figure D.1 ci-dessous en donne une représentation dans une phase ferro-électrique. En fait, la croissance du cristal¹ se fait à des températures supérieures à sa température de Curie (comprise entre 1130 et 1200°C) pour lesquelles le cris-

¹Le niobate de lithium est un cristal de synthèse qui s'obtient par la méthode de Czochralski bien connues des gens qui font de la micro-électronique.

tal est dans la phase para-électrique : les ions lithium sont contenus dans les plans d'oxygènes tandis que les ions niobium se trouvent entre deux plans d'oxygènes à égale distance de ceux-ci. Le cristal est ainsi localement neutre. Lorsque celui-ci refroidit, les forces d'interactions au sein du cristal deviennent prédominantes devant l'électro-neutralité si bien que les ions lithium et niobium sont déplacés par rapport au plans d'oxygènes induisant une polarisation spontanée P_s au sein du cristal.

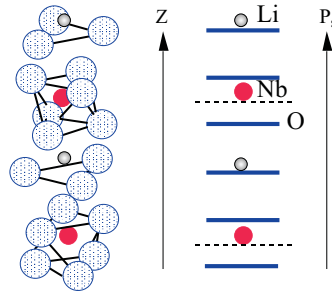


FIG. D.1 – La structure cristalline du niobate de lithium.

On comprend alors que deux polarisations spontanées opposées sont possibles du fait de l'équiprobabilité qu'ont les ions de basculer d'un côté ou de l'autre des plans d'oxygènes, comme le montre la figure D.2 ci-dessous. Or, l'orientation de la polarisation fixe le signe du coefficient non-linéaire (voir chapitre 6 de la référence [96]). C'est pourquoi, afin d'obtenir les conditions de *Quasi-accord de phase*, il suffira d'inverser périodiquement cette polarisation spontanée dont découlera finalement celle du signe du coefficient non-linéaire d_{33} . On parle alors de *l'inversion des domaines ferro-électriques*.

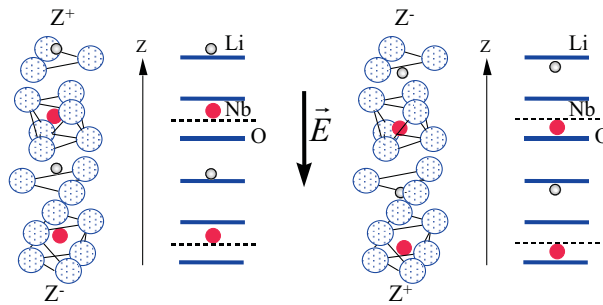


FIG. D.2 – Phases ferro-électriques du niobate de lithium.

Une dernière précision concerne la coupe du wafer à utiliser pour nos substrats. En effet, il existe trois différentes coupes possibles et qui se font

selon les trois axes cristallins X, Y et Z . Comme nous l'avons représenté sur les figures D.1 et D.2 ci-dessus et D.3 ci-dessous, seule la coupe selon l'axe Z porteur de l'indice extraordinaire propose la ferro-électricité dont nous avons besoin pour la polarisation périodique. Les autres coupes pourront permettre la réalisation de guides d'ondes mais pas de guides d'ondes non-linéaires puisque ni l'accord de phase par biréfringence (beaucoup trop complexe à obtenir) ni le QAP ne seront réalisables sur ces puces.

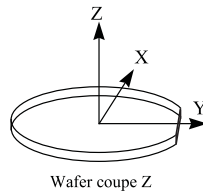


FIG. D.3 – Représentation d'un wafer de niobate de lithium en coupe Z .

D.2 Le retournement des domaines (ou « poling »)

Comme nous l'avons dit précédemment, de nombreuses méthodes existent (exodiffusion du lithium, diffusion titane, bombardement par faisceau d'électrons) pour polariser périodiquement le niobate. Toutefois, l'application d'un champ électrique parallèlement à l'axe de la polarisation spontanée reste la méthode la plus directe et la plus simple que l'on connaisse actuellement.

L'idée est donc ici d'appliquer périodiquement sur l'échantillon un champ électrique dont la valeur soit supérieure à celle du champ coercitif qui vaut environ 20 kV/mm à température ambiante. De cette façon, partant d'un échantillon mono-domaine de $500 \mu\text{m}$ d'épaisseur, on obtient une inversion de la polarisation spontanée au travers d'un masque diélectrique de période convenablement choisie et en rapport avec l'interaction désirée. Notons enfin que les interactions qui nous intéressent requièrent un pas d'inversion du signe du $\chi^{(2)}$ dont la valeur se situe autour de $13,6 \mu\text{m}$.

D.2.1 Le dispositif expérimental

Dans notre cas, le champ électrique est appliqué au cristal par le biais d'électrodes liquides. Pour réussir à appliquer le champ électrique à des endroits précis, nous délimitons spatialement les zones de contact de l'électrode

liquide grâce à un masque de résine isolante disposant d'ouvertures aux endroits voulus.

La première étape consiste donc à déposer le masque isolant sur la surface du wafer. Pour cela, nous utilisons la résine *SHIPLEY-1818* qui après cuisson se révèle être une résine positive² d'environ $1,3\mu m$ d'épaisseur. C'est à ce stade que nous choisissons la période d'inversion Λ du cristal. À l'aide du masque de lithographie approprié, nous venons graver, par photolithographie, les zones qui laisseront place à des ouvertures où l'électrode liquide viendra appliquer le champ électrique. Une représentation schématique d'un masque est donnée sur la figure D.4 sur lequel nous pouvons distinguer de petites ouvertures rectangulaires où la polarisation sera retournée. Ces ouvertures périodiques définissent des bandes de PPLN qui accueilleront par la suite les guides d'ondes.

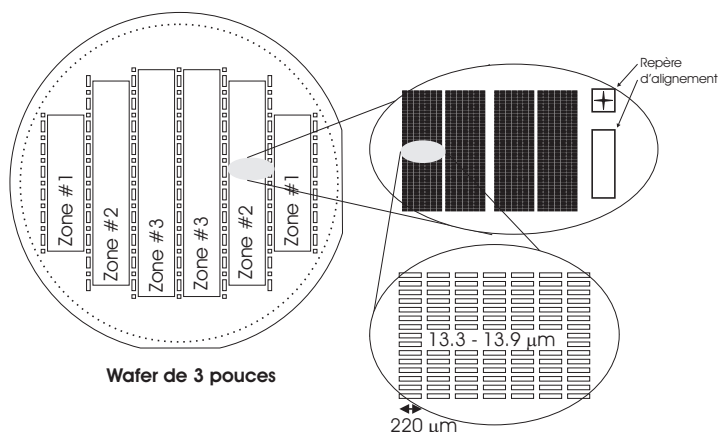


FIG. D.4 – Représentation d'un masque de poling des wafers de niobate de lithium en coupe Z.

L'échantillon est ensuite enserré entre deux coques de plexiglas creuses qui, une fois remplies par l'électrode liquide ($LiCl$), constituent la *cellule de poling* représentée sur la figure D.5 :

² « Positive » signifie qu'à l'attaque chimique, ce sont les zones qui ont subi l'insolation qui disparaissent.

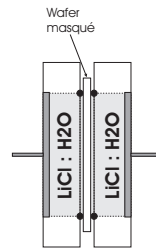
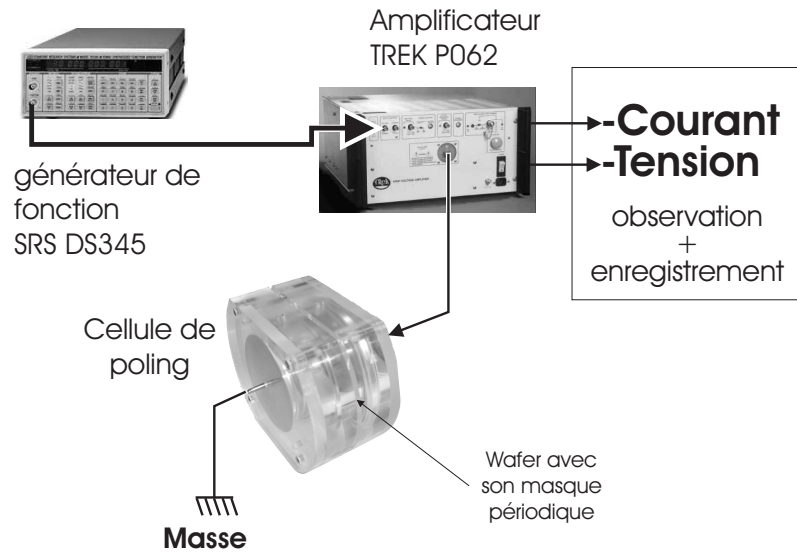


FIG. D.5 – Cellule de poling pour wafer 3 pouces.

