



THÈSE

présentée pour obtenir

le grade de DOCTEUR EN SCIENCES

spécialité

Physique Atomique et Moléculaire

par

ANTOINE MONMAYRANT

FAÇONNAGE ET CARACTÉRISATION D'IMPULSIONS ULTRACOURTES. CONTRÔLE COHÉRENT DE SYSTÈMES SIMPLES.

Soutenance prévue le 27 juin 2005 devant la commission d'examen :

M.	Manuel	JOFFRE	Rapporteur
M.	Jean	OBERLÉ	Rapporteur
Mme.	Béatrice	CHATEL	Directrice de thèse
M.	Jean-Michel	GÉRARD	
M.	Bertrand	GIRARD	Invité
M.	Daniel	KAPLAN	
M.	Xavier	MARIE	Invité
M.	Christoph	MEIER	

REMERCIEMENTS

Je souhaite ici saluer et remercier tous ceux qui m'ont aidé à aller au bout de cette aventure. Malheureusement, comme ils ont été très nombreux à m'accompagner durant ce voyage et comme ces remerciements ne peuvent décemment pas occuper 1000 pages, je sais d'avance que je vais omettre ici un bon nombre d'entre eux. Disons que je vais essayer de faire de mon mieux pour remercier le plus grand nombre. Ne cherchez pas ici la moindre logique dans l'ordre d'apparition des heureux nommés. Il s'agit sûrement de la seule partie de cette thèse où je peux me permettre d'enchaîner les idées comme elles viennent, sans me préoccuper de la logique, alors j'en profite...

Merci d'abord à tous ceux que j'oublie ici. J'espère que cette place de choix dans les remerciements, bien qu'assez impersonnelle, pourra me faire pardonner...

Merci Jérôme (oui c'est lui, non il n'a pas changé) pour avoir été le mentor du petit scarabée débutant sa vie de thésard et de moniteur. Sans ton humour moitié Sud-Ouest / moitié Carambar, la salle de manip et les salles de TP auraient été bien plus sombres.

Un grand merci à ma famille pour son soutien durant cette thèse ainsi que les années d'étude (parfois difficiles) qui l'ont précédée.

Merci Valérie pour ta disponibilité sans faille. Ta générosité et ton calme face aux questions bêtes à répétition m'ont aidé à apprivoiser certaines bestioles utilisées au cours de ma thèse (je pense en particulier à LabView...). Merci aussi pour ta franchise et tes rires sonores...

Merci Sebas, pour tes rayures, tes Doc Martens vertes et pour m'avoir laissé quelques spirales de Cornu et quelques fonctions erf pour que je puisse m'amuser un peu durant cette thèse (lâche ce crapaud tout de suite!).

Merci à Michel, Malika, Arnaud (pas d'élec., mais des idées), Renaud, Eric, et tous ceux avec qui j'ai eu la chance d'enseigner à l'UPS... Le monitorat a vraiment été une expérience passionnante et un bon moyen de sortir le nez du guidon de temps en temps.

Merci au lecteurs de ne pas prêté à tension au éventuels fôtes d'hortagrafe!

Merci à Albert Stolorow ainsi qu'à toute la "Dream Team" d'Ottawa (Ben, Rune, Athony-Lee-, Marc, Jérôme,...) pour leur accueil si chaleureux. Mon séjour au Canada a été un vrai plaisir grâce à eux. Merci aussi de m'avoir initié à la Poutine, la diététique selon les bûcherons québécois.

Merci Jean-Christophe pour m'avoir accompagné durant ce voyage de trois ans : quitte à ramer, autant être plusieurs. Bon courage pour la suite, le tofu et le sake...

Merci à moi de m'être supporté malgré mon mauvais caractère. J'ai souvent cru craquer (et croc!). Mais comment est-ce que les autres ont fait ?

Merci aussi à tous les oiseaux de passages : Barbara, Jean-Baptiste, Marie-Emilie, Giorgio, Peter, Julien, ...

Merci Béatrice pour avoir dirigé cette thèse d'une poigne de fer... non, je rigole ! Merci pour ton enthousiasme communicatif, ta bonne humeur et ta patience. Merci pour ton calme durant mes accès de colère ou de panique, ton assurance dans mes moments de doute... Ce

fut un plaisir de travailler avec toi.

Merci à l'ensemble des techniciens (Daniel, Philippe, Michel, Laurent, William) pour leur aide sans laquelle aucun travail expérimental ne pourrait se faire. Un merci tout particulier aux mécaniciens pour leur précieuse contribution lors de la conception du façonneur à haute résolution.

Merci aussi à Gérard pour ses conseils précieux sur la vivisection des platines de translation et à Gilles pour l'amélioration de l'électronique d'acquisition.

Merci Bertrand d'avoir su prendre du recul quand nous étions tous penchés sur l'expérience, le nez dans le guidon, à ne pas voir plus loin que le bout de notre ligne 4f : "On n'y arrive déjà pas avec une impulsion, pourquoi veux-tu qu'on en mette deux ? On le fait si tu insistes, mais de toute manière, ça ne marchera pas..." Merci pour ta disponibilité et tes conseils.

Merci à Aziz d'avoir partagé sa cage avec un autre fauve. Merci pour tes conseils, tes sifflotements, "Las ketchup song" et merci surtout d'avoir fermé la porte... ;-)

Un grand merci aux membres du jury extérieurs au laboratoire : Manuel Joffre, Jean Oberlé, Jean-Michel Gérard, Daniel Kaplan et Xavier Marie pour leurs remarques constructives concernant mon manuscrit. Merci aussi pour leur bonne volonté, leur flexibilité et leur sens de l'improvisation sans lesquels ma soutenance n'aurait pas été possible.

Merci Wendel pour avoir mis à ma disposition tes programmes de contrôle du façonneur, avec en plus une assistance technique toujours disponible. Ce fut une source d'inspiration très importante pour l'écriture de mes programmes.

Merci à Gregor Strobawa pour ses précieux conseils sur la calibration des cristaux liquides.

Merci à Agnès et Marie-France pour le secours qu'elles ont apporté à l'inadapté de la paperaise administrative que je suis...

Merci aux fourmis de l'improvisation, de la logistique et de la restauration pour tout le travail fourni le jour de la soutenance.

Merci à Daniel, Thomas, Richard ainsi que les autres membres de la société Fastlite pour leur collaboration chaleureuse.

Merci Chris pour tes réponses patientes lors de mes nombreuses incursions express dans ton antre du paquet d'onde et du canard en plastique...

Merci Arnaud (non pas celui-là, l'autre) d'être venu gonfler les rangs des Arnaud du LCAR. Ta bonne humeur et tes déclarations implacables ont été très agréables durant la fin de cette thèse.

Merci au hasard, et aux différents enseignants de mécanique quantique, pour m'avoir détourné de mes premières amours en dévoilant les charmes du pays des chats mi-morts/mi-vivants, des particules qui en seraient vraiment si seulement elles n'étaient pas aussi des ondes, ...

Merci Gribouille pour ton soutien de tous les jours...

Bref, merci, merci, merci...

Table des matières

INTRODUCTION	1
I Dispositif expérimental	5
I.1 Présentation générale	5
I.2 Oscillateur	5
I.3 Amplificateur régénératif	6
I.4 NOPA	7
I.5 Dispositifs de caractérisation	10
I.5.1 Autocorrélation - Corrélation	10
I.5.2 Cross-corrélation résolue spectralement : XFROG	12
I.5.3 SPIDER	13
II Façonneur à cristaux liquides	15
II.1 Mise en forme : état de l'art	15
II.2 Choix de la ligne 4f	18
II.2.1 Géométries non retenues	19
II.2.2 Géométrie retenue : géométrie X	22
II.2.3 Evolutions possibles	23
II.3 Choix du masque	24
II.3.1 Masques possibles	24
II.3.2 Double masque à cristaux liquides	26
II.4 Dimensionnement	27
II.4.1 Paramètres pertinents	27
II.4.2 Application particulière	29
II.5 Construction et alignement de la ligne	30
II.5.1 Télescope cylindrique	31
II.5.2 Réseau d'entrée	32
II.5.3 Réseau de sortie	33
II.5.4 Miroirs de repli	33
II.6 Mise en place du masque et des polariseurs	34
II.7 Calibration du déphasage	35
II.7.1 Fonctionnement du masque	35
II.7.2 Simplifier pour mieux calibrer	38
II.7.3 Mise en oeuvre expérimentale	38
II.7.4 Méthode mieux adaptée	40

II.8	Calibration de la dispersion	43
II.9	Résultats et limitations	45
II.9.1	Approche naïve	45
II.9.2	Pixellisation	47
II.9.3	Couplage spatio-temporel	54
II.9.4	Utilisation d'algorithmes d'optimisation	60
II.9.5	Possibilités de compensation	61
II.10	Applications	64
II.10.1	Expériences de Démonstration	65
II.10.2	Application au contrôle cohérent	69
II.10.3	Compensation de dispersion résonnante	69
II.10.4	Façonneur replié	73
III	Autres dispositifs de mise en forme	75
III.1	Miroirs déformables	76
III.1.1	Rappels sur le NOPA	76
III.1.2	Compresseur à miroir déformable	77
III.2	AOPDF : recompression et mise en forme	80
III.2.1	Principe de fonctionnement	80
III.2.2	Compresseur accordable	85
III.2.3	Capacités de mise en forme	89
III.2.4	Vers un NOPA compact et programmable?	90
III.2.5	Bilan	93
IV	Contrôle cohérent dans le Rubidium	95
IV.1	Formalisme	96
IV.1.1	Notations et définitions	96
IV.1.2	Population dans l'état final - fluorescence	97
IV.1.3	Transitoires cohérents - Spirale de Cornu	99
IV.2	Observation expérimentale	102
IV.3	Contrôle des transitoires	104
IV.3.1	Quelques expériences.	104
IV.3.2	Déphasage des transitoires	106
IV.4	Sensibilité au champ	112
IV.5	Spirographe atomique	114
IV.5.1	Des transitoires cohérents à l'atome	115
IV.5.2	Prérequis expérimentaux	124
IV.5.3	Spirographe expérimental	128
IV.5.4	Premiers résultats	128
IV.5.5	On peut jouer avec le spirographe	133
IV.5.6	Mais attention à la casse	133
IV.6	Caractérisation par transitoires	135
IV.6.1	Signal de reconstruction	136
IV.6.2	Reconstructions expérimentales	138
IV.6.3	Extensions et perspectives	145

IV.6.4 Conclusion	147
CONCLUSION	149
Annexe	153
A Notations pour expérimentateur	153
A.1 Transformée de Fourier	153
A.1.1 Définitions	153
A.1.2 Propriétés et transformées de Fourier usuelles	153
A.2 Champ électrique	154
A.2.1 Définitions - Notations	154
A.2.2 Durée - largeur spectrale	154
B Couplage spatiotemporel	157
B.1 Conventions et notations	157
B.1.1 Champs en entrée	157
B.2 Champs au plan de Fourier	157
B.3 Champ en sortie	159
B.4 Filtrage spatial	160
B.5 Exemples	161
B.5.1 Cas d'un délai	161
B.5.2 Cas d'une phase quadratique	162
Bibliographie	165

INTRODUCTION

Les fondements de la physique quantique reposent sur l'étude d'interactions entre la matière et le rayonnement et l'observation d'effets que la physique classique ne pouvait expliquer. L'interaction lumière-matière reste le cœur de la physique quantique. Elle permet à la fois l'étude et le contrôle de la matière.

Dans le domaine de la spectroscopie ultra-rapide, l'idée de contrôler un système quantique ou une réaction chimique par le biais de son interaction avec une ou plusieurs impulsions lumineuses n'est pas récente. De nombreux schémas de contrôle utilisent l'interaction cohérente entre une ou plusieurs impulsions et le système étudié. Ces schémas reposent sur des phénomènes d'interférences quantiques que l'on peut influencer à l'aide d'un petit nombre de paramètres, comme le délai ou la phase entre deux impulsions excitatrices. De nombreuses avancées tant théoriques[1, 2, 3] qu'expérimentales[4, 5] ont été réalisées dans ce domaine au cours des deux dernières décades. Cette approche est parfois limitée car on ne dispose que d'un nombre très restreint de paramètres.

En 1992, Rabitz propose une approche novatrice de *contrôle optimal*[6] qui va s'avérer très prolifique pour la manipulation de systèmes complexes pour lesquels il existe peu ou pas de modèle théorique. L'idée est de reboucler l'expérience sur elle-même par l'utilisation d'algorithme d'optimisation. On parle souvent d'expériences en *boucle fermée*. Pour pousser le système à produire un résultat donné, comme maximiser un signal de fluorescence ou modifier un ratio de produits d'une réaction chimique, on essaye un grand nombre d'impulsions lumineuses fortement mises en forme. En fonction de la pertinence des différentes formes, l'algorithme d'optimisation génère de nouvelles formes que l'on teste à nouveau. En quelques sortes, dans le contrôle optimal, on laisse le système trouver la voie qui répond à notre objectif. La mise en œuvre pratique a demandé de gros efforts de développement. Il a fallu entre autre trouver un moyen de générer ces impulsions fortement façonnées et ce de façon suffisamment rapide. La première démonstration expérimentale a été faite 6 ans plus tard, par l'équipe de Gerber[7]. Cette approche a montré des résultats spectaculaires, en permettant notamment de modifier les produits d'une réaction chimique de façon importante. De nombreux résultats ont ainsi été obtenus entre autres dans des molécules[7, 8, 9], des agrégats[10] et bien d'autres systèmes complexes[11, 12, 13, 14]. Le contrôle optimal a ainsi prouvé qu'il était possible de manipuler des systèmes complexes dont nous ne connaissions presque rien. Pour autant, cette manipulation ne nous apprend pas forcément beaucoup sur la physique qui gouverne le système manipulé.

Parallèlement, une approche différente s'est développée pour l'étude et le contrôle de systèmes plus simples. Cette technique de *contrôle cohérent*, repose sur un schéma dit de *boucle ouverte*. Pour ces systèmes suffisamment bien modélisables, on peut déterminer à l'avance quelle forme d'impulsion appliquer pour obtenir un résultat donné. Contrairement au contrôle en boucle fermée, il n'y a ici aucune itération. Par contre, les deux approches partagent en commun l'utilisation d'impulsions fortement façonnées. De nombreux résultats dans des systèmes atomiques ont ainsi été obtenus entre autres par l'équipe de Silberberg[15, 16] ainsi que notre équipe[17, 18, 19]. D'autres groupes ont appliqué cette approche avec succès à des petites molécules[20, 21]. Cette technique de boucle ouverte

nécessite une connaissance parfaite du système et de ses interactions et cela limite son champ d'application aux systèmes que l'on est capable de modéliser avec précision.

Ces deux approches de contrôle, en boucle ouverte et boucle fermée, se sont influencées l'une - l'autre. Dans un mouvement assez récent, expérimentateurs et théoriciens se sont ainsi penchés sur l'information que l'on peut retirer des expériences de contrôle optimal[22, 9, 23, 10] : est-ce que la forme des impulsions optimales, ou celle des impulsions inadéquates, ou encore l'évolution des impulsions au cours de l'optimisation peut nous renseigner sur le système et ses couplages ? Des approches mixtes, mêlant les deux approches ont ainsi vu le jour[24, 15, 10]. Il s'agit d'un domaine de recherche très actif et encore très jeune : autant l'efficacité du contrôle optimal pour la manipulation de systèmes complexes est une chose entendue, autant sa pertinence pour la compréhension de ces systèmes est encore à prouver.

Ces différentes approches de contrôle, ont été utilisées avec succès pour l'étude et le contrôle de systèmes atomiques[25], moléculaires[24, 26] ou biologiques[12]. Elles n'ont été possibles que grâce au développement des techniques de mise en forme d'impulsions. Les façonneurs actifs¹[27, 28], permettant de générer des impulsions arbitrairement mises en forme, ont ainsi joué un rôle clef. Ils ont permis de diversifier le type d'impulsions utilisables dans des expériences de contrôle cohérent. De plus, c'est leur capacité à générer rapidement des formes très variées qui a rendu possible la réalisation expérimentale du contrôle optimal. Leur développement a suivi de très près celui des laser ultra-courts. L'absence de modulateurs suffisamment rapides pour façonner directement le profil temporel des impulsions à l'échelle femtoseconde a naturellement conduit à l'utilisation de filtres spectraux. Différentes approches ont été suivies pour réaliser ces filtres.

La plus ancienne repose sur l'étalement spatial des composantes spectrales au sein d'un compresseur et l'utilisation de masques spatiaux. Des masques tout d'abord fixes[29] puis programmables[30, 31, 32] ont été employés. Différents types de compresseurs ont été imaginés pour réaliser l'étalement des composantes spectrales[33]. La géométrie la plus répandue à l'heure actuelle est la ligne à dispersion nulle ou ligne 4f introduite initialement par Froehly[34]. Les façonneurs basés sur ce principe permettent aujourd'hui un contrôle de la phase et de l'amplitude des composantes spectrales de l'impulsion mise en forme et ce avec une grande résolution[35]. Cela permet de générer des impulsions au profil temporel extrêmement complexe. Le contrôle de la polarisation a aussi été réalisé[36] récemment avec ce type de dispositifs .

Une deuxième approche plus récente utilise l'interaction colinéaire entre l'impulsion à mettre en forme et une onde acoustique façonnée au sein d'un cristal biréfringent[37, 27]. Ces différentes techniques ont permis la génération d'impulsions dans le proche infrarouge[38, 39] et plus récemment dans le visible[40, 41]. L'extension de ces techniques à l'infrarouge[42, 43, 44] et à l'ultraviolet[43, 45] est un domaine de recherche très actif et des résultats prometteurs ont été obtenus. Cela devrait permettre d'élargir rapidement le champ d'application des techniques de contrôle cohérent ou de contrôle optimal à des gammes spectrales

¹Cette appellation de "façonneurs actifs" qui a été consacrée par la communauté scientifique, est assez malheureuse. En effet, presque tous ces façonneurs se comportent comme des filtres spectraux passifs, c'est-à-dire sans amplification. L'appellation "façonneurs programmables" serait moins ambiguë. Nous utiliserons dans la suite de cette thèse l'appellation consacrée en gardant à l'esprit cette ambiguïté.

jusqu'ici hors d'atteinte.

Suite aux différentes études de contrôle cohérent réalisées par l'équipe[17, 18, 19, 46, 47], il est apparu opportun de se doter d'un appareil de mise en forme. Nous souhaitions disposer d'un façonneur permettant un contrôle total des impulsions infrarouges issues de notre chaîne : c'est-à-dire contrôlant à la fois l'amplitude et la phase des différentes composantes spectrales de ces impulsions. Nous voulions aussi obtenir la meilleure résolution spectrale possible. La conception et la construction de celui-ci, qui mêlent des solutions inédites[35] à des techniques largement éprouvées[48, 49], a été mon principal travail durant la première moitié de cette thèse.

Par la suite, et dans la continuation d'études précédemment menées au sein de l'équipe[18, 17], nous avons réalisé des expériences de contrôle cohérent dans l'atome de rubidium[47, 50, 51, 52] en exploitant les capacités de ce façonneur.

Un autre façonneur[41] permettant la mise en forme d'impulsions visibles accordables a aussi été étudié durant cette thèse.

Ce rapport est organisé de la façon suivante :

- Le chapitre I présente brièvement les différents éléments constitutifs du système laser que nous avons utilisé, dont plusieurs dispositifs de caractérisation d'impulsions femtoseconde.
- Le chapitre II détaille tout d'abord la conception, la construction et les réglages du façonneur haute résolution : cette première partie se veut une base de départ explicite pour qui souhaite développer ce genre de dispositif. Dans une deuxième partie, les résultats obtenus avec cet appareil, ainsi que ses limitations sont présentées.
- Le chapitre III montre les différents dispositifs de façonnage d'impulsions visibles que nous avons employés. Nous insisterons en particulier sur un dispositif acousto-optique spécialement optimisé pour la recompression et la mise en forme d'impulsions accordables dans toute la gamme visible.
- Enfin, le chapitre IV est consacré aux expériences de contrôle cohérent dans le rubidium. Ces expériences reposent sur le contrôle du comportement transitoire d'un atome lors de son interaction avec une impulsion résonnante. Nous rappellerons tout d'abord l'origine de ce comportement transitoire, nommé *transitoire cohérent*[53], ainsi que sa première mise en évidence expérimentale[17]. Nous présenterons aussi différentes expériences visant à manipuler ce transitoire cohérent en modifiant la forme des impulsions interagissant avec l'atome[16, 18, 47]. Puis, nous verrons comment manipuler différents transitoires cohérents pour reconstruire la fonction d'onde atomique transitoire[51], c'est-à-dire au cours de l'interaction de l'atome avec une impulsion ultra-courte. Cette technique dite de *spirographe atomique* a permis la reconstruction de la fonction d'onde de l'atome soumis à différentes impulsions de contrôle. Nous présenterons ces reconstructions et discuterons de leur domaine de validité ainsi que de leur lien avec d'autres techniques de mesure de fonctions d'onde. Nous présenterons ensuite une technique de caractérisation complète d'impulsions ultra-courtes directement dérivée du spirographe atomique. Cette caractérisation par transitoires cohérents permet de mesurer l'im-

pulsion (les impulsions en fait) interagissant avec l'atome[50, 52]. Elle présente des caractéristiques assez inhabituelles que nous détaillerons. Nous présenterons les différentes mesures expérimentales réalisées et discuterons des perspectives qu'ouvre cette technique de mesure.

Chapitre I

DISPOSITIF EXPÉRIMENTAL

I.1 Présentation générale

La chaîne laser utilisée au sein de l'équipe FEMTO a déjà été décrite très en détail[54, 55] et ses caractéristiques sont très classiques pour une chaîne de spectroscopie ultrarapide. Nous allons donc la présenter de façon succincte. La figure I.1 est une vue schématique des divers éléments de cette chaîne.

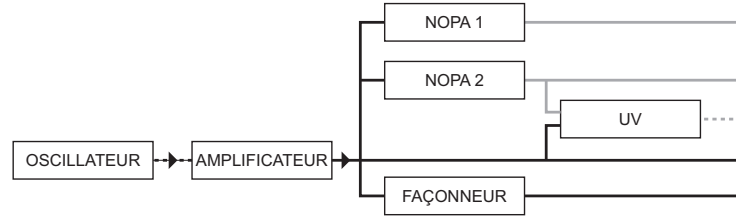


FIG. I.1 – Présentation schématique de la chaîne laser de l'équipe.

L'oscillateur (section I.2) fournit des trains d'impulsions infrarouges (centrées à 800 nm) d'une centaine de femtosecondes, à haute cadence (76 MHz) et de basse intensité crête (5 nJ). Celles-ci sont amplifiées dans un amplificateur régénératif (section I.3 page suivante) : en sortie de ce dernier on obtient une plus faible cadence (1kHz), mais avec des intensités crêtes élevées (1 mJ). Ces impulsions amplifiées sont séparées sur plusieurs voies. Deux d'entre elles alimentent des amplificateurs paramétriques optiques non-colinéaires (section I.4 page 7) ou NOPA (pour l'acronyme anglais Non-colinear Optical Parametric Amplifier). Ceux-ci fournissent des impulsions très courtes (25 fs) et accordables dans le visible (de 500 à 700 nm). Les deux voies restantes peuvent être utilisées telles quelles ou mise en forme à l'aide d'un façonneur décrit en détail dans le chapitre II page 15. De plus, par somme de fréquence entre les impulsions visibles et infrarouges ou par doublage ou triplage de fréquence des impulsions infrarouges, des impulsions UV accordables peuvent être générées. Enfin, différents dispositifs sont utilisés pour caractériser les différentes impulsions générées dans cette chaîne.

I.2 Oscillateur

Le point de départ de toutes les impulsions disponibles dans notre chaîne est l'oscillateur. Il s'agit d'un oscillateur titane-saphir développé et construit dans le laboratoire.

Le principe général de ce type d'oscillateur a fait l'objet de nombreuses descriptions, voir notamment la référence [56]. Les particularités de notre oscillateur ont été décrites dans les thèses de C. Nicole [54] et J. Degert [55]. Il repose sur un verrouillage de modes passif par effet Kerr optique. Il délivre des impulsions de 95 fs de durée¹, avec un taux de répétition de 76 MHz et 5 nJ par impulsion. La longueur d'onde des impulsions est accordable entre 760 et 840 nm.

Le tableau I.1 résume les principales caractéristiques de la sortie de l'oscillateur, en mode continu et impulsionnel.

Propriété	Valeur
Puissance moyenne en régime continu	840 mW
Puissance moyenne en régime impulsionnel	380 mW
Taux de répétition	76 MHz
Énergie par impulsion	5 nJ
Longueur d'onde centrale	760-840 nm
Largeur spectrale	10 nm
Durée	95 fs

TAB. I.1 – Principales caractéristiques de notre oscillateur titane-saphir.

I.3 Amplificateur régénératif

Les impulsions produites par l'oscillateur permettent de réaliser un grand nombre d'expériences. Leur intensité crête est suffisante pour observer un certain nombre d'effets non-linéaires.

Cependant, les expériences en champ fort nécessitent de hautes intensités. De même, de nombreuses expériences requièrent une certaine accordabilité et donc des mélanges non-linéaires.

On amplifie donc les impulsions issues de l'oscillateur grâce à un amplificateur régénératif (ou RGA pour ReGenerative Amplifier) commercial (Spitfire, Spectra Physics). L'amplificateur proprement dit (une cavité avec un barreau de titane-saphir) est placé entre un étireur et un compresseur. Cela permet d'atteindre des intensités crêtes élevées². On obtient ainsi en sortie de l'amplificateur des impulsions de l'ordre de 120 fs, avec un taux de répétition de 1 kHz et 1 mJ par impulsion. Cela représente un gain de plus de 5 ordres de grandeur sur l'énergie par impulsion. En sortie de celui-ci, la cadence des impulsions est abaissée à 1 kHz. Le principe de cet amplificateur est de piéger une impulsion à amplifier dans une cavité possédant un milieu amplificateur (barreau de titane-saphir pompé

¹Toutes les durées, sauf mention contraire, sont les largeurs à mi-hauteur en intensité (FWHM en anglais).

²En particulier, beaucoup plus élevées que le seuil de dommage de la partie amplificateur.

optiquement par un laser kHz commercial). L'impulsion s'intensifie à chaque tour de cavité jusqu'à saturation. Elle est alors éjectée. Une explication détaillée du fonctionnement de cet amplificateur, notamment les procédures d'injection et d'éjection de l'impulsion, se trouve dans la thèse de Jérôme Degert [55].

Les principales caractéristiques des impulsions issues de l'amplificateur régénératif sont rassemblées dans le tableau I.2.

Propriété	Valeur
Puissance moyenne	950 mW
Taux de répétition	1 kHz
Énergie par impulsion	950 μ J
Longueur d'onde centrale	790-810 nm
Largeur spectrale	$\sim 8,5$ nm
Durée	~ 120 fs

TAB. I.2 – Caractéristiques de l'amplificateur régénératif.

I.4 NOPA

Nous disposons, en sortie de l'amplificateur régénératif, d'impulsions intenses (plusieurs centaines de μ J) dans le proche infrarouge (entre 810 et 790 nm). Pour avoir accès à une plus grande variété de systèmes atomiques ou moléculaires, il nous faut des impulsions femtosecondes accordables.

Nous disposons pour cela de deux amplificateurs paramétriques optiques non colinéaires ou NOPA. Ceux-ci sont un type particulier d'oscillateurs paramétriques optiques ou OPA[57], permettant d'obtenir des impulsions de quelques dizaines de femtosecondes, accordables sur l'ensemble du spectre visible. L'étude du fonctionnement des NOPA[58, 59, 60] dépasse le cadre de cette thèse et a déjà été décrit en détail dans la référence [55]. Nous nous bornerons ici à rappeler brièvement le principe de l'amplification paramétrique optique et la spécificité de la géométrie non-colinéaire.

L'amplification paramétrique optique repose sur le principe de la différence de fréquence dans un cristal non-linéaire du second ordre. Ce principe est rappelé sur la figure I.2 page suivante. Un champ de pompe intense, de pulsation ω_p ainsi qu'un champ signal très faible, de pulsation ω_i sont envoyés dans un cristal non-linéaire ($\chi^{(2)}$). Par différence de fréquence, une partie des photons du champ pompe, est convertie en paires de photons de pulsations ω_s et ω_i . Ces pulsations vérifient la relation de conservation de l'énergie $\omega_p = \omega_s + \omega_i$. L'interaction au sein du cristal du champ de pompe avec le champ signal se traduit donc par une amplification de ce dernier et l'apparition d'un champ complémentaire (ou idler) de pulsation $\omega_i = \omega_p - \omega_s$.

Dans le cas d'impulsions courtes, plusieurs processus limitent cette amplification :

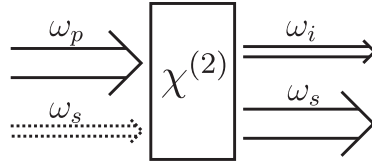


FIG. I.2 – Principe de la différence de fréquence utilisée dans l'amplification paramétrique optique.

- Les impulsions pompe et signal se propagent à des vitesses différentes dans le cristal (désaccord de vitesse de groupe ou GDM pour "Group Delay Mismatch") ce qui réduit leur interaction.
- Les impulsions sont étirées lors de leur propagation (dispersion de la vitesse de groupe ou GDD pour "Group Delay Dispersion") et donc moins intenses ce qui diminue l'amplification.
- En plus de la condition de conservation de l'énergie, les trois champs présents dans le cristal doivent vérifier la conservation de l'impulsion ou accord de phase : $\vec{k}_p = \vec{k}_s + \vec{k}_i$. Cet accord de phase est en général très sélectif en fréquence et ne permet pas d'amplifier de large bande spectrale et donc des impulsions courtes.

La configuration non-colinéaire, développée tout d'abord pour des oscillateurs paramétriques optiques[61, 62], permet de lever ces limitations. L'idée est d'introduire un angle α entre l'impulsion de pompe et l'impulsion signal. Pour un angle bien choisi, l'accord de phase est vérifié sur une plage spectrale qui couvre une grande partie du spectre visible. La figure I.3 montre l'évolution de l'angle d'accord de phase θ en fonction de la longueur d'onde du champ signal, pour différents angles α entre le champ pompe et le signal. Ceci

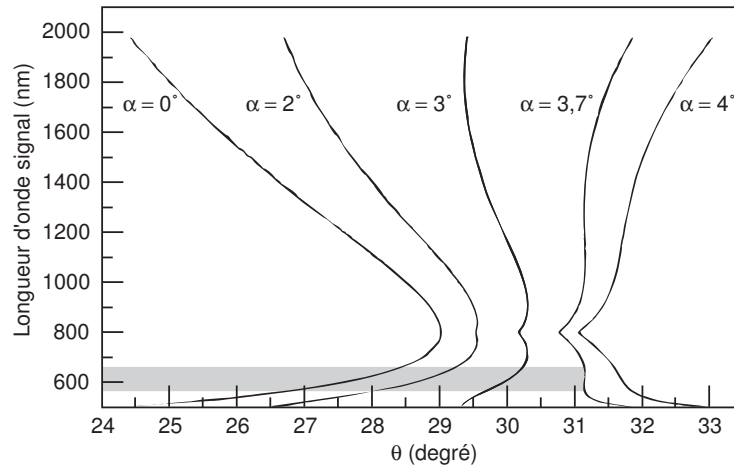


FIG. I.3 – Angle d'accord de phase θ (cristal de BBO type I) pour différents angles α entre les champs pompe et signal. Le champ pompe est centré à 400 nm. Pour un angle de $3,7^\circ$, l'accord de phase est constant sur une plage de 100 nm centrée à 580 nm (rectangle gris).

est tracé pour un cristal de BBO type I et une impulsion pompe centrée à 400 nm. Pour $\alpha \simeq 3,7$ degrés, l'accord de phase est constant sur une plage allant de 530 à 630 nm.

La figure I.4 présente les NOPA tels qu'ils sont développés dans notre laboratoire. Un tiers du faisceau issu de l'amplificateur alimente le NOPA (800 nm, 350 μ J). Une très

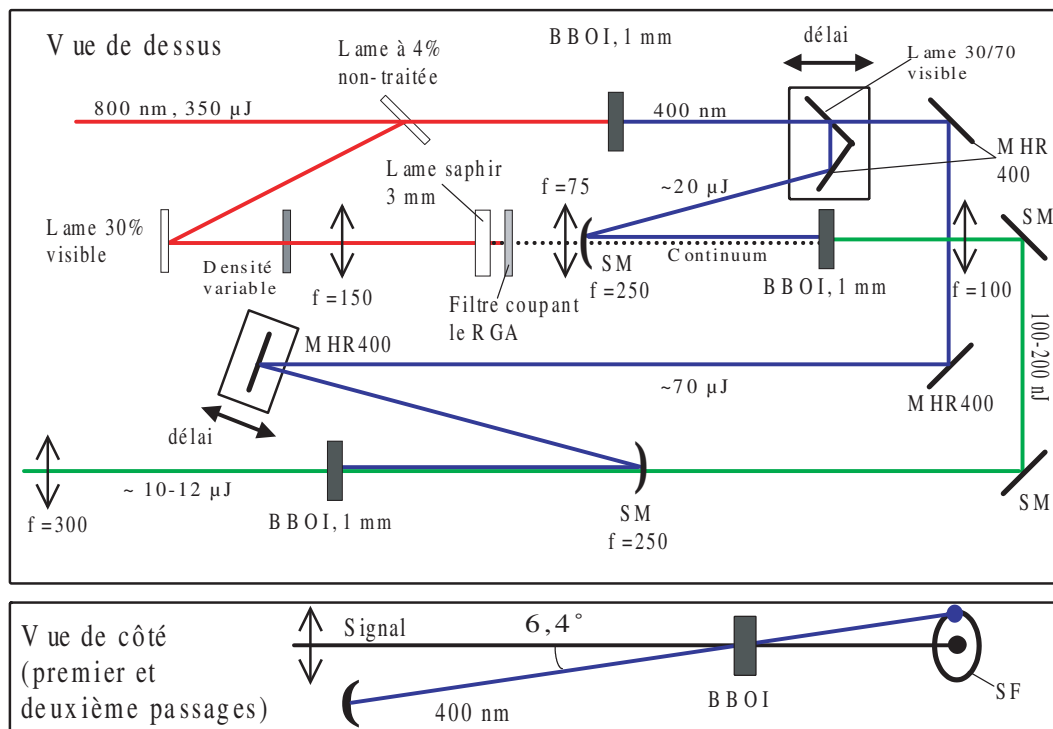


FIG. I.4 – **Figure tirée de la référence [55]** Schéma de l'amplificateur paramétrique optique.

faible partie est prélevée et focalisée dans une lame de saphir (3 mm d'épaisseur). Cela génère un continuum de lumière blanche à forte dérive de fréquence (300 fs de durée). Le reste du faisceau infrarouge est envoyé dans un cristal de BBO type I (1 mm d'épaisseur) afin d'obtenir une impulsion de pompe à 400 nm. Une part de cette pompe (20 μ J) et le continuum sont focalisés dans un cristal de BBO type I (1 mm d'épaisseur). L'amplification paramétrique au sein de ce cristal permet d'obtenir un champ signal de quelques centaines de nJ. La longueur d'onde centrale de ce signal est accordée en changeant le délai entre la pompe et le continuum. En effet, la pompe est plus courte que le continuum et elle n'amplifie qu'une "tranche" temporelle du continuum; compte tenu de la dérive de fréquence du continuum, cette "tranche" a une longueur d'onde qui dépend du délai.

Afin d'obtenir une dizaine de μJ de signal, un deuxième étage est utilisé. Il est identique au premier, mis à part la pompe qui est plus intense ($70 \mu\text{J}$).

La sortie du NOPA est recomprimée par un compresseur à prismes non représenté sur la figure. Le tableau I.3 page suivante récapitule les propriétés de sortie des deux NOPA.

Propriété	Valeur
Énergie par impulsion avant compression	10-14 μJ
Taux de répétition	1 kHz
Longueur d'onde centrale	510-700 nm
Largeur spectrale	15-60 nm
Durée	20-30 fs

TAB. I.3 – Principales caractéristiques des impulsions en sortie des deux NOPA.

I.5 Dispositifs de caractérisation

La caractérisation des impulsions ultracourtes est un problème à part entière. En effet, celle-ci nécessite des techniques tout à fait originales, qui répondent à la question suivante : comment mesurer l'événement le plus court que l'on sache produire ? En effet, les méthodes de mesure habituelles reposent sur l'utilisation d'une référence de temps plus courte que l'événement à mesurer. Dans le cas des impulsions femtoseconde, une telle référence n'existe pas (à moins bien sûr de disposer d'une impulsion attoseconde [63]). Il existe cependant de nombreuses techniques qui permettent une mesure partielle ou complète d'une impulsion femtoseconde. Nous présentons ici trois techniques que nous avons utilisées pour caractériser les impulsions à 800 nm et les impulsions visibles disponibles.

I.5.1 Autocorrélation - Corrélation

I.5.1.1 Autocorrélation

Une première technique est l'autocorrélation intensimétrique. A l'aide d'un interféromètre de Michelson, l'impulsion à caractériser $E(t)$ est séparée en deux parties identiques, décalées temporellement d'un délai τ , comme cela est schématisé sur la figure I.5. Ces deux

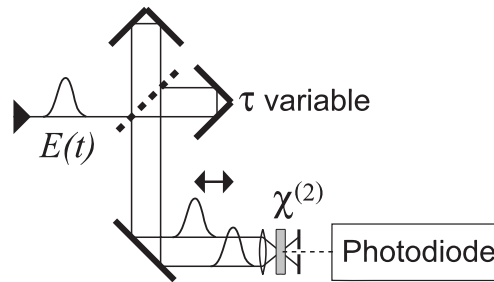


FIG. I.5 – Schéma de principe de l'autocorrélation intensimétrique.

répliques sont superposées dans un cristal de BBO type I, avec un léger angle. Avec cette configuration, l'intensité de la deuxième harmonique générée $I_2(t)$ s'écrit en fonction du

délai τ :

$$I_2(t, \tau) \propto I(t)I(t - \tau) \quad (\text{I.1})$$

où $I(t) = |E(t)|^2$ est l'intensité de l'impulsion à caractériser. Cette intensité $I_2(t, \tau)$ est détectée par une photodiode. Cette dernière ayant un temps de réponse très grand devant la durée de $I_2(t, \tau)$, elle donne un signal $S_{AC}(\tau)$ proportionnel à l'intégrale de cette intensité :

$$S_{AC}(\tau) \propto \int I_2(t, \tau) dt \propto \int I(t)I(t - \tau) dt \quad (\text{I.2})$$

Ce signal n'est autre que la fonction d'autocorrélation de l'intensité $I(t)$. Il est clair que ce signal ne contient pas toute l'information sur $E(t)$: cette mesure n'est pas une mesure complète. En particulier, elle ne donne aucune information sur la phase temporelle et n'apporte qu'une information incomplète sur l'amplitude. Elle permet cependant d'avoir une estimation rapide de la durée d'une impulsion ayant un profil simple. En supposant que $E(t)$ a un profil gaussien, on a un lien simple entre la largeur à mi-hauteur de $S_{AC}(\tau)$, notée $\Delta\tau_{AC}$ et la largeur à mi-hauteur en intensité de $E(t)$, notée Δt :

$$\Delta\tau_{AC} = \Delta t\sqrt{2} \quad (\text{I.3})$$

L'autocorrélateur dont nous disposons permet de mesurer les impulsions à 800 nm issues de l'amplificateur ou les impulsions visibles issues des NOPA. Pour passer d'une longueur d'onde à une autre, on ajuste l'accord de phase dans le cristal de BBO en inclinant ce dernier. Selon la durée des impulsions à mesurer, deux cristaux différents sont utilisables : 200 μm pour les impulsions les plus longues ou 20 μm pour les impulsions courtes. Bien que partielle, cette mesure est très utile, pour la recompression des impulsions issues du NOPA ou pour le réglage du compresseur de l'amplificateur régénératif.

I.5.1.2 Corrélation

La corrélation intensimétrique (aussi appelée "cross-corrélation") est une autre technique de mesure incomplète. Le principe est tout à fait identique à celui de l'autocorrélation mis à part que l'on superpose dans le cristal non-linéaire deux impulsions différentes $E(t)$ et $E_{ref}(t)$, comme cela est schématisé sur la figure I.6. En fonction du délai τ entre les

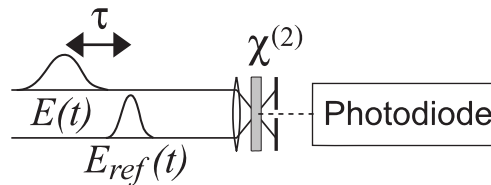


FIG. I.6 – Schéma de principe de la cross-corrélation intensimétrique.

deux impulsions, la photodiode fournit un signal $S_{XC}(\tau)$:

$$S_{XC}(\tau) \propto \int I_2(t, \tau) dt \propto \int I_{ref}(t)I(t - \tau) dt \quad (\text{I.4})$$

où $I(t) = |E(t)|^2$ et $I_{ref}(t) = |E_{ref}(t)|^2$. Les deux impulsions peuvent être dans des domaines spectraux différents. Si les deux impulsions ont un profil temporel gaussien de largeurs à mi-hauteur en intensité respectives de Δt et Δt_{ref} , le signal de cross-corrélation sera aussi gaussien, de largeur à mi-hauteur $\Delta\tau_{XC}$:

$$\Delta\tau_{XC} = \sqrt{\Delta t^2 + \Delta t_{ref}^2} \quad (\text{I.5})$$

En particulier, si l'impulsion de référence $E_{ref}(t)$ est très courte devant $E(t)$, la cross-corrélation donne une bonne estimation du profil d'intensité $I(t) = |E(t)|^2$.