Une fois le masque de résine posé, vient le moment d'appliquer le champ électrique à l'aide du montage représenté sur la figure D.6 :

FIG. D.6 – Montage pour le poling de wafer trois pouces de $LiNbO_3$.

Un générateur de fonction fournit une rampe de tension comprise entre 0 et 3 V qui est ensuite amplifiée par un facteur 3000. Nous obtenons ainsi des champs du même ordre de grandeur que le champ coercitif de notre wafer de $500\ \mu m$, c'est-à-dire $\sim 10\ kV$. La forme de la rampe de tension est discutée plus en détail dans les références [94, 95].

L'amplificateur dispose d'une première sortie (haute tension) qui est dirigée vers la cellule de poling et de deux autres (basse tension) qui permettent d'observer le courant et la tension qui est appliquée en temps réel sur l'échantillon.

Dès lors, l'inversion se fait en trois étapes selon la figure D.7 [94, 95] :

- (i) *La formation et la nucléation* : il s'agit de l'apparition de domaines microscopiques sur les zones non-masquées, en des endroits appelés sites de nucléation.
- (ii) *La propagation sous la forme d'une aiguille* : les micro-domaines créés vont ensuite se propager rapidement à travers toute l'épaisseur du cristal sous la forme d'une aiguille, tandis que leur base s'agrandit.
- (iii) *La propagation des murs* : les parois de l'aiguille se redressent jusqu'à devenir parallèles. Alors, les micro-domaines créés sur chaque face fusionnent dès qu'ils entrent en contact. Au final ils n'en forment plus qu'un qui remplit tout l'espace compris entre les électrodes.

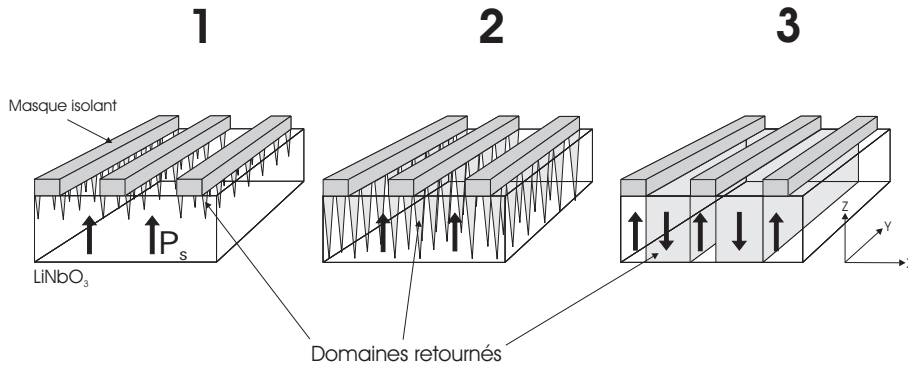


FIG. D.7 – Étapes principales de l'inversion de polarisation d'un cristal monodomaine avec un masque isolant et une électrode liquide.

Grâce à cette technique, on obtient une inversion du coefficient d_{33} dans toute la masse du matériau et non pas en surface uniquement comme c'est le cas avec des méthodes de diffusion titane par exemple [94].

D.2.2 Le contrôle des domaines

Pour obtenir l'interaction non linéaire la plus efficace, il est important que le duty-cycle (coefficient noté a page 72 du chapitre 2) qui correspond au rapport de la largeur des domaines retournés sur la période d'inversion, vaille $1/2$. En effet, ce dernier intervient dans la valeur du coefficient non-linéaire effectif utilisé pour le quasi-accord de phase, selon la formule suivante :

$$\chi_n = \frac{2d_{33}}{n\pi} \sin(n\pi a) \quad (\text{D.1})$$

La mesure de la largeur des domaines inversés est donc importante et devrait idéalement être faite, en temps réel, durant le processus de poling afin de pouvoir ajuster le duty-cycle.

Notons qu'il est difficile de voir les zones inversées à l'œil nu. Il est par contre possible de les observer en microscopie optique en transmission, les murs entre les domaines étant soumis à des contraintes que l'optique permet de mettre en évidence. Toutefois, le meilleur moyen de repérer les zones inversées et de vérifier que nos paramètres de poling sont bons consiste à réaliser une attaque chimique des faces de l'échantillon grâce à un mélange d'acide fluorhydrique (HF) et d'acide nitrique (HNO_3). La cinétique de l'attaque étant beaucoup plus forte sur les éléments polarisés négativement que sur ceux polarisés positivement, il va se produire le gravure d'un réseau en surface qui correspond à celui de l'inversion des domaines ferro-électriques. Bien sûr, après coup, le substrat PPLN testé n'est plus utilisable et il nous faut recommencer... d'où l'intérêt de développer une méthode de contrôle non-destructive.

Nous avons élaboré au LPMC une méthode qui permet de contrôler l'élargissement des domaines sans avoir recours à l'attaque acide.

En fonction du masque de poling utilisé, nous connaissons S_{totale} la surface totale des bandes de poling nommées *zone#1*, *zone#2* et *zone#3* sur le schéma D.4. Cette surface correspond une quantité totale de charges qu'il est possible de retourner en utilisant l'équation suivante

$$Q_{total} = 2P_s \times S_{totale} \quad (D.2)$$

où $P_s = 76 \mu C.cm^{-2}$ est la polarisation spontanée du niobate de lithium.

Évidemment notre objectif est de retourner seulement la moitié de la surface totale afin d'obtenir un *duty-cycle* de 1/2. Pour nous en assurer, nous avons expérimentalement accès à la quantité de charge qui est passé à travers le cristal grâce aux sorties auxiliaires de l'amplificateur. Voici typiquement le relevé que nous obtenons pour le courant I_{poling} après l'application de la rampe de tension U_{poling} .

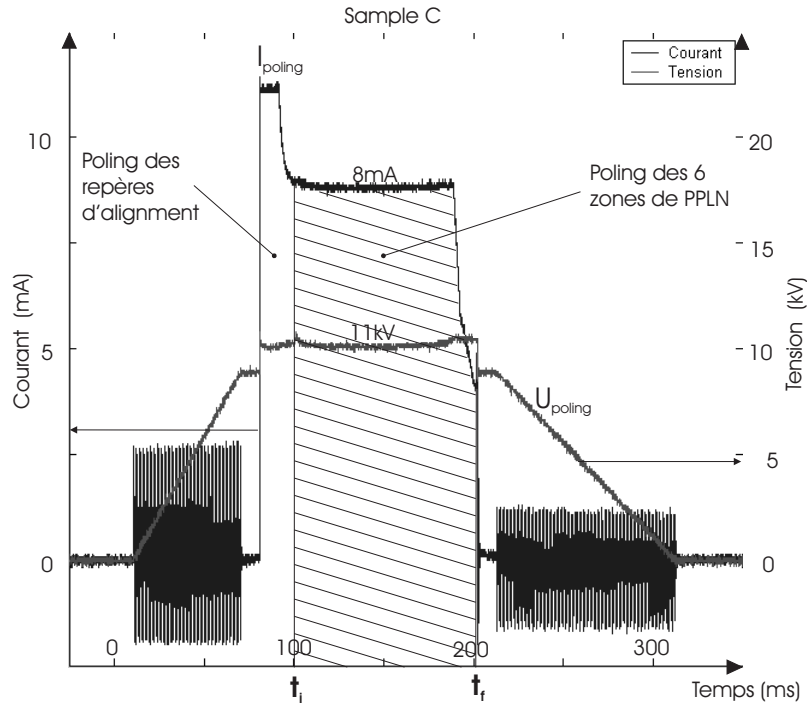


FIG. D.8 – Observation du courant et de la tension appliquée au wafer de $LiNbO_3$. La partie hachurée correspond au retournement effectif des zones de PPLN.

Nous pouvons distinguer deux zones distinctes sur la courbe du courant :

- la première correspond au poling des repères d'alignement de nos masques de lithographie,
- la seconde (hachurée) au poling des zones de PPLN qui nous intéressent.

La raison pour laquelle, les repères d'alignement sont retournés les premiers constitue actuellement le cœur de notre travail, mais nous avons identifié un petit pic tension qui correspond au début du poling des zones de PPLN. Alors la quantité de charge qui a traversé l'échantillon correspond à la surface hachurée qui peut être calculée simplement par :

$$Q = \int_{t_i}^{t_f} I_{poling}(t) dt \quad (D.3)$$

Cette quantité de charge peut alors être assimilée à une surface effective retournée par identification avec la formule D.2 :

$$Q = 2P_s S_{retournée}$$

Le *duty-cycle* a peut être calculé en faisant le rapport de la surface retournée sur la surface totale :

$$a = \frac{S_{\text{retournée}}}{S_{\text{totale}}} = \frac{\int_{t_i}^{t_f} I_{\text{poling}}(t) dt}{2P_s \times S_{\text{totale}}} \quad (\text{D.4})$$

La technique a été validée en sacrifiant quelques wafers. Nous avons comparé le *duty-cycle* moyen calculé avec le *duty-cycle* moyen mesuré après attaque chimique du wafer.

échantillon	$a_{\text{calculé}} (\%)$	$a_{\text{mesuré}} (\%)$
A	57 ± 3	57 ± 2
B	61 ± 3	60 ± 2
C	62 ± 3	66 ± 2
D	65 ± 3	68 ± 2

TAB. D.1 – Comparaison entre le *duty-cycle* moyen calculé et celui mesuré après attaque chimique.

La méthode s'avère être assez précise avec une erreur de 3% sur l'estimation de a . Évidemment, la principale limitation de cette méthode est qu'elle repose sur l'hypothèse que le poling et le *duty-cycle* sont homogènes sur toute la surface du wafer...

Cette méthode permet donc de simplement identifier les wafers qui ont une forte probabilité d'avoir été correctement retournés et ceux qui ne l'ont pas été. Le seul moyen de s'assurer de la qualité du poling, sans attaquer de façon irréversible la surface de l'échantillon, sera d'effectuer, par exemple, une mesure de la puissance de second harmonique générée par le cristal [97, 98, 99].

D.3 Les guides d'ondes *SPE*

Il s'agit désormais de réaliser une structure guidante qui respecte l'inversion périodique du signe du coefficient non-linéaire dont nous venons de décrire le protocole de fabrication. La méthode employée au laboratoire est l'échange protonique qui permet d'accroître l'indice selon l'axe extraordinaire et de le diminuer selon les axes ordinaires. C'est la raison pour laquelle seuls les modes *TM* sont supportés par la structure. Cette façon de procéder consiste à plonger le cristal dans un bain acide fondu à une température comprise entre 160 et 350°C. Il en résulte un échange ionique entre les atomes

de lithium situés à la surface du cristal et les protons présents dans le bain. L'opération peut alors se résumer selon l'équation-bilan suivante :



où x représente le taux de substitution qui dépend de l'acidité du bain, de la coupe du substrat ainsi que de la température de travail. Précisons que les échanges protoniques réalisés dans des bains fortement acides peuvent entraîner la destruction du coefficient non-linéaire du cristal. Toute l'originalité de notre groupe est d'avoir développé un protocole à acidité affaiblie tout en gardant un bon confinement des ondes : *l'échange protonique doux (SPE)*, couplé à l'utilisation d'un masque de protection muni d'ouvertures définissant des canaux où l'échange aura lieu (cf.figure D.9).

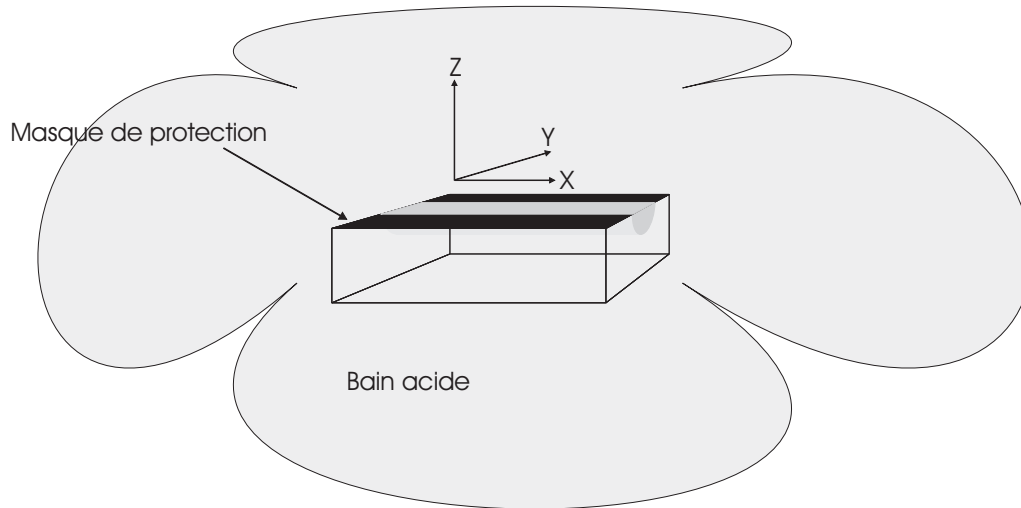


FIG. D.9 – Cristal de $LiNbO_3$ avec son masque pour réaliser l'échange protonique sur une bande de 4, 5, 6, 7 ou 8 μm de large. Ce canal définira le microguide.

D.3.1 Le protocole expérimental

La première étape du protocole consiste une fois encore à déposer une couche de protection munie d'ouvertures qui serviront à définir les canaux où l'échange protonique pourra avoir lieu. Contrairement au poling, une simple couche de résine isolante ne résiste pas à une attaque acide à 300° C, il nous faut donc déposer un matériau bien plus résistant qui va permettre de protéger le $LiNbO_3$.

La technique consiste à déposer sur le wafer 150 nm de *silice* (SiO_2) puis 80 nm de *tantale* (Ta) dans lesquels nous viendrons graver des canaux correspondants aux guides. Le rôle de la silice est de protéger le niobate de lithium du tantale qui protège l'ensemble $\text{SiO}_2\text{--LiNbO}_3$ de l'échange protonique. Ces deux couches sont déposées par sputtering puis masquées par photolithographie classique et enfin gravées par *gravure ionique réactive*.

Nous utilisons à Nice un protocole à acidité affaiblie : l'échange protonique doux ou SPE (pour Soft Proton Exchange). Typiquement la source de protons que nous utilisons est l'acide benzoïque (AB) dont le taux d'acidité est contrôlé par l'adjonction d'une faible quantité de benzoate de lithium (BL). La composition du mélange est alors référencée par le titre massique de BL que l'on désigne généralement par :

$$\rho = \frac{m_{BL}}{m_{BL} + m_{AB}} \times 100 \quad (\text{D.6})$$

où les variables m représentent les masses respectives d' AB et de BL en phase solide (poudre) mises en jeu pour réaliser le bain. Pour les guides dont nous avons besoin, ρ vaut typiquement entre 2,65 et 2,8%. Notons que le lecteur intéressé pourra trouver une description complète de ce protocole dans les références [69] et [25].

Afin de réaliser un guide "idéal", nous entendons un guide qui ne change pas les propriétés quadratiques du cristal, qui soit de bonne qualité optique et qui présente un confinement suffisant, il faut être capable de "maîtriser" un certain nombre de paramètres. Citons notamment :

- (i) *La température d'échange*. Elle contrôle la cinétique de l'échange ainsi que la qualité optique du guide final. D'après de nombreux travaux effectués dans notre groupe, la cinétique optimale s'obtient aux environs des 300°C . Cette température de travail permet non seulement de faire fondre l' AB mais aussi le BL dont le point de fusion est plus haut.

A ce titre, notons que l'échange est réalisé dans une ampoule de verre scellée munie d'un rétrécissement en son milieu comme le montre la figure D.10 ci-dessous.

Cette particularité permet de placer l'échantillon de PPLN masqué³ ainsi que le mélange $AB + BL$ en phase solide respectivement dans deux compartiments différents. L'ampoule, elle-même encapsulée dans un tube en fonte, est ensuite placée dans un four à thermostat réglé sur la température de 300°C désirée. Au moment de l'enfournement, le

³Les motifs du masque correspondent cette fois à des ouvertures perpendiculaires au réseau d'inversion des domaines.