Elle est ainsi tout à fait adaptée à la caractérisation partielle d'impulsions fortement mises en forme ou des trains d'impulsion. Nous l'avons utilisé pour mesurer les impulsions infra-rouges mises en forme par le façonneur haute résolution (voir chapitre II) et les impulsions visibles mises en forme par un autre dispositif de façonnage (voir chapitre III).

I.5.2 Cross-corrélation résolue spectralement : XFROG

Afin d'obtenir une caractérisation complète des impulsions femtoseconde, nous avons réalisé un dispositif de cross-corrélation résolue spectralement ou XFROG[64]. Ce dispositif est une variante du "Frequency Resolved Optical Gating" ou FROG[65, 66]. L'idée est de remplacer la photodiode par un spectromètre afin d'obtenir la répartition spectrale de l'intensité pour chaque délai τ , comme cela est présenté sur la figure I.7. On obtient ainsi

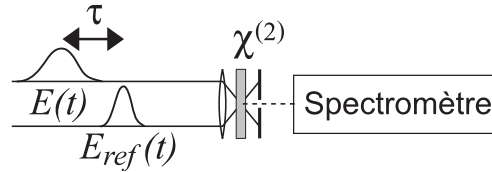


FIG. I.7 – Schéma de principe de la cross-corrélation intensimétrique résolue spectralement (XFROG).

une carte en deux dimensions $S_{XF}(\tau, \omega)$. Celle-ci contient suffisamment d'informations pour pouvoir reconstruire $E_{ref}(t)$ et $E(t)$, en phase et en amplitude. Cette reconstruction fait appel à une procédure itérative. La plus performante à l'heure actuelle repose sur la décomposition de cette trace en composants principaux[67] : "Principal Component Generalized Projects Algorithm" ou PCGPA. $S_{XF}(\tau, \omega)$ contient en fait beaucoup plus d'informations que le strict minimum. Cette redondance assure une certaine robustesse au bruit.

Cette méthode de caractérisation est relativement bien adaptée aux impulsions mises en forme. En effet, elle permet la caractérisation d'impulsions complexes, composées de plusieurs parties distinctes étalées dans le temps. On peut quantifier la complexité d'une impulsion mise en forme[28] en définissant η comme le rapport entre le support temporel ΔT de l'impulsion et la plus petite structure temporelle δt de cette dernière³. Cette définition

³ η peut aussi être défini dans le domaine spectral et le résultat revient au même.

permet de cerner la principale limitation du XFROG concernant les impulsions mises en forme : pour caractériser convenablement une impulsion de complexité η , il faut disposer d'une reconstruction sur un nombre de points de l'ordre de η . Dans le cas du XFROG, cela nécessite une trace $S_{XF}(\tau, \omega)$ contenant η^2 points de mesure. La quantité de données à traiter croît donc de façon quadratique avec la complexité des impulsions et on peut facilement atteindre des temps de calculs rédhibitoires pour des impulsions couramment utilisées durant cette thèse.

Malgré ce défaut, cette technique apparaît comme la plus adaptée pour la caractérisation des impulsions mises en forme. Malheureusement, nous ne l'avons mise en œuvre que tardivement. Autant la réalisation expérimentale s'est avérée très simple, autant l'écriture et l'utilisation de l'algorithme de PCGPA ont été difficiles⁴. C'est pourquoi la plupart des mesures présentées dans cette thèse sont des cross-corrélations.

I.5.3 SPIDER

Le deuxième appareil de caractérisation complète dont nous disposons repose sur la technique de SPIDER[68] : "Spectral Phase Interferometry for Direct Electric Field Reconstruction". Cette technique reprend le principe de l'interférométrie spatiale à décalage[69], mais transposée au domaine spectral. A partir de l'impulsion à caractériser E , on génère deux copies E_{d1} et E_{d2} décalées temporellement de τ et spectralement de $\delta\omega$. Les franges d'interférences spectrales entre ces deux impulsions sont mesurées par un spectromètre. Un traitement adapté de ces franges permet de reconstruire la phase :

$$\theta(\omega) = \phi(\omega + \delta\omega) - \phi(\omega) + \omega\tau \quad (\text{I.6})$$

où $\phi(\omega)$ est la phase spectrale du champ inconnu $\tilde{E}(\omega)$. A partir de cette quantité, connaissant τ et pour $\delta\omega$ suffisamment petit, on peut reconstruire $\phi(\omega)$. L'amplitude spectrale $|\tilde{E}(\omega)|$ peut être déduite du même signal spectral[70] ou mesurée indépendamment. La connaissance de la phase $\phi(\omega)$ et de l'amplitude spectrale $\tilde{E}(\omega)$ assure une caractérisation complète du champ E .

La mise en œuvre pratique de cette technique se heurte au problème du décalage spectral. La réalisation de ce décalage est hors de portée des modulateurs temporels actuels (ou presque, voir la référence[71]). On doit donc utiliser un moyen détourné pour réaliser ce décalage. Pour cela, on utilise le dispositif présenté sur la figure I.8 page suivante.

L'impulsion à caractériser $E(t)$ traverse une lame de microscope de 200 μm . Par réflexion sur chacune des faces de la lame, deux copies de $E(t)$ séparées de $\tau \simeq 1,5$ ps sont générées. Nous les notons $E_1(t)$ et $E_2(t)$. La partie transmise par la lame est envoyée dans un étireur à réseaux en double passage, qui introduit une forte phase quadratique $\phi^{(2)} = 8,5 \text{ fs}^2$. On obtient ainsi une impulsion à forte dérive de fréquence $E_c(t)$. Les deux copies de l'impulsion à caractériser $E_1(t)$ et $E_2(t)$ et l'impulsion à dérive de fréquence sont recombinaées dans un cristal de KDP type II de 200 μm , dans une géométrie non colinéaire. Lors de son interaction dans le cristal avec $E_1(t)$ et $E_2(t)$, $E_c(t)$ se comporte comme une

⁴L'algorithme impose entre autres des contraintes très strictes sur l'échantillonnage de la carte $S_{XF}(\tau, \omega)$. Ces dernières sont difficiles à vérifier et une étape complexe de pré-traitement des données est nécessaire.

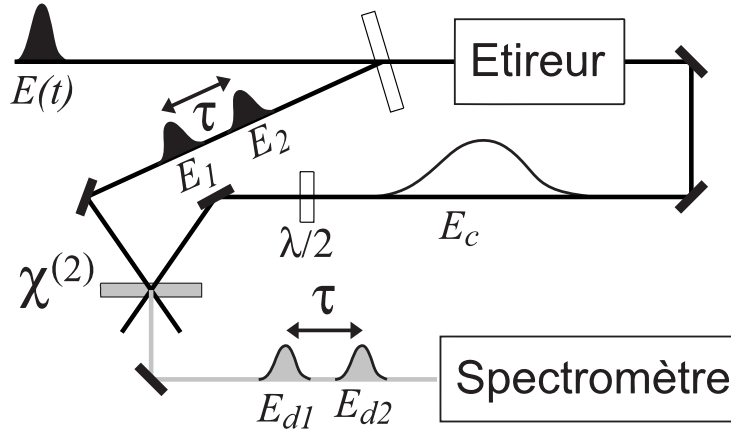


FIG. I.8 – Schéma de principe du SPIDER.

impulsion monochromatique. On appelle ω_1 sa pulsation au moment de son interaction avec E_1 . Par somme de fréquence l'interaction entre E_1 et E_c produit une impulsion notée E_{d1} qui est une copie de E_1 , mais décalée spectralement de ω_1 . A cause de la dérive de fréquence, la pulsation apparente de $E_c(t)$ n'est pas la même au moment de son interaction avec $E_2(t)$: elle vaut $\omega_1 + \delta\omega$. Comme précédemment, l'interaction entre E_2 et E_c produit une impulsion notée E_{d2} qui est une copie de E_2 , décalée spectralement de $\omega_1 + \delta\omega$.

Au bout du compte, on obtient bien deux copies de $E(t)$ (E_{d1} et E_{d2}) qui sont décalées temporellement et spectralement. Ces deux impulsions sont envoyées dans un spectromètre et l'analyse du spectrogramme obtenu permet la reconstruction de la phase spectrale associée à $\tilde{E}(\omega)$. Une mesure indépendante du spectre $|\tilde{E}(\omega)|^2$ de l'impulsion inconnue permet la caractérisation complète.

Cette technique de mesure complète présente l'avantage d'être rapide car elle ne repose sur aucun algorithme itératif. Elle est tout indiquée pour l'optimisation en temps réel de la durée des impulsions issues de l'amplificateur à 800 nm. Cependant, elle est très limitée pour la caractérisation d'impulsions fortement mises en forme. En effet, la validité de la reconstruction repose sur un choix judicieux des deux paramètres que sont le délai τ et la dérive de fréquence $\phi^{(2)}$. Pour des impulsions fortement étirées ou des séquences d'impulsions, il est parfois impossible de trouver des paramètres convenables tout en assurant une quantité de signal suffisante.

Il existe toutefois différentes variantes du SPIDER qui permettent d'obtenir un signal plus important[72, 73] ou de disposer de plus de marge de manœuvre dans le choix des paramètres[74]. Mais ces méthodes sont plus adaptées aux impulsions peu intenses qu'aux impulsions fortement mises en forme. En particulier, elles offrent une résolution relativement faible, de l'ordre du nm à 800 nm, alors que la caractérisation des impulsions issues de notre façonneur (voir chapitre II) réclame une résolution de l'ordre de 0,06 nm à 800 nm. De plus, aucune de ces méthodes de SPIDER n'arrive à reconstruire correctement la phase spectrale lorsque le spectre est composé de différents lobes, comme c'est le cas lors de la génération d'une séquence d'impulsions.

Chapitre II

FAÇONNEUR À CRISTAUX LIQUIDES

Mise en garde : ce chapitre peut paraître un peu technique. Il l'est et se revendique comme tel. Nous souhaitons donner ici les informations nécessaires à la conception et à la construction d'un façonneur à cristaux liquides. En effet, une proportion importante de cette thèse a été dévolue à la conception de notre façonneur. De plus, à l'heure actuelle, de nombreux laboratoires travaillant dans des domaines variés se lancent dans la construction de tels dispositifs. Il nous semble donc opportun de détailler les étapes clefs ainsi que les difficultés d'une telle réalisation.

Après un aperçu des techniques actuelles de mise en forme, nous donnerons les différentes étapes de la conception d'un façonneur à cristaux liquides. Nous verrons tout d'abord comment déterminer la géométrie de la ligne optique (dite ligne à dispersion nulle ou ligne 4f) qui entoure le masque à cristaux liquides. Puis nous présenterons les choix possibles pour le masque. Nous présenterons ensuite comment dimensionner les différents éléments du façonneur en fonction des performances souhaitées. Viendront ensuite la construction proprement dite ainsi que la calibration du masque. Nous présenterons enfin les résultats obtenus avec notre façonneur ainsi que plusieurs applications à des expériences de contrôle cohérent.

II.1 Mise en forme : état de l'art

La mise en forme d'impulsions ultra-courtes a suivi de très près l'évolution des sources laser ultra-rapides. La plupart des systèmes utilisés aujourd'hui pour façonner les impulsions femtoseconde dérivent des appareils développés il y a plus de vingt ans pour les impulsions picosecondes[34]. Une vaste majorité de ces façonneurs combinent un masque spatial programmable et un dispositif optique servant à imager le spectre de l'impulsion sur le masque.

Très souvent le dispositif optique utilisée est une ligne à dispersion nulle ou ligne 4f présentée sur la figure II.1 page suivante. Elle est composée d'une paire de réseaux (R1 et R2) et d'une paire de lentilles de focale f (L1 et L2), le tout disposé dans un arrangement "4f". Pour des impulsions très courtes, les lentilles sont remplacées par des miroirs, afin d'éviter toute dispersion et toute aberration chromatique. Les composantes spectrales de l'impulsion d'entrée sont séparées spatialement et imagées dans le plan de symétrie du dispositif appelé plan de Fourier (PF). Dans ce plan, la pulsation et la position transverse sont couplées : toute modification du profil spatial affectera le profil spectral. Le masque spatial programmable est disposé dans ce plan : il va se comporter comme un masque

spectral. Différents types de masques sont utilisables : barrettes de cristaux liquides[49],

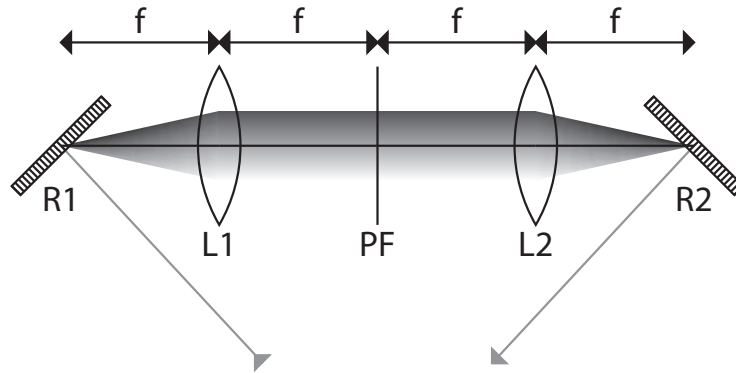


FIG. II.1 – Ligne à dispersion nulle (ligne 4f).

écrans à cristaux liquides[75, 76], modulateurs acousto-optiques transverses[31], miroirs à membrane déformable[32], écran à micro-miroirs électromécaniques[45], valves optiques[77]. Pour les masques réfléchifs (membranes et micro-miroirs), la ligne 4f est repliée comme

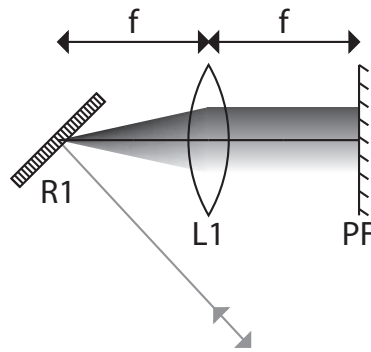


FIG. II.2 – Ligne à dispersion nulle repliée (ligne 2f) pour l'utilisation de masques réfléchifs.

cela est schématisé sur la figure II.2. On a alors un dispositif tout à fait analogue à un spectromètre à réseaux où le masque prend la place de la caméra CCD.

Selon le type de masque utilisé, ces façonneurs contrôlent la phase[28], la phase et l'amplitude[78], la polarisation[36] mais aussi le profil spatial transverse[76]. Ils ont prouvé leur efficacité dans le proche infrarouge et une partie du visible (de 430 nm à 1600 nm pour les masques à cristaux liquides).

Un autre façonneur appelé AOPDF (pour Acousto-Optic Programmable Dispersive Filter) est aussi largement répandu. Contrairement aux autres façonneurs évoqués plus haut, il n'utilise pas d'optique dispersive mais utilise un concept développé récemment[27]. Ce filtre acousto-optique repose sur l'interaction longitudinale entre une impulsion lumineuse et une onde acoustique dans un cristal de TeO_2 . Il permet la mise en forme en phase et en amplitude d'impulsions dans l'infrarouge[79, 27] (entre 800 nm et 2700 nm)

et récemment (voir le chapitre III page 75 de cette thèse) dans le visible[41](de 480 nm à 800 nm).

Ces différents façonneurs possèdent leurs inconditionnels et leurs détracteurs. A la vérité, ils ont tous des caractéristiques très différentes les uns des autres, en particulier en terme de résolution spectrale, de rendement, de contraste ou encore d'énergie disponible en sortie. De même, la facilité d'utilisation, d'alignement et le taux de rafraîchissement sont très variables. Tout dépend du type d'application à laquelle on destine le façonneur. En effet, les caractéristiques pertinentes ne sont pas les mêmes selon que l'on souhaite nettoyer une impulsion courte pour se rapprocher au mieux de la limite par transformée de Fourier ou que l'on veut obtenir les mises en formes les plus complexes possibles, sur de très larges fenêtres temporelles. De même, pour l'utilisation d'algorithme d'optimisation, c'est à dire une utilisation en boucle fermée, un taux de rafraîchissement élevé est souhaitable. Cela est moins important pour une utilisation "directe" en boucle ouverte. L'énergie en sortie (qui dépend à la fois du seuil de dommage et de l'efficacité de transmission) peut aussi être dans certains cas un paramètre important, notamment pour l'étude de phénomènes non-linéaires.

A l'heure actuelle, les façonneurs à barrettes de cristaux liquides permettent d'atteindre les mises en forme les plus complexes. Ils supportent aussi des énergies relativement élevées (plusieurs centaines de μJ par impulsions) avec des rendements importants. Par contre, leur taux de rafraîchissement (de l'ordre du Hz) est faible et les impulsions qu'ils délivrent n'offrent pas le meilleur contraste : différents types de répliques parasites accompagnent les impulsions mise en forme et les profils temporel et spatial sont couplés. Ces répliques se retrouvent sur les autres types de masques pixellisés. Les masques acousto-optiques ou les valves optiques offrent de meilleurs contrastes (absence de répliques notamment) mais avec des rendements et des résolutions spectrales plus faibles. Les acousto-optiques permettent des taux de rafraîchissement plus élevés (jusqu'au kHz). L'AOPDF quant à lui est beaucoup plus simple d'alignement et d'utilisation que les façonneurs à base de ligne 4f. Il offre un contraste et un taux de rafraîchissement élevés. Il ne supporte par contre que des énergies relativement faibles (de l'ordre de 20 μJ à 800 nm). Son rendement varie grandement en fonction de la taille et de la coupe du cristal mais aussi des propriétés des impulsions à mettre en forme. Sa compacité permet toutefois de le placer facilement avant des éventuels étages d'amplification. Récemment, l'AOPDF a montré des performances très intéressante dans le domaine visible[41] (voir aussi la partie III.2 page 80).

Pour le "nettoyage" d'impulsions et l'obtention d'impulsions extrêmement courtes, les miroirs à membranes sont tout à fait indiqués[80] du moment que les impulsions sont très peu intenses. Le traitement de ces derniers est en effet relativement fragile. L'AOPDF est aussi couramment utilisé à cet effet en amont des étages d'amplification.

Il existe plusieurs difficultés qui limitent l'extension de ces différents façonneurs à l'ultraviolet et l'infrarouge. La première est la faible transmission des matériaux utilisés pour les masques (ou pour l'AOPDF) à ces longueurs d'onde. Des problèmes d'absorption se posent aussi et dégradent de façon dramatique le rendement du dispositif. De plus, pour l'ultraviolet, compte tenu des courtes longueurs d'onde, la planéité du masque peut devenir insuffisant pour conserver l'intégrité de l'impulsion à mettre en forme (déformation du front

d'onde) et l'absorption peut détériorer le masque.

Dans l'infrarouge, il est difficile de trouver des masques capables de générer des déphasages d'au moins 2π , condition nécessaire pour contrôler pleinement la phase et l'amplitude de l'impulsion à mettre en forme¹.

Différentes approches sont suivies pour pouvoir contourner ces difficultés. La première consiste à mettre en forme des impulsions visibles ou infrarouges à l'aide des dispositifs présentés plus hauts et de transférer la mise en forme par somme ou différence de fréquence[42, 43, 44]. Cette approche ne permet pas pour l'instant d'obtenir des énergies importantes dans les gammes spectrales visées à cause du faible rendement des processus de conversion de fréquence. De plus, cette efficacité de conversion est fortement couplée au type de mise en forme programmée, ce qui rend l'utilisation de ces systèmes assez difficile. La deuxième est d'adapter les dispositifs actuels pour les rendre directement utilisables dans ces nouvelles gammes spectrales. Récemment, deux méthodes directes de mise en forme d'impulsions à 400 nm ont été démontrées : l'une utilise un masque réflectif à micro-miroirs[45], l'autre un modulateur acousto-optique en silice[81]. Dans l'infrarouge, Nelson et Vaughan ont montré que la mise en forme en phase et amplitude peut être obtenue à l'aide d'écrans à cristaux liquides ne contrôlant que la phase[75]. Leur technique ne nécessite des déphasage que de π et non pas 2π ce qui permet de mettre en forme des impulsions de plus grande longueur d'onde (jusqu'à $3\text{ }\mu\text{m}$ au lieu de $1,5\text{ }\mu\text{m}$ pour les masques standards).

II.2 Choix de la ligne 4f

Nous allons présenter ici le type de ligne 4f que nous avons retenu pour la construction de notre façonneur. La géométrie très simple, présentée sur la figure II.1 page 16, n'est pas optimale pour des impulsions femtosecondes. En effet, comme nous l'avons dit précédemment, le spectre très large des impulsions femtosecondes est sensible à la dispersion non linéaire dans les lentilles ainsi qu'aux aberrations chromatiques. Il nous faut donc remplacer ces optiques réfractives par des optiques réflectives ce qui complique la géométrie puisque la propagation du faisceau ne se fait plus sur un axe unique. Plusieurs solutions sont envisageables et sont liées au choix des éléments réflectifs de la ligne. A première vue, les miroirs paraboliques semblent les meilleurs candidats pour permettre de sortir de l'axe sans introduire d'aberrations. Malheureusement, les miroirs paraboliques disponibles à l'heure actuelle ne proposent pas des états de surface nécessaires à la réalisation d'une ligne 4f. Les deux autres possibilités sont les miroirs sphériques et les miroirs cylindriques. Les miroirs sphériques sont peu onéreux et disponibles dans une grande gamme de distances focales. Cependant, lorsqu'ils sont utilisés hors d'axe, ils introduisent des aberrations dites d'excentrement. De plus, ils focalisent complètement les différentes composantes spectrales dans le plan de Fourier, ce qui entraîne des densités de puissance élevées. Or pour notre façonneur, qui utilise une barrette horizontale de cristaux liquides, il est inutile de focaliser dans la direction verticale. Nous allons donc utiliser des miroirs cylindriques en place des lentilles.

¹Du moins dans le cas très courant de ligne 4f avec des masques à une dimension (des barrettes). Pour des masque à deux dimensions (des écrans), cette condition semble pouvoir être contournée[75].

Il reste cependant à choisir l'agencement des différents éléments de la ligne 4f pour obtenir le meilleur résultat possible. De nombreuses géométries sont disponibles et ont été mises en oeuvre. Comme nous allons le voir, certaines entraînent des aberrations qui nuisent fortement à la qualité du faisceau mis en forme. Nous passerons en revue les géométries le plus souvent employées en mettant en avant leurs qualités et défauts. Nous présenterons ensuite celle que nous avons retenue ainsi qu'une évolution de cette dernière.

II.2.1 Géométries non retenues

Nous présentons ici rapidement différentes géométries couramment employées pour la fabrication de façonneurs. Nous nous attacherons à en donner les limitations.

II.2.1.1 Géométrie Ω

La géométrie en Ω [82, 83] est relativement simple comme on peut le voir sur la figure II.3. Elle est assez intuitive lorsque l'on part du schéma de ligne 4f à lentilles (voir la figure II.1 page 16). De plus, sa symétrie par rapport au plan de Fourier la rend facile à aligner. Par contre, les deux miroirs cylindriques introduisent d'importantes aberrations d'excentrement[84] et la divergence du faisceau de sortie est mal définie².

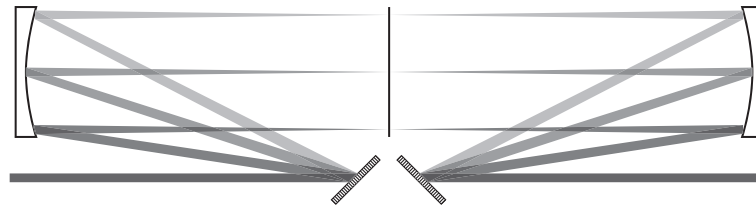


FIG. II.3 – Géométrie en Ω . Aberrations géométriques d'excentrement au niveau des miroirs, importantes pour les courtes longueurs d'onde (les plus déviées).

Ces aberrations apparaissent car le réseau n'est pas en face du miroir cylindrique³, comme cela est schématisé sur la figure II.4 page suivante. Les différents rayons qui en sont issus voient des rayons de courbure effectifs différents : les longueurs d'ondes les moins déviées verront une focale plus longue, les plus déviées une focale plus courte.

Cela a deux conséquences majeures :

1. La paire de miroirs cylindriques se comporte comme un télescope différent pour chaque longueur d'onde. On aura donc une géométrie et une divergence en sortie de la ligne différente pour chaque composante spectrale.
2. Les différentes composantes spectrales ne seront pas toutes focalisées dans le plan de Fourier de la ligne. Cela détériore la résolution de l'appareil et crée des couplages spatio-temporels.

²Chaque composante spectrale a une divergence différente : couplage spatio-spectral.

³Plus précisément car les rayons incidents proviennent d'un point n'appartenant pas au plan vertical de symétrie du miroir.

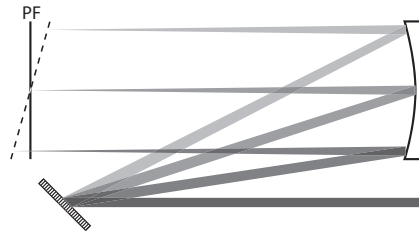


FIG. II.4 – Aberrations d'excentrement sur les miroirs cylindriques. Les différentes longueurs d'ondes voient des focales effectives différentes et ne sont pas focalisées dans le plan de Fourier (PF).

II.2.1.2 Géométrie S

La géométrie en S [49, 83, 38, 40] (voir figure II.5) est une évolution de la géométrie précédente tendant à diminuer l'effet des aberrations d'excentrement. L'idée est d'équilibrer les aberrations du premier et du second miroir : une longueur d'onde subissant de fortes aberrations sur le premier en subira de plus faibles sur le second et inversement. Dans ce cas, la géométrie du faisceau de sortie dépend de la longueur d'onde⁴.

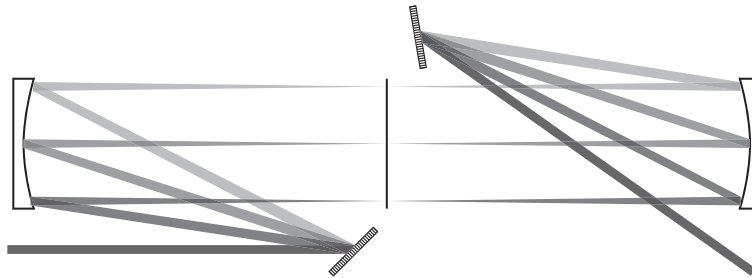


FIG. II.5 – Géométrie en S . Aberrations géométriques d'excentrement au niveau des miroirs, réparties entre les différentes longueurs d'onde.

II.2.1.3 Géométrie T

Pour éviter toute aberration d'excentrement, la géométrie en T [85] peut être choisie. Comme on le voit sur la figure II.6 page ci-contre, on introduit un miroir de repli à 45° entre le réseau et le miroir cylindrique afin d'avoir une incidence dans l'axe. Cette géométrie est symétrique, elle est donc facile à régler. Par contre, les miroirs de repli compliquent le dispositif et rajoutent des degrés de liberté à ajuster.

De plus, les miroirs cylindriques sont placés plus en hauteur que les autres éléments. Il est donc nécessaire de changer la hauteur du faisceau soit au niveau des réseaux, soit au niveau des miroirs de repli. Quelque soit l'option choisie, on introduira une inclinaison des taches focales dans le plan de Fourier.

⁴Chaque longueur d'onde voit un télescope cylindrique confocal avec un grandissement différent.

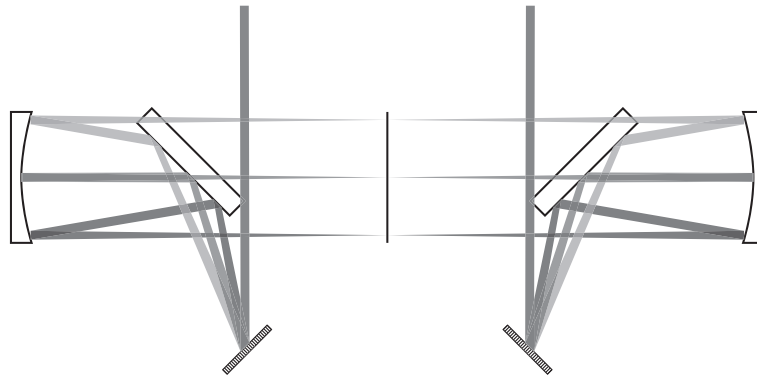


FIG. II.6 – Géométrie en T . Pas d'aberrations d'excentrement. Inclinaison des taches focales due au réseau ou au miroir de repli.

En effet, les différentes composantes ne font pas des trajets de même longueur après le miroir de repli : si on élève le faisceau à l'aide de ce miroir, les composantes spectrales arriveront sur le miroir cylindrique à des hauteurs différentes. De même, si on change la hauteur au niveau du réseau, les différentes longueurs d'ondes sont diffractées dans un plan non pas horizontal, mais oblique (voir la figure II.7).

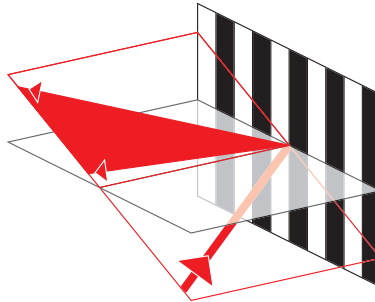


FIG. II.7 – Incidence non horizontale sur le réseau de diffraction. Le faisceau diffracté est oblique.

Dans les deux cas, cela change la façon dont les différentes longueurs d'onde sont dispersées et focalisées dans le plan de Fourier. Deux cas de figure se présentent, selon l'orientation choisie pour le miroir cylindrique qui collimate le faisceau diffracté et focalise les différentes composantes spectrales :

- On ne tient pas compte de l'obliquité du faisceau incident et on règle le miroir cylindrique de sorte que son axe de révolution soit vertical. Dans ce cas, seule la dispersion dans le plan horizontal est compensée. Les taches focales des différentes longueurs d'onde sont alors verticales et réparties selon un axe oblique dans le plan de Fourier, comme cela est présenté sur la figure II.8 page suivante. En sortie de la ligne, le faisceau a une forme et une divergence différente selon l'axe horizontal et l'axe vertical. De plus, la dérive de fréquence selon l'axe vertical introduite par le réseau n'est pas

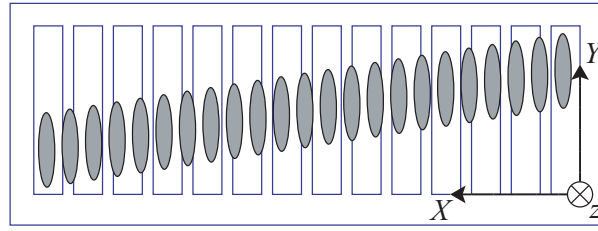


FIG. II.8 – Inclinaison des taches focales dans le plan de Fourier. Cas d’une inclinaison sur le réseau de diffraction non compensée par les miroirs cylindriques.

compensée : l’impulsion de sortie possède une dérive de fréquence spatiale, ce qui pose de nombreux problèmes.

- On tourne le miroir cylindrique afin que son axe soit orthogonal au plan de diffraction. Cela permet de bien recollimater le faisceau en compensant la dérive de fréquence spatiale selon les deux axes. Dans le plan de Fourier, les taches focales sont toujours étalées selon un axe oblique mais elles sont maintenant inclinées, comme cela est schématisé sur la figure II.9. Le problème majeur dans ce cas est qu’une même

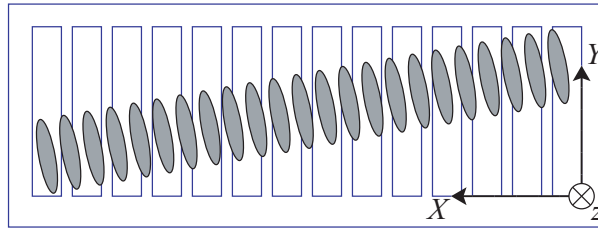


FIG. II.9 – Inclinaison des taches focales dans le plan de Fourier. Cas d’une inclinaison sur le réseau de diffraction compensée par les miroirs cylindriques.

longueur d’onde s’étale sur plusieurs pixels du fait de l’inclinaison de l’axe de dispersion. Cela tend à réduire la limite de résolution spectrale du façonneur. On peut toutefois remédier à ce problème en pivotant le masque dans le plan de Fourier pour faire coïncider les axes de dispersion des pixels et des longueurs d’onde. On a alors un faisceau de sortie exempt d’aberrations, mais la construction et l’alignement de la ligne sont plus complexes.

II.2.2 Géométrie retenue : géométrie X

Pour éviter les aberrations d’excentrement et l’inclinaison des tâche focales, la géométrie à utiliser doit vérifier les conditions suivantes⁵ :

⁵Ces conditions sont suffisantes mais sûrement pas nécessaires.

- Le réseau (ou son image par un éventuel miroir de repli) doit se trouver dans le plan vertical de symétrie du miroir cylindrique.
- L'incidence sur le réseau doit se faire dans le plan orthogonal aux traits du réseau : il ne faut pas qu'il y ait "d'angle vertical".
- Un miroir de repli placé entre le réseau et le miroir cylindrique ne doit pas changer à la fois l'inclinaison du faisceau par rapport à l'horizontale et sa direction dans le plan horizontal. En d'autre terme le faisceau peut être défléchi horizontalement ou verticalement, mais pas les deux à la fois.

La première condition évite les aberrations d'excentrement, les deux suivantes l'inclinaison des taches focales dans le plan de Fourier.

La géométrie que nous avons retenue[83, 35] pour notre façonneur est schématisée sur la figure II.10. Elle ne présente ni aberration d'excentrement, ni inclinaison des taches focales. Sa structure est très semblable à celle de la géométrie précédente, à ceci près que les miroirs de repli font face aux miroirs cylindriques. C'est cette différence qui permet de s'affranchir de l'inclinaison des taches focales.

La disposition particulière des faisceaux d'entrée et de sortie permet de les faire passer par un diaphragme commun, ce qui est extrêmement pratique pour l'alignement.

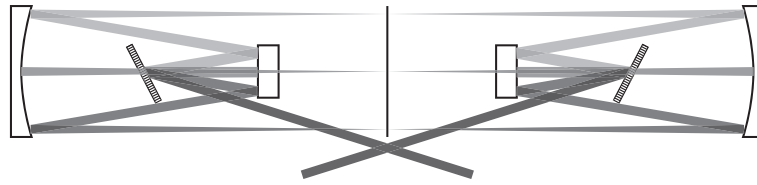


FIG. II.10 – Géométrie en X. Pas d'aberration d'excentrement ou d'inclinaison dans le plan de Fourier. Alignement facilité par le croisement des faisceaux d'entrée et de sortie.

II.2.3 Evolutions possibles

Toutes les géométries présentées ici peuvent être modifiées pour obtenir un système plus compact. Il est en effet possible de placer un miroir après le masque⁶ et de n'utiliser qu'une demi-ligne[35].

Cette configuration offre deux avantages principaux :

- une réduction importante de l'encombrement et du coût de la ligne (un seul réseau, un seul miroir cylindrique).
- une plus grande simplicité d'alignement. En particulier, il est difficile avec la ligne complète d'avoir des angles identiques sur les deux réseaux alors que cela est essentiel à la qualité du faisceau de sortie⁷. Avec la ligne repliée, cette difficulté disparaît car il n'y a qu'un seul réseau.

⁶Certaines compagnies proposent même des masques phase/amplitude dont une face est réfléchissante (traitement métallique ou diélectrique).

⁷Si les angles sont différents, la dispersion angulaire du premier réseau n'est pas parfaitement compensée par le second et le faisceau de sortie présente une dérive de fréquence spatiale ("spatial chirp" en anglais).

Elle nécessite par contre l'utilisation du masque en double passage, ce qui n'est pas toujours un avantage car cela divise sa dynamique par deux. En effet, la tension commandant le déphasage est digitalisée et la plus petit incrément de phase correspond à un digit de tension : en double passage il est deux fois plus grand qu'en simple passage. La figure II.11 compare la ligne à dispersion nulle complète à sa version repliée.

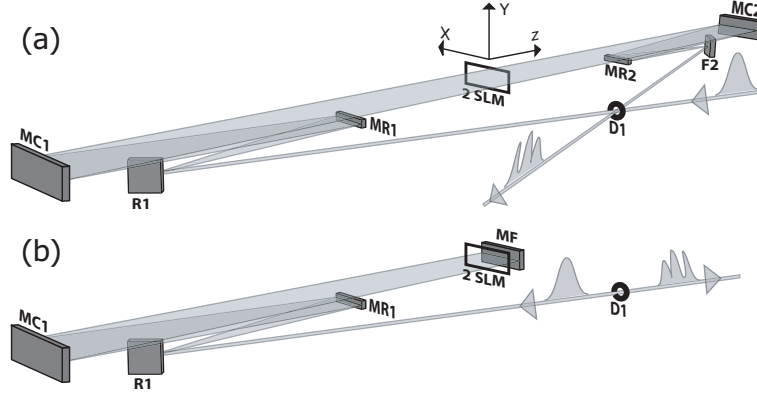


FIG. II.11 – Ligne à dispersion nulle utilisée au laboratoire (a) et sa version repliée (b). Dans la version repliée, un miroir MF est placé après le masque SLM. L'alignement est simplifié et le masque est utilisé en double passage.

II.3 Choix du masque

II.3.1 Masques possibles

A l'heure actuelle il existe un vaste éventail de masques possibles. Ils offrent des caractéristiques très différentes et tout dépend du type d'utilisation de l'on souhaite faire. Lors de la conception de notre façonneur nous avons un objectif simple : contrôler à la fois la phase et l'amplitude avec une haute résolution et une énergie par impulsion en sortie relativement élevée.

Cet objectif élimine d'entrée de jeu les masques bidimensionnels (certains miroirs déformables, valves optiques et écrans à cristaux liquides). En effet, la deuxième dimension ne nous sert pas et ces masques ne contrôlent que la phase. Ils peuvent par contre se montrer très utiles pour le contrôle spatio-temporel des impulsions[76]. Il semble aussi possible de contrôler la phase et l'amplitude, mais c'est une technique qui vient juste d'être démontrée et qui n'est pas encore utilisable de manière efficace[75].

Il reste à choisir parmi les masques uni-dimensionnels : barrettes à cristaux liquides (LC pour Liquid Crystal), modulateurs acousto-optiques transverses (AOM pour Acousto-Optic Modulators), les miroirs déformables à membrane (LMDM pour Linear Membrane Deformable Mirrors) et les miroirs déformables micro-usinés (MEMM pour MicroElectro-Mechanical Mirrors) développés très récemment. Le tableau II.1 page ci-contre compare les principales caractéristiques de ces différents masques. Seuls les cristaux liquides (LC) et les

modulateurs acousto-optiques (AOM) offrent une haute résolution (c'est à dire un grand nombre de paramètres de contrôle). Les modulateurs acousto-optiques ont par contre un rendement assez faible et nécessitent la construction d'une ligne 4f spécifique⁸. De plus, ils ne peuvent être utilisés efficacement qu'avec des lasers ayant des taux de répétition inférieurs au MHz. Ils ne peuvent donc pas être utilisés sur notre oscillateur (voir la partie I.1 page 6). Nous nous sommes donc tournés vers les masques à cristaux liquides qui offrent un meilleur rendement et n'imposent aucune contrainte sur le taux de répétition du laser mis en forme.

	LC	AOM	MEMM	LMDM
Paramètres de contrôle	128, 640	$\simeq 1000$	240	19
Pixellisé ?	oui	non	oui	non
Ouverture (mm ²)	13x5, 64x10	?	3x3	11x39
Rafraîchissement (Hz)	10	>1000	$\simeq 1000$	5000
Contrôle	amplitude phase	amplitude phase	phase	phase
Rendement à 800 nm	>70%	$\simeq 30\%$	$\simeq 80\%$	$\simeq 80\%$
Énergie en entrée (μJ)	>300	?	2	< 2
Gamme spectrale (nm)	430-1600	600-900 ?	200-900	550-1000, 266, 350
Principale limitation	pixelisation	faible efficacité	technologie très récente	résolution faible

TAB. II.1 – Comparatif des différents masques spatiaux unidimensionnels.

Lorsque nous avons conçu notre façonneur, il existait peu de masques à cristaux liquides utilisables pour réaliser un façonneur. La compagnie CRI proposait des masques contrôlant la phase et l'amplitude sur 128 pixels. Il s'agit en fait de deux barrettes de 128 pixels collées à un même substrat de verre. La compagnie Jenoptik proposait quant à elle des masques simples ne contrôlant que la phase mais sur 640 pixels. Dans notre recherche d'une haute résolution en phase et amplitude, nous avons opté pour une solution hybride : nous avons combiné deux masques Jenoptik pour pouvoir contrôler phase et amplitude sur

⁸Le modulateur acousto-optique dévie le faisceau et il faut donc décaler la deuxième moitié de la ligne 4f par rapport à l'axe de propagation libre. Cela rend la construction et l'alignement plus délicat qu'avec les masques à cristaux liquides et introduit des aberrations.