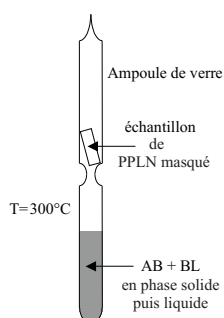


FIG. D.10 – Protocole de l'échange protonique dans une ampoule de verre.

mélange acide est encore en phase solide et à l'état de poudre. Après quelques minutes à 300°C , celui-ci est complètement liquéfié et un retournement de l'ampoule lui permet de couler jusqu'à l'échantillon et à l'échange protonique de débuter.

- (ii) *La durée de l'échange.* Dans le cas d'un mélange $AB + BL$ homogène, celle-ci n'intervient à priori que sur la profondeur du guide final. Il faut savoir que plus le guide est profond, plus le nombre de modes guidés est important, et ce d'autant plus que la longueur d'onde considérée est basse. Dans notre cas, la configuration modale qui nous intéresse est *unique* au signal et à l'idler (1310 et 1550 nm). En effet, nous désirons assurer un recouvrement maximal entre le mode guidé pour le signal de fluorescence et le mode de la fibre optique télécom que l'on placera en sortie du guide pour la récolte des paires de photons (voir le chapitre 2). Cette précieuse caractéristique peut être obtenue grâce à une durée d'échange moyenne de *72 heures* environ. Toutefois, comprenons bien que cette durée est un paramètre délicat dans le sens où, si elle est rendue trop courte pour limiter par exemple le nombre de modes à la pompe, le mode fondamental escompté pour l'onde signal ou idler pourrait ne pas être supporté par la structure.
- (iii) *Le titre massique de benzoate de lithium.* Il représente typiquement le paramètre le "plus libre" du protocole d'échange. Il est d'ailleurs d'une importance capitale puisque de sa valeur dépendent l'accroissement d'indice, la qualité optique mais aussi la structure modale du guide dont nous venons de parler. On sait en effet qu'il existe un ρ_{seuil} ($\sim 2,6\%$ en coupe Z) en dessous duquel les guides obtenus sont "archi-multimodes" à la pompe, de qualité optique médiocre et présentant de surcroît une annulation du coefficient non-linéaire (voir la référence [69]). Le titre massique de BL permettant de protéger l'inversion périodique des domaines, d'obtenir de faibles pertes à la propagation ($< 0,5\text{ dB/cm}$)

ainsi que des guides unimodaux à 1550 nm est compris entre 2,65 et 2,8%. Il faut tout de même préciser que plus ce pourcentage est proche du seuil, plus fort est l'accroissement d'indice qui offre alors un meilleur confinement des ondes guidées. En revanche pour les raisons que nous venons de citer, le risque de rater l'échantillon est plus grand.

Des travaux de thèses précédentes [68, 69] ont permis d'identifier une valeur de ρ pour laquelle toutes ces conditions sont remplies. Pour le guide dont nous avons besoin, ρ vaudra typiquement 2,75% et le cristal subira l'échange à 300°C pendant 72 h .

Annexe E

Liste des symboles utilisés

Symbole	Signification	Unité
Grandeurs « classiques »		
t	Temps	s
λ_j	Longueur d'onde du champ optique j	nm
ω_j	Pulsation du champ optique j	s^{-1}
$\vec{k}_j$	Vecteur d'onde du champ optique j	m^{-1}
c	Vitesse de la lumière dans le vide	$m.s^{-1}$
ϵ_0	Permittivité du vide	$S.I.$
$\vec{P}$	Polarisation du milieu diélectrique	$V.m^{-1}$
$\vec{E}$	Champ électrique	$V.m^{-1}$
$\vec{D}$	Induction électrique	$V.m^{-1}$
$\vec{H}$	Champ magnétique	T
$I(t)$	Intensité d'un rayonnement optique à l'instant t	W
Grandeurs « quantiques »		
$ e\rangle$	Niveau excité d'un système atomique	—
$ g\rangle$	Niveau fondamental d'un système atomique	—
$\hbar$	Constante de Planck	$J.s$
$\hat{H}_0$	Hamiltonien du champ électrique non perturbé	—
$\hat{H}_{int}$	Hamiltonien d'interaction de trois champs par le tenseur $\chi^{(2)}$	—
$\hat{a}_j^\dagger$	Opérateurs création associé au mode j	—
$\hat{a}_j$	Opérateurs annihilation associé au mode j	—
n_j	Nombre de photons dans le mode j	—
v_p	Amplitude du champ de pompe traité « classiquement »	$cm.s^{-1}$
<i>A suivre...</i>		

Symbole (suite)	Signification (suite)	Unité (suite)
APD	Détecteurs de photons (Photodiode à Avalanche)	
Ge	<i>Germanium</i>	—
η_{Ge}	Efficacité de la photodiode Ge	—
D_{Ge}^c	Taux de coups sombres dans la photodiode Ge	s^{-1}
S_{Ge}^{brut}	Taux de comptage brut dans la photodiode Ge	s^{-1}
S_{Ge}^{net}	Taux de comptage net dans la photodiode Ge (= <i>Taux de coups bruts</i> — <i>Taux de coups sombres</i>)	s^{-1}
$InGaAs$	<i>Arseniure de Gallium</i> sur substrat de <i>Phosphure d'Indium</i>	—
η_{In}	Efficacité de la photodiode $InGaAs$	—
D_{In}^c	Taux de coups sombres dans la photodiode $InGaAs$	s^{-1}
S_{In}^{brut}	Taux de comptage brut dans la photodiode $InGaAs$	s^{-1}
S_{In}^{net}	Taux de comptage net dans la photodiode $InGaAs$ (= <i>Taux de coups bruts</i> — <i>Taux de coups sombres</i>)	s^{-1}
ΔT	Durée d'ouverture de la fenêtre de détection associée au photon annoncé	ns
Si	<i>Silicium</i>	—
η_{In}	Efficacité de la photodiode Si	—
D_{In}^c	Taux de coups sombres dans la photodiode Si	s^{-1}
S_{In}^{brut}	Taux de comptage brut dans la photodiode Si	s^{-1}
S_{In}^{net}	Taux de comptage net dans la photodiode Si (= <i>Taux de coups bruts</i> — <i>Taux de coups sombres</i>)	s^{-1}
	Probabilités	
P_i	Probabilité d'avoir i photons ($i = 0, 1, 2, \dots$)	—
P_{2+}	Probabilité d'avoir plus de deux photons = $P_2 + P_3 + \dots$	—
$g^{(2)}(\tau)$	Fonction d'autocorrélation d'ordre deux à l'instant τ	—
μ_p	Taux moyen de paires de photons créées	s^{-1}
μ_s	Taux moyen de photons <i>signal</i> créées	s^{-1}
μ_i	Taux moyen de photons <i>idler</i> créées	s^{-1}
γ_s	Facteur de collection associé aux photons <i>signal</i> (<i>pertes + récolte</i>)	—
γ_i	Facteur de collection associé aux photons <i>idler</i> (<i>pertes + récolte</i>)	—
	Matériaux	
$LiNbO_3$	Niobate de Lithium	
<i>A suivre...</i>		

Symbole (suite)	Signification (suite)	Unité (suite)
<i>PPLN</i>	Niobate de Lithium Périodiquement Polarisé	
<i>KNbO₃</i>	Niobate de Potassium	
<i>LiCl</i>	Électrode liquide à base de chlorure de lithium	
<i>QAP</i>	Quasi-Accord de Phase	
$\chi^{(2)}$	Tenseur susceptibilité d'ordre deux du matériau	$pm.V^{-1}$
n_e	Indice optique extraordinaire d'un matériau biréfringent	–
n_o	Indice optique ordinaire d'un matériau biréfringent	–
P_s	Polarisation spontanée du matériau	$\mu C.cm^{-2}$
Q	Quantité de charges électriques	C
S	Surface	cm^2
Optique intégrée		
<i>SPE</i>	Échange protonique doux (Soft Proton Exchange)	
n_j^{guide}	Indice optique « vu » par les l'onde associée au mode j	–
δn	Accroissement d'indice à la surface du microguide	–
Λ	Période d'inversion du signe du coefficient non-linéaire	μm
I_j	Intégrale de recouvrement entre les ondes guidés pour le mode j	–
ΔK	Désaccord de phase entre les ondes guidées	m^{-1}
α_j	Pertes dans le guide « vues » par le mode j	cm^{-1}
d_{klm}	Coefficient non-linéaires constituant le tenseur $\chi^{(2)}$	$pm.V^{-1}$
g	Terme de couplage contenant la susceptibilité d'ordre deux $\chi^{(2)}$ et les intégrales de recouvrement I_j	cm^{-1}
l	Longueur du guide	cm
$\delta\lambda_j$	Largeur spectrale du mode j issu de la conversion paramétrique dans le microguide	cm
β_j	Constante de propagation associée au mode j	m^{-1}
Optique guidée		
<i>WDM</i>	Coupleur directionnel en longueur d'onde (Wavelength De-Multiplexer)	
R	Coefficient de réflexion de <i>Fresnel</i>	
γ_{pertes}^λ	Pertes associées aux composants optiques à la longueur d'onde λ	dB
Électronique de comptage		
<i>A suivre...</i>		

Symbole (suite)	Signification (suite)	Unité (suite)
NIM	Standard d'électronique logique $0\text{ V} / -0,8\text{ V}$	
ECL	Standard d'électronique logique $-1,8\text{ V} / 0,1\text{ V}$	
TTL	Standard d'électronique logique $0\text{ V} / 5\text{ V}$	
TAC	Convertisseur Temps-Amplitude	
ATM	Mode de Transfert Asynchrone	
R_c	Taux de coïncidences sur le TAC	s^{-1}
N_T	Taux de triggers émis par la source de photons annoncés	s^{-1}
Sécurité de l'échange quantique de clé		
$N_{Shannon}$	Nombre de bits que doivent sacrifier Alice et Bob	—
e	Taux d'erreur dans la clé échangée	—
D	Distance de l'échange de la clé	km
α	Pertes par kilomètre lors de l'échange	$dB.km^{-1}$
p^{net}	Probabilité que Bob détecte un vrai photon	—
p^{dc}	Probabilité que Bob détecte coup accidentel	—

TAB. E.1: Récapitulatifs des symboles utilisés

Table des figures

1	Schéma d'un système quantique à deux niveaux.	23
2	Centre coloré dans le diamant	25
3	Microcavité autour d'une boîte quantique	25
4	Utilisation des ondes acoustiques de surface pour la génération de photons uniques	27
5	Principe de la post-sélection des impulsions contenant un photon	28
6	Atomes en cavité pour la génération de photons uniques . . .	30
7	Utilisation d'un condensat d'atomes froids pour la génération de photons uniques	31
1.1	Interaction paramétrique à trois photons	35
1.2	Montage pour effectuer la mesure de la séparation temporelle entre les photons signal et l'idler	43
1.3	Graphique de la séparation temporelle entre les photons signal et idler	43
1.4	Représentation de paires créées aléatoirement	45
1.5	Probabilité d'avoir un intervalle vide entre deux photons successifs	46
1.6	Schéma de principe sur la probabilité d'avoir un ou plusieurs photons dans une fenêtre tem	
1.7	Schéma de la source de photons uniques annoncés	49
1.8	Montage de type <i>Hanbury-Brown & Twiss</i>	53
1.9	Évolutions de P_1 et $g^{(2)}(0)$ en fonction du coefficient de collection des paires γ_i et γ_s	56
1.10	Évolutions de P_1 et $g^{(2)}(0)$ en fonction de γ_i	57
1.11	Évolutions de P_1 et $g^{(2)}(0)$ en fonction de γ_s	58
1.12	Évolutions de P_1 et $g^{(2)}(0)$ en fonction de μ_p	59
1.13	Évolutions de P_1 et $g^{(2)}(0)$ en fonction de D_{Ge}^c	60
1.14	Évolutions de P_1 et $g^{(2)}(0)$ en fonction de ΔT	61
1.15	Taux de bits utilisables par impulsion en fonction de la distance	63
2.1	Guide d'onde intégré sur niobate de lithium	67
2.2	Accord de phase dans un guide d'onde	68
2.3	Puissance de l'onde signal pour le désaccord de phase et 2 types d'accord de phase	69
2.4	Définition des axes cristallins du Niobate de Lithium	70
2.5	Guide d'onde intégré sur substrat PPLN	72
2.6	Schéma indicatif de l'élargissement spectral	75
2.7	Guide PPLN final	79
2.8	Montage expérimental pour la caractérisation de la fluorescence paramétrique	80

2.9	Spectres de fluorescence expérimental	82
2.10	Courbe de quasi-accord de phase	84
2.11	Spectres de fluorescence	86
2.12	Schéma microguide asymétrique dans le plan vertical	88
2.13	Représentation de la densité de puissance $3D$ associés aux 3 modes à l'intérieur d'un	
2.14	Représentation de la densité de puissance $2D$ associés aux 3 modes à l'intérieur d'un	
2.15	Représentation de la densité de puissance $2D$ associés aux 3 modes à l'intérieur d'un	
2.16	Schéma complet de la source de photons uniques annoncés . . .	93
2.17	Synoptique d'une fibre « pigtailed »	96
2.18	Connecteur optique standard (FC-PC), couramment appelé <i>I-optique</i> . 97	
2.19	Le U-bench OFR connectorisé.	97
2.20	Transmission du filtre <i>DJ1160a</i>	98
2.21	WDM 1310/1550 nm.	99
2.22	Photodiode à avalanche fibrée en <i>germanium</i>	100
2.23	Montage expérimental pour la détermination de mesurer γ_i . .	103
2.24	Diagramme de coïncidences signaux électriques–photons annoncés	104
2.25	Source de photon uniques annoncés « Niçoise ».	108
3.1	Représentation globale de la source de photons uniques annoncés	112
3.2	Principe du miroir semi-réfléchissant	113
3.3	Montage de type « Hanbury-Brown & Twiss » utilisé à Nice .	114
3.4	Évolution de la probabilité P_1 et de la valeur de $g^{(2)}(0)$ en fonction de la puissance.	
3.5	Montage de type « Hanbury-Brown & Twiss »	132
3.6	Courbe d'anti-coïncidences en régime pulsé pour une source poissonnienne	134
3.7	Courbe d'anti-coïncidences en régime asynchrone avec montage classique	135
3.8	Principe de la post-sélection des impulsions contenant au moins un photon. Les imp	
3.9	Histogramme expérimental pour une source typiquement poissonnienne.	137
3.10	Schéma source de photon uniques annoncés utilisée à Genève .	138
3.11	Montage de type Hanbury-Brown & Twiss amélioré pour une source asynchrone	139
3.12	Diagramme logique de la carte d'acquisition « High speed digital I/O device »	140
3.13	Histogramme d'anti-coïncidences issue de la carte d'acquisition	142
3.14	Histogramme d'anti-coïncidences et son ajustement poissonnien	144
3.15	Histogramme d'anti-coïncidences pour le bruit	146
3.16	Schéma d'ensembles joints	147
4.1	Évolutions des largeurs spectrales des photons annoncés en fonction de la longueur	
4.2	Probabilité d'obtenir le photon idler à la sortie du guide en fonction de la longueur	
A.1	Taux de bits utilisables par impulsion en fonction de la distance	176
B.1	Caractéristique schématique du principe d'une jonction PN en mode inversé	178
B.2	Circuit de polarisation de type <i>extinction passive</i>	179

B.3	Principe de caractérisation d'une photodiode en mode passif .	181
B.4	Courbe étalonnage de l'efficacité de l'APD- <i>Ge</i>	182
B.5	Principe de caractérisation d'une photodiode en mode déclenché	183
B.6	Courbe étalonnage de l'efficacité des APD- <i>InGaAs</i>	185
C.1	Probabilité qu'un intervalle de temps supérieur ou égal à ΔT sépare deux événements succe	
C.2	Schéma expérimental pour mesurer l'intervalle de temps entre deux détections successives	18
C.3	Distributions des intervalles de temps expérimentaux	190
D.1	La structure cristalline du niobate de lithium.	195
D.2	Phases ferro-électriques du niobate de lithium.	195
D.3	Représentation d'un wafer de niobate de lithium en coupe Z. .	196
D.4	Représentation d'un masque de poling des wafers de niobate de lithium en coupe Z.	197
D.5	Cellule de poling pour wafer 3 pouces.	198
D.6	Montage pour le poling de wafer trois pouces de <i>LiNbO₃</i>	198
D.7	Étapes principales de l'inversion de polarisation d'un cristal monodomaine avec un masque	
D.8	Observation du courant et de la tension appliquée au wafer de <i>LiNbO₃</i> . La partie hachurée	
D.9	Cristal de <i>LiNbO₃</i> avec son masque	203
D.10	Protocole de l'échange protonique dans une ampoule de verre. .	205