640 pixels[35]. Ces masques sont décrits en détails dans la référence[49]. Nous ne donnerons ici que les caractéristiques essentielles.

Suite à notre démonstration la compagnie Jenoptik a développé un masque double couche proposant permettant le contrôle de la phase et de l'amplitude sur 640 pixels. Depuis peu, la compagnie CRI propose un produit similaire. A notre connaissance, il n'existe pas d'expérience à l'heure actuelle démontrant les possibilités de ces nouveaux masques.

II.3.2 Double masque à cristaux liquides

Les deux masques que nous avons utilisés sont composés d'une fine couche ($9\ \mu\text{m}$) de cristaux liquides de type nématique prise entre deux substrats de verre. L'un des substrats porte 640 électrodes transparentes d'oxyde d'indium-étain (ITO pour Indium Tin Oxyde) permettant d'appliquer une tension différente à chacun des pixels. L'autre substrat porte une seule électrode d'ITO couvrant toute la surface et servant de masse. Les pixels font $10\ \text{mm}$ de haut par $97\ \mu\text{m}$ de large et sont séparés par des interstices (des "gaps") de $3\ \mu\text{m}$ (voir la figure II.12). La zone active est de $10\ \text{mm}$ par $64\ \text{mm}$.

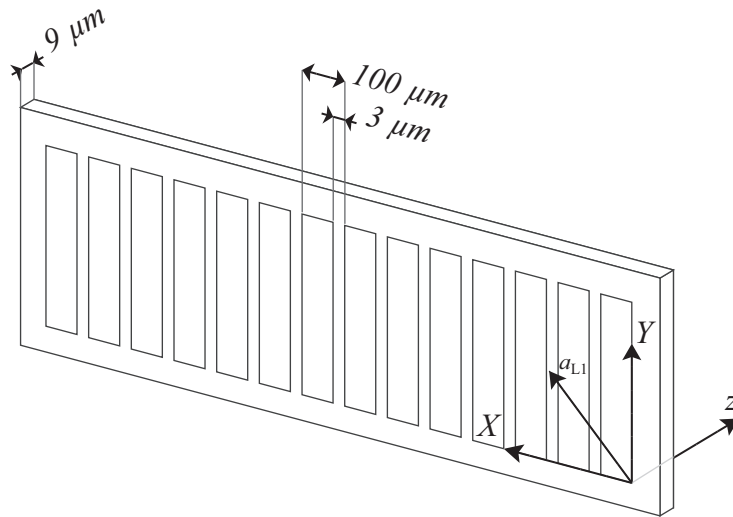


FIG. II.12 – Masque à cristaux liquides à 640 pixels utilisé dans le façonneur.

Les cristaux liquides nématiques utilisés ont une forme de bâtonnets qui sont orientés parallèlement aux substrats en absence de tension, comme cela est représenté dans le haut de la figure II.13 page suivante. Leur direction d'ancrage est imposée par brossage des électrodes (vecteur $\vec{a}_{L1}$). Les cristaux présentent une anisotropie diélectrique et ils pivotent lorsqu'on applique une tension entre l'électrode d'un pixel et la masse (voir le bas de la figure II.13). Cela modifie le chemin optique vu par une onde polarisée selon l'axe d'ancrage $\vec{a}_{L1}$. En combinant deux masques agissant sur des polarisations différentes et en utilisant un polariseur, on peut ainsi contrôler la phase et l'amplitude au niveau de chaque pixel du double masque.

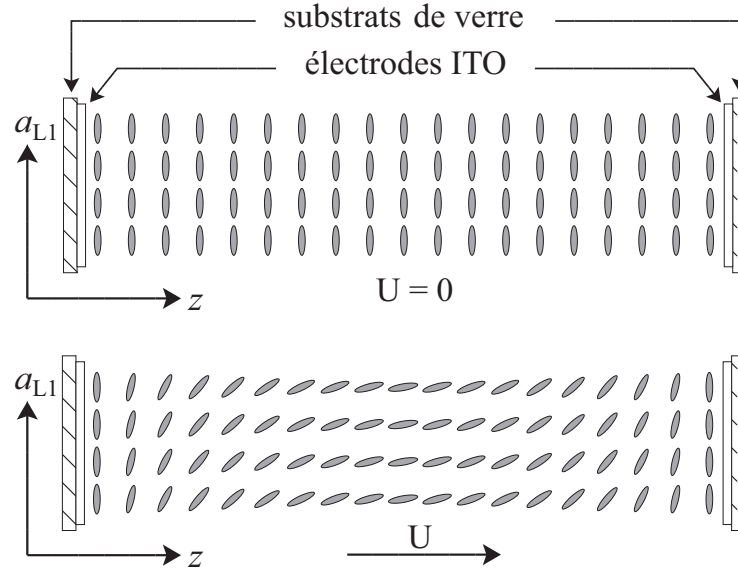


FIG. II.13 – Rotation des cristaux liquides par application d’une tension. Le chemin optique change.

II.4 Dimensionnement

Maintenant que le choix du masque est arrêté, il nous faut dimensionner les différents éléments de la ligne à dispersion nulle, à savoir les réseaux et les miroirs. C’est une étape essentielle pour obtenir un façonneur répondant à nos exigences. Pour y parvenir il faut tenir compte d’un grand nombre de paramètres :

- Les dimensions du masque ; en particulier la taille de la zone active, mais aussi celle des pixels et interstices.
- Les propriétés spectrales de l’impulsion à mettre en forme à savoir longueur d’onde centrale et largeur spectrale.
- Les dimensions spatiales et l’énergie de ces mêmes impulsions.
- Les objectifs que nous nous sommes fixés pour notre façonneur : une haute résolution spectrale.

Nous allons voir tout d’abord comment résumer le problème du dimensionnement dans le cas général à quelques équations simples et quelques paramètres pertinents. Nous appliquerons ensuite cela au cas particulier de notre façonneur.

II.4.1 Paramètres pertinents

Notre objectif est d’obtenir la meilleure résolution spectrale pour notre façonneur. Il nous faut donc étaler au maximum le spectre des impulsions à mettre en forme sur le masque à cristaux liquides, sans toutefois le tronquer. Dans un premier temps, nous devons donc déterminer la dispersion introduite par le réseau.

Nous appelons λ_0 et $\Delta\lambda_L$ la longueur d’onde centrale et la largeur spectrale à mi-

hauteur en intensité (FWHM) des impulsions. Ces impulsions arrivent sur le réseau avec un angle d'incidence θ_i . On note θ_{d0} l'angle de diffraction pour λ_0 et $\Delta\theta_d$ la séparation angulaire correspondant à $\Delta\lambda_L$. On suppose que les angles d'incidence et de diffraction sont proches de l'angle de Littrow à λ_0 . Le pas du réseau est noté d , la focale de la lentille, f et la largeur du masque L . La figure II.14 résume ces notations et présente un schéma simplifié de la situation (le miroir cylindrique y est remplacé par une lentille par commodité).

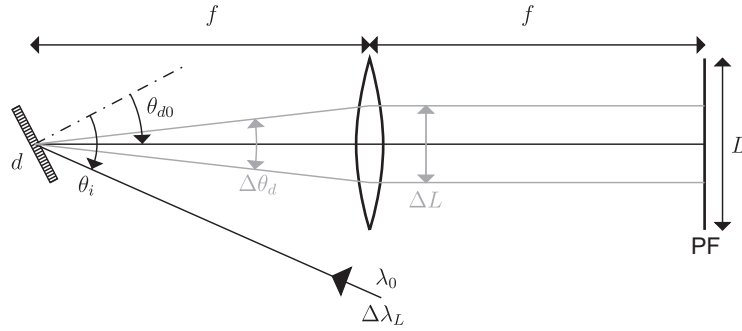


FIG. II.14 – Détermination de l'étalement spatial du spectre des impulsions.

Les propriétés de base des réseaux nous permettent d'écrire :

$$\frac{\Delta\theta_d}{\Delta\lambda_L} = \frac{1}{d \cos(\theta_{d0})} \quad (\text{II.1})$$

Comme le montre la figure II.14, au niveau de la lentille, on obtient un étalement spatial ΔL :

$$\Delta L = f \tan(\Delta\theta_d) = f \tan\left(\frac{\Delta\lambda_L}{d \cos(\theta_{d0})}\right) \quad (\text{II.2})$$

Le réseau étant dans le plan focal objet de la lentille, cette dernière rattrape la dispersion angulaire et on retrouve l'étalement ΔL au niveau du masque (dans le plan de Fourier PF).

Pour éviter de tronquer le spectre, il faut que ΔL soit de l'ordre d'un tiers de L . Cela revient à dire que le masque doit avoir une largeur spectrale $\Delta\lambda_M$ trois fois plus grande que celle de l'impulsion à mettre en forme. Ce facteur trois permet d'avoir une intensité spectrale aux bords du masque de moins de 0,2% de l'intensité maximale. La largeur spectrale du masque peut s'exprimer en fonction de la lentille et du réseau choisis :

$$\Delta\lambda_M = \arctan\left(\frac{L}{f}\right) d \cos(\theta_{d0}) \quad (\text{II.3})$$

On peut ainsi chercher les couples lentille-réseau permettant d'avoir $\Delta\lambda_M \simeq 3\Delta\lambda_L$. Il en existe de toute évidence une infinité mais parmi ceux-ci, seuls certains conviendront.

En effet, ces couples doivent répondre à une deuxième contrainte pour obtenir une grande résolution. Celle-ci concerne la taille des taches focales au niveau du masque. Chaque longueur d'onde de l'impulsion d'entrée donne une image au niveau du masque dont la taille doit être comparable à celle des pixels du masque. Effectivement, si ces images sont trop grandes par rapport aux pixels, on va perdre en résolution car une même longueur d'onde

va être étalée sur un grand nombre de pixels. A l’opposé, il faut éviter des images trop piquées car cela exalte l’effet des interstices. Pour expliciter cette contrainte, on suppose que l’impulsion a un profil spatial gaussien au niveau du réseau et on note Δx la largeur à mi-hauteur en intensité de ce profil. On peut alors montrer que les taches focales ont un profil gaussien de largeur à mi-hauteur ΔX donné par[86]⁹ :

$$\Delta X = \frac{2 \ln(2) \lambda_0 f \cos(\theta_i)}{\pi \Delta x \cos(\theta_{d0})} \quad (\text{II.4})$$

Il nous faut donc trouver un couple lentille-réseau (miroir-réseau en fait !) qui permet d’avoir à la fois le bon étalement $\Delta \lambda_L \simeq \Delta \lambda_M / 3$ et la bonne taille pour les taches focales $\Delta X \simeq l/3$ où l est la distance entre deux pixels.

II.4.2 Application particulière

Dans le cas particulier de notre façonneur, les impulsions ont une largeur spectrale $\Delta \lambda_L = 10$ nm et une longueur d’onde centrale $\lambda_0 = 795$ nm. Le diamètre du faisceau laser est de $\Delta x = 2,3$ mm. Le masque a une largeur $L = 64$ mm et une distance inter-pixel de $l = 100$ μm (pixel de 97 μm et interstice de 3 μm).

Il nous faut donc trouver un couple (f, d) pour lequel la bande spectrale correspondant au masque $\Delta \lambda_M$ sera de l’ordre de 30 nm et la largeur des taches focales ΔX de l’ordre de 30 μm . Les miroirs cylindriques que l’on peut utiliser ont des focales comprises entre 200 mm et 800 mm. Pour les focales plus courtes des aberrations sphériques apparaissent et la zone de Rayleigh devient très courte. Les grandes focales quant à elles sont difficilement réalisables avec des planéités suffisantes et augmentent considérablement l’encombrement de l’appareil.

De même, il n’existe pour des réseaux efficaces à 795 nm qu’une dizaine de densités de traits possibles comprises entre 300 et 2400 traits par millimètre.

La figure II.15 page suivante présente le résultat des simulations effectuées en tenant compte de ces contraintes. Nous avons choisi pour ces calculs un angle d’incidence θ_i aussi proche que possible de l’angle de Littrow θ_l , tout en permettant de séparer angulairement le faisceau incident et celui diffracté ($\theta_i = \theta_l + 10^\circ$). L’évolution de la bande spectrale du masque $\Delta \lambda_M$ (partie haute, en nm) et de la taille des taches focales ΔX (partie basse, en μm) y est représentée en fonction de la distance focale des miroirs et de la densité de traits des réseaux.

Comme on peut le constater, les conditions sur la bande spectrale au niveau du masque ne sont vérifiées que sur une zone assez étroite. La taille des taches focales apparaît moins critique. En recoupant ces résultats avec les miroirs et réseaux disponibles, nous avons finalement retenus des réseaux de 2000 traits par millimètre et des miroirs cylindriques de 600 mm de distance focale (figure II.15, traits noirs).

Le tableau II.4.2 page suivante récapitule les caractéristiques retenues pour les différents constituant de notre façonneur.

⁹Cette largeur est celle donnée par un simple calcul d’optique gaussienne auquel on doit appliquer un facteur d’agrandissement dû au réseau.

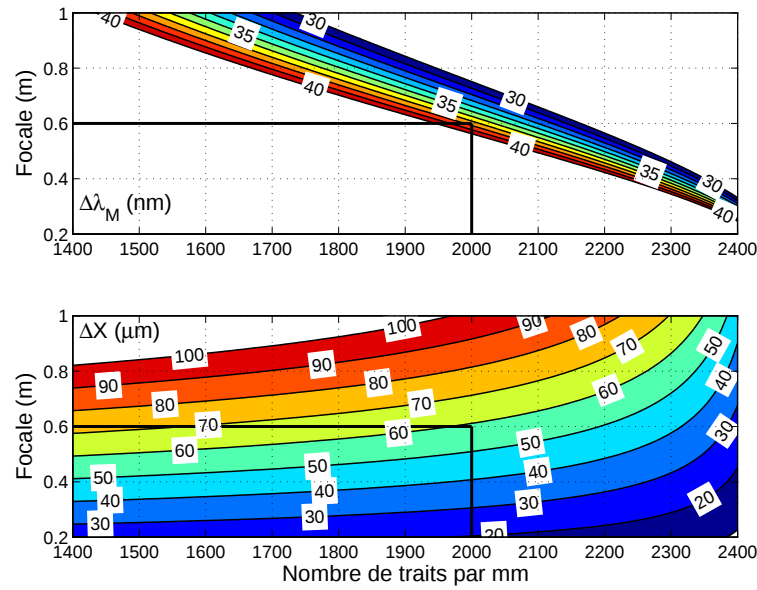


FIG. II.15 – Bande spectrale en nm (haut) et largeur des taches focales en μm (bas) au niveau du masque, en fonction de la focale et de la densité de traits. Les segments noirs indiquent le couple (miroir, réseau) finalement choisi.

Élément	Caractéristiques
Masques	2 masques de phase, 640 pixels, zone active de 64mm
Miroirs cylindriques	focale 600mm, traités argent
Réseaux	2000 traits/mm, traité or, haute énergie

TAB. II.2 – Caractéristiques des différents élément composant le façonneur.

II.5 Construction et alignement de la ligne

L'alignement de la ligne à dispersion nulle entourant le masque est une étape très critique. En effet, comme son nom l'indique, cette ligne ne doit pas introduire de dispersion ou de déformation sur les impulsions d'entrée alors même qu'elle contient des éléments extrêmement dispersifs que sont les réseaux. Cette ligne est en fait un compresseur à réseaux de compression nulle et elle réclame donc les mêmes attentions lors des réglages. En particulier, il est crucial d'avoir exactement les mêmes angles sur les réseaux d'entrée et de sortie faute de quoi les impulsions de sortie présenteront une forte dérive de fréquence spatiale ainsi qu'une phase spectrale cubique. De même, les quatre distances f entre les réseaux, les miroirs cylindriques et le plan de Fourier doivent être respectées. Tout écart introduit une phase quadratique importante et peut aussi entraîner des déformations spatiales du faisceau. Pour plus de détails sur les effets de désalignement de la ligne à dispersion nulle se reporter à la référence[87]. L'alignement de la ligne à dispersion nulle se décompose en

quatre étapes que nous décrirons ici. La première est l'alignement du télescope formé par la paire de miroirs cylindriques (MC1 et MC2, voir la figure II.16). La deuxième est le réglage de l'orientation et des angles d'incidence sur le réseau (R1). Vient ensuite l'ajout du réseau R2. Enfin, on insère les miroirs de repli (MR1 et MR2).

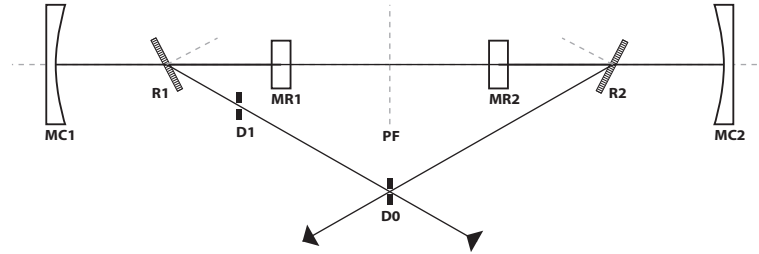


FIG. II.16 – Schéma de la ligne à dispersion nulle.

II.5.1 Télescope cylindrique

On place tout d'abord les deux miroirs cylindriques à une distance $2f$ l'un de l'autre. Les supports des deux miroirs doivent permettre un réglage fin de tous les degrés de rotation ainsi qu'une translation selon l'axe de la ligne à dispersion nulle. Il est important de repérer clairement l'axe choisi pour la propagation dans la ligne à dispersion nulle.

On fait ensuite passer un faisceau de grand diamètre¹⁰, parfaitement collimaté, dans le télescope formé par les deux miroirs (voir figure II.17). Il est très important d'avoir un faisceau bien collimaté et d'attaquer le télescope dans l'axe de la ligne pour éviter toute aberration d'excentrement. Si ces précautions sont prises, on peut utiliser les déformations du profil en sortie du télescope pour régler ce dernier.

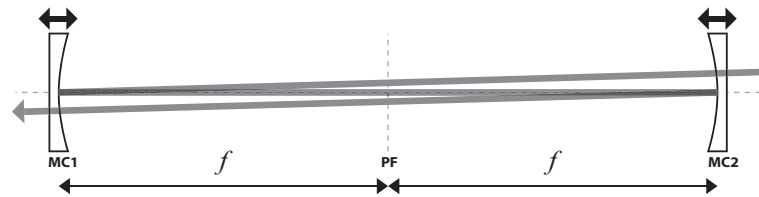


FIG. II.17 – Alignement des miroirs cylindriques à l'aide d'un faisceau de grand diamètre. Pour plus de clarté, on a introduit un angle entre le faisceau incident et l'axe de la ligne. En réalité, cet angle doit être nul.

En effet, le profil en sortie reste circulaire uniquement si la distance et l'orientation relative des deux miroirs sont parfaitement réglées. Il faut bien évidemment veiller en premier lieu à ce que le faisceau de sortie suive toujours l'axe de la ligne. On utilise ensuite, les règles suivantes, de manière itérative, pour corriger le profil du faisceau :

¹⁰Un laser HeNe, avec un diamètre grossi à 1 cm peut très bien faire l'affaire.

- Si le profil se déforme en fonction de la distance d’observation en prenant des formes d’ellipses inclinées, c’est que les axes de révolution des deux cylindres ne sont pas parallèles. Il faut alors tourner l’un des deux miroirs sur lui-même.
- Si le faisceau diverge dans le plan horizontal, c’est qu’il faut augmenter la distance entre les deux miroirs. S’il converge, il faut la diminuer.

Une fois que ces réglages sont faits, on ne doit normalement plus toucher qu’à un seul degré de liberté des miroirs cylindriques : l’inclinaison par rapport à la verticale (afin de régler la hauteur du faisceau réfléchi). Il est important de relever la position du plan de symétrie de ce télescope : il correspond au plan de Fourier (PF) où l’on placera le masque à cristaux liquides.

II.5.2 Réseau d’entrée

Pour pouvoir régler les réseaux, on doit disposer d’un faisceau centré à la longueur d’onde centrale λ_0 des impulsions à mettre en forme. Il est plus simple de travailler avec un faisceau monochromatique à cette longueur d’onde¹¹. Nous supposons que c’est le cas. Il est cependant possible de travailler directement avec les impulsions à mettre en forme. Dans ce cas, on peut réaliser les alignements avec l’extrémité d’un spectromètre à fibre (sans sphère intégrante) : cela permet de discriminer λ_0 au milieu de la tache de diffraction.

Tout comme pour les miroirs cylindriques, les supports des réseaux doivent permettre de régler finement tous les degrés de liberté, notamment la rotation du réseau autour de sa normale. Il faut aussi une rotation graduée et micrométrique pour la rotation autour de l’axe z (verticale).

Il faut effectuer deux réglages sur le réseau d’entrée. En premier, il faut s’assurer que les traits du réseau sont bien verticaux pour que la diffraction d’un faisceau horizontal n’entraîne pas un faisceau diffracté oblique (voir la partie II.2.1.3 page 20). Il faut ensuite régler l’incidence sur le réseau pour renvoyer la longueur d’onde centrale λ_0 dans l’axe de la ligne.

On place aussi précisément que possible le réseau R1 ainsi que deux diaphragmes D0 et D1 qui définissent le faisceau d’entrée dans la ligne¹² (voir figure II.18 page suivante). Le diaphragme D0 doit se trouver dans le plan de Fourier car il servira aussi pour l’alignement du faisceau de sortie.

On règle ensuite l’orientation des traits du réseau. Pour cela on commence par renvoyer l’ordre 0 en autocollimation. Puis on tourne le réseau afin de renvoyer maintenant l’ordre 1. On règle sa hauteur en tournant le réseau autour de sa normale. On répète cette procédure jusqu’à ce qu’elle converge.

On règle alors l’angle d’incidence sur le réseau en tournant ce dernier jusqu’à ce que la longueur d’onde centrale soit diffractée dans l’axe de la ligne. Il faut noter de quel angle on a tourné car il faudra utiliser le même angle pour R2.

¹¹A 800nm, on peut utiliser par exemple un oscillateur en mode continu.

¹²L’angle entre le faisceau entrant et l’axe de la ligne doit correspondre à la différence entre l’angle d’incidence choisi et l’angle diffracté correspondant.

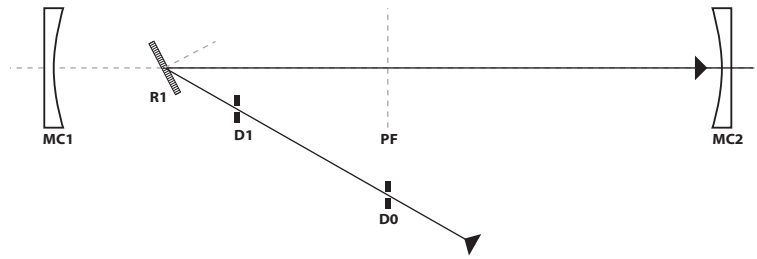


FIG. II.18 – Alignement du premier réseau. La longueur d'onde centrale doit être diffractée dans l'axe de la ligne.

II.5.3 Réseau de sortie

On règle maintenant le deuxième réseau (R2). On le place aussi précisément que possible dans la ligne (voir figure II.19) et on utilise le faisceau diffracté par le premier réseau pour l'alignement. La procédure est tout à fait similaire à celle utilisée pour le premier réseau. Il est primordial que les angles pour R1 et R2 soient les mêmes pour éviter d'introduire une phase spectrale cubique et une dérive de fréquence spatiale. Si tout se passe bien, on doit alors envoyer le faisceau de sortie dans le diaphragme D0.

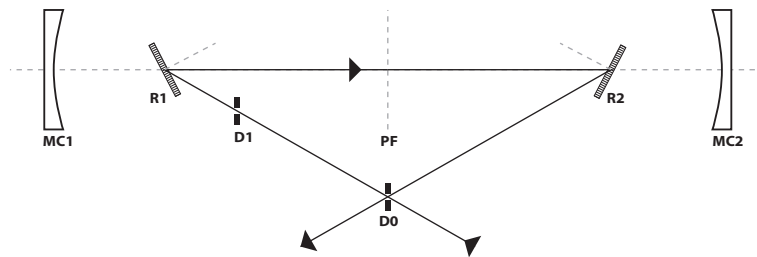


FIG. II.19 – Alignement du deuxième réseau. On utilise le faisceau diffracté par le premier réseau.

II.5.4 Miroirs de repli

Pour finir la construction de la ligne à dispersion nulle, il ne reste plus qu'à insérer les deux miroirs de repli MR1 et MR2 (voir la figure II.20 page suivante). Les supports de ces derniers doivent permettre une translation fine dans l'axe de la ligne. On les place de façon à avoir une distance f entre le réseau et le miroir cylindrique. On ajuste alors l'inclinaison des miroirs de repli et des miroirs cylindriques afin de faire sortir le faisceau de la ligne par le diaphragme D0.

Pour ajuster finement la position des miroirs de repli, on fait passer par la ligne des impulsions parfaitement compressées et on mesure la durée des impulsions en sortie. En déplaçant les miroirs de repli, on doit pouvoir obtenir des impulsions presque aussi courtes que celle d'entrée. Il faut veiller à déplacer les deux miroirs de façon symétrique.

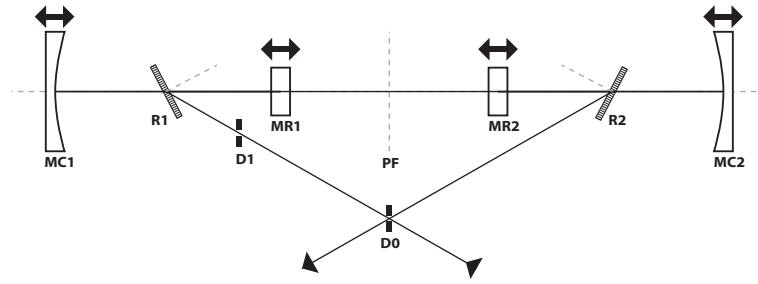


FIG. II.20 – Ligne complète après insertion des miroirs de repli.

II.6 Mise en place du masque et des polariseurs

La ligne étant parfaitement alignée, il faut maintenant placer le masque au plan de Fourier. Cela ne pose normalement pas de difficulté particulière. Le masque à cristaux liquides tend à allonger légèrement le chemin optique entre les miroirs cylindriques. On doit ainsi observer une légère augmentation de la durée des impulsions en sortie de la ligne. On compense celle-ci en ré-ajustant les miroirs MR1 et MR2.

Dans le cas de notre façonneur, l'association de deux masques de phase réclame en plus une étape assez délicate d'alignement d'un masque avec l'autre.

Il faut tout d'abord s'assurer que les surfaces des deux masques sont parfaitement parallèles, ce qui peut être fait par autocollimation. Il faut ensuite que les pixels des deux masques soient bien verticaux : c'est à dire que l'on doit contrôler la rotation des deux masques autour de l'axe de propagation (axe z sur la figure II.21 page 36). Pour ce faire, on envoie sur les deux masques un faisceau laser HeNe : chaque masque produit une figure de diffraction qui est horizontale pour des pixels verticaux. En superposant les deux figures de diffraction, on peut aligner les pixels avec une précision de l'ordre de 2 mrad. La dernière étape est de faire en sorte que les pixels des deux masques tombent bien en face les uns des autres : le décalage transverse (translation selon l'axe X sur la figure II.21) entre les deux masques doit être nul, avec une précision de l'ordre de la largeur d'un interstice ($3 \mu\text{m}$). Pour vérifier cela, on programme le premier masque avec une phase de π sur les 320 premiers pixels et de 0 sur les suivants. Le deuxième masque a une programmation inversée : 0 sur les 320 premiers pixels et π ensuite. Si le déplacement latéral est nul, cette programmation donne une transmission nulle (comme cela sera démontré plus loin). Sinon, on observe une bande lumineuse dont la largeur est proportionnelle au décalage transverse. En translatant un des deux masques pour annuler cette bande, on aligne les pixels avec la précision nécessaire. Ceci n'est qu'un résumé succinct de cet alignement, la référence[35] contient plus de détails.

Les nouveaux masques à cristaux liquides contrôlant phase et amplitude avec un grand nombre de pixels évitent ces ennuis.

Pour finir la construction du façonneur, on dispose un polariseur sur le faisceau d'entrée et un sur le faisceau de sortie, tous deux étant horizontaux. Ils sont tout à fait essentiels au bon fonctionnement de l'appareil, tous les calculs de la partie II.7.1.2 page 36 reposant sur des polarisations bien définies en entrée et en sortie de la ligne.

Pour pouvoir utiliser correctement le façonneur, il reste deux calibrations à faire :

- Tout d'abord, il faut connaître la tension à appliquer à un pixel donné pour y introduire un déphasage donné. Il faut donc déterminer la phase ϕ introduite par un pixel en fonction de la tension U qui lui est appliquée et de la longueur d'onde incidente λ .
- Ensuite, pour utiliser le masque à cristaux liquides comme un masque spectral, il faut déterminer la dispersion dans le plan de Fourier. Cela permet de savoir quel pixel agit sur quelle composante spectrale.

Ces deux calibrations, souvent mal faites, sont essentielles à un fonctionnement correct du façonneur. Elles sont décrites en détail dans les deux parties suivantes.

II.7 Calibration du déphasage

Pour réaliser des mises en formes données à l'aide des deux masques de phase placés dans la ligne 4f, nous appliquons des tensions données à chacun des 1280 pixels (2x640). Nous n'avons pour l'instant pas explicité le lien entre la tension que l'on applique au niveau d'un pixel et le déphasage introduit par celui-ci. Nous allons le détailler dans une première sous-partie et nous verrons que ce déphasage est fonction de la longueur d'onde incidente. Ensuite, nous verrons comment simplifier le problème de la calibration. Puis nous présenterons deux techniques de calibration permettant de mesurer expérimentalement le déphasage introduit pour une tension donnée.

II.7.1 Fonctionnement du masque

II.7.1.1 Généralités et notations

On note pour toute la suite $\vec{z}$ la direction de propagation de la lumière, qui est normale aux faces des barrettes de cristaux liquides, $\vec{X}$ la direction selon laquelle se répartissent les pixels et $\vec{Y}$ la direction d'orientation de chaque pixel (voir la Figure II.21 page suivante). Les directions d'ancrage (voir la partie II.3.2 page 26) des cristaux liquides sont notées $\vec{a}_{L1}$ pour la première barrette et $\vec{a}_{L2}$ pour la deuxième. Sur chaque barrette, les pixels se comportent comme des lames d'onde commandées en tension ayant toutes les mêmes lignes neutres. Pour une longueur d'onde λ_0 , le déphasage ϕ introduit par la lame vaut

$$\phi(\lambda_0, U) = \frac{2\pi\Delta n(\lambda_0, U)e_{CL}}{\lambda_0} \quad (\text{II.5})$$

où $\Delta n(\lambda_0, U)$ est la différence entre les indices de l'axe lent et de l'axe rapide, U la tension appliquée, et e_{CL} l'épaisseur du cristal liquide.

L'axe lent de chaque barrette correspond à la direction d'ancrage des cristaux liquides soit $\vec{a}_{L1}$ pour la première et $\vec{a}_{L2}$ pour la deuxième. On note $\vec{a}_{R1}$ et $\vec{a}_{R2}$ les axes rapides des deux barrettes (voir la figure II.22 page suivante). La conception des barrettes est telle que l'axe lent de la première est incliné de $+45^\circ$ et celui de la deuxième de -45° par rapport à l'axe $\vec{X}$. On a donc :

$$\vec{a}_{L1} = \frac{1}{\sqrt{2}}(\vec{X} + \vec{Y}) \quad \text{et} \quad \vec{a}_{L2} = \frac{1}{\sqrt{2}}(\vec{X} - \vec{Y}) \quad (\text{II.6})$$

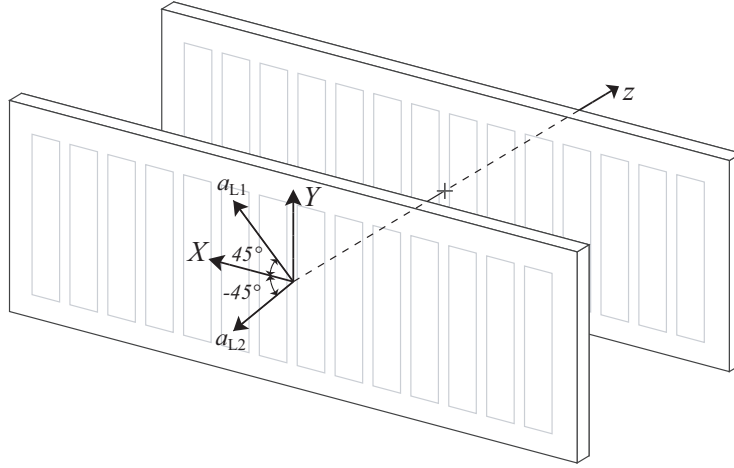


FIG. II.21 – Deux barrettes de cristaux liquides en phase constituant le masque de notre façonneur.

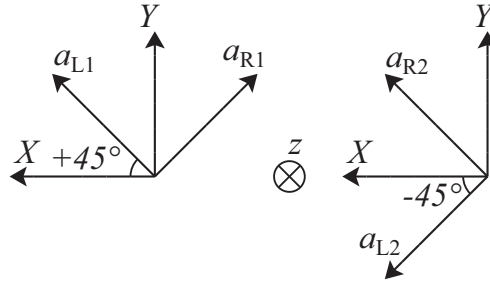


FIG. II.22 – Axes rapides et lents des deux barrettes de cristaux liquides

II.7.1.2 Effet des deux barrettes

On considère une onde incidente $\vec{E}_0$ polarisée selon $\vec{X}$ et se propageant selon $\vec{z}$ qui va traverser les deux barrettes de cristaux liquides. On note $\vec{E}_1$ l'onde sortant de la deuxième barrette. On s'intéresse ici à l'effet d'un seul pixel : on ne tient pas compte de l'extension de l'onde selon $\vec{X}$. On note aussi $\phi_1(U_1)$ (respectivement $\phi_2(U_2)$) le déphasage introduit par la première barrette (respectivement par la deuxième). U_1 (respectivement U_2) est la tension appliquée au pixel de la première barrette (respectivement de la deuxième). La dépendance avec la longueur d'onde est implicite.

L'onde incidente s'écrit :

$$\vec{E}_0 = E_0 \vec{X} = \frac{E_0}{\sqrt{2}} (\vec{a}_{L1} - \vec{a}_{R1}) = \frac{E_0}{\sqrt{2}} (\vec{a}_{R2} + \vec{a}_{L2}) = \frac{E_0}{\sqrt{2}} (\vec{a}_{L1} + \vec{a}_{L2}) \quad (\text{II.7})$$

Compte tenu des déphasages introduits, on a en sortie

$$\vec{E}_1 = \frac{E_0}{\sqrt{2}} (\vec{a}_{L1} e^{i\phi_1(U_1)} + \vec{a}_{L2} e^{i\phi_2(U_2)}) \quad (\text{II.8})$$

qui peut encore se réécrire :

$$\vec{E}_1 = \frac{E_0}{2} \left[\vec{X} (e^{i\phi_1(U_1)} + e^{i\phi_2(U_2)}) + \vec{Y} (e^{i\phi_1(U_1)} - e^{i\phi_2(U_2)}) \right] \quad (\text{II.9})$$

Finalement, on obtient

$$\vec{E}_1 = E_0 \exp \left(i \frac{\phi_1 + \phi_2}{2} \right) \left[\vec{X} \cos \left(\frac{\phi_1 - \phi_2}{2} \right) + i \vec{Y} \sin \left(\frac{\phi_1 - \phi_2}{2} \right) \right] \quad (\text{II.10})$$

la dépendance en tension des déphasages étant implicite pour alléger l'expression. Si l'on s'en tient là, on réalise une mise en forme en phase et en polarisation de l'onde incidente $\vec{E}_0$.

En polarisant $\vec{E}_1$ selon l'axe $\vec{X}$, on obtient en sortie une mise en forme en phase et en amplitude :

$$\vec{E}_1 \cdot \vec{X} = \exp \left(i \frac{\phi_1 + \phi_2}{2} \right) \cos \left(\frac{\phi_1 - \phi_2}{2} \right) E_0 = B(U_1, U_2) E_0 \quad (\text{II.11})$$

où B est la transmission complexe associée aux tensions appliquées.

En conclusion, pour réaliser une mise en forme donnée au niveau d'un pixel, avec une configuration de polariseur donnée, il existe un couple unique¹³ de déphasages (ϕ_1, ϕ_2) . Le déphasage étant une fonction monotone de la tension appliquée, cela revient à dire qu'il existe un couple unique de tensions (U_1, U_2) .

II.7.1.3 Masque introduit par les barrettes

On tient compte à présent de l'extension des barrettes et du champ incident selon $\vec{X}$. Elles sont orientées de telle sorte que le pixel numéro n de la première soit en face du pixel numéro n de la seconde. Dans toute la suite, on nommera pixel l'association des pixels des deux barrettes. On néglige l'interstice qui sépare chaque pixel des pixels voisins¹⁴. On suppose qu'un polariseur orienté selon $\vec{X}$ est placé en sortie. On note $B^{(n)} = B^{(n)}(U_1^{(n)}, U_2^{(n)})$ la transmission complexe au niveau du pixel numéro n . Dans la suite, la dépendance en tension sera implicite. Le filtre spatial équivalent aux deux barrettes s'écrit :

$$M_X(X) = \sum_{n=-N/2}^{N/2-1} B^{(n)} \delta(X - nl) \otimes \text{rect}(X/l) \quad (\text{II.12})$$

où N est le nombre de pixels (le pixel 0 étant au centre), X est la position transverse dans le plan de Fourier, l est la distance entre deux pixels voisins ($100 \mu m$), $\delta(u)$ représente la fonction de Dirac et la fonction carré $\text{rect}(u)$ est définie par :

$$\text{rect}(u) = \begin{cases} 1 & \text{si } |u| \leq 1/2, \\ 0 & \text{si } |u| > 1/2. \end{cases} \quad (\text{II.13})$$

¹³Il est toujours possible d'inverser les rôles joués par les deux barrettes ou d'ajouter un multiple de 2π à ϕ_1 ou ϕ_2 .

¹⁴L'effet principal de ces interstices est de créer un réseau périodique. Cela entraîne une diffraction de l'onde incidente et se traduit par l'apparition de répliques temporelles, comme cela est expliqué dans la partie II.9.2 page 47.

II.7.2 Simplifier pour mieux calibrer

Nous souhaitons montrer ici comment nous pouvons simplifier quelque peu la calibration du déphasage. En effet, il faudrait en toute rigueur mesurer le déphasage introduit par chaque pixel de chaque barrette (soit 1280 pixels), pour toutes les tensions possibles (il y en a 4096) et pour les différentes longueurs d'onde présentes dans le spectre de l'impulsion à mettre en forme. Cela reviendrait à faire plus de 5 millions de mesures pour chaque longueur d'onde ce qui prendrait beaucoup trop de temps.

Nous pouvons tout d'abord calibrer en une seule fois le déphasage pour tous les pixels d'une même barrette. En effet, il a été vérifié expérimentalement que les pixels sont relativement homogènes et agissent de la même façon pour des tensions suffisamment élevées. Il nous suffit donc de travailler avec les tensions élevées pour pouvoir calibrer chaque barrette en une étape. Par contre, il existe des différences significatives d'une barrette à l'autre et il nous faudra les calibrer indépendamment.

Pour ce qui est de la dépendance en longueur d'onde, nous allons voir que nous pouvons calibrer à une longueur d'onde donnée et extrapoler ce résultat à d'autres longueurs d'onde.

Nous considérons un seul masque de phase. Cela revient soit à retirer le second (ce qui est faisable dans notre cas), soit à l'éteindre. Nous appliquons la même tension U à tous les pixels du masque de phase.

Nous rappelons ici l'expression du déphasage $\phi(\lambda, U)$ introduit (donnée précédemment dans l'équation II.5 page 35) :

$$\phi(\lambda, U) = \frac{2\pi\Delta n(\lambda, U)e_{CL}}{\lambda} \quad (\text{II.14})$$

La différence d'indice $\Delta n(\lambda, U)$ peut s'écrire comme

$$\Delta n(\lambda, U) = \Delta n(\lambda, 0).g(U) = \Delta n(\lambda).g(U) \quad (\text{II.15})$$

$\Delta n(\lambda)$ étant la biréfringence du cristal liquide, fonction connue¹⁵, et $g(U)$ étant à déterminer.

On peut alors écrire :

$$\phi(U, \lambda) = \phi(U, \lambda_{cal}) \cdot \frac{\lambda_{cal}}{\lambda} \cdot \frac{\Delta n(\lambda)}{\Delta n(\lambda_{cal})} \quad (\text{II.16})$$

$\Delta n(\lambda)$ étant connu, il suffit de mesurer le déphasage $\phi(U, \lambda_{cal})$ en fonction de la tension U pour une longueur d'onde λ_{cal} fixée pour connaître $\phi(U, \lambda)$ pour toute longueur d'onde et toute tension.