Liste des tableaux

2.1	Caractéristiques générales de la fibre optique SMF-28.	94
2.2	Comparaison des pertes d'une fibre « pigtailed » au guide . . .	96
2.3	Données expérimentales associées à la mesure du coefficient de couplage γ_i .	106
3.1	Tableau récapitulatif des mesures expérimentales. Les efficacités des APDs sont extraites de	
3.2	Tableau récapitulatif des performances expérimentales.	123
3.3	Tableau comparatif des performances théoriques et expérimentales.	125
3.4	Comparaison des performances pour différents types de sources de photons uniques.	128
3.5	Tableau récapitulatif des mesures expérimentales faites à Genève.	149
3.6	Tableau récapitulatif des performances de la source de Genève à partir du programme développ	
D.1	Comparaison entre le <i>duty-cycle</i> moyen calculé et celui mesuré après attaque chimique.	202
E.1	Récapitulatifs des symboles utilisés	210

Bibliographie

- [1] J S Bell. On the Einstein-Podolsky-Rosen Paradox. *Physics (Long Island City, N.Y.)*, 1 :195–200, 1964.
- [2] A Aspect, P Grangier, and G Roger. Experimental realization of Einstein-Podolsky-Rosen-Bohm Gedankenexperiment : A new violation of Bell's inequalities. *Phys. Rev. Lett.*, 49 :91–94, 1982.
- [3] J G Rarity and P R Tapster. Experimental violation of Bell's inequality based on phase and momentum. *Phys. Rev. Lett.*, 64 :2495–2498, 1990.
- [4] W Tittel, J Brendel, H Zbinden, and N Gisin. Violation of Bell inequalities by photons more than 10 km apart. *Phys. Rev. Lett.*, 81 :3563–3566, 1998.
- [5] C. H. Bennett and D. P. DiVincenzo. Quantum information and computation. *Nature*, 404(247), 2000.
- [6] C H Bennett and S Wiesner. Communication via one and two particle operations on Einstein-Podolsky-Rosen states. *Phys. Rev. Lett.*, 69 :2881–2884, 1992.
- [7] Izo Abram and Philippe Grangier. From quantum optics to quantum communications. *C. R. Physique*, 4, 2003.
- [8] J.-Ph. Poizat and R. Mosseri. Introduction à l'information quantique. *GdR : Information et communication quantique*, 2000.
- [9] J.G. Rarity. Quantum communication and beyond. *Phil. Trans. R. Soc. Lond. A*, 361 :1507–1518, 2003.
- [10] C H Bennett and G Brassard. Quantum cryptography : Public key distribution and coin tossing. *Proceedings of the IEEE International Conference on Computers, Systems and Signal Processing, Bangalore, India*, page 175, 1984.
- [11] P Shor. Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer. *Proc. 35th IEEE Symp. on Foundations of Computer Science*, Santa Fe, ed S. Goldwasser :124, 1994.

- [12] Charles H. Bennett, Francois Bessette, Gilles Brassard, Louis Salvail, and John Smolin. Experimental quantum cryptography. *Journal of Cryptology*, 5(1) :3, 1992.
- [13] W. T. Buttler, R. J. Hughes, P. G. Kwiat, S. K. Lamoreaux, G. G. Luther, G. L. Morgan, J. E. Nordholt, C. G. Peterson, and C. M. Simmons. Practical free-space quantum key distribution over 1 km. *Phys. Rev. Lett.*, 81 :3283, 1998.
- [14] J. Rarity, P. Taptser, and P. Gorman. Secure free-space key exchange to 1.9 km and beyond. *J. Mod. Opt.*, 48 :1887, 2001.
- [15] C. Kurtsiefer, P. Zarda, M. Handler, H. Weinfurter, P. M. Gorman, P. R. Tapster, and J. G. Rarity. Quantum cryptography : A step towards global key distribution. *Nature*, 419 :450, 2002.
- [16] Richard J Hughes, Jane E Nordholt, Derek Derkacs, and Charles G Peterson. Practical free-space quantum key distribution over 10 km in daylight and at night. *new J. Phys.*, 4(43), 2002.
- [17] A. Muller, H. Zbinden, and N. Gisin. Quantum cryptography over 23 km in installed under-lake telecom fibre. *Europhys. Lett.*, 33 :335, 1996.
- [18] Sébastien Tanzilli, Wolfgang Tittel, Hugues De Riedmatten, Hugo Zbinden, Pascal Baldi, Marc De Micheli, Daniel Barry Ostrowsky, and Nicolas Gisin. PPLN waveguide for quantum communication. *Eur. Phys. J. D*, 18 :155–160, 2002.
- [19] C.K. Hong and Leonard Mandel. Experimental realization of a localized one-photon state. *Phys. Rev. Lett.*, 56(1) :58, 1986.
- [20] I. Marcikic, H. de Riedmatten, W. Tittel, H. Zbinden, and N. Gisin. Long-distance teleportation of qubits at telecommunication wavelengths. *Nature*, 421 :509, 2003.
- [21] D. Stucki, N. Gisin, O. Guinnard, G. Ribordy, and H. Zbinden. Quantum key distribution over 67 km with a plug-and-play system. *J. Phys.*, 4(41), 2002.
- [22] C. Gobby, Z. L. Yuan, and A. J. Shields. Quantum key distribution over 122 km of standard telecom fiber. *App. Phys. Lett.*, 84(19) :3762, 2004.
- [23] Hideo Kosaka, Akihisa Tomita, Yoshihiro Nambu, Tadamasa Kimura, and Kazuo Nakamura. Single-photon interference experiment over 100 km for quantum cryptography system using a balanced gated-mode photon detector. *Electron. Lett.*, 39(16) :1199, 2003.
- [24] Tadamasa Kimura, Yoshihiro Nambu, Takaaki Hatanaka, Akihisa Tomita, Hideo Kosaka, and Kazuo Nakamura. Single-photon interference

- over 150 km transmission using silica-based integrated-optic interferometers for quantum cryptography. *arXiv :quant-ph/0403104*, 2004.
- [25] Loic Chanvillard, Pierre Aschiéri, Pascal Baldi, Daniel Barry Ostrowsky, L. Huang, and Marc De Micheli. Soft proton exchange on PPLN : a simple waveguide fabrication process for highly efficient non-linear interactions.
 - [26] Sébastien Tanzilli, Hugues De Riedmatten, Wolfgang Tittel, Hugo Zbinden, Pascal Baldi, Marc De Micheli, Daniel Barry Ostrowsky, and Nicolas Gisin. Highly efficient photon-pair source using periodically poled lithium niobate waveguide. *Elec. Lett.*, 37(1), 2001.
 - [27] H. Kimble, M. Dagenais, and L. Mandel. Photon antibunching in resonance fluorescence. *Phys. Rev. Lett.*, 39 :691, 1977.
 - [28] F. Diedrich and H. Walter. Nonclassical radiation of a single stored ion. *Phys. Rev. Lett.*, 58 :203, 1987.
 - [29] T. Bashe, W. Moerner, M. Orrit, and H. Talon. Photon antibunching in the fluorescence of a single dye molecule trapped in a solid. *Phys. Rev. Lett.*, 69 :1516, 1992.
 - [30] C. Brunel, B. Lounis, P. Tamarat, and M. Orrit. Triggered source of single photons based on controlled single molecule fluorescence. *Phys. Rev. Lett.*, 83 :2722, 1999.
 - [31] François Treussart, Romain Alléaume, V. Le Floch, L.T. Xiao, J.-M. Courty, and J.-F. Roch. Direct measurement of the photon statistics of a triggered single photon source. *Phys. Rev. Lett.*, 89(9) :093601–1, 2002.
 - [32] B. Lounis and W.E. Moerner. Single photons on demand from a single molecule at room temperature. *Nature*, 407 :491–493, 2000.
 - [33] Alexios Beveratos, Rosa Brouri, Thierry Gacoin, André Villing, Jean-Philippe Poizat, and Philippe Grangier. Single photon quantum cryptography. *Phys. Rev. Lett.*, 89 :187901, 2002.
 - [34] J. P. Grangier, G. Roger, and A. Aspect. Experimental evidence for a photon anticorrelation effect on a beam splitter : a new light on single-photon interferences. *Europhys. Lett.*, 1(4) :173, 1986.
 - [35] C. Kurtsiefer, S. Mayer, P. Zarda, and Harald Weinfurter. Stable solid-state source of single photons. *Phys. Rev. Lett.*, 85(2) :290–293, 2000.
 - [36] A. Beveratos, S. Kuhn, R. Gacoin, J.-P. Poizat, and P. Grangier. Room temperature stable single-photon source. *Eur. Phys. J. D*, 18 :191–196, 2002.

- [37] E. Moreau, I. Robert, J. M. Gérard, I. Abram, L. Manin, and V. Thierry-Mieg. Single-mode solid-state single photon source based on isolated quantum dots in pillar microcavities. *Appl. Phys. Lett.*, 79(18) :2865, 2001.
- [38] M. Pelton, Charles Santori, Jelena Vuckovic, B. Zhang, Glenn S. Solomon, J. Plant, and Yoshihisa Yamamoto. Efficient source of single photons : A single quantum dot in a micropost microcavity. *Phys. Rev. Lett.*, 89(23) :233602–1, 2002.
- [39] D. Fattal, E. Diamanti, K. Inoue, and Y. Yamamoto. Quantum teleportation with a quantum dot single photon source. *Phys. Rev. Lett.*, 92(3) :037904, 2004.
- [40] G. Messin, J. P. Hermier, E. Giacobino, P. Desbiolles, and M. Dahan. Bunching and antibunching in the fluorescence of semiconductor nanocrystals. *Opt. Lett.*, 26(23) :1891, 2001.
- [41] X. Brokmann, G. Messin, P. Desbiolles, E. Giacobino, M. Dahan, and J. P. Hermier. Colloidal *CdSe/ZnS* quantum dots as single-photon sources. *New J. Phys.*, 6(99), 2004.
- [42] C. L. Foden, V. I. Talyanskii, G. J. Milburn, M. L. Leadbeater¹, and M. Pepper. High-frequency acousto-electric single-photon source. *Phys. Rev. A*, 62 :011803, 200.
- [43] J. Cunningham, V. I. Talyanskii, J. M. Shilton, M. Pepper, M. Y. Simmons, and D. A. Ritchie. Single-electron acoustic charge transport by two counterpropagating surface acoustic wave beams. *Phys. Rev. B*, 60 :4850, 1999.
- [44] S. Takeuchi. Beamlike twin-photon generation by use of type II parametric downconversion. *Opt. Lett.*, 26(11) :843, 2001.
- [45] C.H. Monken, P.H. Souto Ribeiro, and S. Pàdua. Optimizing the photon pair collection efficiency : A step toward a loophole-free bell's inequalities experiment. *Phys. Rev. A*, 57(4) :2267, 1998.
- [46] John Rarity, P.R. Tapster, and E. Jakeman. Observation of sub-poissonian light in parametric downconversion. *Optics Comm.*, 62(3) :201, 1987.
- [47] Alfred B. U'Ren, Christine Silberhorn, Konrad Banaszek, and Ian A. Walmsley. Conditional preparation of single photons for scalable quantum-optical networking. *Phys. Rev. Lett.*, 93 :093601, 2004.
- [48] Shigeki Takeuchi, Ryo Okamoto, and Keiji Sasaki. High-yield single-photon source using gated spontaneous parametric downconversion. volume 43, page 5708, 2004.

- [49] Alfred B. U'Ren, Christine Silberhorn, Konrad Banaszek, and Ian A. Walmsley. Conditional preparation of single photons for scalable quantum-optical networking. *arXiv :quant-ph/0312118*, 2004.
- [50] T.B. Pittman, B.C Jacobs, and J.D. Franson. Heralding single photons from pulsed parametric down-conversion. *arXiv :quant-ph/0408093*, 2004.
- [51] Sylvain Fasel, Olivier Alibart, Sebastien Tanzilli, Pascal Baldi, Alexios Beveratos, Nicolas Gisin, and Hugo Zbinden. High quality asynchronous heralded single photon source at telecom wavelength. *New J. Phys.*, 6(163), 2004.
- [52] T. B. Pittman, B. C. Jacobs, and J. D. Franson. Single photons on pseudodemand from stored parametric down-conversion. *Phys. Rev. A*, 66 :042303, 2002.
- [53] Evan Jeffrey, Nicholas A Peters, and Paul G Kwiat. Towards a periodic deterministic source of arbitrary single-photon states. *New J. Phys.*, 6 :100, 2004.
- [54] M. Hennrich, T. Legero, A. Kuhn, and G. Rempe. Photon statistics of a non-stationary periodically driven single-photon source. *arXiv :quant-ph/0406034*, 2004.
- [55] C. W. Chou, S. V. Polyakov, A. Kuzmich, and H. J. Kimble. Single-photon generation from stored excitation in an atomic ensemble. *arXiv :quant-ph/0401147*, 2004.
- [56] L. M. Duan, M. D. Lukin, J. I. Cirac, and P. Zoller. Long-distance quantum communication with atomic ensembles and linear optics. *Nature*, 414 :413, 2001.
- [57] J.A. Armstrong, N. Bloembergen, J. Ducuing, and P.S. Pershan. *Phys. Rev. Lett.*, 127 :1918–1939, 1962.
- [58] N Bloembergen and P S Pershan. Light waves at the boundary of non-linear media. *Phys. Rev.*, 128 :606, 1962.
- [59] Amnon Yariv. *Optical Electronics*. harcourt brace jovanovich college publisher, fourth edition edition, 1991.
- [60] Leonard Mandel and Emil Wolf. *Optical coherence and quantum optics*, chapter 22.4.7, pages 1084–1088. Cambridge university press, 1995.
- [61] D. F. Walls and G. J. Milburn. *Quantum Optics*. Springer - Verlag, 1994.
- [62] C. K. Hong, Z. Y. Ou, and L. Mandel. measurement of subpicosecond time intervals between two photons by interference. *Phys. Rev. Lett.*, 59(18) :2044–2046, 1987.

- [63] <http://www.granddictionnaire.com>.
- [64] Roy J. Glauber. The quantum theory of optical coherence. *Phys. Rev.*, 130(6) :2529, 1963.
- [65] Roy J. Glauber. Coherent and incoherent states of the radiation field. *Phys. Rev.*, 131(6) :2766, 1963.
- [66] R. Hanbury brown and R. Q. Twiss. *Nature*, 178 :1046, 1956.
- [67] Pascal Baldi. *Génération de fluorescence paramétrique guidée sur niobate de lithium polarisés périodiquement. Etude préliminaire d'un oscillateur paramétrique optique intégré*. PhD thesis, Université de Nice Sophia-Antipolis, France, 1994.
- [68] Sébastien Tanzilli. *Optique intégrée pour les communications quantiques*. PhD thesis, Université de Nice Sophia-Antipolis, France, 2002.
- [69] Loïc Chanvillard. *Interactions paramétriques guidées de grande efficacité : Utilisation de l'échange protonique doux sur du niobate de lithium inversé périodiquement*. PhD thesis, Université de Nice Sophia-Antipolis, France, 1999.
- [70] Vladislav Beskrovnyy and Pascal Baldi. Optical parametric fluorescence spectra in periodically poled media. *Optics Express*, 10(19) :990, 2002.
- [71] G. J. Edwards and M. Lawrence. A temperature dependent dispersion equation for congruently grown lithium niobate. *Opt. Quant. Electron.*, 16 :373–374, 1984.
- [72] K. V. D. Velde, M. Thienpont, and R. V. Green. Extending the effective index method for arbitrarily shaped inhomogeneous optical waveguides. *IEEE JLT*, 6 :1153–1159, 1988.
- [73] Ajoy Ghatak and K. Thyagarajan. *Introduction to fiber optics*. Cambridge University Press, 1999.
- [74] S. Cova, M. Ghioni, A. Lacaita, and F. Zappa. Avalanche photodiodes and quenching circuits for single-photon detection. *Appl. Opt.*, 35 :1956–1976, 1996.
- [75] T.B. Pittman, B.C. Jacobs, and J.D. Franson. Demonstration of non-deterministic quantum logic operation using linear optical elements. *Phys. Rev. Lett.*, 88(25) :257902, 2002.
- [76] J. D. Franson, M. M. Donegan, M. J. Fitch, B. C. Jacobs, and T. B. Pittman. High-fidelity quantum logic operations using linear optical elements. *Phys. Rev. Lett.*, 89(13) :137901, 2002.
- [77] Ch. Wang, Ch. Braig, P. Zarda, H. Weinfurter, and Ch. Kurtsiefer. Solid-state source of single photons for quantum cryptography and student-labs. In *CLEO/IQEC*, IW2, page 125. OSA, 2004.