II.7.3 Mise en oeuvre expérimentale

Pour mesurer l'évolution du déphasage avec la tension, on place le masque de phase entre deux polariseurs horizontaux. On éclaire le tout avec un laser HeNe ($\lambda_{cal} = 632,8$ nm) de fort diamètre, comme cela est schématisé sur la figure II.23 page suivante. On éclaire

¹⁵Elle fait partie des données techniques fournies avec les barrettes de cristaux liquides.

ainsi un grand nombre de pixels et on applique une même tension U à tous les pixels de la barrette. L'intensité transmise est mesurée à l'aide d'un électrophotomultiplicateur.

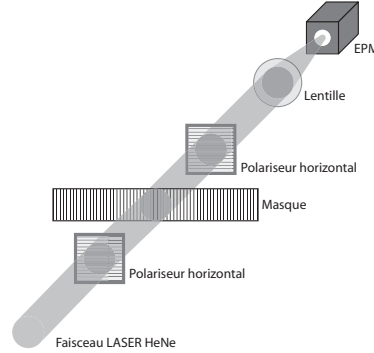


FIG. II.23 – Schéma de principe de la calibration du déphasage en fonction de la tension appliquée.

On obtient ainsi un signal $S(\lambda_{cal}, U)$ qui dépend de la tension appliquée à la barrette. En simplifiant l'équation II.11 page 37 au cas d'un seul masque de phase, on trouve :

$$S(\lambda_{cal}, U) \propto \left[\cos \left(\frac{\phi(\lambda_{cal}, U)}{2} \right) \right]^2 \quad (\text{II.17})$$

Les signaux obtenus lors de la calibration, normalisés à 1, sont présentés sur la figure II.24. La tension y est représentée en digit : la tension appliquée aux pixels varie entre 0 et 8 volts et est codée sur 12 bits, c'est à dire 4096 niveaux que nous appelons digits. Le trait en pointillés indique la tension minimale que nous utiliserons dans la pratique (environ 130 digits) pour des raisons que nous détaillerons plus loin.

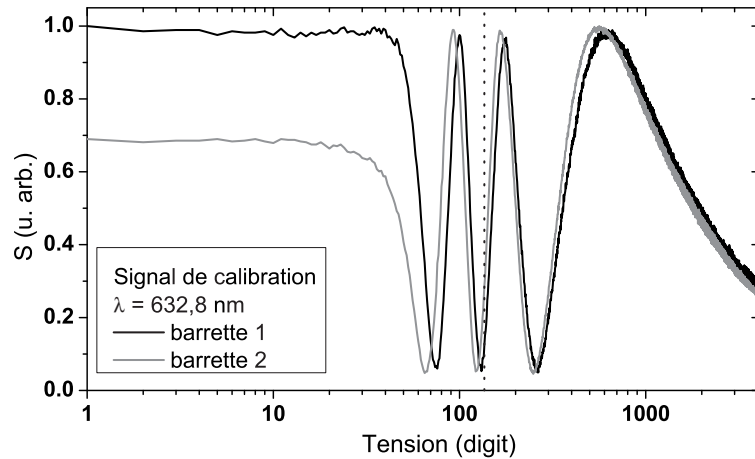


FIG. II.24 – Transmission des barrettes de cristaux liquides à 632,8 nm en fonction de la tension

A l'aide de l'équation II.17, on peut remonter au déphasage introduit $\phi(U)$ pour une tension U donnée. La figure II.25 page suivante présente le déphasage ainsi obtenu pour les

deux barrettes. Il s'agit en fait de courbes d'ajustement qui reproduisent très fidèlement l'évolution du déphasage en fonction de la tension en s'affranchissant du bruit de mesure. Les ajustements sont de la forme

$$\phi(\lambda_{cal}, U) = e^{P(U)} + C \quad (\text{II.18})$$

où $P(x)$ est un polynôme de degrés 9 et C une constante.

Nous n'avons reproduit le déphasage que pour les tensions élevées (supérieures à 130 digits environ). En effet, il n'est pas conseillé d'utiliser des tensions trop faibles car le déphasage introduit pour une tension donnée varie alors beaucoup d'un pixel à l'autre à cause des défauts de surface. De plus, compte tenu de la relation non linéaire entre la tension et la phase, la sensibilité en phase est très réduite pour des tensions basses (la pente de $\phi(\lambda_{cal}, U)$ est très forte pour U faible). On se limite donc à un intervalle de tensions élevées en faisant en sorte de pouvoir contrôler le déphasage sur une gamme de 2π pour toutes les longueurs d'onde du spectre des impulsions à mettre en forme.

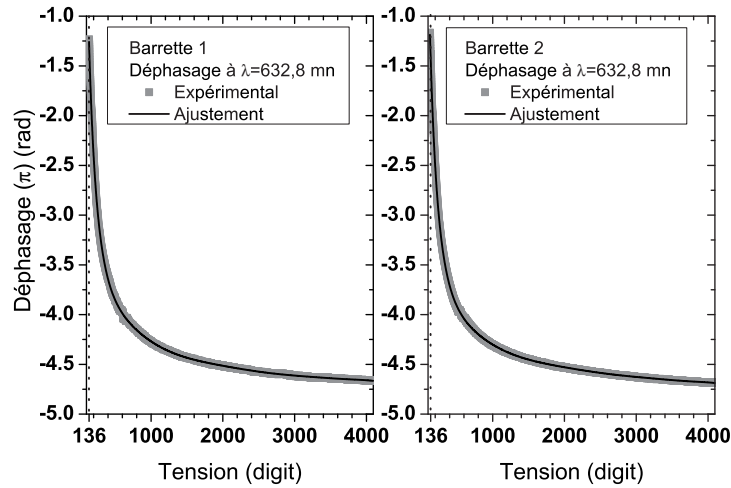


FIG. II.25 – Déphasages introduits par les barrettes de cristaux liquides pour une onde incidente à 632,8 nm en fonction de la tension appliquée.

II.7.4 Méthode mieux adaptée

La méthode de calibration présentée ci-dessus présente quelques défauts et peut être améliorée.

Tout d'abord, elle doit être faite avant d'introduire le masque dans la ligne à dispersion nulle. Pour recalibrer le masque avec cette méthode, il faudrait le sortir de la ligne.

Ensuite, elle a été faite à une longueur d'onde unique et la dépendance en longueur d'onde est déduite d'une formule d'ajustement fournie avec le masque et non vérifiée.

Enfin et surtout, le déphasage obtenu à partir des données expérimentales (figure II.25) par inversion de l'équation II.17 page précédente présente des défauts. En effet, les différents

maxima et minima ne sont pas tous au même niveau¹⁶ et cela introduit des déformations importantes dans le déphasage à proximité des extrema.

Nous avons développé durant la thèse une autre technique de calibration, assez proche de la précédente. Nous l'avons mise en oeuvre lors de la construction d'un façonneur, dans l'équipe d'Albert Stolow au Steacie Institute For Molecular Sciences. Elle a servi à la calibration d'un masque de phase seule à 128 pixels de chez CRI. Elle s'affranchit des défauts énumérés ci-dessus. L'idée est de calibrer le masque à l'intérieur même de la ligne à dispersion nulle, comme cela est schématisé sur la figure II.26. On place avant le masque (M), une fente (F) et après, une lame d'onde (Lo), dont les lignes neutres sont alignées avec celles du masque. En translatant la fente, on choisit la longueur d'onde de calibration parmi celles présentes dans le spectre de l'impulsion. La lame d'onde introduit quant à elle un déphasage supplémentaire qui dépend de son inclinaison par rapport à l'axe de propagation de la lumière. En jouant sur l'inclinaison de la lame d'onde, on peut déplacer la position (en tension) des extrema.

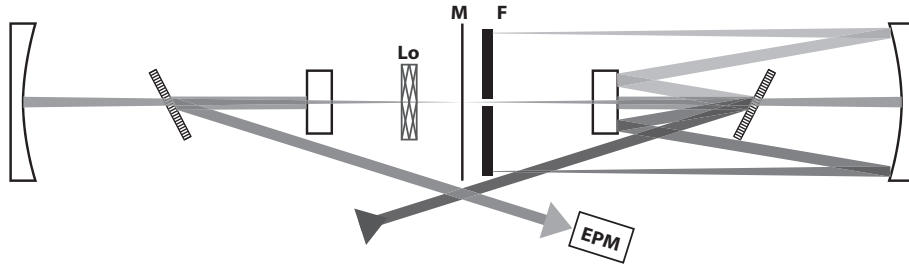


FIG. II.26 – Autre méthode de calibration du déphasage. Une fente F et une lame d'onde Lo sont introduite dans la ligne 4f. Cette technique a été mise en oeuvre à Ottawa pour calibrer un masque CRI de 128 pixels.

On obtient ainsi, pour les différentes longueurs d'onde λ , un ensemble de signaux de transmission

$$S(\lambda, U, \phi_{Lo}) \propto \left[\cos \left(\frac{\phi(\lambda, U) + \phi_{Lo}}{2} \right) \right]^2 \quad (\text{II.19})$$

où ϕ_{Lo} est le déphasage ajouté par la lame d'onde. La figure II.27 page suivante montre une série de signaux de transmission pour différentes inclinaisons de la lame d'onde pour une longueur d'onde de 795 nm. Comme on le voit, les extrema et donc les zones déformées se déplacent en fonction de l'inclinaison de la lame d'onde. En traitant ces courbes de transmission comme dans la précédente méthode (c'est à dire en inversant l'équation II.19), on obtient plusieurs reconstruction du déphasage introduit par le masque. Chaque reconstruction est décalée d'un certain ϕ_{Lo} (qui n'a pas besoin d'être connu) et présente des déformations autour de tensions différentes. En combinant toutes ces reconstructions (par exemple en les moyennant), on obtient une reconstruction du déphasage très précise et exempte des imperfections que présente la méthode précédente.

¹⁶Cela peut être dû aux interstices entre les pixels. Ils laissent passer une partie de l'onde sans introduire de déphasage et l'intensité totale résulte de l'interférence entre ces fuites et l'onde déphasée par les pixels.

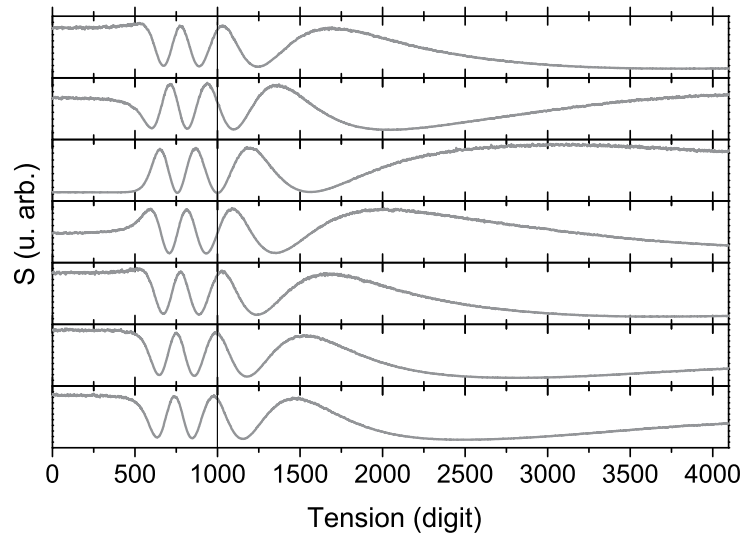


FIG. II.27 – Transmission de la barrette (CRI 128 pixels) de cristaux liquides à 795 nm en fonction de la tension, pour différentes orientation de la lame d'onde.

La figure II.28 montre le déphasage reconstruit sur l'ensemble de la gamme de tension. Il ne s'agit pas ici d'une courbe d'ajustement, la phase reconstruite étant suffisamment peu déformée pour être utilisée telle quelle. Comme pour le façonneur installé à Toulouse, nous n'avons utilisé au final que les tensions les plus élevées.

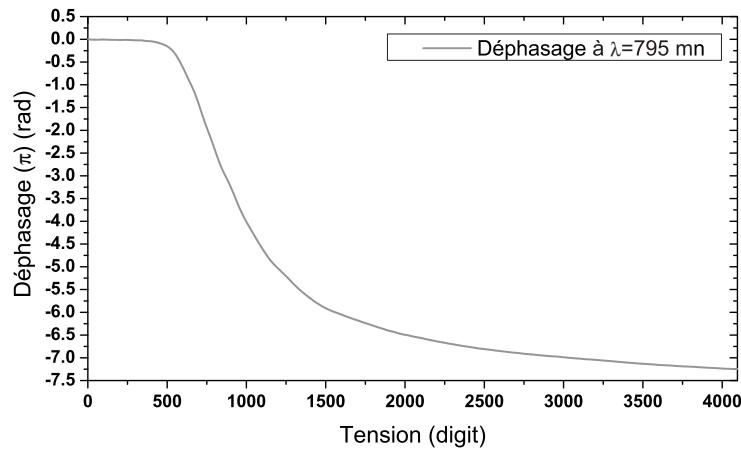


FIG. II.28 – Déphasage introduit par le masque (CRI 128 pixels) pour une onde incidente à 795 nm en fonction de la tension appliquée.

II.8 Calibration de la dispersion

Nous venons de déterminer la calibration du déphasage en fonction de la tension appliquée à un pixel et de la longueur d'onde de la lumière incidente. Forts de l'équation II.11 page 37, nous savons donc quel couple de tensions appliquer au niveau d'un pixel du masque pour contrôler la phase et l'amplitude d'une onde de longueur d'onde λ . Pour pouvoir utiliser notre masque comme un masque spectral, il nous faut aussi savoir la relation de dispersion au plan de Fourier, c'est à dire le lien entre la position transverse X et la pulsation ω (voir la figure II.29). En d'autres termes, il nous faut savoir quel pixel contrôle quelle composante spectrale de l'impulsion à mettre en forme.

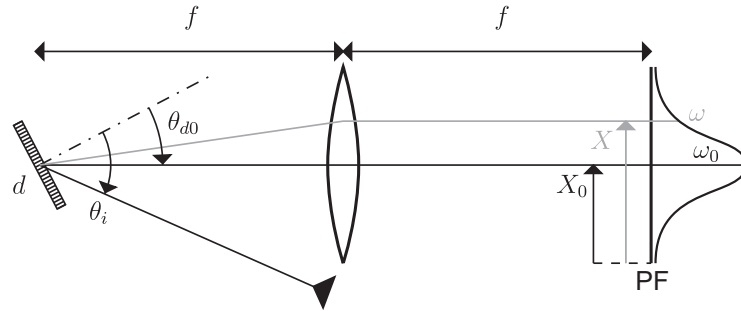


FIG. II.29 – La calibration de la dispersion donne le lien entre la position transverse X et la pulsation ω dans le plan de Fourier PF.

Pour cela, on programme le masque de façon à avoir une transmission nulle pour un certain nombre de pixels, à des positions connues. La figure II.30 page suivante présente un spectre obtenu lors d'une calibration. Au niveau de chacun des pixels éteints, on observe un minimum d'intensité. Le contraste obtenu est plutôt mauvais (on est loin d'une extinction totale). Ce contraste n'est absolument pas représentatif des capacités du façonneur. En effet, ne connaissant pas à l'avance la longueur d'onde correspondant aux pixels que l'on veut éteindre, on programme les tensions correspondant à une transmission nulle à la longueur d'onde centrale, ce qui n'est pas idéal. De plus, une partie de l'intensité incidente est transmise par les interstices du masque et dégrade le contraste. L'effet de ces interstices peut être compensé, ce qui n'est pas le cas ici. Enfin, la résolution du spectromètre utilisé (de l'ordre de 0,1 nm) ne permet pas de bien résoudre les creux au niveau des pixels éteints.

Pour cette calibration 20 pixels ont été éteints, espacés les uns des autres de 10 pixels. On obtient ainsi la relation entre la position sur le masque et la longueur d'onde pour une vingtaine de points¹⁷.

Ces points sont alors utilisés pour obtenir une formule de dispersion liant la position transverse X à la pulsation ω . Un ajustement linéaire de la dispersion est souvent suffisant. Il est donné par :

$$X(\omega) - X_0 \simeq \alpha(\omega - \omega_0) \quad (\text{II.20})$$

¹⁷Lorsqu'on utilise un masque de phase seule, on peut obtenir un résultat similaire en programmant des sauts de phase de π : au niveau de ces derniers, on observe une encoche dans le spectre[28].

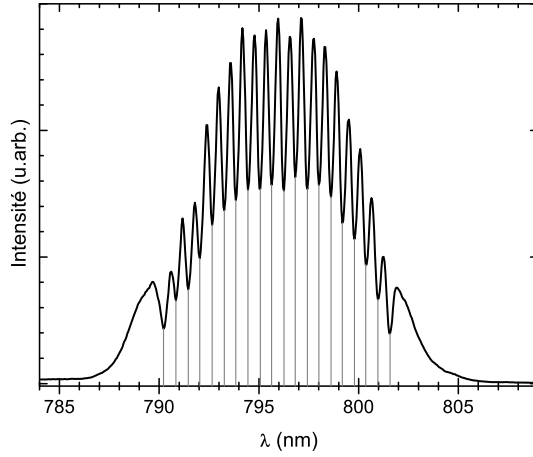


FIG. II.30 – Spectre typique lors d’une calibration de la dispersion. Le spectre transmis (trait noir) présente des minima au niveau des différents pixels choisis pour la calibration (indiqués par les traits gris).

où X_0 est la position correspondant à la pulsation centrale ω_0 , et α est donné par les formules de base des réseaux :

$$\alpha = -\frac{2\pi cf}{d\omega_0^2 \cos(\theta_{d0})} \quad (\text{II.21})$$

θ_{d0} est l’angle de diffraction pour la longueur d’onde centrale de l’impulsion :

$$\theta_{d0} = \arcsin\left(\frac{\lambda_0}{d} - \sin(\theta_i)\right) \quad (\text{II.22})$$

Dans le cas de notre façonneur, qui est très dispersif, il est essentiel de tenir compte des ordres supérieurs¹⁸. En effet, nous avons utilisé dans un premier temps l’ajustement linéaire présenté ci-dessus et nous avons constaté des déformations sur les impulsions mises en forme, comme cela a déjà été reporté[88, 49, 85]. Entre autres, l’introduction d’un délai entraînait une dérive de fréquence sur les impulsions. Ces problèmes ont été résolu en utilisant la formule de dispersion complète :

$$X(\omega) - X_0 = f \cdot \tan\left[\arcsin\left(\frac{2\pi c}{\omega d} - \sin(\theta_i)\right) - \arcsin\left(\frac{2\pi c}{\omega_0 d} - \sin(\theta_i)\right)\right] \quad (\text{II.23})$$

La focale f , le pas du réseau d et la pulsation centrale ω_0 étant fixé, on ajuste la position de la longueur d’onde centrale X_0 et l’angle d’incidence θ_i à l’aide des points mesurés. Dans la pratique, d’un jour à l’autre, la position varie au plus d’un pixel et l’angle d’incidence de quelques dizaines de milliradians.

¹⁸La dispersion non-linéaire a aussi pour effet de déformer fortement les répliques dues à la pixellisation (voir la partie II.9.2 page 47).

II.9 Résultats et limitations

Nous possédons maintenant tous les éléments pour utiliser notre façonneur. En effet, nous savons quel pixel agit sur quelle longueur d'onde (voir la partie II.8 page 43) et pour une longueur d'onde donnée, nous savons quelle paire de tensions appliquer pour obtenir une phase et une amplitude souhaitées (voir l'équation II.11 page 37 et la partie II.7 page 35).

De façon plus précise, nous pouvons contrôler la phase et l'amplitude sur $N = 640$ pixels, de largeur spectrale $\delta\omega$ et centrés en des pulsations ω_n . Nous avons donc un masque spectral de la forme

$$M_w(\omega) = \sum_{n=-N/2}^{N/2-1} B^{(n)} \delta(\omega - \omega_n) \otimes \text{rect}(\omega/\delta\omega) \quad (\text{II.24})$$

où $B^{(n)}$ est la transmission complexe du $n^{\text{ième}}$ pixel que l'on contrôle par le biais du couple de tensions appliquées au pixel.

Nous allons voir qu'une utilisation naïve de ce masque fait apparaître de nombreuses différences entre la mise en forme programmée et celle réellement obtenue. Nous verrons que ces différences sont dues à deux limitations inhérentes à ce type de façonneur. La première est la pixellisation du masque utilisé et nous traiterons de ses effets en détail. La deuxième est liée à la ligne à dispersion nulle qui introduit un couplage spatiotemporel que nous présenterons brièvement. Nous verrons enfin quelles précautions et quelles stratégies nous pouvons appliquer pour corriger ces défauts et obtenir les mises en formes souhaitées.

II.9.1 Approche naïve

Nous allons ici donner deux exemples simples de mise en forme que nous appliquerons de manière directe. La démarche est la suivante. Nous souhaitons réaliser une mise en forme spectrale donnée par la fonction de transfert $H_i(\omega)$. Nous allons tout simplement échantillonner cette fonction sur les pixels du masques, c'est à dire que nous allons programmer

$$M_w(\omega) = \sum_{n=-N/2}^{N/2-1} H_{i,n} \delta(\omega - \omega_n) \otimes \text{rect}(\omega/\delta\omega) \quad (\text{II.25})$$

où $H_{i,n} = H_i(\omega_n)$.

II.9.1.1 Cas d'un délai

Appliquons la démarche précédente au cas d'un délai $\phi^{(1)}$. Cela correspond à une phase spectrale linéaire et on a la fonction de transfert suivante :

$$H_i(\omega) = e^{i\phi^{(1)}(\omega - \omega_0)} \quad (\text{II.26})$$

où ω_0 est la pulsation centrale de l'impulsion à mettre en forme.

Nous avons programmé de manière directe ce type de mise en forme pour différents délais pour des impulsions à 800 nm de 100 fs de durée. La figure II.31 présente les intensités de corrélation obtenues pour trois délais : 15 ps en trait gris clair, 20 ps en trait gris et 30 ps en trait noir. Une intensité de 1 correspond à l'intensité de l'impulsion obtenue en demandant au façonneur un délai nul. On retrouve bien dans les trois cas une copie de l'impulsion à mettre en forme autour du délai voulu. Cependant, l'intensité obtenue dépend fortement du délai programmé, alors que la mise en forme programmée est une simple phase (voir l'équation II.26 page précédente) et qu'elle devrait donc toujours donner la même intensité en sortie. De plus on constate dans les trois cas une réplique non voulue à délai nul dont l'intensité ne dépend pas du délai programmé. Enfin, une réplique étirée précède toujours l'impulsion mise en forme de 35 ps environ. Elle est d'autant plus intense que le délai demandé est important et elle peut même devenir nettement plus intense que l'impulsion voulue. C'est le cas pour un délai de 30 ps où la pré-réplique à -5 ps est 7 fois plus intense que l'impulsion programmée. En terme d'énergie (donnée par l'aire sous l'impulsion), l'écart entre la pré-réplique et l'impulsion voulue est encore plus grand.

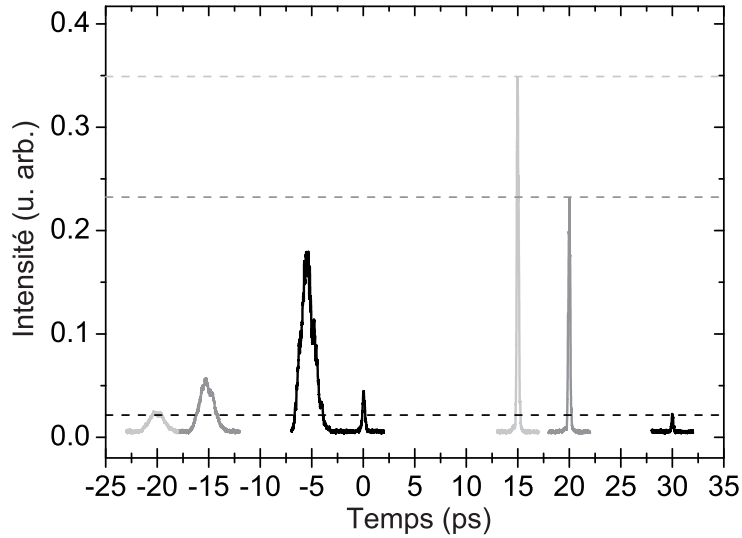


FIG. II.31 – Programmation directe de trois délais : 15 ps (trait gris clair), 20 ps (trait gris) et 30 ps (trait noir).

II.9.1.2 Cas d'une dérive de fréquence

Nous persévérons en appliquant la même démarche pour une dérive de fréquence $\phi^{(2)}$. Cela correspond à une phase spectrale parabolique et on a la fonction de transfert suivante :

$$H_i(\omega) = e^{i\frac{1}{2}\phi^{(2)}(\omega-\omega_0)^2} \quad (\text{II.27})$$

La figure II.32 page suivante présente l'intensité de corrélation pour une phase spectrale quadratique $\phi^{(2)} = 7.10^5 \text{ fs}^2$. On constate un pic étroit à délai nul en plus de l'impulsion à dérive de fréquence. De plus, l'impulsion à dérive de fréquence ainsi obtenue est nettement

plus courte que ce que prévoit la théorie : 12 ps de largeur à mi-hauteur en intensité au lieu des 19 ps attendues¹⁹.

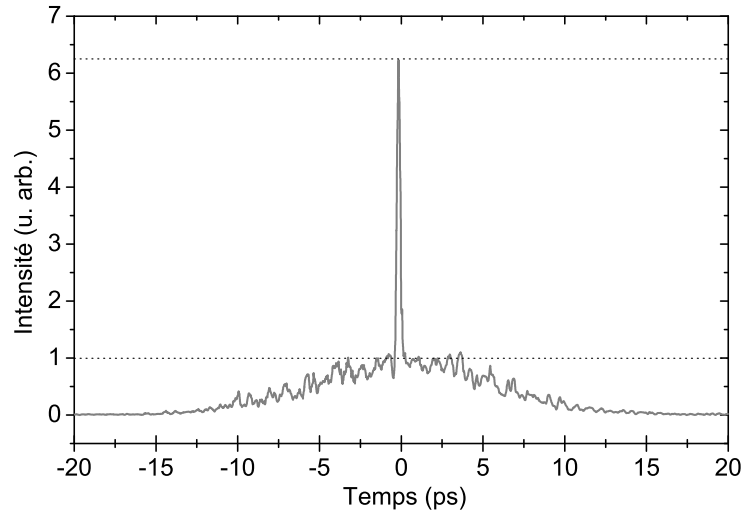


FIG. II.32 – Programmation directe d'une forte phase quadratique $\phi^{(2)} = 7.10^5 \text{ fs}^2$.

Les résultats présentés ici sont quelques peu caricaturaux. En effet, pour les besoins de la cause, nous nous sommes placés dans des situations très défavorables. Néanmoins, ces résultats montrent que l'approche naïve présente des dangers et qu'elle limite grandement la qualité des mises en forme obtenues. Il nous faut donc pousser un peu plus loin notre analyse du fonctionnement du façonneur. Nous allons en particulier détailler les effets de la pixellisation du masque. Nous évoquerons ensuite les problèmes de couplage spatio-temporels introduits par la ligne à dispersion nulle. Bien évidemment, cette analyse est nécessaire lorsque l'on veut travailler en boucle ouverte (c'est à dire lorsque l'on souhaite obtenir une mise en forme donnée directement, sans utiliser ni rétroaction, ni itérations). Dans le cas d'une utilisation en boucle fermée, où on utilise un algorithme d'optimisation pour converger vers une forme de sortie donnée, cette analyse apparaît moins utile. Pourtant, nous allons montrer que même dans ce cas, il faut connaître précisément les limitations de son façonneur.

II.9.2 Pixellisation

Le masque que nous utilisons est composé de 640 pixels séparés par des interstices (des "gaps" en anglais). Cette pixellisation entraîne deux conséquences.

Tout d'abord, comme nous l'avons vu dans l'approche naïve, nous ne pouvons pas reproduire réellement une fonction de transfert donnée $H_i(\omega)$ (voir notamment l'équation II.25 page 45). Nous pouvons simplement échantillonner cette dernière, comme cela

¹⁹Par contre, les oscillations sur le profil temporel ont déjà été observées sur des corrélations d'impulsions à forte dérive de fréquence, obtenues par d'autres moyens que le façonneur (étireur à réseaux en double passage).

est schématisé sur la figure II.33 : on programme sur un pixel centré en ω_n la transmission $H_i(\omega_n)$. Comme tout échantillonnage, cette étape introduit une distortion entre la mise en forme souhaitée et la mise en forme réellement appliquée. Le principal effet de cet échantillonnage spectral est d'introduire dans le domaine temporel des répliques de l'impulsion mise en forme.

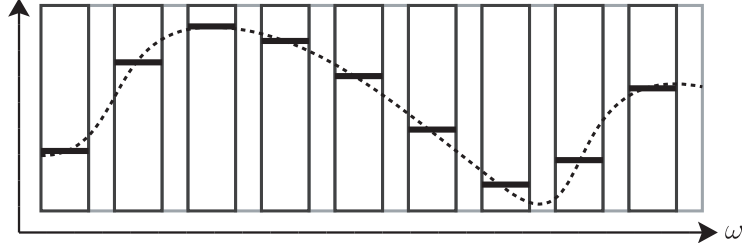


FIG. II.33 – La fonction de transfert à reproduire (pointillés) est échantillonnée sur les pixels du masque (traits noirs).

Ensuite, nous ne contrôlons pas la mise en forme au sein des interstices. Ces derniers vont eux aussi produire un jeu de répliques temporelles qui seront liées non pas à l'impulsion mise en forme, mais à l'impulsion entrant dans le façonneur.

Nous allons dans un premier temps étudier ces deux effets de pixellisation en supposant que la dispersion des composantes spectrales est linéaire au niveau du masque. Nous étudierons ensuite l'effet de la dispersion non-linéaire.

II.9.2.1 Pixellisation pour une dispersion linéaire

Nous considérons ici une dispersion linéaire dans le plan du masque pour les pulsations de l'impulsion à mettre en forme. Nous noterons ω l'écart à la pulsation centrale de l'impulsion incidente, X la position dans le plan de Fourier, l'origine étant la position de la pulsation centrale. La dispersion linéaire est notée α (voir l'équation II.21 page 44 pour sa valeur) et on a ainsi la relation $X = \alpha\omega$ entre la position et la pulsation. On suppose que l'on a programmé le masque de façon à reproduire une fonction de transfert idéale $H_i(\omega)$. La figure II.34 page suivante regroupe les différentes notations utilisées.

En s'inspirant de ce qui a été présenté sur le masque à cristaux liquides (voir l'équation II.12 page 37), et compte tenu de la dispersion linéaire, on peut écrire pour le masque :

$$M_w(\omega) = \underbrace{\sum_{n=-N/2}^{N/2-1} H_{i,n} \text{rect} \left[\frac{\omega - n\delta\omega}{r\delta\omega} \right]}_{\text{contribution des pixels}} + \underbrace{\sum_{n=-N/2}^{N/2-1} G_n \text{rect} \left[\frac{\omega - (n+1/2)\delta\omega}{(1-r)\delta\omega} \right]}_{\text{contribution des interstices}} \quad (\text{II.28})$$

où N est le nombre de pixels, $\delta\omega$ la largeur spectrale d'un ensemble pixel-interstice, et $\text{rect}()$ est la fonction carré définie par l'équation II.13 page 37 et qui vaut 1 dans $[-\frac{1}{2}, \frac{1}{2}]$ et nulle partout ailleurs. $\delta\omega$ est reliée à la distance l entre deux pixels par $\delta\omega = l/\alpha$. $H_{i,n}$ est la transmission programmée sur le $n^{\text{ième}}$ pixel, elle vaut $H_{i,n} = H_i(n\delta\omega)$. De même, G_n

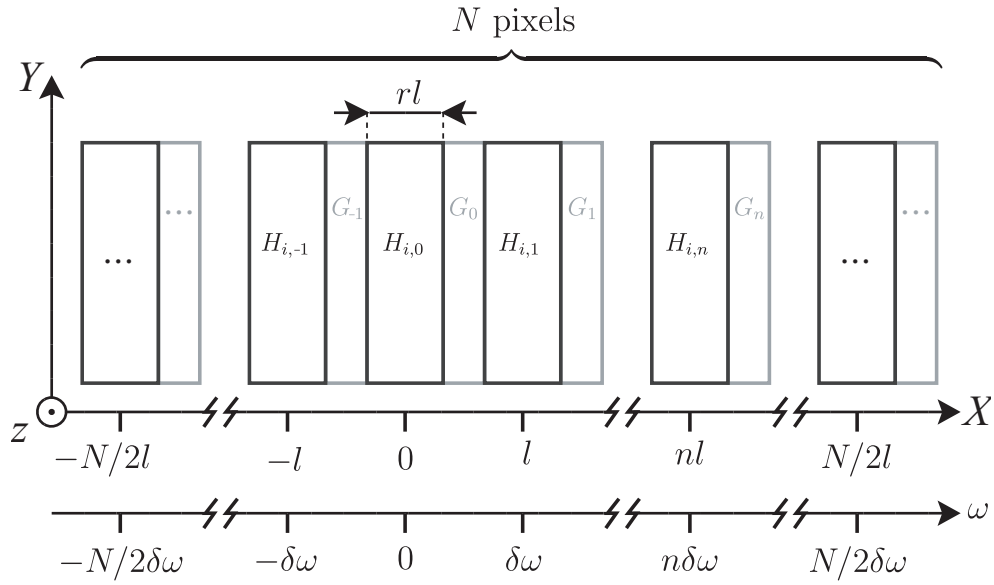


FIG. II.34 – Programmation du masque pixellisé. $H_{i,n}$ est la transmission du $n^{\text{ième}}$ pixel et G_n celle du $n^{\text{ième}}$ interstice.

est la transmission du $n^{\text{ième}}$ interstice. Cette transmission est assez mal caractérisée et on suppose le plus souvent qu'elle correspond simplement à la transmission d'un pixel soumis à une tension nulle, soit dans notre cas une transmission de 1. Cette supposition reproduit assez bien les observations expérimentales.

Enfin, r est un ratio qui donne la part de l occupée par le pixel : le pixel a une largeur spatiale (respectivement spectrale) de rl (respectivement $r\delta\omega$) et l'interstice de $(1-r)l$ (respectivement de $(1-r)\delta\omega$). Pour notre façonneur, on a $r = 0,97$.

Nous pouvons modifier légèrement cette équation en repoussant les bornes de la sommation à $-\infty$ et $+\infty$. Cela revient à dire que le masque est infini. Cette modification est justifiée par le fait que le spectre des impulsions à mettre en forme est normalement plus étroit que le spectre transmis par le masque et que cela ne fera donc pas de différence sur l'impulsion de sortie.²⁰ On obtient ainsi l'expression suivante pour le masque spectral[89] :

$$M_w(\omega) = \sum_{n=-\infty}^{\infty} H_{i,n} \text{rect} \left[\frac{\omega - n\delta\omega}{r\delta\omega} \right] + \sum_{n=-\infty}^{\infty} G_n \text{rect} \left[\frac{\omega - (n + 1/2)\delta\omega}{(1-r)\delta\omega} \right] \quad (\text{II.29})$$

Dans le premier terme de cette expression, on reconnaît la fonction de transfert idéale H_i , échantillonnée spectralement avec un pas $\delta\omega$. Ce sous échantillonnage va se traduire dans le domaine temporel par l'apparition de répliques espacées d'un délai $T_0 = 2\pi/\delta\omega$. Pour notre façonneur, cette fenêtre temporelle T_0 est de l'ordre de 35,8 ps. Une autre façon de voir l'effet de ce sous-échantillonnage est de dire que l'on ne contrôle le champ mis en forme que sur une fenêtre de largeur T_0 .

²⁰Même dans le cas d'un spectre tronqué par le masque, on peut faire cette approximation remplaçant le spectre d'entrée par le spectre tronqué.

De plus, comme chaque échantillon a un profil carré de largeur $r\delta\omega$, on va voir apparaître une enveloppe en sinus cardinal, de pseudo-période $T_0/r \simeq 36,9$ ps. Il en est de même pour la contribution des interstices qui va se traduire par des répliques espacées de T_0 , modulées par une enveloppe en sinus cardinal de pseudo-période $T_0/(1-r)r \simeq 1,19$ ns.

Plus précisément, la transformée de Fourier de l'équation II.29 page précédente, donne pour la réponse impulsionnelle $m(t)$:

$$m(t) = r \operatorname{sinc}\left(\frac{r\pi t}{T_0}\right) \sum_{n=-\infty}^{\infty} H_{i,n} e^{in\delta\omega t} + (1-r) \operatorname{sinc}\left(\frac{(1-r)\pi t}{T_0}\right) \sum_{n=-\infty}^{\infty} G_n e^{i(n+1/2)\delta\omega t} \quad (\text{II.30})$$

On voit ainsi bien apparaître les enveloppes en sinus cardinal (fonction notée $\operatorname{sinc}(\cdot)$), ainsi que le poids relatif des répliques dues aux pixels et de celles dues aux interstices. On a ainsi un rapport en intensité théorique de $(r/[1-r])^2 \simeq 1000$ entre la contribution des pixels et celle des interstices.

Pour illustrer ceci, nous allons prendre le cas d'un simple délai : $H_i(\omega) = \exp(i\omega\tau)$. Nous supposons que la transmission des interstices et la même partout et qu'elle est égale à celle d'un pixel soumis à une tension nulle. Cela correspond à une transmission unitaire $G_n = 1$. Dans ce cas, la réponse impulsionnelle voulue est une fonction de Dirac $\delta(t - \tau)$ et celle obtenue est donnée par²¹ :

$$m(t) = r \operatorname{sinc}\left(\frac{r\pi t}{T_0}\right) \sum_{n=-\infty}^{\infty} \delta(t - \tau - nT_0) + (1-r) \operatorname{sinc}\left(\frac{(1-r)\pi t}{T_0}\right) \sum_{n=-\infty}^{\infty} \delta(t - nT_0) e^{\frac{i\delta\omega t}{2}} \quad (\text{II.31})$$

Dans ce cas très simple on obtient une infinité de copies exactes de la mise en forme souhaitée, séparées les unes des autres de T_0 , avec une amplitude donnée par l'enveloppe en sinus cardinal. On obtient aussi une infinité de répliques non-mises en forme, dues aux interstices.

La figure II.35 page suivante présente une simulation pour un délai de 12 ps dans le cas de notre façonneur ($T_0 = 35,8$ ps, $r = 0,97$, durée des impulsions à l'entrée de 100 fs). L'intensité du champ sortant (trait noir) ainsi que l'enveloppe des répliques dues aux pixels (trait gris) et de celles dues aux interstices (pointillés gris) sont représentées. Le temps est gradué en multiples de la fenêtre T_0 et l'échelle d'intensité est logarithmique.

On observe bien que les répliques dues aux pixels suivent l'impulsion principale en restant sous l'enveloppe en sinus cardinal [89], alors que celles dues aux interstices sont fixes.

Pour des délais τ de l'ordre de $T_0/2$, la première réplique due aux pixels (celle en $\tau - T_0$) aura une intensité comparable à celle de l'impulsion mise en forme. On ne peut donc plus faire de distinction entre une l'impulsion principale et sa réplique : on aurait un résultat identique en demandant un délai de $-T_0/2$. Nous obtenons cette singularité car nous atteignons la limite de Nyquist. En effet, nous échantillonnons notre mise en forme avec un pas spectral $\delta\omega$. Pour éviter des ambiguïtés dues au repliement temporel, la mise en forme spectrale doit avoir une période plus grande que $2\delta\omega$. Dans le cas d'une phase,

²¹On a utilisé ici la relation de fermeture $\sum e^{in\delta\omega t} = \sum \delta(t - nT_0)$.

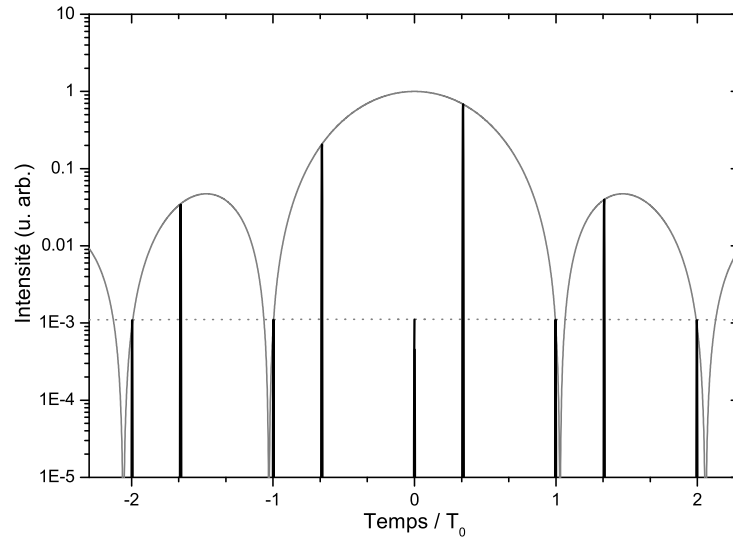


FIG. II.35 – Effet de la pixellisation pour un retard de 12 ps, pour une dispersion linéaire. Outre la mise en forme demandée, l'intensité en sortie (trait noir) présente deux séries de répliques espacées de $T_0 = 35,8$ ps. Les répliques dues aux pixels suivent la mise en forme et sont modulées par une enveloppe en sinus cardinal (trait gris) de pseudo-période T_0/r . Celles dues aux interstices sont fixes, modulées par un sinus cardinal (pointillés gris) de pseudo-période $T_0/(1-r)$ très grande devant T_0 .

cela revient à dire que l'accroissement de phase ne doit pas dépasser 2π sur deux pixels soit π/pixel . Pour une phase linéaire de type $\omega\tau$, cela fixe une limite supérieure à τ :

$$\tau < \pi/\delta\omega = T_0/2$$

Il s'agit donc là d'un exemple extrême. Pour éviter des répliques dues aux pixels trop importantes, il faut restreindre la fenêtre temporelle dans laquelle on réalise les mises en forme.

Cette approche simple, supposant une dispersion linéaire, nous a permis d'expliquer la présence des différentes répliques (dues aux pixels ou dues aux interstices) dans les figures II.31 page 46 et II.32 page 47. Elle explique aussi en partie la diminution d'amplitude observée sur les impulsions fortement décalées de la figure II.31.

Par contre, elle prévoit une réplique à délai nul (due aux interstices) nettement plus faible que celle observée expérimentalement (0,1% en théorie, de l'ordre de 2,5% dans l'expérience). De plus cette réplique ne dépend pas de la mise en forme programmée alors que dans certains cas nous avons pu observer expérimentalement des variations significatives de celle-ci. Ces deux différences sont explicables par le comportement des cristaux liquides au sein du masque. En effet, le masque de transmission est en fait lié à la rotation des cristaux liquides sous l'action des deux masques de tensions qui sont appliqués aux deux

barrettes. A la frontière entre un pixel (soumis à une tension non nulle) et l'interstice voisin, la tension appliquée change brutalement, mais les cristaux liquides présentent quant à eux une transition douce. Ils vont ainsi passer continûment de la rotation imposée au niveau du pixel à la position de repos qui prévaut au sein de l'interstice. Les phases appliquées par chacune des deux barrettes vont elles aussi présenter de telles transitions continues. On peut ainsi reproduire les observations expérimentales en introduisant une fonction de réponse spatiale du masque qui vient lisser les créneaux de phase appliqués pour réaliser une mise en forme donnée.

Autre limite de cette approche, elle n'explique pas l'étalement des répliques de pixellisation observé sur la figure II.31 page 46. Pour expliquer cet étalement, il faut tenir compte de la dispersion non-linéaire dans le plan de Fourier.

II.9.2.2 Effet de la dispersion non-linéaire

Dans notre façonneur, il nous faut tenir compte de la dispersion non-linéaire, comme nous l'avons mentionné lors de la calibration de la dispersion. En particulier, si l'on utilise l'approximation linéaire de la dispersion pour programmer le masque, on observe une distorsion des impulsions mises en forme[88].

Le principal effet de cette dispersion non-linéaire est que le masque spectral $M_w(\omega)$ n'est plus régulier. En effet, le masque spatial $M_X(X)$ est composé de pixels spatialement périodiques mais comme la position X et la pulsation ω ne sont plus liées par une simple relation linéaire, ces pixels sont de largeur variable dans le domaine spectral.

Or les répliques étaient directement liées à la périodicité de l'échantillonnage dans le domaine spectral. Maintenant que cette périodicité est mise en défaut, les différentes répliques vont être déformées[85]. La figure II.36 page suivante présente la même simulation que celle de la figure II.35 page précédente mais en prenant la formule complète pour la dispersion spatiale au niveau du masque. Les mêmes paramètres et notations sont utilisées que pour la figure II.35. On constate que les différentes répliques montrent un fort élargissement et une intensité crête très inférieure à celle obtenue dans l'approximation linéaire (seule la réplique à délai nulle reste inchangée). En fait, ces répliques ont une forte dérive de fréquence, d'autant plus grande qu'elles sont loin du délai nul. Ce type de déformation a déjà été rencontré lors de l'étude de la dispersion non-linéaire dans les spectromètre et de son effet sur l'interférométrie spectrale[90]. L'effet de la dispersion non-linéaire est plutôt positif car il tend à diminuer l'intensité des répliques, qu'elles soient dues aux pixels ou aux gaps²². Dans le cas, fréquent, où l'impulsion mise en forme participe à un processus non-linéaire, cette diminution de l'intensité crête de plus d'un ordre de grandeur est très profitable.

En fait, l'étirement de ces répliques est l'analogue de l'étalement d'un paquet d'onde vibrationnel dans un puits de potentiel anharmonique. En effet, tout comme un paquet d'onde, l'impulsion mise en forme s'écrit comme la superposition d'un nombre fini de niveaux. Ces niveaux sont les N pixels composant le masque, centrés à des pulsations ω_n qui dépendent de la dispersion dans le plan du masque. La non linéarité de la dispersion est l'analogue de l'anharmonicité des niveaux qui constituent le paquet d'onde.

²²Cependant, l'énergie des différentes répliques n'est pas modifiée par la dispersion non-linéaire.

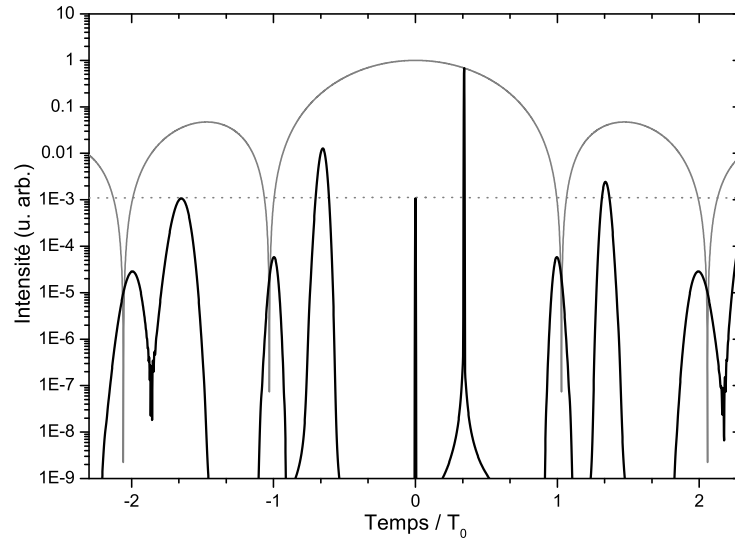


FIG. II.36 – Effet de la dispersion non-linéaire sur les répliques dues à la pixelisation pour un retard de 12 ps. Les répliques présentes sur l'intensité de sortie (trait noir) présentent de fortes dérives de fréquence. Leur intensité crête est ainsi inférieure à l'enveloppe prévue (trait gris pour les répliques des pixels et pointillés gris pour celles des interstices) par l'approximation linéaire de la dispersion.

Si on crée un paquet d'onde dans un puits anharmonique, celui-ci va osciller en se déphasant ; il va revenir périodiquement à sa position d'origine en étant à chaque fois plus dispersé. C'est exactement ce que l'on observe dans la mise en forme obtenue : périodiquement, on voit réapparaître l'impulsion mise en forme, à chaque fois plus étalée. Il existe cependant une différence quantitative qui est le nombre de niveaux : nous avons 640 pixels, ce qui est 10 à 50 fois supérieur au nombre d'états mis en jeu habituellement dans les expériences de dynamique de paquet d'onde[91, 92]. Cela nous rend extrêmement sensible à l'anharmonicité.

De plus dans le cas de notre façonneur, on a deux "paquets d'onde" : l'un est dû aux interstices et l'autre aux pixels. En jouant sur les transmissions complexes des pixels, on change le paquet d'onde initial dû aux pixels (c'est à dire la réplique principale mise en forme). On ne peut par contre pas agir sur le paquet d'onde dû aux interstices.

Une application intéressante de cette analogie est la possibilité de refocaliser une des répliques dues aux interstices : c'est l'analogue de la refocalisation de paquet d'onde étudiée théoriquement[93] et observée expérimentalement[94]. Dans ces expériences, on prépare le paquet d'onde initial avec une relation de phase donnée. Celle-ci est l'opposée de la phase non linéaire accumulée par le paquet d'onde après un certain temps de propagation τ . Au bout d'un temps τ , on obtient un paquet d'onde focalisé temporellement.

Transposé au cas de notre façonneur, cela veut dire que l'on peut recomprimer n'importe quelle réplique due aux interstices. Il suffit pour cela de programmer sur la réplique principale la phase compensant la dispersion de la réplique visée.

Une telle refocalisation a été simulée sur la figure II.37. Le trait noir représente l'intensité temporelle obtenue (en échelle logarithmique) en programmant une impulsion retardée de 12 ps (soit de l'ordre de $T_0/3$) avec une phase quadratique $\phi^{(2)} = -6.10^4 \text{ fs}^2$. Pour comparaison, les tirets gris montrent l'intensité déjà obtenue en absence de phase quadratique. On observe la recompression presque parfaite de la réplique en $t = -23 \text{ ps}$ (c'est à dire $-2/3T_0$). L'intensité crête de cette réplique est 5 fois plus élevée que celle de l'impulsion principale étirée.

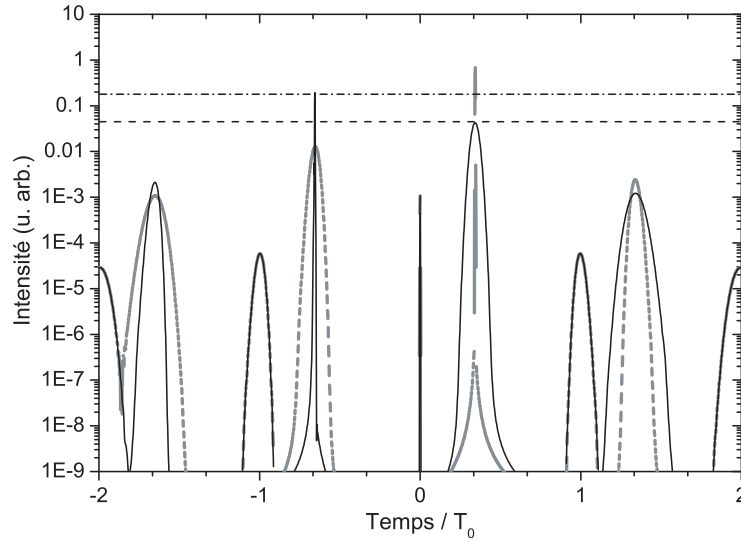


FIG. II.37 – Effet de la dispersion non-linéaire et de la pixellisation pour un retard de 12 ps : sans dérive de fréquence (tirets gris) et avec une phase quadratique $\phi^{(2)} = -6.10^4 \text{ fs}^2$ (trait noir).

Il est ainsi très important de tenir compte des effets de pixellisation et de dispersion non linéaire lors de la programmation du façonneur, faute de quoi on peut se retrouver dans une situation où l'impulsion principale est nettement moins intense que certaines répliques de pixellisation.

Tout ce qui a été présenté ici se transpose à tout autre masque pixellisé, comme les masques à cristaux liquides à deux dimensions[76] et les miroirs MEMS[45]. Par contre, les masques non pixellisés comme les modulateurs acousto-optiques, les miroirs à membrane déformable ou les valves optiques ne présentent pas de telles répliques²³.

II.9.3 Couplage spatio-temporel

Aux effets de la pixellisation du masque, s'ajoute une limitation qui est propre aux lignes à dispersion nulle. C'est le défaut le plus connu de l'ensemble des façonneurs reposant

²³Ils ont par contre d'autres inconvénients comme le manque d'efficacité pour les modulateurs acousto-optiques ou la complexité pour les valves optiques qui nécessitent un faisceau laser de contrôle mis en forme à l'aide d'un masque à cristaux liquides. De plus, ces dispositifs non pixellisés ne peuvent pas reproduire des variations très abruptes de la phase ou de l'amplitude comme des sauts de phase ou des créneaux d'amplitude

sur une ligne à dispersion nulle : le couplage spatio-temporel. On peut se faire une idée de la physique que cache une telle appellation assez simplement. Il suffit pour cela de supposer que la dispersion au niveau du masque est linéaire et de prendre l'exemple d'un délai. L'introduction d'un délai correspond à une phase spectrale linéaire. Spatialement, le masque applique donc une phase linéaire. Il se comporte donc comme un prisme et l'impulsion est déviée. En sortie, l'impulsion sera donc décalée temporellement et spatialement.

A travers cet exemple, on comprend bien l'origine du couplage spatio-temporel : on agit sur les composantes spectrales alors qu'elles sont couplées à la position transverse dans le plan de Fourier. On agit ainsi sur les deux à la fois.

Nous ne détaillerons pas ici complètement les calculs permettant d'obtenir le champ de sortie en fonction de la mise en forme programmée. Ceux-ci ont été détaillés dans plusieurs articles, notamment les références [86, 95]. Nous donnons les étapes essentielles de ces calculs en annexe B page 157. Dans cette partie, nous nous contenterons de donner l'expression du champ en sortie du façonneur et d'en montrer les implications. Nous verrons en particulier que ce champ se présente le plus souvent comme une fonction où les variables spatiales et temporelles ne sont pas séparables. Nous donnerons enfin l'expression analytique simplifiée du champ de sortie dans le cas d'un filtrage spatial de la sortie du façonneur. Nous illustrerons ces explications par l'exemple de séquences d'impulsions.

II.9.3.1 Expression du champ de sortie

Nous souhaitons donner ici le lien entre le champ en entrée du façonneur et celui obtenu en sortie, pour un masque programmé $M_w(\omega)$. Du point de vue spatial, seule nous importe la coordonnée transverse horizontale x selon laquelle se fait l'étalement des pulsations au sein de la ligne. Selon la verticale y , il n'y a pas de changement et nous omettrons cette coordonnée. Nous supposerons que le champ en entrée possède un profil transverse gaussien de largeur à mi-hauteur en intensité Δx . Nous supposerons que les dépendances spatiales et spectrales sont factorisables²⁴. Nous l'écrivons donc sous la forme :

$$E_E(\omega, x) = E_E(\omega) e^{-2 \ln(2) \left(\frac{x}{\Delta x}\right)^2} = E_E(\omega) 4^{-\left(\frac{x}{\Delta x}\right)^2} \quad (\text{II.32})$$

La pulsation ω est définie par rapport à la pulsation centrale ω_0 .

Le champ en sortie du façonneur peut s'écrire comme :

$$E_S(\omega, x) = E_E(\omega) \int e^{i \frac{x}{v} \omega'} M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta \omega'}\right)^2} d\omega' \quad (\text{II.33})$$

où on a introduit les deux paramètres v et $\Delta \omega'$ dont les expressions sont détaillées dans l'annexe B page 157. Nous nous contenterons ici de donner une idée qualitative de leur signification physique. $\Delta \omega'$ est ainsi lié à la taille finie ΔX des taches focales dans le plan du masque : celle-ci "lissent" l'effet des cristaux liquides. La partie intégrale de l'équation II.33 traduit entre autre cette convolution du masque par le profil spatial de chacune des composantes spectrales. La vitesse de couplage spatio-temporel v traduit quant à elle le

²⁴Il n'y a aucun couplage entre les deux : ni dérive spatiale de fréquence, ni aberration chromatique...

double rôle joué par le masque à cristaux liquides : masque spectral (grâce à la dispersion des composantes spectrales) mais aussi masque spatial qui agit sur fréquences spatiales du faisceau de sortie. Elle est directement proportionnelle au pas des réseaux : elle sera d'autant plus faible que les réseaux seront dispersifs. Plus cette vitesse est faible et moins les effets de couplage spatio-temporel se feront sentir. Bien qu'implicite, ce couplage apparaît très clairement dans cette équation : la position transverse est, à un facteur $1/v$ près, la variable conjuguée de la pulsation. De plus, cette équation permet de calculer analytiquement le champ en sortie pour certains masques simples²⁵.

II.9.3.2 Couplage : cas d'un délai.

Nous allons ici intégrer analytiquement l'équation II.33 page précédente dans le cas simple d'un décalage temporel. Cela nous permettra de faire apparaître le décalage spatial associé par couplage et d'expliquer de façon rigoureuse l'exemple donné page 55.

Pour retarder d'un délai τ l'impulsion traversant la façonneur, le masque doit introduire une phase spectrale linéaire de la forme $\phi(\omega) = \omega\tau$. Cela correspond à un masque :

$$M_w(\omega) = e^{i\omega\tau} \quad (\text{II.34})$$

Le champ en sortie s'écrit donc :

$$E_S(\omega, x) = e^{i\omega\tau} E_E(\omega) \int e^{i\frac{x}{v}\omega'} e^{-i\omega'\tau} 4^{-\left(\frac{\omega'}{\Delta\omega}\right)^2} d\omega' \quad (\text{II.35})$$

On voit ainsi apparaître la mise en forme spectrale souhaitée ainsi qu'une intégrale dépendant de la position et du délai introduit. Cette intégrale n'est autre que la transformée de Fourier d'une gaussienne, multipliée par une phase linéaire. On a donc

$$E_S(\omega, x) = e^{i\omega\tau} E_E(\omega) 4^{-\left(\frac{x/v-\tau}{\Delta T}\right)^2} \quad (\text{II.36})$$

où l'on posé $\Delta T = 4 \ln(2)/\Delta\omega'$. En utilisant la définition de $E_E(\omega, x)$ ainsi que le lien entre $\Delta\omega'$ et Δx (voir l'équation B.17 page 161) on peut écrire

$$E_S(\omega, x) = e^{i\omega\tau} E_E(\omega) 4^{-\left(\frac{x-v\tau}{\Delta x}\right)^2} = e^{i\omega\tau} E_E(\omega, x - v\tau) \quad (\text{II.37})$$

En définissant $E_S(t, x)$ comme la transformé de Fourier inverse de $E_S(\omega, x)$ on peut exprimer ce même résultat dans le domaine temporel :

$$E_S(t, x) = E_E(t - \tau, x - v\tau) \quad (\text{II.38})$$

On obtient ainsi bien un champ décalé temporellement de τ , mais aussi décalé spatialement de $v\tau$.

La figure II.38 page suivante est une simulation présentant les impulsions obtenues en programmant une séquence de trois impulsions, décalées de 5 ps. Nous avons utilisé les

²⁵Autre intérêt de cette écriture, I. Wamsley et C. Dorrer ont montré que l'on peut calculer simplement un paramètre de le couplage spatio-temporel à partir de $E(x, \omega)$. Ils ont aussi démontré la possibilité de mesurer ce paramètre expérimentalement de façon linéaire[96].

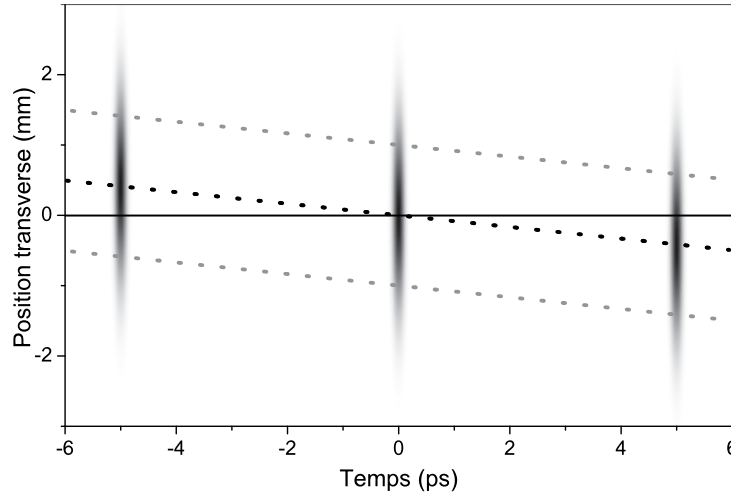


FIG. II.38 – Effet du couplage spatio-temporel sur une séquence de trois impulsions séparées de 5 ps. L'intensité des trois répliques (dégradé du blanc au noir, noir pour le plus intense) est représentée en fonction du temps et de la position transverse. Les pointillés noirs relient les maxima des répliques et les pointillés gris les mi-hauteurs.

caractéristiques de notre façonneur et pris des impulsions d'entrée de largeur à mi-hauteur en intensité de $\Delta t = 150$ fs pour le profil temporel et de $\Delta x = 2$ mm pour le profil spatial transverse. L'intensité du champ en sortie du façonneur (dégradé du blanc au noir, noir pour le plus intense) est représentée en fonction de la position transverse et du temps. Les pointillés noirs relient les maxima des répliques, les pointillés gris les mi-hauteurs. On observe bien un décalage spatial proportionnel au décalage temporel. Ces deux décalages sont liés par la vitesse de couplage[48] introduite plus haut : v est la pente de la droite en pointillés noir. On voit bien sur cette figure que plus v sera faible, moins les répliques seront décalées hors de l'axe de propagation. Pour notre façonneur, particulièrement dispersif, ce couplage vaut en valeur absolue $83 \mu\text{m}/\text{ps}$. Ceci est relativement faible comparé à des dispositifs standards ($145 \mu\text{m}/\text{ps}$ dans la référence [95]).

Cet effet n'est pas négligeable pour autant et il faut en tenir compte lorsque la mise en forme programmée s'étend sur un intervalle temporel dépassant la dizaine de picosecondes.

Dans le cas d'une phase spectrale quadratique, on peut aussi écrire le champ en sortie de façon analytique. Ce calcul est mené dans l'annexe B page 157 et il montre que la dérive de fréquence temporelle est accompagnée d'une dérive de fréquence spatiale ("spatial chirp" en anglais), les deux étant liées par le coefficient v .

II.9.3.3 Filtrage spatial et couplage spatio-temporel

Pour limiter les effets de couplage spatio-temporel, il est possible de filtrer spatialement le faisceau de sortie. Différents types de filtrage sont possibles, mais le plus souvent un simple diaphragme est placé sur le faisceau de sortie. Dans notre façonneur, ce filtrage est réalisé par le diaphragme commun aux faisceaux d'entrée et de sortie (D0 sur la figure II.16

page 31).

Après le filtrage, on peut écrire pour le champ de sortie[86, 48, 95, 28] :

$$E_S(\omega, x) = E_E(\omega, x) \int M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' \quad (\text{II.39})$$

Cela revient en fait à décomposer l'équation II.33 page 55 sur les différents modes hermito-gaussiens[29] et à n'en garder que le mode gaussien.

On voit apparaître la fonction de transfert $H(\omega)$ du façonneur qui vaut :

$$H(\omega) = \int M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' = \int M_w(\omega') 4^{-\left(\frac{\omega' - \omega}{\Delta\omega'}\right)^2} d\omega' \quad (\text{II.40})$$

Le masque programmé est ainsi convolué par une gaussienne de largeur à mi-hauteur en intensité $\Delta\omega'$. Cela revient à dire que la mise en forme temporelle correspondant à ce masque sera multipliée par une enveloppe gaussienne de largeur à mi-hauteur :

$$\Delta T = \frac{4 \ln(2)}{\Delta\omega'} = \frac{\Delta x}{|v|} \quad (\text{II.41})$$

Même après le filtrage, on voit encore un effet du couplage spatio-temporel : la mise en forme temporelle se fait sous une enveloppe gaussienne directement proportionnelle au profil spatial de l'impulsion d'entrée.

Une explication simple à cette gaussienne temporelle et à son lien avec le profil spatial peut être donnée en prenant l'exemple d'une séquence d'impulsions. On reprend pour cela ce qui a été fait dans la partie précédente, notamment l'expression d'une impulsion décalée temporellement (voir l'équation II.38 page 56). On suppose toujours que les impulsions d'entrée ont une durée de $\Delta t = 150$ fs et une largeur de $\Delta x = 2$ mm et on utilise la vitesse de couplage spatio-temporel v de notre façonneur ($v = -83 \mu\text{m/ps}$).

On considère une séquence de treize impulsions séparées de 5 ps réalisée avec le façonneur. Elle est représentée, en intensité, en fonction de la position transverse et du temps, dans le haut de la figure II.39 page ci-contre. Comme dans la figure II.38 page précédente, les pointillés noirs relient les maxima des répliques, les pointillés gris les mi-hauteurs. L'approximation gaussienne, ou le filtrage spatial, revient à ne garder de tout ceci que ce qui est le long de l'axe $x = 0$. C'est ce qui a été fait dans le bas de la figure II.39 : le trait noir représente l'intensité en fonction du temps le long de la coupe et le trait gris est l'enveloppe temporelle gaussienne théorique.

On voit ici l'origine physique de l'enveloppe temporelle gaussienne : plus les impulsions sont décalées temporellement, plus elles s'éloignent de l'axe $x = 0$ et plus l'intensité qui reste au niveau de l'axe est faible. En fait, En décalant temporellement une impulsion, on parcourt son profil transverse : plus l'impulsion sera large spatialement, plus il faudra la décaler temporellement pour obtenir une atténuation donnée. Ceci explique la proportionnalité entre ΔT et Δx (équation II.41).

Enfin, on retrouve la largeur à mi-hauteur de la gaussienne temporelle en regardant l'intersection de l'axe $x = 0$ avec les deux droites en pointillés gris qui délimitent la largeur à mi-hauteur spatiale.

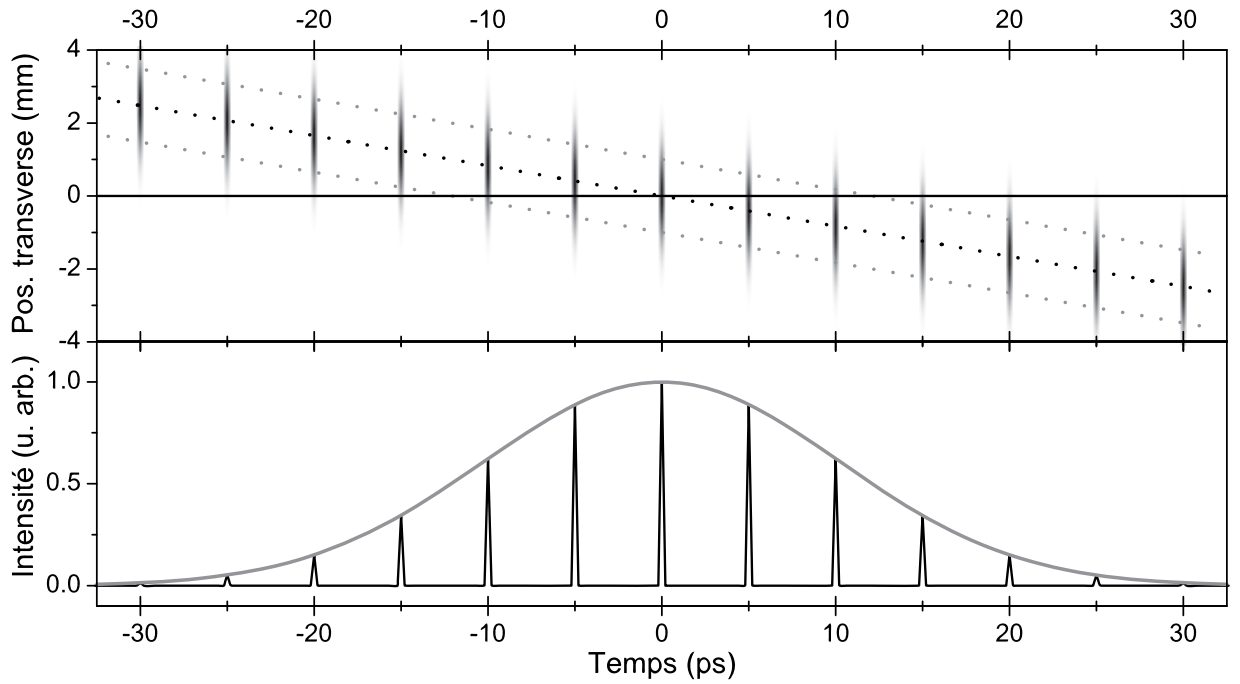


FIG. II.39 – **Haut** : 13 impulsions séparées de 5 ps. Intensité (blanc au noir, noir pour le plus intense) en fonction du temps et de la position transverse. Les pointillés noirs relient les maxima, les gris les mi-hauteurs. **Bas** : même séquence dans le cas d'un filtrage spatial. L'intensité de sortie (trait noir) ainsi que l'enveloppe temporelle (trait gris) sont représentées en fonction du temps.

Cette enveloppe gaussienne a été mesurée expérimentalement, pour différentes largeurs spatiales, en appliquant divers délais à l'aide du façonneur. La figure II.40 page suivante présente les résultats obtenus pour un faisceau incident de 2,0 mm de diamètre (largeur à mi-hauteur en intensité Δx).

Les carrés noirs correspondent aux intensités de corrélation de l'impulsions retardée avec une impulsion de référence. En trait gris est représenté un ajustement de l'enveloppe gaussienne. On trouve une largeur temporelle de 23 ps en bon accord avec la théorie qui donne 24 ps. Une deuxième expérience avec un diamètre de faisceau de 1,0 mm donne une largeur de 13 ps toujours en bon accord avec la théorie.

On remarque que la gaussienne temporelle n'est pas tout à fait centrée sur $t = 0$, mais légèrement décalée vers les délais positifs. Ceci est simplement dû à un diaphragme de filtrage spatial légèrement hors d'axe : la coupe du profil spatio-temporel se fait ainsi selon un axe $x = x_0$ avec $x_0 \neq 0$ (voir la figure II.39). Le décalage temporel observé ici correspond à un décalage spatial de l'ordre de $150 \mu\text{m}$, ce qui est tout à fait raisonnable. On peut en fait choisir le centre de la gaussienne temporelle en décalant latéralement le diaphragme de filtrage.

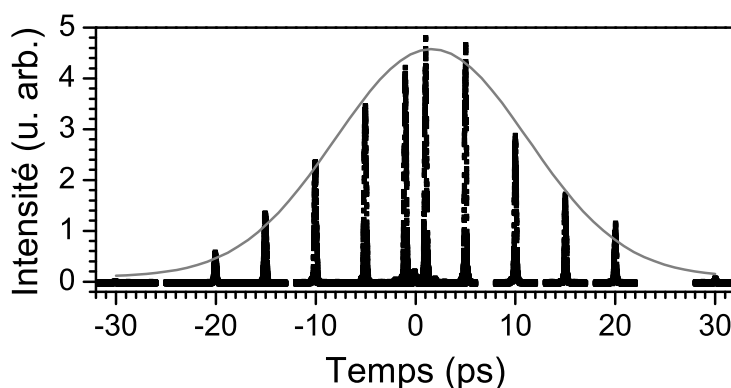


FIG. II.40 – Mise en évidence de l’enveloppe temporelle. Intensité de corrélation (carrés noirs) pour différents retards introduits par le façonneur. Une enveloppe temporelle gaussienne (trait gris) de 23 ps de durée est clairement visible.

II.9.3.4 Effet du filtrage sur les répliques de pixelisation

Le filtrage que nous avons présenté ci-dessus a un autre effet bénéfique : il tend à réduire l’importance des répliques dues aux pixels et aux interstices. Grace à l’enveloppe temporelle gaussienne, les répliques sont plus faibles et ce d’autant plus qu’elles sont éloignées de $t=0$. Cela est clairement visible sur la figure II.41 page suivante où l’on a comparé l’intensité en sortie de la ligne sans tenir compte du couplage spatiotemporel (pointillés gris) et en tenant compte de celui-ci (trait noir). Nous avons repris les mêmes paramètres que pour la figure II.41 page ci-contre, en utilisant pour le couplage spatiotemporel un diamètre à mi-hauteur en entrée de la ligne de 2 mm. Cela correspond à une gaussienne temporelle de 23 ps (voir page 58).

II.9.4 Utilisation d’algorithmes d’optimisation

Les différentes limitations que nous venons de détailler sont évidemment très importante pour les expérimentateurs qui souhaitent obtenir directement une mise en forme donnée (travail en boucle ouverte). De nombreuses expérimentateurs adoptent une approche un peu différente en laissant la programmation de leur façonneur évoluer de manière itérative afin d’obtenir une mise en forme, ou un résultat expérimental souhaité. Il s’agit alors d’une utilisation en boucle fermée.

Même pour ce genre d’utilisation du façonneur, il est extrêmement important d’en connaître les limites. En effet, nous avons vu au travers des exemples de l’approche naïve (voir la partie II.9.1 page 45) que les mises en forme programmées peuvent fortement différer de celle obtenues si on ne prend pas un minimum de précautions. Il ne faut donc pas analyser la solution vers laquelle une optimisation a convergé en regardant uniquement la programmation du façonneur. Il est important de caractériser de manière indépendante l’impulsion obtenue en sortie.

De plus, il faut être très prudent lors de cette caractérisation. Nous avons vu en

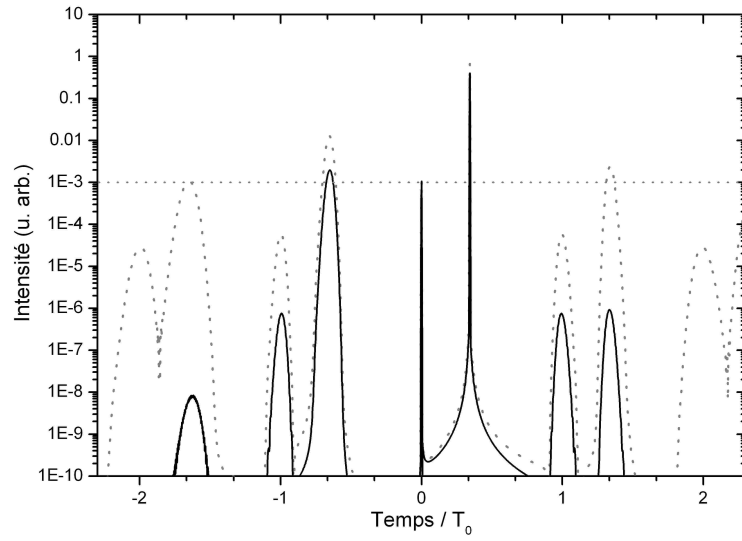


FIG. II.41 – Effet du couplage spatiotemporel sur les répliques. Intensité des répliques en tenant compte du couplage spatiotemporel (trait noir) et en le négligeant (pointillés gris). Lorsque l'on tient compte de tous les effets cumulés on constate une nette diminution des répliques.

effet lors de l'étude de la pixellisation que sous certaines conditions, l'impulsion principale (celle à laquelle on s'attend dans une approche naïve) est nettement moins intense que les répliques de pixellisation (voir la partie II.9.2.2 page 53). Or ces dernières sont loin de l'impulsion principale et on peut facilement les rater lors de la caractérisation si on n'a pas en tête la largeur de la fenêtre temporelle.

De même, de nombreux travaux sont effectués avec des façonneurs qui ne contrôlent que la phase spectrale des impulsions, pas l'amplitude. Ces façonneurs ont l'inconvénient de ne pas permettre la réalisation de mises en forme arbitraires, celles-ci nécessitant aussi le contrôle de l'amplitude. Ils permettent par contre de simplifier l'analyse des résultats car l'amplitude spectrale est constante, quelle que soit la mise en forme. Cela permet de comparer différents résultats d'optimisation en supposant l'énergie constante. Dans ce genre d'analyse, il est très important de se préoccuper du couplage spatio-temporel. En effet, ce dernier introduit une sorte de mise en forme en amplitude qui dépend de la mise en forme en phase programmée. Nous avons ainsi vu que des impulsions décalées temporellement ont une amplitude qui dépend fortement du délai introduit et de la taille spatiale du faisceau d'entrée. Pour pouvoir prétendre travailler à intensité constante, il faut restreindre les mises en forme réalisées par l'algorithme à une plage temporelle faible devant l'enveloppe temporelle gaussienne, ou encore introduire des corrections d'amplitude.

II.9.5 Possibilités de compensation

Les deux principales limitations présentées précédemment ne sont pas aussi rédhibitoires qu'il n'y paraît.

Tout d'abord, dans certains cas, la fonction de transfert du façonneur, obtenue sans la moindre correction, est tout à fait suffisante. Prenons par exemple une mise en forme s'étendant sur une fenêtre de quelques picosecondes, afin que l'atténuation temporelle gaussienne et sinus cardinal soit négligeable, et supposons que cette fenêtre ne contient aucune réplique de pixels ou d'interstices. Si on fait une somme de fréquence entre cette impulsion mise en forme et une impulsion de durée comparable, on obtiendra un signal somme ne présentant ni réplique, ni déformation temporelle. On peut remplacer dans cet exemple la somme de fréquence par n'importe quel système ayant un "temps de réponse" court devant la durée séparant l'impulsion mise en forme de ses plus proches répliques.

Ensuite, dans les cas où le système qui interagit avec l'impulsion mise en forme "intègre" sur un temps long (par exemple une interaction résonnante avec une transition à deux niveaux vers un niveau métastable), il est possible de jouer sur les paramètres du façonneur et ceux de l'impulsion d'entrée pour minimiser les répliques et les déformations temporelles.

Nous allons ainsi présenter comment éliminer les répliques dues aux interstices et comment tirer profit du couplage spatio-temporel pour atténuer celles dues aux pixels.

II.9.5.1 Compensation des répliques dues aux interstices

Dans la partie II.9.2 page 47, nous avons supposé que la transmission complexe G_n des différents pixels était la même pour tous et indépendante de la mise en forme programmée. On suppose le plus souvent que cette transmission équivaut à celle d'un pixel pour une tension d'adressage nulle. Cette supposition implique que les répliques dues aux interstices sont fixes à des délais multiples de T_0 . On peut alors calculer le masque à appliquer pour obtenir une mise en forme donnée tout en supprimant ces répliques non voulues[48]. Il est aussi possible de compenser "à la main" ces répliques en ajoutant à la mise en forme voulue une mise en forme correspondant à une impulsion à délai nul, avec une phase ϕ et une amplitude A :

$$M_w^{(1)}(\omega) = M_w(\omega) + Ae^{i\phi} \quad (\text{II.42})$$

Avec cette modification, on ajoute à chaque réplique due aux interstices, une réplique de pixellisation liée à $Ae^{i\phi}$. En ajustant ϕ et A , on peut annuler presque complètement les répliques des interstices. C'est cette méthode que nous avons employée le plus souvent lorsqu'il s'est avéré nécessaire de compenser l'effet des interstices.

La figure II.42 page ci-contre montre une telle compensation dans le cas d'une impulsion décalée temporellement de -2 ps. Le trait gris (partie haute) est l'intensité de corrélation avec une impulsion courte, en l'absence de compensation des répliques dues aux interstices. Le trait noir (partie basse), est la même intensité, avec compensation. La compensation a été faite avec une amplitude $A = 0,18$ et une phase $\phi = \pi$. On constate une disparition totale de la réplique due aux interstices.

Sur la figure II.43 page suivante la compensation est faite sur une impulsion à forte dérive de fréquence : $\phi^{(2)} = 7.10^5 \text{ fs}^2$. Le trait gris (partie haute) est l'intensité en l'absence de compensation et le trait noir (partie basse), avec compensation. La compensation a été réalisée avec une amplitude $A = 0,20$ et une phase $\phi = \pi$. Ici encore, la réplique a été totalement annulée.

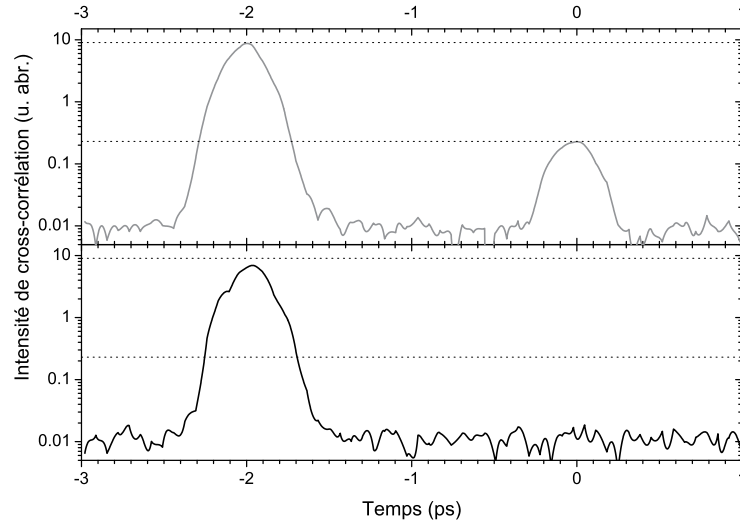


FIG. II.42 – Compensation des répliques dues aux interstices pour un délai de -2 ps. Sans compensation (trait gris), on observe sur l'intensité un pic à $t=0$. Avec compensation (amplitude $A = 0,18$ et phase $\phi = \pi$) (trait noir), ce pic a totalement disparu.

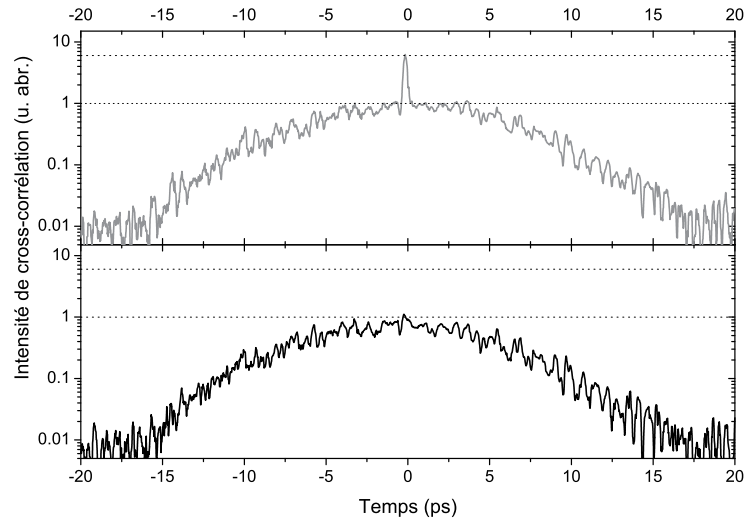


FIG. II.43 – Compensation des répliques dues aux interstices pour une phase quadratique $\phi^{(2)} = 7.10^5 \text{ fs}^2$. Sans compensation (trait gris), on observe à $t=0$ un pic étroit six fois plus intense que l'impulsion mis en forme. Avec compensation (amplitude $A = 0,2$ et phase $\phi = \pi$) (trait noir), ce pic a totalement disparu.