- [78] Jelena Vuckovic, David Fattal, Charles Santori, B. Zhang, Glenn S. Solomon, and Yoshihisa Yamamoto. Enhanced single-photon emission from a quantum dot in a micropost microcavity. *Appl. Phys. Lett.*, 82(21) :3596, 2003.
- [79] Andrea Fiore, J. X. Chenb, and M. Ilegems. Scaling quantum-dot light-emitting diodes to submicrometer sizes. *Appl. Phys. Lett.*, 81(10) :1756, 2002.
- [80] Richard P. Mirin. Photon antibunching at high temperature from a single InGaAs/GaAs quantum dot. *Appl. Phys. Lett.*, 84(8) :1260, 2004.
- [81] Grégoire Ribordy, Jurgen Brendel, Jean-Daniel Gautier, Nicolas Gisin, and Hugo Zbinden. Long-distance entanglement-based quantum key distribution. *Phys. Rev. A*, 63 :012309, 2000.
- [82] Xinhua Gu, Roman Y. Korotkov, , Yujie J. Ding, Jin U. Kang, and Jacob B. Khurgin. Backward second-harmonic generation in periodically poled lithium niobate. *J. Opt. Soc. Am. B*, 15(5), 1998.
- [83] <http://www.ino.qc.ca>.
- [84] S. Fasel, N. Gisin, G. Ribordy, and H. Zbinden. Quantum key distribution over 30 km of standard fiber using energy-time entangled photon pairs : a comparison of two chromatic dispersion reduction methods. *Eur. Phys. J. D*, 30 :143–148, 2004.
- [85] Olivier Alibart, Sebastien Tanzilli, Daniel Barry Ostrowsky, and Pascal Baldi. Guided wave technology for a telecom wavelength heralded single photon source. *arXiv :quant-ph/0405075*, 2004.
- [86] Zhi Zhao, Yu-Ao Chen, An-Ning Zhang, Tao Yang, Hans J. Briegel, and Jian-Wei Pan. Experimental demonstration of five-photon entanglement and open-destination teleportation. *Nature*, 430 :54, 2004.
- [87] M. Eibl, S. Gaertner, M. Bourennane, Ch. Kurtsiefer, M. Zukowski, and H. Weinfurter. Experimental observation of four-photon entanglement from down-conversion. *arxiv :quant-ph/0302042*, 2003.
- [88] Jian-Wei Pan, Matthew Daniell, Sara Gasparoni, Gregor Weihs, and Anton Zeilinger. Experimental demonstration of four-photon entanglement and high-fidelity teleportation. *Phys. Rev. Lett.*, 86(20) :4435, 2001.
- [89] Dik Bouwmeester, Jian-Wei Pan, Matthew Daniell, Harald Weinfurter, and Anton Zeilinger. Observation of three-photon greenberger-horne-zeilinger entanglement. *Phys. Rev. Lett.*, 82(7) :1345, 1999.
- [90] N. Lutkenhaus. Security against individual attacks for realistic quantum key distribution. *Phys. Rev. A*, 61 :052304, 2000.

- [91] C. Shannon. A mathematical theory of communication. *Bell System Technical Journal*, 27(379), 1948.
- [92] G Ribordy, J D Gautier, H Zbinden, and N Gisin. Performance of InGaAs/InP avalanche photodiodes as gated-mode photon counters. *Appl. Opt.*, 37 :2272–2277, 1998.
- [93] D Stucki, G Ribordy, A Stefanov, and H Zbinden. Photon counting for quantum key distribution with peltier cooled InGaAs/InP APD's. *Journal of Modern Optics*, 43, 2001.
- [94] Katia Gallo. *Guides enterrés polarisés périodiquement et étude théorique d'un amplificateur paramétrique contrapropagatif*. PhD thesis, Université de Nice Sophia-Antipolis, France, 2001.
- [95] Irene Aboud. *polarisation périodique et échange protonique dans le niobate de lithium*. PhD thesis, Université de Nice Sophia-Antipolis, France, 2000.
- [96] Amnon Yariv. *Quantum Electronics*. Wiley, New-York, third edition, 1989.
- [97] Steffen Kjaer Johansen and Pascal Baldi. Characterization of quasi-phase-matching gratings in quadratic media through double-pass second-harmonic power measurements. *J. Opt. Soc. Am. B*, 21(6) :1137, 2004.
- [98] Sunao Kurimura and Yoshiaki Uesu. Application of the second harmonic generation microscope to nondestructive observation of periodically poled ferroelectric domains in quasi-phase-matched wavelength converters. *J. Appl. Phys.*, 871(1) :369, 1997.
- [99] D. Rouede, Y. Le Grand, L. Leroy, J.R. Duclere, and M. Guilloux-Viry. Nonlinear optical properties and domain microstructure of epitaxial $SrBi_2Nb_2O_9$ thin films on $SrTiO_3$ and on MgO substrates studied by second-harmonic generation. *Opt. Comm.*, 222 :289, 2003.

Résumé :

Ce manuscrit présente l'étude et la réalisation expérimentale d'une source de photons uniques annoncés, basée sur l'utilisation de paires de photons issues d'un guide d'onde intégré sur un substrat de niobate de lithium polarisé périodiquement (PPLN). Le principe repose sur la séparation des photons d'une paire et sur la détection de l'un pour annoncer la présence de l'autre. Aussi, le guide d'onde nous permet de venir récolter les paires de photons simplement à l'aide d'une fibre optique monomode, gage de compacité et stabilité de la source. L'intervalle de temps entre deux paires successives n'étant pas défini, cette source présente un fonctionnement dit « asynchrone ». Afin de caractériser les performances de ce type de source, nous proposons deux méthodes expérimentales originales. La première repose sur un modèle d'analyse des statistiques des détections dans un montage de type « Hanbury-Brown & Twiss » pour remonter aux probabilités d'avoir 0, 1 ou 2 photons, tandis que la seconde est une « version asynchrone » du montage original de « Hanbury-Brown & Twiss » pour tracer la fonction de corrélation croisée du second-ordre. Les performances de cette première source de photons uniques aux longueurs d'ondes télécom se situent parmi les meilleures au monde avec une probabilité d'avoir un photon unique à 1550 nm dans une fibre optique monomode de 0,37 tandis que les événements à deux photons sont réduits d'un facteur 12 par rapport à une source poissonnienne équivalente.

Mots-clés : Optique Intégrée, Guides PPLN, Optique Quantique, Sources de Photons Uniques Annoncés.

Abstract :

This thesis reports the realization of a heralded single photon source (HSPS) based on a periodically poled lithium niobate (PPLN) waveguide. The HSPS relies on photon pairs generated by spontaneous parametric down-conversion and the idea is to use one of the photons to herald the arrival of the second one. Taking advantage of the guided structure, the photon pairs are collected by a single mode telecom fiber attached to the output of the waveguide. This demonstrates the potential of waveguide technologies for building efficient, stable, and compact sources. The creation time of two successive photons pairs is unknown, essentially, this is a quantum equivalent of the classical « asynchronous transfer mode ». We point out two means to characterize the efficiency of this type of source. We first investigate an analysis model that allows us to infer the probability of having 0, 1 or 2 photons from the detection on a « Hanbury-Brown & Twiss » type setup, while we build an asynchronous equivalent to the « Hanbury-Brown & Twiss » setup in order to measure the second-order cross-correlation function. This work has lead to the demonstration of having a single photon at 1550 nm into a single-mode optical fiber with a probability of 0.37, whereas the multi-photon emission probability is reduced by a factor of 12 compared to weak laser poissonian light sources at equal P_1 .

Key-words : Integrated Optics, PPLN waveguides, Quantum Optics, Heralded Single Photon Sources.