Comme ces exemples le montrent, cette approche "manuelle" fonctionne relativement bien. La compensation à introduire semble conforter l'idée que la transmission des interstices correspond à la mise en forme pour une tension nulle. En effet, pour une tension nulle, on obtient une mise en forme proche de l'identité et ici, la compensation est optimale est toujours obtenue pour $\phi = \pi$.

Cependant, l'amplitude de la compensation à apporter varie en fonction de la mise en forme programmée. Ceci est dû au comportement des cristaux liquides à l'interface entre un pixel et l'interstice voisin (déjà évoqué dans la partie II.9.2.1 page 51). Ceux-ci passent continûment de l'orientation imposée au niveau du pixel à l'orientation de repos au niveau de l'interstice. Cette transition dépend fortement de la tension qui est appliquée au pixel et donc du type de mise en forme. Il n'existe pour l'instant pas de modèle permettant de calculer l'importance de la réplique due aux interstices pour une mise en forme donnée.

Pour compenser plus précisément cette réplique, ainsi que les autres déformations, il est possible d'avoir recours à des algorithmes d'optimisation itératifs. Si un tel modèle pour la transition à l'interface pixel/interstice existait, nous pourrions tout à fait compenser la réplique sans recourir à des algorithmes itératifs.

II.9.5.2 Atténuation des répliques de pixellisation

Nous venons de voir qu'il était relativement facile de compenser les répliques dues aux interstices. Il serait illusoire de croire que l'on peut faire de même avec les répliques dues à la pixelisation²⁶. Nous avons vu en effet qu'à cause de la pixellisation nous ne contrôlons le champ mis en forme qu'à l'intérieur d'une fenêtre de $T_0 \simeq 35,8$ ps et que justement les répliques sont à l'extérieur de cette fenêtre.

Il est cependant possible de jouer sur le couplage spatio-temporel pour atténuer ces répliques. En effet, le couplage spatio-temporel dans l'approximation gaussienne se traduit par une multiplication du champ de sortie par une gaussienne temporelle. On peut varier à volonté la largeur et la position temporelle de cette gaussienne en jouant respectivement sur le diamètre du faisceau d'entrée (voir l'équation II.41 page 58) et le filtrage spatial en sortie (voir la remarque page 59). Expérimentalement, il suffit donc de changer l'ouverture du diaphragme d'entrée et la position du diaphragme de sortie pour contrôler cette enveloppe gaussienne. L'idée est de centrer cette gaussienne sur la mise en forme souhaitée en ajustant au mieux sa largeur pour écraser les répliques sans toutefois trop déformer la mise en forme.

II.10 Applications

Nous avons utilisé le façonneur que nous venons de décrire en détail pour de nombreuses expériences. Les objectifs que nous nous étions fixés lors de sa construction étaient une haute résolution spectrale et un contrôle complet du champ (phase et amplitude). La grande résolution spectrale se traduit par une fenêtre temporelle disponible pour la mise en forme très large (de l'ordre de 35 ps). Elle se traduit aussi par la possibilité de réaliser des mises en forme très complexes.

²⁶La solution existe mais n'est pas simple : remplacer le masque par un masque non pixellisé !

Nous allons présenter ici différents résultats qui sans être exhaustifs nous semblent caractéristiques des possibilités de notre façonneur. Nous montrerons tout d'abord les premières expériences visant à déterminer les capacités et les limites de notre façonneur. Nous présenterons ensuite des mises en forme utilisées pour du contrôle cohérent dans l'atome de rubidium (détaillé dans le chapitre IV page 95). Puis nous décrirons succinctement les résultats obtenus sur la compensation de dispersion résonnante (traitée en détail dans la thèse de Jean-Christophe Delagnes[97] ainsi que dans la référence[98]). Enfin, nous montrerons les possibilités de la ligne repliée II.2.3 page 23.

II.10.1 Expériences de Démonstration

Lors des expériences de caractérisation du façonneur, nous avons pu montrer une transmission en puissance de plus de 60% pour l'ensemble du dispositif avec un facteur d'extinction²⁷ de 20 dB. Avec des corrections appropriées (voir la partie II.9.2 page 47), il est même possible d'atteindre 30 dB.

Pour montrer l'intervalle temporel disponible pour réaliser les mise en formes, nous avons programmé différents délais. Cela correspond à des fonctions de transfert du type $H(\omega) = e^{-i\phi^{(1)}(\omega-\omega_0)}$, où ω_0 est la pulsation centrale du laser et $\phi^{(1)}$ le délai souhaité. La figure II.44 rappelle les résultats obtenus pour différents délais compris entre -30 ps et 30 ps. Comme cela a été expliqué dans la partie II.9.3.3 page 57, on voit apparaître une

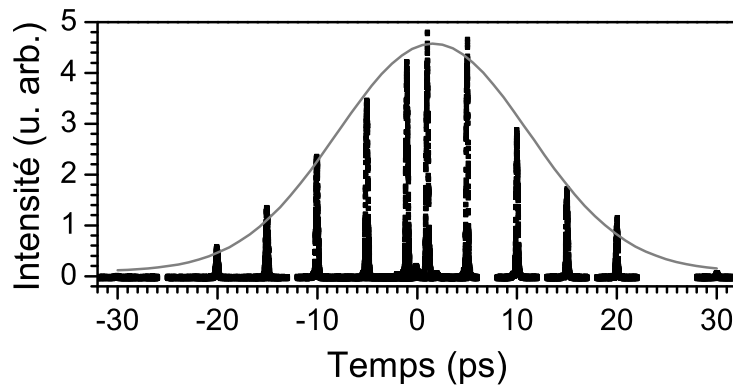


FIG. II.44 – Mise en évidence de l'enveloppe temporelle. Intensité de cross-corrélation (carrés noirs) pour différents délais introduits par le façonneur. Une enveloppe temporelle gaussienne (trait gris) est clairement visible (voir la partie II.9.3.3 page 57).

enveloppe gaussienne de largeur à mi-hauteur en intensité de 23 ps. Cette durée dépend fortement du diamètre du faisceau incident et peut être augmentée en utilisant un faisceau plus gros. Avec le faisceau utilisé ici (2 mm de diamètre), on peut donc réaliser des mises en forme s'étalant sur une vingtaine de picosecondes sans trop perdre en intensité. Cet intervalle temporel est plus faible que la fenêtre temporelle maximale $T_0 \simeq 35,8$ ps

²⁷Il s'agit du rapport entre l'intensité transmise maximale et l'intensité transmise minimale.

de notre façonneur en l'absence de couplage spatio-temporel. Il reste cependant très grand devant ce qui est réalisable avec les façonneurs que l'on trouve dans la littérature[28, 95, 48].

A titre d'exemple, la figure II.45 donne les tensions appliquées aux deux barrettes de cristaux liquides (SLM1 et SLM2) pour obtenir un délai de -1 ps.

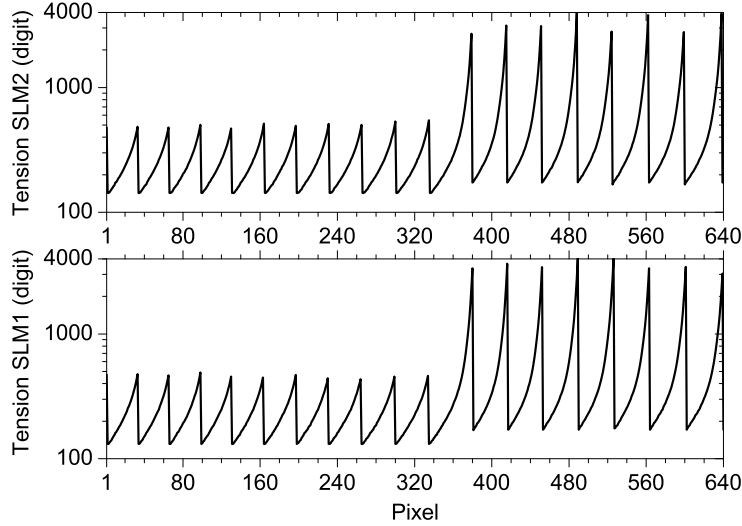


FIG. II.45 – Tensions (en digit : 4096 niveaux de gris) appliquées aux deux barrettes de cristaux liquides (SLM1 et SLM2) pour l'obtention d'un délai de -1 ps.

Afin d'illustrer la réalisation de mise en formes complexe, nous avons programmé les mises en formes suivantes.

- (a) Trois paires d'impulsions de même amplitude. Les intensités relatives entre les paires sont de $1/2$, 1 et $1/3$. Ceci correspond à une fonction de transfert :

$$H_{(a)}(\omega) \propto \frac{e^{-i\phi_{11}^{(1)}(\omega-\omega_0)} + e^{-i\phi_{12}^{(1)}(\omega-\omega_0)}}{2} + e^{-i\phi_{21}^{(1)}(\omega-\omega_0)} + e^{-i\phi_{22}^{(1)}(\omega-\omega_0)} \\ + \frac{e^{-i\phi_{31}^{(1)}(\omega-\omega_0)} + e^{-i\phi_{32}^{(1)}(\omega-\omega_0)}}{3} + A_{(a)} e^{i\phi_{(a)}}$$

où les τ_{ij} sont les positions des différentes impulsions de la séquence, ω_0 la pulsation centrale du laser et $A_{(a)} e^{i\phi_{(a)}}$ la compensation des répliques dues aux interstices, ajustée à la main.

- (b) Une impulsion carré de 700 fs de durée, soit une fonction de transfert :

$$H_{(b)}(\omega) \propto \text{sinc}((\omega - \omega_0)T/2) e^{-i\phi^{(1)}(\omega-\omega_0)} + A_{(b)} e^{i\phi_{(b)}}$$

avec T la durée de l'impulsion carré et $\phi^{(1)}$ sa position.

(c) Trois impulsions, la première avec une phase cubique, la seconde une phase plate et la dernière une phase quadratique. Soit :

$$H_{(c)}(\omega) \propto e^{-i\phi_1^{(1)}(\omega-\omega_0)-i\frac{\phi^{(3)}}{6}(\omega-\omega_0)^3} + \mathbb{1} + e^{-i\phi_2^{(1)}(\omega-\omega_0)-i\frac{\phi^{(2)}}{2}(\omega-\omega_0)^2} + A_{(c)}e^{i\phi_{(c)}}$$

La figure II.46 présente l'intensité de corrélation obtenue pour ces trois exemples. Les

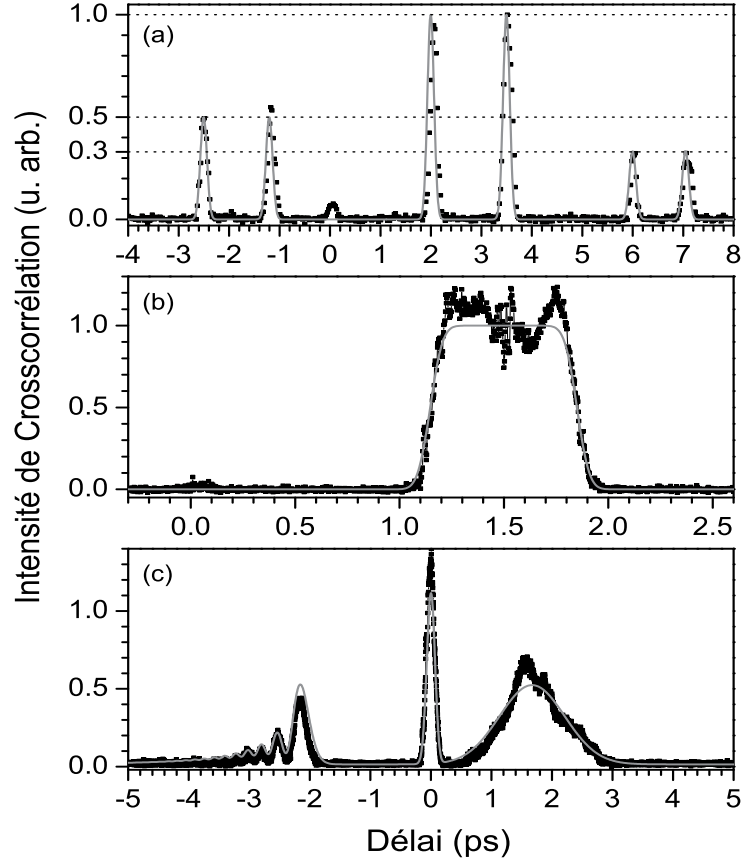


FIG. II.46 – Cross-corrélations théoriques (traits gris) et expérimentales (carrés noirs). (a) 3 paires d'impulsions d'intensités relatives 1/2, 1 et 1/3, (b) impulsion carré de 700 fs, et (c) trois impulsions avec une phase cubique (9.10^7 fs^3), une phase plate et une phase quadratique (5.10^4 fs^2).

carrés noirs correspondent aux résultats expérimentaux et les traits gris aux mises en forme demandées²⁸. Dans tous les cas, l'accord est excellent, bien que la compensation des défauts inhérents à ce type de façonneur soit ici très basique. En effet, nous n'avons pas cherché à compenser l'effet de l'enveloppe temporelle gaussienne (couplage spatio-temporel) et nous avons compensé la réplique à délai nul simplement en ajoutant une impulsion à délai nul, en opposition de phase (voir la partie II.9.5.1 page 62). Pour de meilleurs résultats, les

²⁸C'est à dire aux fonctions de transfert présentées plus haut. Nous ne tenons absolument pas compte des défauts de la ligne, à savoir la pixellisation et le couplage spatio-temporel pour ces courbes théoriques.

effets de pixellisation et de couplage spatio-temporel devraient être compensés de manière plus fine. Cela peut être fait en utilisant une procédure itérative[48].

De plus, on observe que la plage temporelle de mise en forme est particulièrement large. Des mises en formes s'étalant sur plus de 10 ps sont facilement réalisables. Des mises en forme plus longues peuvent être générées mais il faut alors corriger l'amplitude temporelle pour tenir compte de l'enveloppe gaussienne.

Dans le cas de l'impulsion carré (figure II.46 (b)), on observe des oscillations sur le plateau qui sont entièrement dues à la faible largeur spectrale de l'impulsion à mettre en forme. La mise en forme spectrale est donc tronquée. Pour des impulsions carrées plus longues, on observe des oscillations moindres.

A titre d'exemple, les tensions appliquées aux deux barrettes de cristaux liquides (SLM1 et SLM2) pour la réalisation de cette impulsion carrée sont données sur la figure II.47. Les tensions sont données en digit, c'est à dire en niveaux de gris compris entre 1 et 4096 (4096 digits correspondent à une tension maximale de 8V).

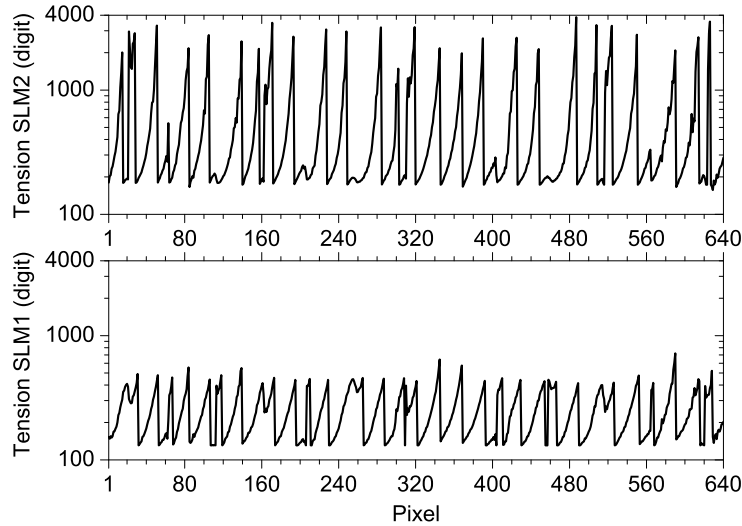


FIG. II.47 – Tensions (en digit : 4096 niveaux de gris) appliquées aux deux barrettes de cristaux liquides (SLM1 et SLM2) pour l'obtention de l'impulsion carré de la figure II.46(b).

La figure II.48 page suivante présente un autre exemple où seule la mise en forme en amplitude a été utilisée. Le spectre de notre amplificateur régénératif (ronds gris) est modifié pour obtenir un spectre parfaitement gaussien (carré noir). Pour ce faire, nous avons tenu compte du spectre d'entrée dans la fonction de transfert programmée. L'accord avec la mise en forme demandée (trait gris) est excellent. Aucune itération n'a été utilisée pour cette mise en forme et aucune correction des défauts n'a été nécessaire. Il est tout à fait possible de combiner cette mise en forme en amplitude avec une phase arbitraire. Cela peut permettre d'avoir une impulsion modèle (spectre gaussien avec phase polynômiale par exemple) qui correspondent tout à fait au type d'impulsions utilisées dans de nombreuses simulations.

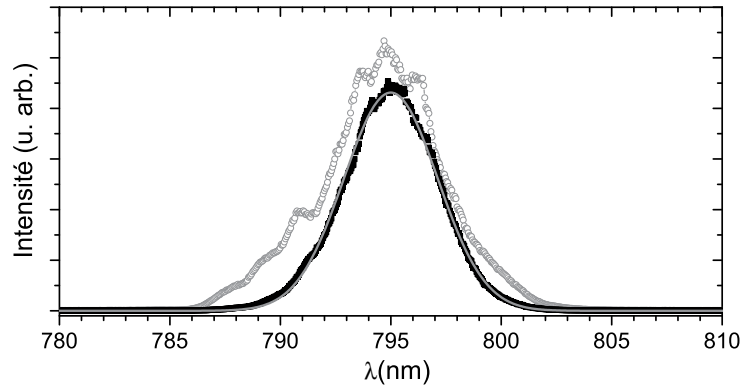


FIG. II.48 – Correction d'intensité spectrale : spectre non corrigé (ronds gris), spectre gaussien expérimental (carrés noirs) et spectre gaussien théorique (trait gris).

II.10.2 Application au contrôle cohérent

D'autres résultats ont été obtenus avec cet appareil, notamment pour des expériences de contrôle cohérent dans le rubidium[47] présentées dans le chapitre IV page 95. Des impulsions à forte dérive de fréquence avec des sauts de phase de π ont ainsi été synthétisées. Des phases quadratiques de plus de $2 \cdot 10^5 \text{ fs}^2$ ont été obtenues sans déformation du profil temporel. Le critère de Nyquist (variation de la phase inférieure à π d'un pixel à son voisin) permet de donner une phase quadratique maximale : $\phi_L^{(2)} = T_0^2 / [2\pi N] \simeq 3,2 \cdot 10^5 \text{ fs}^2$ (T_0 étant la fenêtre temporelle accessible et N le nombre de pixels). En dessous de cette limite, les profils temporels sont très bien reproduits. Des phases quadratiques plus importantes (jusqu'à $9 \cdot 10^5 \text{ fs}^2$), ont été aussi obtenues, mais avec des durées pour les impulsions étirées plus faibles que les durées théoriques. Cette réduction de la durée est entre autres causée par l'enveloppe temporelle gaussienne due au couplage spatio-temporel.

Dernièrement, nous avons généré des séquences d'impulsions[50, 51] (l'une limitée par transformée de Fourier, l'autre à forte dérive de fréquence) séparées de plusieurs picosecondes avec un contrôle interférométrique du délai. Ces résultats sont présentés en détail dans le chapitre IV.

II.10.3 Compensation de dispersion résonnante

Nous avons réalisé à l'aide du façonneur des expériences de compensation de dispersion résonnante[98]. Nous allons ici présenter brièvement le principe de cette expérience ainsi que les résultats observés. Tout ceci est traité plus en détail dans la thèse de Jean Christophe Delagnes[97] qui s'est intéressé tout particulièrement aux effets de propagation d'impulsions ultracourtes.

Les problèmes de dispersion sont très bien connus dans le domaine de l'optique ultrarapide. Habituellement, on s'intéresse à la propagation d'une impulsion courte dans un milieu transparent. La pulsation centrale de cette impulsion est donc loin de toute résonance du milieu. On observe alors en sortie du milieu une impulsion déformée à cause de la

phase spectrale non linéaire $\phi_{nl}(\omega)$ accumulée lors de la propagation. Le principal effet de cette phase est la dispersion de la vitesse de groupe (GVD) qui introduit une dérive de fréquence temporelle dans l'impulsion de sortie (chirp). De nombreux dispositifs simples ont été développés pour compenser cette dispersion de la vitesse de groupe[99, 100]. Dans la plupart des cas, la compensation des ordres les plus bas de $\phi_{nl}(\omega)$ est largement suffisante.

Ici, nous nous intéressons à la propagation d'une impulsion ultracourte résonnante avec une transition du milieu. Cette situation est très fréquente dans les expériences de spectroscopie ultra rapide et dans certaines expériences de dynamique de paquets d'onde[101]. Comme l'impulsion est courte, son spectre est beaucoup plus large que la largeur spectrale de la transition. L'absorption ne concerne donc qu'une partie extrêmement étroite du spectre et le principal effet de la propagation est un effet dispersif. On peut montrer que la phase accumulée lors de la propagation s'écrit

$$\phi_{nl}(\omega) = -\frac{\omega}{c}l(n(\omega) - 1) \simeq \frac{\alpha l \Delta_d}{\omega - \omega_{at}} \quad (\text{II.43})$$

où α est le coefficient d'absorption à résonance, l la longueur du milieu, Δ_d la largeur Doppler de la transition et ω_{at} la pulsation de la transition.

Cette phase présente un comportement hyperbolique et on ne peut absolument pas développer en série de Taylor autour de ω_{at} . On ne peut donc plus ici compenser la dispersion à l'aide de dispositifs habituels qui ne permettent d'introduire que des phases quadratiques et cubiques.

Notre façonneur permet lui d'introduire une phase $\phi_{cp}(\omega)$ opposée à celle accumulée lors de la propagation :

$$\phi_{cp}(\omega) = -\phi_{nl}(\omega) \simeq -\frac{\alpha l \Delta_d}{\omega - \omega_{at}} \quad (\text{II.44})$$

Nous avons pour cela utilisé les impulsions délivrées par notre oscillateur (150 fs de durée, longueur d'onde centrale 795 nm, 9,7 nm de largeur spectrale). Celles-ci traversent le façonneur ainsi qu'un four contenant du rubidium, comme cela est schématisé sur la figure II.49.

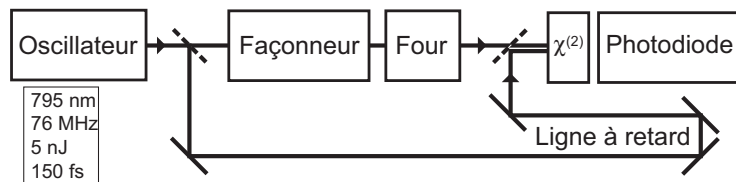


FIG. II.49 – Dispositif expérimental de compensation de dispersion résonnante. L'impulsion déformée par le façonneur et le four est détectée par corrélation avec une impulsion courte.

La transition résonnante est la $5s^2S_{1/2} \mapsto 5p^2S_{1/2}$ à $\lambda_{at} = 794,75$ nm. Le four est porté à une température de 160° ce qui correspond à une pression de rubidium de l'ordre de $7,9 \cdot 10^{-3}$ mbar. A la sortie, l'impulsion transmise est détectée par corrélation avec une impulsion courte²⁹.

²⁹Celle-ci provient d'une partie du faisceau prélevée directement en sortie de l'oscillateur

La figure II.50 présente les résultats obtenus lors de cette expérience. Sur les trois sous-

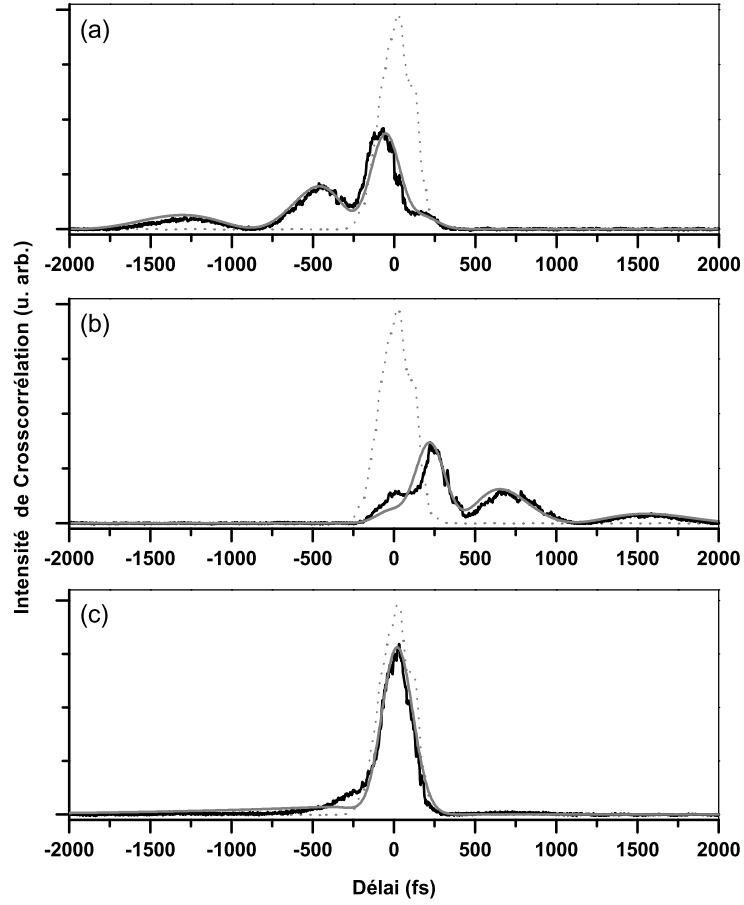


FIG. II.50 – Compensation de la dispersion par une vapeur atomique résonnante et dense : impulsion de départ en pointillés gris, impulsion déformée en noir pour l'expérience, en gris pour la théorie. (a) impulsion déformée par la vapeur, (b) impulsion pré-compensée par le façonneur, (c) impulsion après la traversée du façonneur et de la vapeur.

figures (a), (b) et (c), les pointillés gris représentent la corrélation de l'impulsion n'ayant pas subi de déformation, le trait noir, la même impulsion après déformation et le trait gris correspond à la déformation théorique. Sur la sous-figure (a) l'impulsion est déformée par la traversée du four, le façonneur étant inactif. L'impulsion est fortement étirée et des rebonds sont visibles sur plusieurs picosecondes. Sur la sous-figure (b), le four est vide et seul le façonneur est actif. Ce dernier été programmé de façon à simuler la dispersion opposée à celle du four. On observe une déformation symétrique par renversement du temps par rapport à celle induite par le four. Enfin, sur la sous-figure (c), le façonneur est actif et le four plein. On observe une impulsion très peu déformée par rapport à l'impulsion de départ.

On a donc réussi à compenser la dispersion résonnante introduite par la vapeur ato-

mique, ce qui n'est pas réalisable par des dispositifs habituels. On note toutefois sur l'impulsion compensée un épaulement aux temps négatifs et un léger rebond autour de 700 fs. La compensation n'est pas parfaite et est limitée par la résolution de notre façonneur. En effet, la phase de compensation (voir l'équation II.44 page 70) diverge en ω_{at} : sa pente est d'autant plus raide que l'on s'approche de ω_{at} . On a donc une zone autour de ω_{at} où l'évolution de la phase est trop rapide pour être reproduite par le façonneur. Autour de ω_{at} , le façonneur ne peut reproduire la phase de compensation. On observe dans cette zone des sauts de phase très importants d'un pixel à l'autre. Ces sauts très rapprochés vont introduire un trou en intensité que l'on observe clairement sur le spectre en sortie du façonneur (voir figure II.51 (a)).

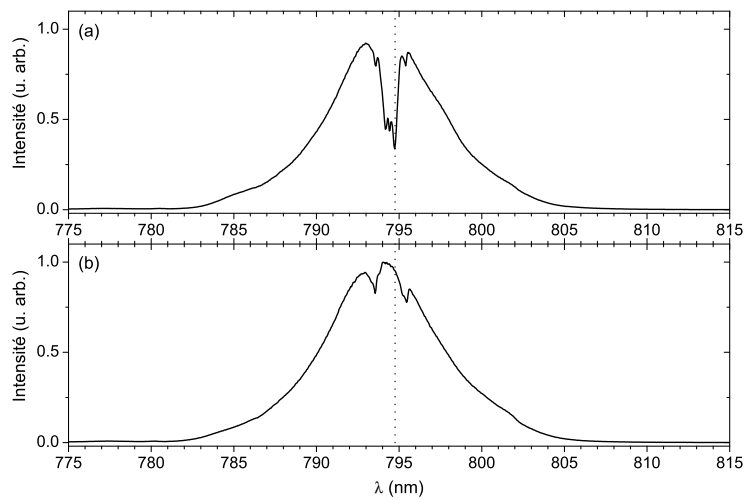


FIG. II.51 – Spectres en sortie du façonneur pour une compensation de la dispersion résonnante : (a) pour toutes les longueurs d'onde ; (b) pour une phase constante autour de λ_{at} . Les pointillés indiquent la position de λ_{at} .

On peut comprendre l'origine de ce trou en se rappelant que le masque programmé est convolué par le profil spatial des composantes spectrales dans le plan de Fourier (voir la partie II.9.3.3 page 57). Le profil spatial de chaque composante spectrale moyenne l'effet de plusieurs pixels : comme ils introduisent ici des phases très différentes le résultat global est une diminution de l'amplitude.

Pour éviter ce problème, on peut programmer sur toute la zone une phase constante. Cela évite la perte de transmission, comme on peut le voir sur la figure II.51 (b), tout en offrant une compensation presque aussi efficace. Dans le cas précédent où on cherche à reproduire la phase de compensation à proximité de ω_{at} , cela se traduit par une forte atténuation des composantes spectrales proches de ω_{at} . Avec cette nouvelle approche, ces composantes ne sont pas atténuées, mais elles ne sont pas non plus compensées. Elles sont donc transmises, mais fortement étalées temporellement.

II.10.4 Façonneur replié

Les résultats présentés ici ont été obtenus avec une demi-ligne à dispersion nulle, comme celle présentée sur la figure II.11(b), page 24 dans la partie II.2.3. Nous avons pour cela modifié notre façonneur en plaçant un miroir juste derrière les deux masques à cristaux liquides. L'intérêt de cette ligne repliée est double. Tout d'abord, cela diminue par deux l'encombrement de la ligne réduisant aussi son coût : un seul réseau et un seul miroir cylindrique sont nécessaires. Ensuite, cela rend le réglage de la ligne beaucoup plus simple. En effet, les degrés de liberté de réglage sont grandement réduits. De plus, le réglage le plus critique et de loin le plus difficile de la ligne standard est l'orientation des deux réseaux : les angles doivent être parfaitement identiques. Ici, on utilise le même réseau et les angles sont automatiquement identiques.

Le double passage double les déphasages qui peuvent être introduits, ce qui n'est pas vraiment intéressant vu que 2π suffisent. Cela réduit la sensibilité d'un facteur deux ce qui peut éventuellement poser problème.

La figure II.52 présente deux mises en formes spectrales obtenues avec le façonneur habituel (trait noir) et le façonneur replié (trait gris). L'exemple II.52(a) correspond à deux impulsions séparées de 1 ps alors que le II.52(b) est un trou spectral de 20 pixels de large (soit 1,2 nm) centré à 794,3 nm. L'accord est très bon dans les deux cas et les petites différences observées sont certainement dues à la modification de l'alignement dans le spectromètre et à une dérive du spectre du laser.

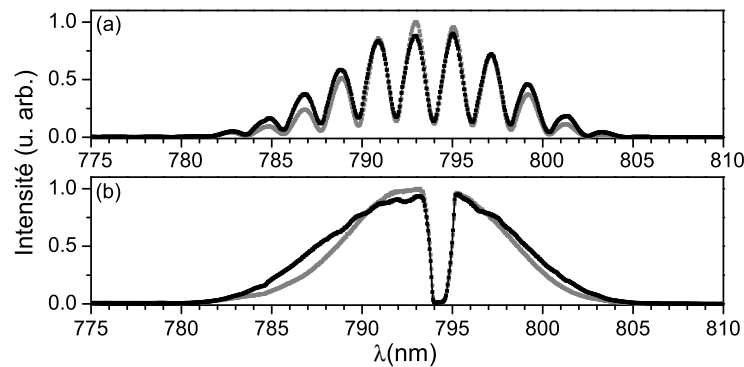


FIG. II.52 – Mises en forme avec le façonneur normal (trait gris) et replié (trait noir) : (a) séquence de deux impulsions séparées de 1 ps, (b) trou spectral de 1,2 nm centré à 794,3 nm.

Nous avons aussi réalisé toute une série de mise en formes temporelles, avec notamment des phases quadratiques et cubiques importantes. Nous avons dans la plupart des cas obtenu exactement les mêmes résultats qu'avec la ligne normale. La figure II.53 page suivante présente le résultat obtenu pour différents délais réalisés avec la ligne repliée. L'intensité de corrélation avec une impulsion courte est représentée par les carrés noirs. On obtient des résultats tout à fait similaires à ceux obtenus avec la ligne standard (voir la figure II.44 page 65). La principale différence est la largeur de l'enveloppe gaussienne qui est ici deux fois plus faible (plus exactement 11,3 ps). Ceci est simplement dû à l'utilisation

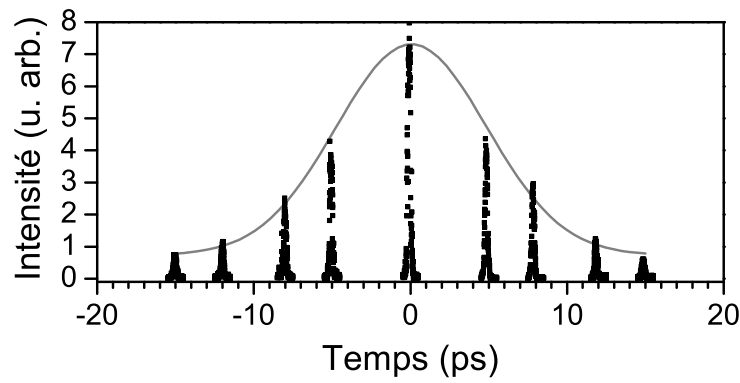


FIG. II.53 – Intensité de corrélation (carrés noirs) pour différents délais introduits avec le façonneur replié. L’enveloppe gaussienne (trait gris) de largeur à mi hauteur de 11,3 ps est due au couplage spatio-temporel. Elle correspond à un diamètre de faisceau d’entrée légèrement inférieur à 1 mm.

d’un faisceau de diamètre deux fois plus petit (1 mm de diamètre en largeur à mi-hauteur en intensité).

La disposition repliée du façonneur est donc particulièrement intéressante car elle permet d’obtenir des performances comparables à celles du façonneur normal, tout en simplifiant le dispositif et en diminuant son encombrement et son coût de fabrication (un seul réseau et un seul miroir cylindrique).

Chapitre III

AUTRES DISPOSITIFS DE MISE EN FORME

Outre le développement et l'utilisation du façonneur haute résolution à 800 nm (chapitre II), j'ai aussi travaillé à la mise en forme d'impulsions visibles accordables. Ces impulsions sont générées par l'un des deux NOPA décrits dans le chapitre I.

Elles présentent de nombreux intérêts. Tout d'abord, elles permettent d'étendre la gamme de longueurs d'ondes accessibles à l'ensemble du spectre visible. Ceci rend possible l'étude et le contrôle de nombreux systèmes atomiques possédant des transitions dans cette gamme. Ensuite, on peut réaliser des expériences pompe-sonde avec des impulsions dans des domaines spectraux complètement différents (pompe-sonde à deux couleurs). Enfin, un intérêt tout particulier des NOPA (qui les distingue notamment des OPA colinéaires) est la durée de ces impulsions visibles[57]. En effet, correctement recomprimées, ces impulsions accordables offrent des durées nettement plus courtes (entre 30 fs et 6 fs selon le type de NOPA) que l'impulsion utilisée pour pomper le NOPA (durée de l'ordre de 130 fs le plus souvent). La brièveté de ces impulsions est un paramètre déterminant pour la résolution temporelle des expériences de pompe-sonde dans le visible.

La recompression de ces impulsions est donc une étape clef sur laquelle de nombreux groupes se sont penchés[102, 103, 104, 80, 105, 60, 106, 58]. Plusieurs techniques sont utilisées, plus ou moins raffinées. Une solution souvent employée est le compresseur à prismes en double passage[60, 104, 107, 58]. Une autre solution plus récente, qui a séduit de nombreux groupes, est l'utilisation de miroirs à dispersion contrôlée ou miroirs "chirpés"[80, 105]. Ces miroirs présentent un revêtement diélectrique multicouche spécialement conçu pour introduire une dispersion négative. L'idée est que les courtes longueurs d'ondes sont réfléchies par les couches externes et les grandes longueurs d'ondes par les couches profondes. Lors d'une réflexion sur un tel miroir, les grandes longueurs d'ondes parcourent un chemin plus grand que les courtes longueurs d'ondes, d'où une dispersion négative.

Au delà de la simple recompression des impulsions, il peut s'avérer très intéressant de contrôler la forme de ces impulsions visibles accordables. Cette capacité de mise en forme peut servir à des expériences de contrôle cohérent[46, 19, 12] dans le visible ou bien à l'optimisation d'impulsions dans l'UV obtenues à partir des impulsions visibles par doublage ou somme de fréquence.

Nous détaillerons ici deux approches différentes mises en œuvre pour réaliser à la fois la recompression et le façonnage des impulsions issues du NOPA. La première repose sur l'utilisation d'un miroir déformable à l'intérieur du compresseur à prismes en double passage. Elle a été développée lors d'une collaboration avec l'équipe d'Albert Stolow au

STEACIE Institute. La deuxième a été développée à Toulouse en partenariat avec Daniel Kaplan¹. Elle utilise un filtre acousto-optique dispersif programmable ou AOPDF (pour "Acousto-Optic Programmable Dispersive Filter") spécialement conçu pour la recompression et la mise en forme d'impulsions visibles accordables[41].

Il n'y a pas eu avec ces deux façonneurs un travail de développement aussi important que pour le façonneur haute résolution présenté au chapitre II. Il s'agit plus ici d'utilisation que de développement.

III.1 Miroirs déformables

Nous présentons ici le compresseur à prismes et à miroir déformable que j'ai installé et utilisé au STEACIE Institute, à Ottawa. Nous rappellerons d'abord les principales caractéristiques du NOPA en insistant sur la nécessité d'un système de compression en sortie. Nous présenterons ensuite le principe du compresseur à miroir déformable en montrant son champ d'applications.

III.1.1 Rappels sur le NOPA

L'amplificateur paramétrique optique ou NOPA (Non-collinear Optical Parametric Amplifier) est une source d'impulsions ultrabrèves, accordable sur l'ensemble de la gamme visible. Des durées de l'ordre de 20 fs sont couramment atteintes avec des dispositifs commerciaux. La gamme spectrale d'accordabilité s'étend de 510 nm à 700 nm. L'énergie typique obtenue est de l'ordre de la dizaine de μJ .

Le principe même des NOPA rend nécessaire l'utilisation d'un compresseur en sortie. En effet, le NOPA repose sur la génération d'un continuum de lumière blanche : une partie de ce continuum va être amplifiée (par différence de fréquence dans un cristal non-linéaire, comme cela est brièvement expliqué dans la partie I.4 page 7) pour donner l'impulsion large bande. On retrouve donc dans les impulsions de sortie la phase spectrale du continuum plus celle introduite par les optiques traversées dans les différents étages d'amplification (lentilles, cristaux,...). La phase spectrale totale est particulièrement importante et complexe. Sans compresseur, les impulsions obtenues ont des durées de l'ordre de quelques centaines de fs.

Plusieurs types de compresseurs sont disponibles. Les miroirs chirpés, les compresseurs à prismes et les compresseurs à réseaux sont les plus courants. Le NOPA avec lequel nous avons travaillé au STEACIE Institute (Clark-MXR NOPA) disposait d'un compresseur à prismes. Son principe est rappelé sur la figure III.1 page suivante. Les différentes composantes spectrales sont séparées par un premier prisme et suivent des chemins optiques différents. Un deuxième prisme tête-bêche, compense la séparation angulaire. Un miroir de repli renvoie les différentes composantes dans le compresseur. On obtient en sortie un faisceau recomposé avec une phase spectrale dépendant de la distance entre les prismes et de l'épaisseur de verre traversée. On peut ainsi introduire une phase spectrale quadratique de signe arbitraire. Il est à noter que les différentes composantes sont séparées au niveau du

¹de la société Fastlite.

miroir de repli. Avec ce type de compresseur, on obtient des durées entre 30 fs et 15 fs,

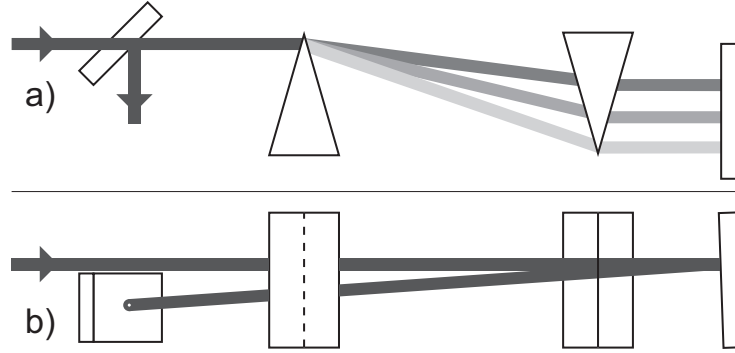


FIG. III.1 – Compresseur à prisme : a) vue de dessus, b) vue de côté.

selon le type de NOPA.

Bien que très courtes, ces impulsions ne sont pas les plus courtes possibles : elles ne sont pas limitées par transformée de Fourier. En effet, les compresseurs à prismes ne compensent que les ordres les plus bas de la phase spectrale : phase quadratique dans la plupart des cas ou phase quadratique et phase cubique pour des compresseurs spécifiques, non accordables. Or les impulsions issues du NOPA présentent des phases relativement complexes et il faut compenser, en plus de la phase quadratique, la phase résiduelle pour se rapprocher des durées les plus courtes. Le plus souvent, on utilise pour cela un miroir déformable placé au sein du compresseur à prismes ou à réseaux. Avec ce genre de compresseur et des NOPA spécialement optimisés, plusieurs équipes ont obtenu des impulsions de quelques femtosecondes[103, 102, 108].

III.1.2 Compresseur à miroir déformable

Le compresseur à miroir déformable que nous avons réalisé est très proche du compresseur à prismes de la figure III.1. La seule différence est que le miroir de repli est remplacé par un miroir à membrane déformable, comme cela est schématisé sur la figure III.2 page suivante. Ce miroir (Oko 19-channel linear micromachined deformable mirrors) est constitué d'une membrane suspendue à 19 actionneurs piézo-électriques. La membrane offre une surface de 11x39 mm. Chaque actionneur est commandé en tension (entre 0 et 280 V codé sur 8 bits) et permet un déplacement local de la membrane (maximum de 6 μm). Pour un jeu de tensions appliqué, le miroir présente un profil $\mathcal{M}(x)$, x étant la direction d'étalement des composantes spectrales (voir la figure III.2 b)). Les différentes composantes spectrales parcourent donc un trajet plus ou moins long selon la position x où elles se réfléchissent. En notant $\omega(x)$ la dispersion au niveau du miroir, on a une phase spectrale introduite par le miroir :

$$\phi_{\mathcal{M}}(\omega) = \frac{2\pi}{\lambda(x)} 2\mathcal{M}(x) = \frac{4\pi}{c} \omega(x) \mathcal{M}(x) \quad (\text{III.1})$$

Grâce au miroir déformable, on peut ajouter à la phase quadratique introduite par le compresseur une phase arbitraire $\phi_{\mathcal{M}}(\omega)$.

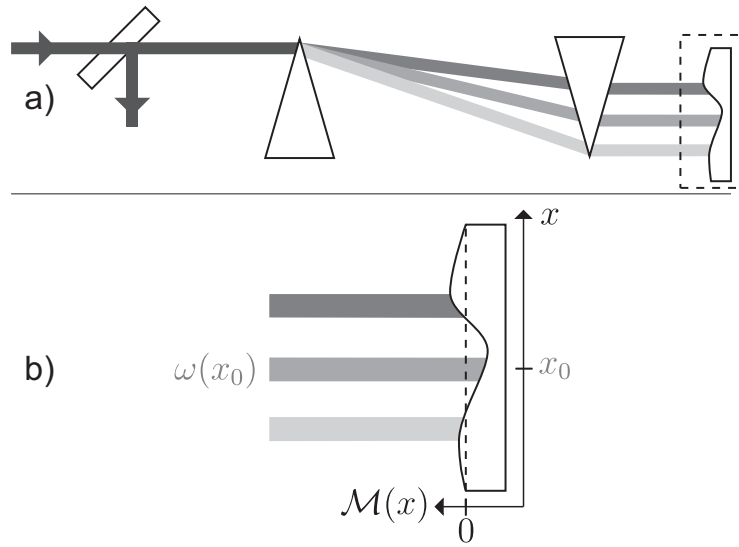


FIG. III.2 – Compresseur à miroir déformable : a) vue de dessus, b) détail du miroir et notations.

Outre le fait qu'il ne contrôle que la phase, le façonneur ainsi obtenu est très différent du façonneur haute résolution. Tout d'abord, sa résolution est très faible. Cela est dû en grande partie au faible nombre de paramètres libres (19 actionneurs seulement). S'ajoute à cela le fait que les différentes composantes spectrales ne sont pas proprement imagées au niveau du miroir déformable : chaque composante spectrale possède au niveau du miroir déformable le même profil spatial que le faisceau incident. Le masque de phase $\phi_{\mathcal{M}}(\omega)$ est donc convolué par un profil spatial relativement large, ce qui peut dégrader la résolution de ce façonneur.

Par contre, la phase $\phi_{\mathcal{M}}(\omega)$ est une fonction continue de la pulsation. En effet la membrane se déforme de façon continue d'un actionneur à l'autre. Cela présente un avantage considérable par rapport à notre façonneur pixellisé : aucune réplique n'est générée.

Cependant, les phases programmables sont beaucoup plus restreintes tant en forme qu'en excursion. Ainsi, contrairement à un masque pixellisé, on ne peut pas reproduire de saut abrupt. Il est donc impossible d'introduire un saut de phase, par exemple. Cela limite aussi la gamme de phases que l'on peut reproduire, même dans le cas d'une phase continue. En effet, avec un masque pixellisé, on peut toujours replier la phase à programmer dans l'intervalle $[0, 2\pi]$: il est ainsi possible de reproduire des phases présentant des excursions importantes. Ici au contraire, la membrane interdit de tels repliements et on est limité par l'excursion maximale que la membrane peut supporter.

Ces deux limitations ne sont pas gênantes. En effet, le but du façonneur à membrane est ici simplement de corriger une phase résiduelle. En effet, le plus gros de la phase est déjà corrigé par les prismes et il n'est pas nécessaire d'introduire des phases très importantes. Enfin, la phase résiduelle à corriger ne présente pas de sauts abrupts. Toutefois, si on souhaite aussi utiliser ce façonneur pour mettre en forme l'impulsion corrigée, on sera fortement limité.

Une dernière différence qui a son importance est le taux de rafraîchissement. Alors

que les cristaux liquides limitent ce dernier à une dizaine de Hz, le miroir à membrane atteint des taux de l'ordre de 5 kHz. Cette rapidité de mise à jour ainsi que le faible nombre de paramètres libres le rendent tout à fait approprié à des optimisations en boucle fermée. Ceci est d'autant plus intéressant qu'on ne connaît habituellement pas à l'avance la phase résiduelle à compenser et qu'une optimisation en boucle fermée est donc tout à fait indiquée.

Plusieurs éléments limitent l'utilisation de ce type de miroir. Tout d'abord, ils ne supportent pas de densité de puissance élevée. Il est apparemment difficile de faire des traitements résistant au flux sur la membrane flexible. Ce défaut est d'autant plus gênant qu'il existe peu (voire même pas du tout) d'informations sur la tenue au flux pour des impulsions courtes. Le plus souvent, les données existantes traitent des lasers continus ou nanosecondes.

Ensuite, le miroir à membrane tend à déformer le profil spatial de l'impulsion. En particulier, le fait de déformer le profil transverse tant à modifier aussi le rayon de courbure selon l'axe vertical. Le miroir déformable a ainsi tendance à modifier la divergence et le pointé du faisceau sortant du compresseur, ce qui est très gênant. Pour limiter ces effets-là, on peut introduire une lentille cylindrique qui focalisera le faisceau selon la dimension verticale, comme cela est schématisé sur la figure III.3. Cela permet de rendre le profil

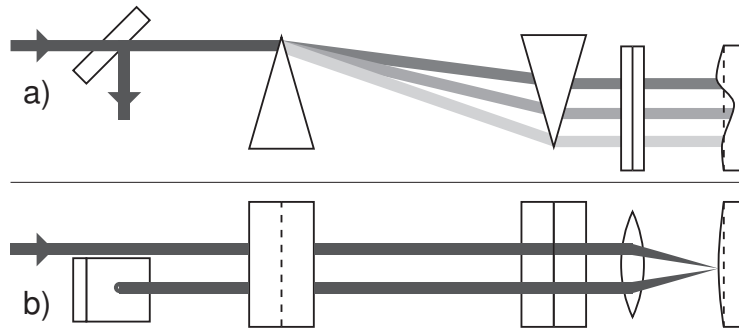


FIG. III.3 – Ajout d'une lentille cylindrique dans le compresseur à miroir déformable : a) vue de dessus, b) vue de côté.

vertical insensible aux déformations de la membrane. Par contre, cela augmente la densité de puissance sur le miroir, ce qui compte tenu de la faible tenue au flux, limite d'autant plus l'énergie des impulsions à mettre en forme. Nous avons ainsi modifié le compresseur en introduisant une lentille cylindrique de 40 cm de focale pour limiter les déformations du faisceau de sortie. Cela nous a permis d'obtenir un pointé et une divergence presque insensible aux déformations programmées, mais nous a limité à des énergies de l'ordre du μJ en sortie du compresseur sur la dizaine de μJ que peut délivrer le NOPA installé au STEACIE Institute.

Malheureusement, je n'ai pas de résultats probants à présenter sur l'amélioration de la compression grâce au miroir déformable. En effet, lors de ce séjour au STEACIE Institute, de nombreux problèmes de stabilité de la chaîne amplifiée nous ont empêchés d'obtenir des résultats convenables. Ces problèmes de stabilité sont depuis réglés et un nouveau séjour

est planifié pour mener à terme ces études.

Comme nous venons de le montrer le miroir à membrane déformable au sein du compresseur permet d'obtenir des impulsions plus courtes et plus propres en sortie du NOPA. Il permet aussi de mettre en forme l'impulsion de sortie. Cette mise en forme est toutefois extrêmement limitée et seules des phases continues avec de faibles excursions sont réalisables. De plus, il ne possède qu'un très petit nombre de degrés de liberté. Au final, le façonneur à membrane déformable présenté ici doit être considéré plus comme un compresseur versatile et performant que comme un réel outil de mise en forme.

Ensuite, le compresseur à prismes, avec ou sans miroir déformable, nécessite un réalignement complet lorsqu'on accorde la longueur d'onde du NOPA. Ce réalignement peut s'avérer fastidieux et complique la procédure d'accord du NOPA. Enfin, l'encombrement de ce compresseur est important. Il rajoute un chemin optique considérable (de l'ordre du mètre), et ce dernier doit être compensé pour la réalisation d'éventuelles expériences pompe-sonde.

Il est possible de réaliser un compresseur beaucoup plus compact et ne nécessitant aucun réalignement à l'aide d'un AOPDF. Il est même possible d'obtenir de véritables capacités de mise en formes directement en sortie du NOPA.

III.2 AOPDF : recompression et mise en forme

Nous avons réalisé à Toulouse deux séries d'expériences où l'on a remplacé le compresseur de sortie du NOPA par un AOPDF. Ces expériences ont été réalisées avec Daniel Kaplan². Le but de ces expériences était d'obtenir un façonneur spécialement optimisé, permettant la compression et la mise en forme des impulsions issues du NOPA, sur toute sa gamme spectrale et ce avec un rendement suffisamment élevé.

Nous rappellerons tout d'abord brièvement le principe de fonctionnement de l'AOPDF, celui-ci étant tout à fait particulier. Nous montrerons ensuite les résultats de compression accordable, obtenus avec deux cristaux différents. Nous verrons ensuite les capacités de mise en forme du deuxième cristal. Enfin, nous montrerons notre tentative de réalisation d'un NOPA programmable à l'aide de l'AOPDF.

III.2.1 Principe de fonctionnement

L'AOPDF est un façonneur d'une extrême compacité (voir la figure III.4 page suivante) et très simple d'utilisation. Les détails intimes de son fonctionnement, eux, le sont beaucoup moins. Nous ne prétendons pas ici expliquer par le menu la physique qui régit l'AOPDF. Nous souhaitons juste donner une idée intuitive de cette physique. Nous invitons les lecteurs qui veulent en savoir plus à consulter les différentes références traitant de l'AOPDF[27, 37, 79, 109, 110, 111].

L'AOPDF repose sur l'interaction colinéaire entre une onde acoustique $v(t)$ et une onde optique $E_i(t)$ au sein d'un cristal de TeO_2 , comme cela est schématisé sur la figure III.5.

²de la société Fastlite qui développe l'AOPDF, dont le nom commercial est DAZZLERTM.

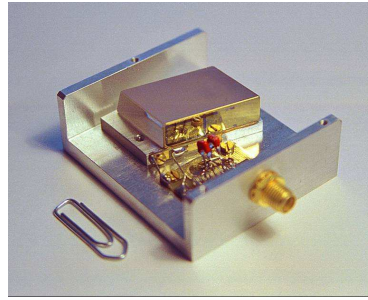
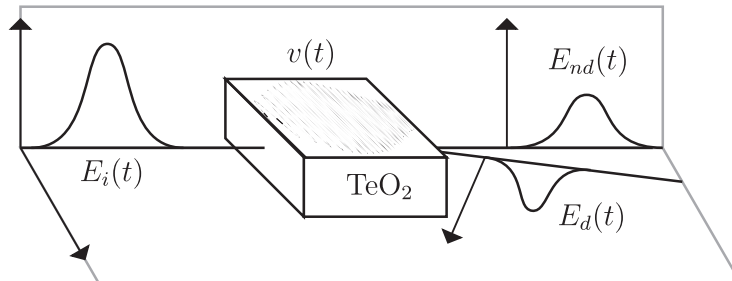


FIG. III.4 – Photographie d'un prototype d'AOPDF.

Ce cristal de paratellurite est fortement biréfringent tant pour l'acoustique que l'optique. L'onde optique incidente (il s'agit de l'impulsion à mettre en forme) se propage selon l'axe ordinaire du cristal. Par diffraction de Bragg assistée par l'onde acoustique, les différentes composantes spectrales de l'onde optique peuvent être basculées sur l'axe extraordinaire du cristal, en des positions différentes du cristal. Or les vitesses de groupes ne sont pas les mêmes sur l'axe ordinaire et sur l'axe extraordinaire. Le temps mis par une composante donnée pour traverser le cristal dépendra donc de l'endroit où elle aura basculé. On obtient ainsi en sortie du cristal deux impulsions lumineuses : l'une non-diffractée $E_{nd}(t)$ de même polarisation que l'impulsion incidente, et une diffractée par l'onde acoustique $E_d(t)$ dont la polarisation est orthogonale à la polarisation d'incidence. Cette onde diffractée $E_d(t)$

FIG. III.5 – Schéma de principe de l'AOPDF. L'interaction colinéaire d'une impulsion $E_i(t)$ avec une onde acoustique façonnée $v(t)$ produit une impulsion mise en forme diffractée $E_d(t)$.

présente une mise en forme directement liée à la position de diffraction dans le cristal des différentes composantes spectrales de $E_i(t)$. La figure III.6 page suivante donne un exemple schématique du processus. Par souci de clarté, l'onde non diffractée n'est pas représentée. On suppose que l'onde incidente $E_i(t)$ présente une dérive de fréquence positive : les courtes longueurs d'onde arrivent en fin de l'impulsion. En diffractant ces courtes longueurs d'onde en début de cristal (traits bleus) et les grandes longueurs d'onde en fin de cristal (traits rouges), il est possible d'annuler la dérive de fréquence et donc de comprimer l'impulsion lumineuse. Pour cela, l'onde acoustique programmée doit présenter une dérive de fréquence opposée à celle de l'impulsion à comprimer.

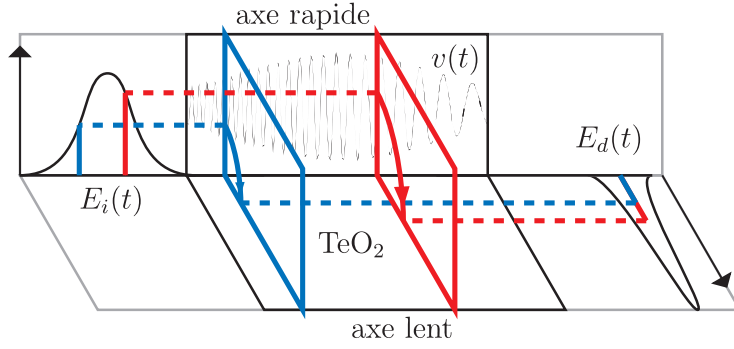


FIG. III.6 – Fonctionnement schématisé de l'AOPDF : les différentes composantes spectrales de l'onde incidente basculent de l'axe rapide à l'axe lent du cristal en des positions différentes. Cela modifie le temps de parcours de chacune d'elle à l'intérieur du cristal.

L'onde diffractée s'exprime simplement comme la convolution de l'impulsion incidente par l'onde acoustique, affectée d'un facteur d'échelle :

$$E_d(t) = E_i(t) \circ v(\alpha t) \quad (\text{III.2})$$

où $\circ$ représente la convolution et avec un facteur d'échelle α valant :

$$\alpha = \frac{c}{\Delta n V} \sim 10^7 \quad (\text{III.3})$$

Δn est l'anisotropie optique du cristal et V la vitesse de propagation de l'onde acoustique. C'est ce facteur d'échelle α qui permet de transférer la forme de l'onde acoustique sur l'impulsion optique. Il est lié à la condition d'accord de Bragg[111] qui donne la fréquence acoustique f_{ac} à appliquer en un point du cristal pour y faire basculer une composante optique de longueur d'onde λ_{op} :

$$f_{ac} = \frac{c}{\alpha \lambda_{op}} = \frac{f_{op}}{\alpha} \quad (\text{III.4})$$

Compte tenu de la valeur du facteur d'échelle α , l'onde acoustique doit avoir une fréquence de quelques dizaines de MHz pour pouvoir mettre en forme des impulsions dans le visible (quelques centaines de THz). L'onde acoustique est en fait une onde progressive qui est synchronisée avec l'impulsion laser à mettre en forme. Celle-ci est générée par un transducteur collé au cristal et commandé par un synthétiseur haute fréquence.

Il est donc possible d'appliquer à l'impulsion incidente $E_i(t)$ une mise en forme temporelle arbitraire en générant l'onde acoustique idoine.

III.2.1.1 Propriétés - fenêtre temporelle

Tout comme les autres façonneurs présentés jusqu'ici, l'AOPDF est un filtre spectral. Contrairement au façonneur haute résolution à cristaux liquides, il n'est pas pixellisé et

n'introduit pas de couplage spatio-temporel. Il permet aussi un taux de rafraîchissement des mises en forme programmées allant jusqu'au kHz (contre une dizaine de Hz pour les cristaux liquides). Il est par contre limité à des densités d'énergie plus faibles et à des taux de répétition inférieurs au MHz pour le laser incident.

Comme pour le façonneur à cristaux liquides, il est possible de définir une fenêtre temporelle dans laquelle il sera possible de réaliser une mise en forme. Il faut faire très attention car cette fenêtre n'a absolument pas la même origine physique que pour le façonneur à cristaux liquides. Dans le cas des cristaux liquides, elle est liée à la pixellisation dans le domaine spectral et des répliques temporelles apparaissent à l'extérieur de cette fenêtre. L'AOPDF quant à lui n'est pas pixellisé et il n'y a donc aucune réplique temporelle.

Par contre, la fenêtre temporelle accessible est limitée par la taille du cristal. En effet, considérons une composante optique donnée : le délai programmable le plus court correspond à sa propagation entièrement sur l'axe ordinaire (elle bascule sur l'axe extraordinaire à la fin du cristal). Le délai le plus long correspond à une propagation selon l'axe extraordinaire (basculement au début du cristal). La différence entre ces deux temps de parcours définit la fenêtre temporelle réalisable par l'AOPDF, pour cette composante spectrale.

Pour une impulsion incidente centrée en λ_0 , cette fenêtre T_{max} vaut :

$$T_{max} = \Delta n_g(\lambda_0) \cos^2(\theta) L / c \quad (\text{III.5})$$

où $\Delta n_g(\lambda_0)$ est la biréfringence de groupe à λ_0 , L la longueur du cristal et θ l'angle entre la direction de propagation dans le cristal et l'axe [110] du cristal. θ est directement lié à la coupe du cristal. T_{max} dépend de la longueur d'onde par le biais de Δn_g dont l'évolution est donnée sur la figure III.7. Cette figure montre que T_{max} augmente pour les courtes longueurs d'onde.

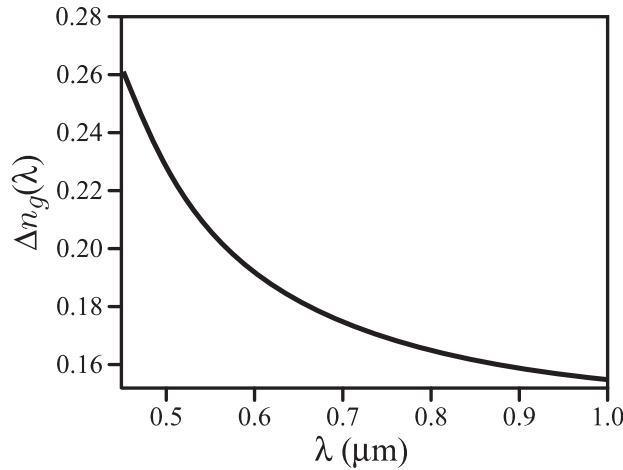


FIG. III.7 – Evolution de la biréfringence de groupe Δn_g en fonction de la longueur d'onde.

Ce délai T_{max} va être utilisé pour mettre en forme l'impulsion de sortie. Cependant, l'AOPDF doit compenser la dispersion due à la propagation dans le cristal de TeO_2 . Une partie de T_{max} - appelée T_{cp} - est ainsi utilisée pour cette *autocompensation*. La dispersion

dans le TeO_2 est d'autant plus grande que le cristal est long et que la longueur d'onde de l'impulsion incidente est courte³. De plus, l'effet de cette dispersion sur l'impulsion traversant le cristal est d'autant plus important que cette dernière est large spectralement.

T_{cp} est donc d'autant plus grand que le cristal est long et que l'impulsion est large spectralement. Il augmente pour les courtes longueurs d'onde.

Il reste donc une fenêtre $T_{eff} = T_{max} - T_{cp}$ pour réaliser les mises en forme souhaitées. L'évolution de T_{eff} en fonction des différents paramètres du cristal et des propriétés de l'impulsion incidente est relativement complexe. Elle résulte d'une compétition entre les évolutions respectives de T_{max} et T_{cp} . Pour une coupe de cristal donnée (c'est-à-dire θ fixé) on peut donner les grandes lignes de cette évolution.

Tout d'abord, T_{eff} est d'autant plus grande que le cristal est long. Ceci se comprend assez facilement : plus le cristal est long et plus la différence entre le temps de parcours sur l'axe rapide et l'axe lent sera grande.

Ensuite, T_{eff} est d'autant plus faible que la longueur d'onde centrale de l'impulsion est courte, comme on peut le constater sur la figure III.8 qui donne l'évolution de T_{eff} (en ps) en fonction de la longueur d'onde pour différentes durées de l'impulsion incidente⁴, pour un cristal de 25 mm de long, taillé à $\theta = 38^\circ$. Cette diminution est due à un accroissement

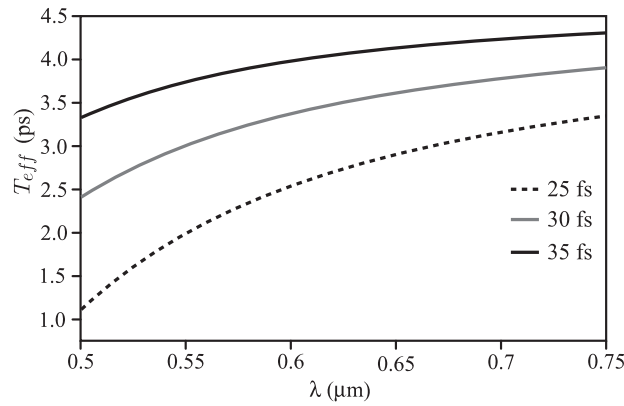


FIG. III.8 – Evolution de T_{eff} en fonction de la longueur d'onde pour différentes largeurs spectrales, correspondant à des durées limitées par transformée de Fourier de 25 fs (pointillés noirs), 30 fs (trait gris) et 35 fs (trait noir).

de T_{cp} plus rapide que celui de T_{max} pour les courtes longueurs d'onde. On disposera donc pour un cristal donné de plus de possibilités de mise en forme dans le rouge que dans le bleu.

Enfin, T_{eff} est d'autant plus faible que la largeur spectrale de l'impulsion incidente est élevée ou, ce qui revient au même, que sa durée limitée par transformée de Fourier est faible. On retrouve cette tendance sur la figure III.8. Cela aussi se comprend simplement : plus l'impulsion incidente est large spectralement, plus elle sera affectée par la dispersion

³Ce qui est le cas de la plupart des matériaux transparents pour les longueurs d'onde visibles.

⁴On entend par durée la largeur à mi-hauteur en intensité du profil temporel de l'impulsion. Les impulsions incidentes sont ici limitées par transformée de Fourier.

naturelle du TeO_2 . Il faudra alors utiliser une plus grande partie du délai disponible T_{max} pour l'autocompensation du cristal.

T_{max} ainsi que l'efficacité de diffraction de l'AOPDF sont fonctions de la coupe du cristal, comme cela est schématisé sur la figure III.9. La capacité de mise en forme varie en

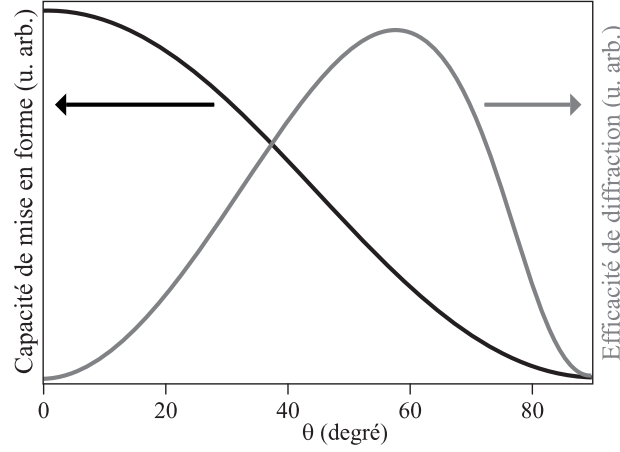


FIG. III.9 – Evolution de l'efficacité de diffraction (trait gris) et de la capacité de mise en forme (c'est-à-dire T_{max} , en trait noir) en fonction de l'angle θ de coupe du cristal.

cosinus carré de θ (voir l'équation III.5) et présente donc un maximum autour de $\theta = 0^\circ$. L'efficacité de diffraction quant à elle passe par un maximum autour de $\theta = 58,5^\circ$. Il est ainsi possible, en choisissant l'angle de coupe θ , de trouver un compromis entre capacité de mise en forme et efficacité de diffraction.

Jusqu'à présent, l'AOPDF a surtout été utilisé autour de 800 nm[27, 39, 112, 113, 114]. Celui-ci est souvent utilisé en amont des étages d'amplification[39, 113, 114]. Il a ainsi permis d'obtenir des impulsions très large bande en compensant le rétrécissement par le gain dû à l'amplification[113]. Des impulsions mises en forme intenses[114] ont aussi été obtenues directement en sortie des amplificateurs. Son extrême compacité et ses capacités de mise en forme ont aussi été utilisées pour la réalisation d'un dispositif de caractérisation d'impulsion simple et versatile[112].

III.2.2 Compresseur accordable

III.2.2.1 Dispositif expérimental

Pour réaliser les différentes expériences de recompression et de mise en forme, nous avons placé l'AOPDF en sortie du NOPA, comme cela est montré sur la figure III.10 page suivante. Celui-ci vient remplacer le compresseur à prismes. En sortie de l'AOPDF, les impulsions sont mesurées par autocorrélation intensimétrique. Le NOPA délivre des impulsions de l'ordre de 11 μJ accordables entre 490 nm et 680 nm. Dans un premier temps, notre but était de recomprimer les impulsions issues du NOPA sur toute sa bande

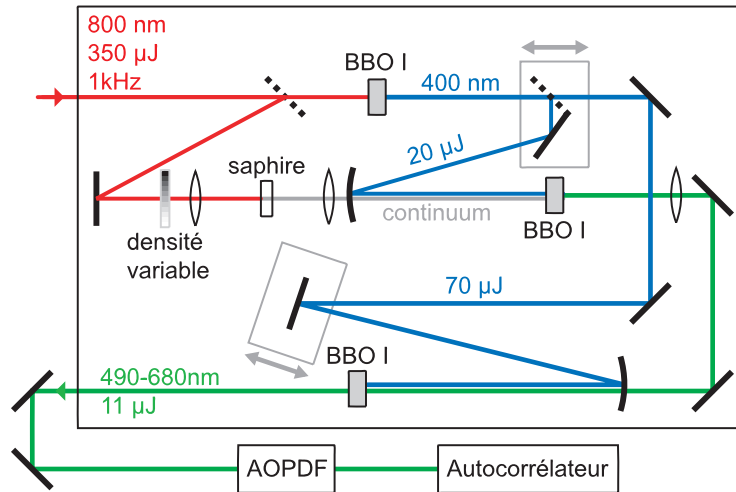


FIG. III.10 – L'AOPDF est directement placé en sortie du NOPA à la place du compresseur à prismes.

spectrale, sans avoir d'élément à déplacer, simplement en changeant la programmation de l'AOPDF.

III.2.2.2 Premier candidat

Nous avons tout d'abord utilisé un premier cristal de 25 mm de long, taillé pour $\theta = 38^\circ$. Cette coupe, comme on peut le voir sur la figure III.9 page précédente, est un compromis entre la capacité de mise en forme et l'efficacité de diffraction. Cette coupe permet entre autre l'autocompensation d'impulsions large bande. Nous avons ainsi pu recomprimer les impulsions issues du NOPA sur une grande partie de la gamme spectrale accessible. La figure III.11 page ci-contre montre des résultats typiques dans le gamme 500-610 nm. La partie haute présente le spectre en entrée (trait gris) et en sortie (carrés noirs) de l'AOPDF. La partie basse représente les intensités d'autocorrélation correspondantes (carrés gris) ainsi que leur ajustement en sech^2 (trait gris). Les durées indiquées correspondent aux largeurs à mi-hauteur en intensité déduites des ajustements. Elles sont tout à fait similaires à celles que donne une recompression par des prismes.

Ces impulsions ne sont pas limitées par transformée de Fourier⁵. Ceci est dû au NOPA utilisé qui est optimisé pour des impulsions intenses (pour un NOPA!) et pour une accordabilité aisée. En contre-partie, les impulsions à sa sortie présentent une phase spectrale importante ainsi que différentes distortions géométriques[115, 116] qui limitent les possibilités de recompression.

Les résultats présentés ici ne montrent pas la partie 610-680 nm de la bande spectrale du NOPA. Nous nous sommes en effet attachés à la partie "courtes longueurs d'onde" car il s'agit de la zone dans laquelle la recompression est plus délicate. En effet, pour mettre en forme des longueurs d'onde courtes, il faut générer des ondes acoustiques de fréquence élevée, pour lesquelles l'absorption dans le cristal est plus importante. On a

⁵Il y a environ un facteur 1,6 entre les durées obtenues et les limites par transformée de Fourier.

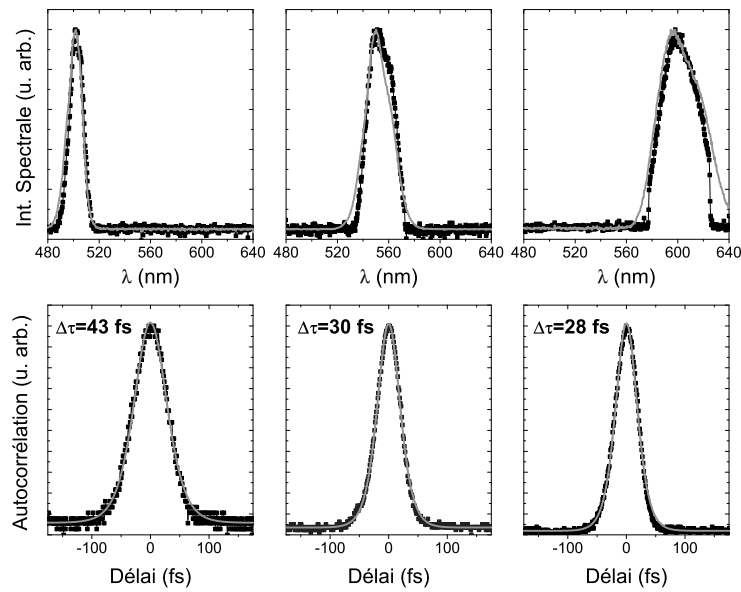


FIG. III.11 – Impulsions en sortie de l’AOPDF. Haut : Intensités spectrales avant (trait gris) et après l’AOPDF (carrés noirs). Bas : intensité d’autocorrélations correspondantes (carrés noirs pour l’expérience, trait gris pour l’ajustement en sech^2).

ainsi observé une très faible énergie pour les impulsions recomprimées autour de 500 nm (quelques centaines de nanoJoules seulement). Malgré la normalisation, la faiblesse du signal se voit clairement sur le bruit de numérisation de l’autocorrélation à 500 nm (figure III.11 en bas à gauche).

Ce premier cristal a permis de montrer que malgré cette absorption, la recompression sur toute la bande spectrale était possible, ce qui est un net avantage par rapport au compresseur à prismes nécessitant un réalignement pour chaque changement de longueur d’onde. Autre avantage non-négligeable, l’AOPDF permet un gain de place conséquent. La figure III.12 montre l’encombrement de l’AOPDF et celui du compresseur à prismes utilisés pour la recompression du NOPA.

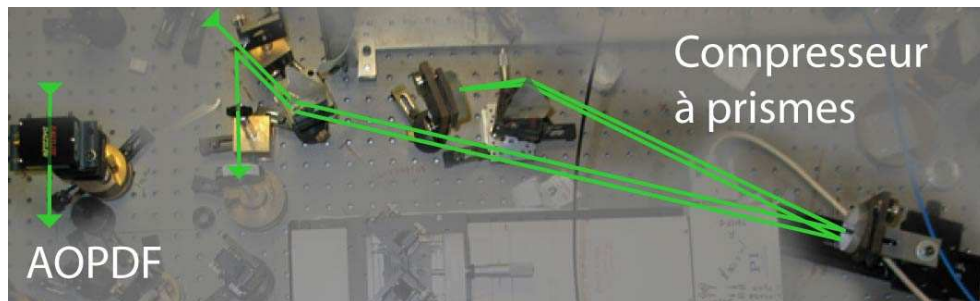


FIG. III.12 – Comparaison entre l’AOPDF et le compresseur à prismes.

Cependant, le rendement de diffraction de l’AOPDF est faible : 20 à 30% maximum

pour des impulsions correctement recomprimées. De plus, l'absorption acoustique limite fortement l'énergie diffractée pour les impulsions autour de 500 nm. Tout ceci peut s'avérer rédhibitoire pour une application concrète. Fort de cette constatation, un deuxième prototype, taillé différemment a été fabriqué et testé.

III.2.2.3 Deuxième candidat

Celui-ci est coupé pour avoir $\theta = 45^\circ$ et est de même longueur que le précédent (25 mm). Cette coupe permet d'augmenter l'efficacité de diffraction au détriment de la fenêtre temporelle T_{max} . Elle réduit la fréquence acoustique utilisée pour diffracter une longueur d'onde optique donnée. Elle diminue donc aussi les effets d'absorption observés pour les courtes longueurs d'onde. Cela nous a permis d'obtenir des efficacités de diffraction bien meilleures ainsi que de réelles capacités de façonnage.

Les impulsions recomprimées par cet AOPDF sont présentées sur la figure III.13. La

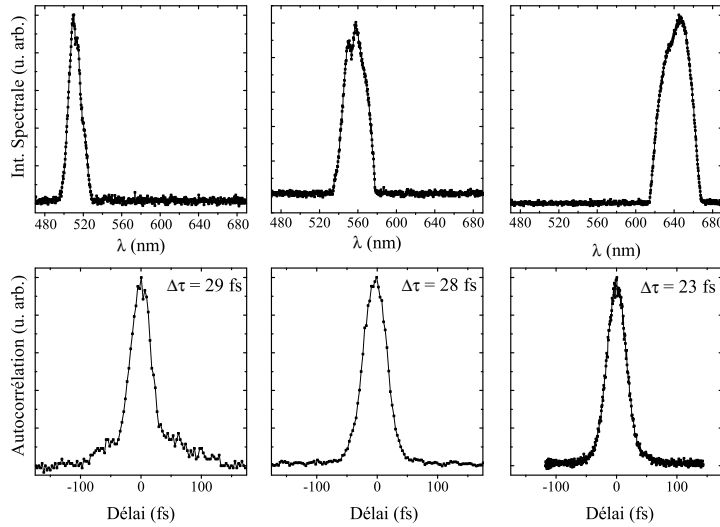


FIG. III.13 – Impulsions en sortie du nouvel AOPDF. Haut : Intensités spectrales (carrés noirs). Bas : intensité d'autocorrélations correspondantes (carrés noirs). Les durées indiquées sont celles d'un ajustement en sech^2 .

partie haute présente le spectre en sortie, la partie basse l'autocorrélation correspondante. Les durées indiquées sont déduites d'un ajustement en sech^2 . On retrouve les mêmes durées qu'avec l'ancien AOPDF, avec une gamme spectrale plus étendue. Celle-ci est maintenant limitée par les capacités du NOPA. Le piédestal nettement visible sur l'autocorrélation à 510 nm semble être dû à notre NOPA. En effet, des expériences similaires réalisées à Leeds par B. Whitaker avec le même AOPDF ont donné de meilleures autocorrélations à cette longueur d'onde[41].

Autre avancée importante avec ce nouvel AOPDF, l'intensité en sortie est nettement plus élevée, notamment aux alentours de 500 nm. En effet, il supporte des densités de puissance plus élevées (de l'ordre de $300 \mu\text{J}/\text{cm}^2$) et offre des efficacités de diffraction plus

importantes. Nous avons ainsi pu obtenir des efficacités de diffraction de l'ordre de 75% sans distortion des impulsions de sortie. Les résultats présentés ici ont été obtenus avec des efficacités de 50% et des impulsions de l'ordre de $4 \mu\text{J}$ en entrée.

Il est donc possible d'obtenir des impulsions courtes, accordables sur tout le visible avec des intensités de plusieurs μJ . Contrairement aux précédents résultats, nous atteignons ici des caractéristiques tout à fait suffisantes pour une application concrète.

III.2.3 Capacités de mise en forme

Outre la recompression accordable, nous avons pu réaliser des mises en forme relativement complexes avec ce deuxième AOPDF. Pour caractériser ses capacités de façonnage, nous avons utilisé le montage expérimental présenté sur la figure III.14. Les impulsions

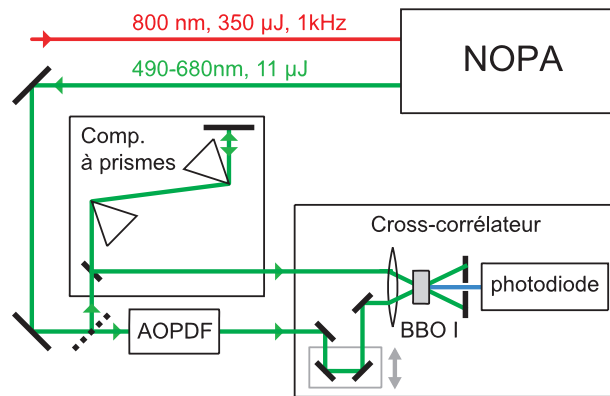


FIG. III.14 – Dispositif expérimental. Une partie de la sortie du NOPA est recomprimée par l'AOPDF, l'autre par le compresseur à prismes. La sortie de l'AOPDF est détectée par crosscorrélacion avec la sortie du compresseur ou par autocorrélacion (non représenté).

de sortie du NOPA sont séparées en deux. Une première partie est recomprimée par un compresseur à prismes standard, la deuxième traverse l'AOPDF. On réalise ensuite la corrélation entre les impulsions issues de l'AOPDF et celles issues du compresseur.

Nous avons pu montrer une fenêtre temporelle disponible pour du façonnage de plus de 4 ps, sur toute la gamme d'accordabilité du NOPA. La figure III.15 page suivante illustre ce résultat avec des intensités de corrélation obtenues à différentes longueurs d'onde :

- (a) Cross-corrélacions pour différents délais programmés dans l'AOPDF, à 510 nm. La fenêtre de 4 ps est clairement visible. Le délai 0 correspond à une impulsion diffractée au milieu du cristal, les délais positifs à une diffraction en début et les négatifs en fin du cristal. La diminution d'intensité des impulsions autour de -1 ps illustre l'absorption de l'onde acoustique à ces courtes longueurs d'onde.
- (b) Séquence de 5 impulsions à 640 nm. Cela montre les mises en formes "complexes" qui peuvent être générées par l'AOPDF.

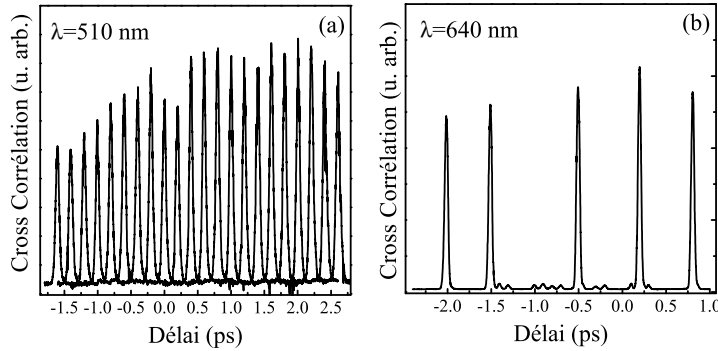


FIG. III.15 – Intensité de corrélation entre la sortie du compresseur à prismes et celle de l'AOPDF. (a) Superposition de plusieurs corrélations pour différents délais obtenues à 510 nm. (b) Séquence de 5 impulsions à 640 nm.

Une autre mise en forme complexe mettant en avant les possibilités de façonnage est montrée sur la figure III.16. Il s'agit d'une séquence de trois impulsions, chacune ayant une

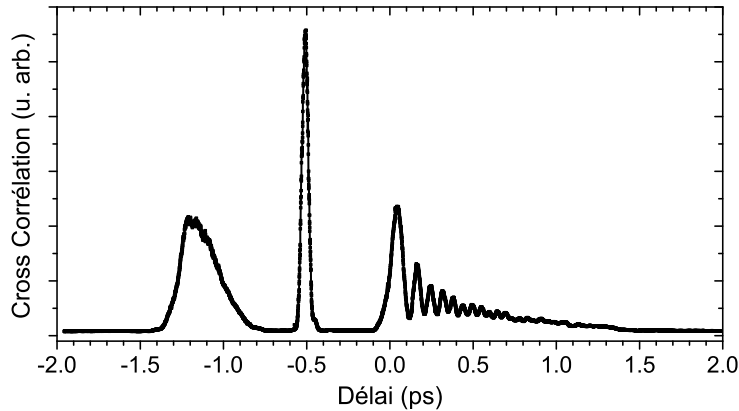


FIG. III.16 – Intensité de corrélation d'une séquence de trois impulsions : avec une phase spectrale quadratique, une plate et une cubique.

phase spectrale différente. La première présente une phase quadratique de $-16,9 \cdot 10^4 \text{ fs}^2$, la seconde une phase nulle et la troisième une phase cubique de $28 \cdot 10^4 \text{ fs}^3$. La mise en forme est globalement bien reproduite avec toutefois un défaut sur la première impulsion. Celle-ci est en effet plus courte qu'elle ne le devrait. Ceci est dû à des effets de saturation de l'onde acoustique et peut être évité en utilisant une efficacité de diffraction plus faible.

III.2.4 Vers un NOPA compact et programmable ?

Nous avons enfin utilisé l'AOPDF à l'intérieur du NOPA. En effet, l'AOPDF est suffisamment petit pour être introduit en amont des étages d'amplification sans trop compliquer

le dispositif. Nous souhaitons ainsi obtenir des impulsions courtes directement en sortie du NOPA. Nous espérons aussi pouvoir accorder le NOPA à l'aide de l'AOPDF. Cette expérience se rapproche de celle présentée dans la référence[117] où une ligne à dispersion nulle avec un masque à cristaux liquides est utilisée pour mettre en forme le continuum de lumière blanche d'un NOPA.

Nous avons ainsi réalisé le montage présenté sur la figure III.17. Il est assez sem-

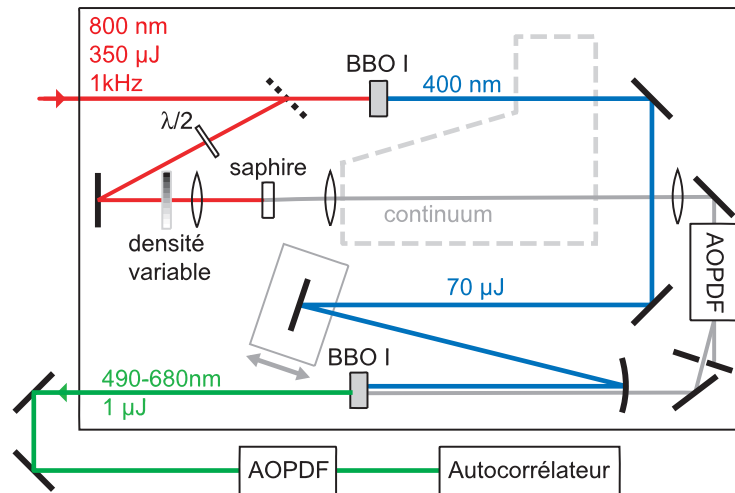


FIG. III.17 – L'AOPDF est placé sur le continuum avant l'amplification. Le premier étage d'amplification a été retiré (voir figure III.10 page 86 pour comparaison).

blable à celui présenté sur la figure III.10 page 86. Nous avons juste retiré le premier étage d'amplification et utilisé l'AOPDF pour mettre en forme le continuum. Nous avons aussi inséré une lame demi-onde pour compenser le changement de polarisation entre l'entrée et la sortie de l'AOPDF. Le faisceau diffracté par l'AOPDF est amplifié dans le cristal qui servait auparavant au deuxième étage d'amplification.

L'accord en longueur d'onde du NOPA se fait en décalant temporellement le continuum par rapport à l'impulsion de pompe (voir la partie I.4 page 7). En effet, le continuum présente une forte dérive de fréquence et le décalage temporel permet de superposer, dans le cristal d'amplification, l'impulsion de pompe avec la partie du continuum que l'on souhaite amplifier. Habituellement, cela se fait par le biais de translations que l'on ajuste à la main, comme on peut le voir sur la figure III.10 (il y a autant de translations que d'étages d'amplification).

Notre objectif est de réaliser ce décalage à l'aide de l'AOPDF en programmant différents retards en fonction de la longueur d'onde voulue. Initialement, nous pensions que les délais à introduire étaient relativement faibles (de quelques dizaines à quelques centaines de fs) et étaient donc facilement réalisables avec l'AOPDF. C'était sans tenir compte de la forte variation de l'indice du TeO_2 dans la gamme visible. Cette variation de l'indice rajoute un délai dû au temps de traversée de l'AOPDF, qui dépend fortement de la longueur d'onde choisie. Ceci va limiter la gamme d'accordabilité accessible.

En plus de ce décalage, l'AOPDF doit aussi recomprimer les impulsions. Il faut pour cela compenser la phase spectrale du continuum et pré-compenser la phase qui sera accumulée lors de la traversée des optiques placées après l'AOPDF.

Les résultats d'accordabilité et de recompression sont présentés sur la figure III.18. Les

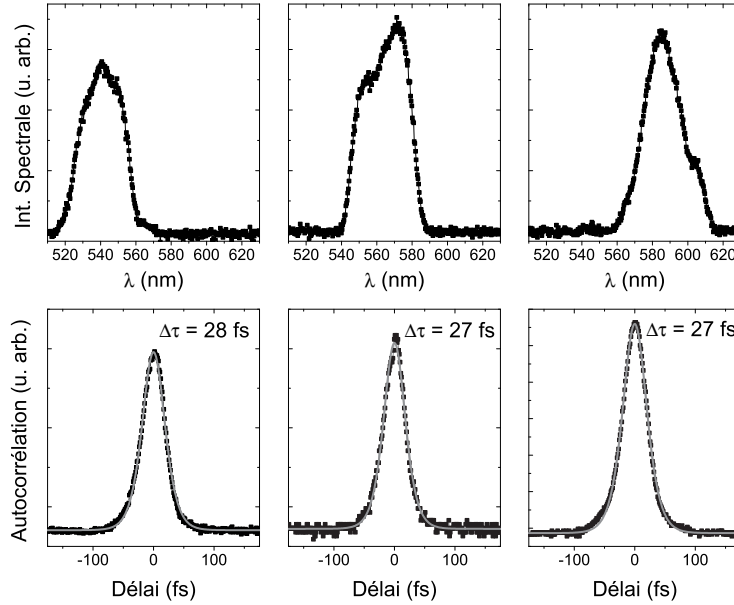


FIG. III.18 – Impulsions en sortie du NOPA avec un AOPDF avant amplification. Haut : Intensités spectrales (carrés noirs). Bas : intensité d'autocorrélations correspondantes (carrés noirs pour l'expérience, trait gris pour l'ajustement en sech^2).

durées obtenues sont tout à fait honorables et comparables à celle obtenues précédemment. Par contre, comme nous le pressentions, la plage d'accordabilité est relativement réduite (530-600 nm) par rapport à celle du NOPA standard (480-680 nm).

Pour pouvoir accorder le NOPA sur une plus large bande, il faut à nouveau utiliser une translation pour rattraper à la main le décalage temporel introduit par la traversée du cristal de TeO_2 .

Une autre limitation de cette approche est l'énergie en sortie du NOPA. Cela ne tient pas au fait de n'utiliser qu'un étage d'amplification. En effet, nous pourrions en mettre deux ; nous avons choisi cette solution car elle minimisait les changements entre les deux dispositions possibles (AOPDF en sortie et AOPDF dans le NOPA). Nous avons remarqué que l'énergie obtenue en sortie était plus faible que pour un NOPA à un étage sans AOPDF : recomprimer le continuum avant l'amplification dégrade le rendement. En effet, l'efficacité d'amplification dépend fortement du recouvrement temporel dans le cristal d'amplification entre l'impulsion de pompe et la partie du continuum à amplifier. Comme la pompe est plutôt longue (de l'ordre de 200 fs) il n'est pas efficace d'avoir un continuum court.

Pour augmenter l'énergie disponible, on peut placer des bouts de verres en sortie du NOPA et demander à l'AOPDF de précompenser la dispersion introduite par ces éléments. Cela permet d'avoir un continuum plus long dans le cristal d'amplification tout en ayant

une impulsion courte après les bouts de verres. Nous avons mis en oeuvre cette modification et cela a permis un net gain en énergie au prix d'une complexité accrue.

III.2.5 Bilan

Nous avons ici montré pour la première fois qu'un AOPDF pouvait servir de système de compression et de mise en forme directement en sortie d'un NOPA et ce sur toute la gamme visible[41]. Il apporte de nombreux avantages par rapport aux systèmes de compression traditionnels (prismes, réseaux, miroirs "chirpés"). Tout d'abord, il ne nécessite aucun réalignement. Il est de plus beaucoup plus simple et compact. Enfin, il permet, en plus de la recompression, une mise en forme de la sortie du NOPA. Cette dernière caractéristique peut être très intéressante pour des expériences de contrôle optimal ou cohérent dans des systèmes moléculaires ou biologiques.

Chapitre IV

CONTRÔLE COHÉRENT DANS LE RUBIDIUM

Le contrôle de systèmes très divers (atomes, molécules, objets biologiques, etc...) au niveau quantique est un domaine très dynamique de la physique moderne. Un grand nombre de schémas de contrôle repose sur l'interaction cohérente entre le système étudié et une ou plusieurs impulsions lumineuses contrôlées[2, 118, 119, 4]. C'est entre autres le cas du contrôle cohérent[120, 25, 121] et du contrôle optimal[6, 24]. Ces vingt dernières années ont vu une explosion du nombre de systèmes étudiés, dans des domaines très variés, grâce notamment à l'apparition de façonneurs actifs¹ (voir en particulier le chapitre II page 15). Ceux-ci permettent la réalisation d'impulsions avec des profils très complexes. Ils permettent aussi de changer très rapidement de mise en forme, rendant ainsi possible des contrôles en boucle fermée[6] (contrôle optimal).

Cette utilisation d'impulsions fortement mises en forme a entraîné un développement rapide des outils et des techniques de caractérisation de ces dernières.

Dans ce chapitre, nous allons nous intéresser à un type particulier de processus cohérent : les transitoires cohérents. Ceux-ci apparaissent lors de l'interaction d'un système atomique avec une impulsion laser ultracourte. Par des techniques de contrôle cohérent qui utilisent le façonneur haute résolution présenté précédemment, nous pouvons contrôler le comportement de ces transitoires. Cela nous a aussi permis d'accéder à la fonction d'onde atomique du système utilisé. Enfin, nous avons exploité cette information sur l'atome pour caractériser l'impulsion ultracourte mise en forme ayant servi à la génération et au contrôle des transitoires cohérents.

Ce travail s'inscrit dans la continuité des précédentes études de transitoires cohérents réalisées au sein de l'équipe[122, 55]. Nous présenterons en premier lieu le formalisme des transitoires cohérents. Ensuite, nous rappellerons succinctement la première mise en évidence expérimentale. Puis, nous donnerons un aperçu des différents schémas de contrôle qui ont été mis en œuvre, avant cette thèse ou pendant.

Nous présenterons ensuite comment il est possible d'obtenir des informations sur la fonction d'onde atomique. Enfin, nous montrerons que les transitoires cohérents peuvent servir de système de caractérisation d'impulsions employées.

¹Cette appellation de "façonneurs actifs" qui a été consacrée par la communauté scientifique, est assez malheureuse. En effet, presque tous ces façonneurs se comportent comme des filtres spectraux passifs, c'est-à-dire sans amplification. L'appellation "façonneurs programmables" serait moins ambiguë. Nous nous conformerons cependant à l'appellation communément admise de "façonneurs actifs" en gardant à l'esprit cette ambiguïté.

IV.1 Formalisme

Les transitoires cohérents se produisent lors de l'interaction d'un système à deux niveaux $|g\rangle$ et $|e\rangle$ avec une impulsion dite de pompe, peu intense, résonnante avec la transition entre ces deux niveaux. Ils se traduisent par une évolution très particulière de la population dans l'état excité $|e\rangle$ au cours de l'interaction avec l'impulsion de pompe. Pour observer cette population, on utilise une impulsion courte dite de sonde qui est séparée de l'impulsion de pompe par un délai variable. Celle-ci permet de transférer une part de la population présente dans l'état excité à un instant donné vers un niveau supérieur $|f\rangle$. Une mesure de la population du niveau supérieur (par observation de la fluorescence émise par exemple) pour différents délais permet de suivre l'évolution des transitoires cohérents. Nous nous intéressons donc à l'interaction entre un système à trois niveaux et deux impulsions courtes.

IV.1.1 Notations et définitions

IV.1.1.1 Niveaux atomiques

On considère le système atomique à trois niveaux présenté à la figure IV.1. On note

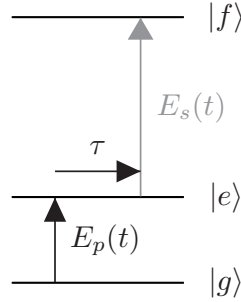


FIG. IV.1 – Système atomique à trois niveaux considéré.

$|g\rangle$, $|e\rangle$ et $|f\rangle$ ces trois niveaux et respectivement $\hbar\omega_g$, $\hbar\omega_e$ et $\hbar\omega_f$ leurs énergies propres. La fonction d'onde de ce système s'écrit :

$$|\Psi(t)\rangle = a_g(t) |g\rangle e^{-i\omega_g t} + a_e(t) |e\rangle e^{-i\omega_e t} + a_f(t) |f\rangle e^{-i\omega_f t} \quad (\text{IV.1})$$

où $a_g(t)$, $a_e(t)$ et $a_f(t)$ sont les amplitudes de probabilités respectives des différents états. Les pulsations associées aux transitions de pompe ($|g\rangle \mapsto |e\rangle$) et de sonde ($|e\rangle \mapsto |f\rangle$) sont notées $\omega_{eg} = \omega_e - \omega_g$ et $\omega_{fe} = \omega_f - \omega_e$, respectivement.

Dans toute la suite, nous considérerons les durées de vie de ces trois états comme infinies² ou du moins excessivement longues devant tous les processus étudiés.

²Expérimentalement, cette approximation est valide car nous allons étudier des dynamiques de quelques picosecondes, les durées de vie des états utilisés étant supérieures à la nanoseconde.

IV.1.1.2 Champs électriques

Les champs de pompe et de sonde sont des quantités réelles. Nous les noterons respectivement $\mathcal{E}_p(t)$ et $\mathcal{E}_s(t)$. Cependant, dans toute la suite, seule la partie de fréquences positives de ces champs nous importera. On définit ainsi le champ de pompe complexe

$$E_p(t) = A_p(t) e^{-i\omega_p t} e^{i\varphi_p(t)} \quad (\text{IV.2})$$

où ω_p est la pulsation centrale du champ. On appelle Δt_p sa largeur à mi-hauteur en intensité. Dans toute la suite et pour tous les autres champs, nous définirons la durée d'une impulsion comme sa largeur à mi-hauteur en intensité.

Ce champ $E_p(t)$ est lié au champ réel $\mathcal{E}_p(t)$ par la relation

$$\mathcal{E}_p(t) = \frac{1}{2} [E_p(t) + E_p^*(t)] \quad (\text{IV.3})$$

Dans le domaine spectral, on définit le champ $\tilde{E}_p(\omega)$

$$\mathcal{F}[E_p(t)] = \tilde{E}_p(\omega) = A_p(\omega) \exp[i\phi_p(\omega)] \quad (\text{IV.4})$$

où $\mathcal{F}$ note la transformée de Fourier directe (voir l'équation A.1 page 153).

On note de façon similaire le champ de sonde $E_s(t)$. En particulier, sa pulsation centrale est notée ω_s , sa durée Δt_s et sa transformée de Fourier $\tilde{E}_s(\omega)$.

On appelle τ le délai pompe-sonde, c'est à dire le délai entre l'impulsion de pompe et celle de sonde. On considère de plus que l'impulsion de pompe (respectivement de sonde) est quasi-résonnante avec la transition de pompe (respectivement de sonde). On note le désaccord $\delta\omega_{eg} = \omega_p - \omega_{eg}$ (respectivement $\delta\omega_{fe} = \omega_s - \omega_{fe}$). Les champs de pompe et de sonde sont tous deux de polarisation linéaire notée $\mathbf{e}$.

IV.1.2 Population dans l'état final - fluorescence

L'objectif de cette section est d'établir la population P_f dans l'état final $|f\rangle$ en fonction du délai pompe-sonde τ . Le signal de fluorescence mesuré expérimentalement et noté $S(\tau)$, est directement proportionnel à cette population.

En utilisant l'approximation de l'onde tournante ainsi que l'approximation perturbative, on trouve[123, 122, 55] pour l'amplitude de probabilité dans l'état excité $|e\rangle$, à un instant t :

$$a_e(t) = -\frac{\mu_{eg}}{2i\hbar} \int_{-\infty}^t E_p(t') e^{i\omega_{eg}t'} dt' \quad (\text{IV.5})$$

avec $\mu_{eg} = \langle e | \mathbf{D} \cdot \mathbf{e} | g \rangle$, $\mathbf{D}$ étant l'opérateur dipolaire électrique et $\mathbf{e}$ la polarisation du champ électrique. La population dans l'état $|e\rangle$ s'écrit quant à elle :

$$P_e(t) = |a_e(t)|^2 = \left| \frac{\mu_{eg}}{2\hbar} \int_{-\infty}^t E_p(t') e^{i\omega_{eg}t'} dt' \right|^2 \quad (\text{IV.6})$$

Les mêmes approximations appliquées à la transition de sonde donnent[122, 55] pour l'amplitude de probabilité dans l'état final $|f\rangle$ à un instant t :

$$a_f(t, \tau) = -\frac{\mu_{fe}}{2i\hbar} \int_{-\infty}^t E_s(t' - \tau) e^{i\omega_{fe}t'} a_e(t') dt' \quad (\text{IV.7})$$

où μ_{fe} est le moment dipolaire de la transition de sonde. L'amplitude de probabilité dans l'état final dépend du délai pompe-sonde τ par le biais de la sonde $E_s(t - \tau)$.

On peut ainsi écrire la population dans l'état final à un instant t sous la forme

$$P_f(t, \tau) = |a_f(t, \tau)|^2 \propto \left| \int_{-\infty}^t dt' E_s(t' - \tau) e^{i\omega_{fe}t'} a_e(t') \right|^2 \quad (\text{IV.8})$$

Le signal que nous observons expérimentalement est en fait une intensité de fluorescence qui est proportionnelle à la population dans l'état final. Cette intensité n'est pas résolue en temps et elle est ainsi liée à la population asymptotique $P_f(\infty, \tau)$ et ne dépend que du délai pompe-sonde.

Il est donc important de regarder le comportement asymptotique de l'amplitude de probabilité $a_f(\infty, \tau)$ ainsi que celui de la population. A partir de l'équation IV.7, on trouve

$$a_f(\infty, \tau) = -\frac{\mu_{fe}}{2i\hbar} \int_{-\infty}^{\infty} dt' E_s(t' - \tau) e^{i\omega_{fe}t'} a_e(t') \quad (\text{IV.9})$$

que nous ré-écrivons sous la forme

$$a_f(\infty, \tau) = e^{i\omega_{fe}\tau} b_f(\tau) \quad (\text{IV.10})$$

en introduisant la quantité $b_f(\tau)$:

$$b_f(\tau) = -\frac{\mu_{fe}}{2i\hbar} \int_{-\infty}^{\infty} dt' E_s(t' - \tau) e^{i\omega_{fe}(t' - \tau)} a_e(t') \quad (\text{IV.11})$$

Cette quantité est proportionnelle à la corrélation³ entre le champ de sonde décalé spectralement de ω_{fe} (c'est à dire $E_s(t) \exp[i\omega_{fe}t]$) et l'amplitude de probabilité dans l'état excité $a_e(t)$. On peut alors écrire le signal de fluorescence $S(\tau)$ comme :

$$S(\tau) \propto P_f(\infty, \tau) = |a_f(\infty, \tau)|^2 = |b_f(\tau)|^2 \quad (\text{IV.12a})$$

$$S(\tau) \propto \left| \int_{-\infty}^{\infty} dt' E_s(t' - \tau) e^{i\omega_{fe}(t' - \tau)} a_e(t') \right|^2 \quad (\text{IV.12b})$$

Tous ces résultats présentés ne sont valables que tant que l'approximation perturbative est vérifiée. Cela se traduit par les conditions suivantes :

$$P_g(t) \simeq 1 \quad \text{et} \quad P_e(t), P_f(t) \ll 1 \quad \text{pour tout instant } t. \quad (\text{IV.13})$$

L'effet principal de la sonde $E_s(t)$ est de moyenner l'amplitude de probabilité $a_e(t)$ sur un temps caractéristique de l'ordre de sa durée Δt_s . Si cette impulsion est suffisamment courte devant les temps caractéristiques d'évolution de $a_e(t)$, elle se résume à un simple

³Plus précisément, $[f(t)*g(t)](\tau)$ désignant le produit de convolution de f par g , on peut écrire $b_f(\tau) = -\mu_{fe}/(2i\hbar) [E_s(-t) \exp(i\omega_{fe}(-t)) * a_e(t)](\tau)$.

delta de Dirac et on obtient ainsi un signal de fluorescence proportionnel à la population dans l'état excité $|e\rangle$:

$$S^\delta(\tau) \propto P_e(\tau) \quad (\text{IV.14})$$

Dans toute la suite, nous omettrons les différents facteurs multiplicatifs afin d'alléger les expressions.

IV.1.3 Transitoires cohérents - Spirale de Cornu

Nous allons pour l'instant étudier le signal de fluorescence $S^\delta(\tau)$ en supposant l'impulsion de sonde infiniment courte (l'effet de cette sonde a été détaillé dans la référence[122] et sera traité en détail dans la partie IV.6 page 135). Le signal $S^\delta(\tau)$ nous donne donc accès directement à $P_e(\tau)$.

Nous allons montrer ici que cette population dans l'état excité présente un régime oscillant transitoire lorsque que la pompe possède une dérive de fréquence[53]. C'est ce régime transitoire que nous appelons *transitoires cohérents*.

Nous supposons tout d'abord que l'impulsion de pompe a un profil gaussien et possède une forte dérive de fréquence :

$$E_p(t) = A_p(t) e^{-i\omega_p t} e^{i\varphi_p(t)} \quad (\text{IV.15a})$$

$$\text{avec } A_p(t) = 4^{-(t/\Delta t_p)^2} \quad (\text{IV.15a})$$

$$\text{et } \varphi_p(t) = -\alpha_p t^2 \quad (\text{IV.15b})$$

Nous supposons ensuite qu'elle est résonnante avec la transition de pompe c'est à dire que $\omega_p = \omega_{eg}$. La pulsation instantanée de cette impulsion (donnée par la dérivée de la phase totale $[\omega_p t - \varphi_p(t)]$) vaut :

$$\omega(t) = \omega_p + 2\alpha_p t = \omega_{eg} + 2\alpha_p t \quad (\text{IV.16})$$

On observe donc bien une dérive de fréquence linéaire en temps. On obtient une telle impulsion en ajoutant dans le domaine spectral une phase quadratique $\phi_p^{(2)}$ à une impulsion limitée par transformée de Fourier de durée initiale Δt_{p0} . On a les relations suivantes (voir par exemple la référence [124]) :

$$\Delta t_p = \Delta t_{p0} \sqrt{1 + \left(\frac{2\phi_p^{(2)}}{\Delta t_{p0}^2} \right)^2} \quad (\text{IV.17a})$$

$$\alpha_p = \frac{2\phi_p^{(2)}}{\Delta t_{p0}^4 + 4\phi_p^{(2)2}} \quad (\text{IV.17b})$$

Fort de ces suppositions, on peut ré-écrire l'équation IV.5 page 97 sous la forme :

$$a_e(t) \propto \int_{-\infty}^t A_p(t') e^{i(\delta\omega_{eg}t' + \alpha_p t'^2)} dt' \quad (\text{IV.18})$$

L'évolution de cette amplitude de probabilité ainsi que celle de la population qui y est associée sont représentées sur la figure IV.2 page suivante. La partie gauche (a) correspond

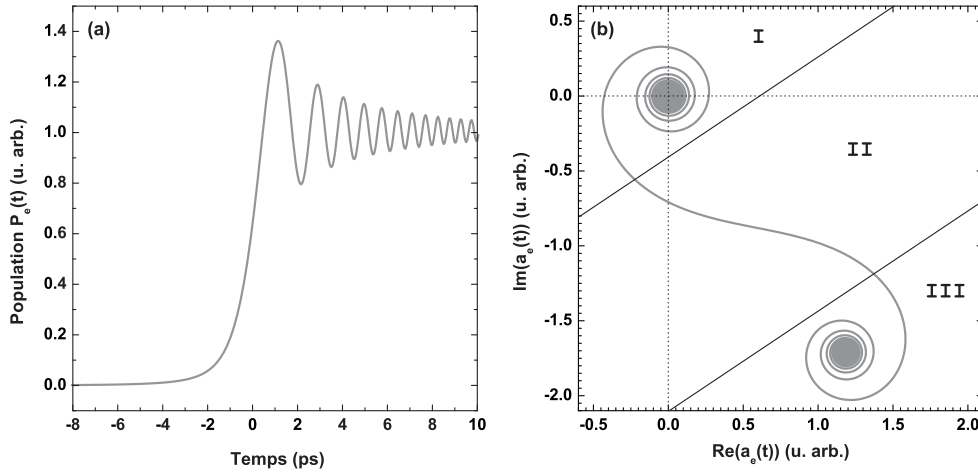


FIG. IV.2 – Population dans l'état excité P_e (a) et amplitude de probabilité a_e (b) en fonction du délai pompe-sonde. La pompe est résonnante avec la transition pompe ($\omega_p = \omega_{eg}$), possède une phase quadratique $\phi_p^{(2)} = 8.10^5 \text{ fs}^2$. Sa durée est de l'ordre de $\Delta t_p = 20 \text{ ps}$ (durée sans dérive de fréquence de 130 fs).

à la population dans l'état excité $P_e(t)$. La partie droite (b) présente l'évolution, dans le plan complexe, de l'amplitude de probabilité $a_e(t)$. Dans ce plan, la population à un délai donné correspond au carré de la distance entre l'origine (0,0) et la position actuelle. Le champ de pompe est choisi à résonance ($\delta\omega_{eg} = 0$), avec une durée initiale de $\Delta t_{p0} = 130 \text{ fs}$ et une phase quadratique spectrale de $\phi_p^{(2)} = 8.10^5 \text{ fs}^2$. La durée de l'impulsion de pompe est ainsi de l'ordre de $\Delta t_p = 20 \text{ ps}$.

On observe sur la population $P_e(t)$ un accroissement rapide autour de $t = 0$ suivi d'oscillations transitoires qui se produisent durant l'interaction avec l'impulsion de pompe. Ce sont ces oscillations que nous appelons transitoires cohérents. L'accroissement à $t = 0$ correspond au passage par la résonance, c'est-à-dire à l'arrivée dans l'impulsion de la pulsation ω_{eg} . A cette contribution principale vient ensuite s'ajouter celles des autres pulsations, non résonnantes, qui arrivent linéairement au cours du temps (voir l'équation IV.16 page précédente). Selon la différence entre la phase accumulée par ces pulsations et celle accumulée par la contribution résonnante, la population va s'accroître (contribution constructive) ou diminuer (contribution destructive).

On peut aussi voir ces oscillations comme des battements entre l'amplitude de probabilité transférée à résonance, qui évolue librement dans l'état $|e\rangle$ en accumulant une phase $\omega_{eg}t$, et la composante spectrale non résonnante qui intervient à un instant t et a accumulé une phase⁴ $\int_{-\infty}^t \omega(t')dt' = \omega_{eg}t + \alpha_p t^2$. On a alors à un instant t un déphasage $\alpha_p t^2$ entre les deux contributions ce qui donne des oscillations qui s'accroissent linéairement avec le temps. On voit donc bien que c'est la phase temporelle quadratique qui est responsable des oscillations transitoires.

⁴en supposant l'impulsion de pompe centrée à résonance : $\omega_p = \omega_{eg}$

L'effet de ce terme de phase se voit mieux sur l'évolution de l'amplitude de probabilité dans le plan complexe (figure IV.2 (b)). L'évolution au cours du temps reproduit une spirale de Cornu similaire à celle rencontrée dans la diffraction d'une onde plane par un bord d'écran, dans l'approximation de Fresnel[125]. Cette évolution peut se décomposer en trois étapes :

- Pour un temps infiniment négatif, il n'y a pas de population dans l'état excité et l'amplitude de probabilité est en (0,0). Pour un délai négatif (zone **I**), la phase temporelle évolue très rapidement alors que le rayon augmente lentement de façon monotone : l'amplitude de probabilité se déroule autour de (0,0). Plus le délai se rapproche du délai nul, plus le rayon augmente et plus l'accroissement de la phase se ralentit.
- Lorsque le délai est suffisamment proche de zéro (zone **II**), on observe le passage par la résonance. Ce passage correspond, dans l'intégrale de l'équation IV.18 page 99, à une condition de phase stationnaire : la phase n'évolue presque plus et le rayon s'accroît fortement.
- La dernière étape d'évolution (zone **III**) correspond aux délais positifs où l'on observe les oscillations sur la population. Dans le plan complexe, cette étape est tout à fait similaire à la première étape. Dans le cas présent (où $\delta\omega_{eg} = 0$), elle est même le symétrique de la première étape par rapport au passage par la résonance. Le terme en $i\alpha_p t^2$ entraîne une évolution de plus en plus rapide de la phase alors que la population atteint progressivement sa valeur asymptotique. La différence par rapport à la première étape est que la spirale s'enroule maintenant autour d'une valeur non nulle. Le rayon, et donc la population, vont ainsi présenter des maxima et des minima.

Lorsque l'interaction est terminée, l'amplitude de probabilité et la population asymptotiques valent :

$$a_e(\infty) = -\frac{\mu_{eg}}{2i\hbar}\sqrt{2\pi}E_p(\omega_{eg}) \propto A_p(\omega_{eg})e^{i\phi_p(\omega_{eg})} \quad (\text{IV.19a})$$

$$P_e(\infty) = |a_e(\infty)|^2 \propto A_p^2(\omega_{eg}) \quad (\text{IV.19b})$$

On a donc une population asymptotique $P_e(\infty)$ totalement indépendante de la phase spectrale à résonance $\phi_p(\omega_{eg})$. Cette population dépend uniquement de l'intensité spectrale à la pulsation de résonance de la transition $|e\rangle \mapsto |f\rangle$. Quel que soit le profil de phase spectrale $\phi_p(\omega)$ de l'impulsion de pompe, le transfert de population sera le même. Pour la population, la phase n'intervient que sur l'évolution transitoire d'où le nom de *transitoires* cohérents.

Par contre, la phase de la probabilité asymptotique $a_e(\infty)$ est directement liée à la phase de l'impulsion à la pulsation de résonance. Dans le plan complexe, la spirale de Cornu converge asymptotiquement vers un point dont la position dépend de $\phi_p(\omega_{eg})$. Cette capacité qu'a l'atome de garder en mémoire la phase du champ de pompe sera un élément clef par la suite. Elle nous permettra de réaliser des mesures de fonctions d'onde atomique ainsi que des caractérisations d'impulsions courtes.

Rappelons enfin que ces transitoires cohérents sont en fait parfaitement analogues à la diffraction d'une onde monochromatique par un bord d'écran, dans l'approximation de

Fresnel[125]. Cette analogie est détaillée dans les références [122, 55, 18] et a été exploitée pour la réalisation de lentilles de Fresnel temporelles [18].

IV.2 Observation expérimentale

La première observation expérimentale des transitoires cohérents a été réalisée au sein de l'équipe, en 2001, durant la thèse de Sébastien Zamith[17]. L'expérience a été réalisée avec la transition $5S_{\frac{1}{2}}$ (état $|g\rangle$), $5P_{\frac{1}{2}}$ (état $|e\rangle$) dans l'atome de rubidium. Cette transition a le double avantage de présenter une force d'oscillateur élevée et une longueur d'onde de résonance autour de 795 nm, proche de la longueur d'onde centrale de notre amplificateur. Les niveaux mis en jeu sont présentés sur la figure IV.3 page suivante. L'impulsion de pompe $E_p(t)$ excite la transition $5S_{\frac{1}{2}} \mapsto 5P_{\frac{1}{2}}$. Elle possède une forte phase quadratique et sa durée étirée est de l'ordre de 20 ps. Des précautions particulières sont prises pour éviter d'exciter le deuxième état du doublet de structure fine $5P_{\frac{1}{2}}$ centré autour de 780 nm. La population transférée dans l'état $5S_{\frac{1}{2}}$ à l'instant τ est mesurée à l'aide d'une impulsion de sonde $E_s(t)$, décalée de τ par rapport à la pompe. Elle est beaucoup plus courte que cette dernière (durée de l'ordre de 30 fs) et que le temps caractéristique des oscillations du transitoire. Nous la supposons aussi limitée par transformée de Fourier. Elle transfère une part de la population vers des niveaux supérieurs ($nS, n'D$). La cascade radiative de ces états vers le fondamental produit un signal de fluorescence à 420 nm. C'est cette fluorescence qui est détectée en fonction du délai de pompe-sonde τ : elle est directement proportionnelle à la population dans l'état excité⁵ au temps τ .

Le schéma du dispositif expérimental est présenté figure IV.4 page suivante. La sortie de l'amplificateur régénératif (795 nm, 130 fs, 1 mJ, 1 kHz) est scindée en deux. Une première partie pompe un NOPA qui produit l'impulsion de sonde (607 nm, 30 fs, 5 μ J, 1 kHz). La deuxième partie est étirée dans un étireur à réseaux en double passage. Celui-ci introduit une phase quadratique importante $\phi_p^{(2)} = -8,4.10^5 \text{ fs}^2$ et donne l'impulsion de pompe (795 nm, 20 ps, 5 μ J, 1 kHz). Le délai de pompe sonde est introduit par une table à retard motorisée et les deux impulsions sont recombinaées et envoyées dans une cellule de rubidium. Cette cellule de 80 mm est fermée par des fenêtres de 3 mm d'épaisseur, taillées à Brewster. Elle est composée d'un queusot contenant du rubidium solide et d'une chambre où ce dernier est sous forme de gaz. La pression de gaz est contrôlée par la température du queusot. Pour éviter la déposition de rubidium sur les fenêtres, la chambre (83° C) est toujours plus chaude que le queusot (69° C) qui fait office de point froid. La pression correspondante est de l'ordre de $2,7.10^{-3} \text{ Pa}$. La détection de la fluorescence est assurée par deux lentilles de grandes ouvertures et un photomultiplicateur (RCA IP28) qui sont accolés à la cellule. La sélection spectrale est réalisée par un filtre coloré. Le photomultiplicateur est porté à une tension de 800V et relié à un intégrateur.

Le signal expérimental obtenu est présenté sur la figure IV.5 page 104. Les carrés noirs correspondent au signal de fluorescence, le trait noir à la théorie, en fonction du délai de pompe-sonde. L'impulsion de pompe dure environ 20 ps et on observe donc l'évolution

⁵Ceci est lié aux suppositions faites sur la sonde. Son effet est étudié en détail dans la partie IV.1.2 page 97 ainsi que dans la référence [122].

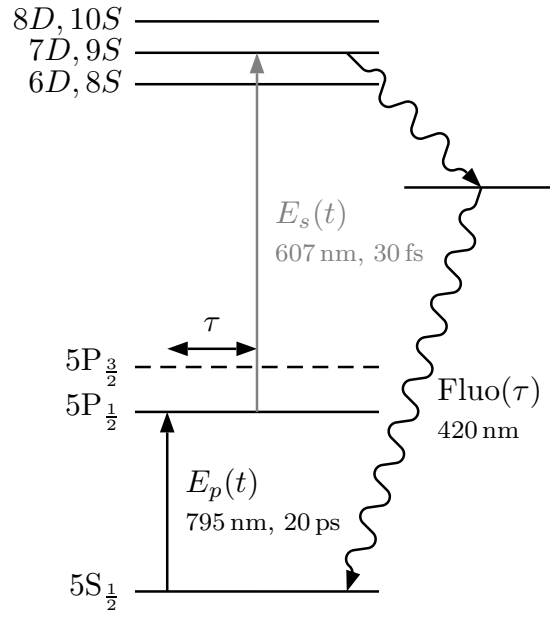


FIG. IV.3 – Niveaux atomiques du rubidium utilisés pour la détection des transitoires cohérents par pompe-sonde.

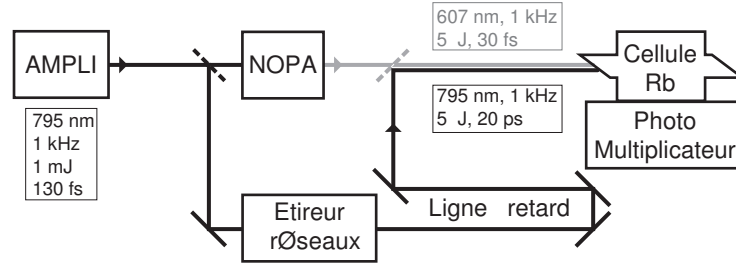


FIG. IV.4 – Dispositif expérimental pour l'observation des transitoires cohérents par pompe-sonde.

de la population dans l'état excité au cours de son interaction avec la pompe. Sur la courbe théorique, l'amplitude, la ligne de base et le délai zéro sont déduits de l'expérience. La phase quadratique est quant à elle ajustée pour reproduire au mieux la période des oscillations. La phase quadratique de l'ajustement est en bon accord avec la géométrie de l'étireur à réseaux.

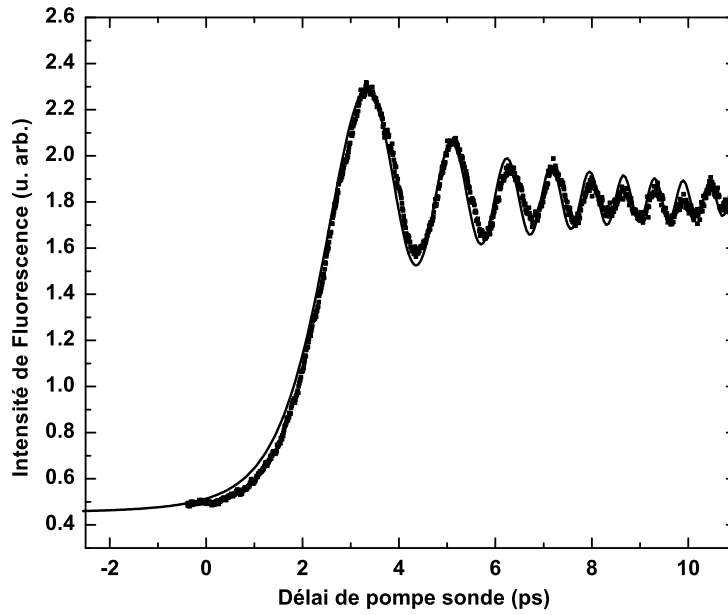


FIG. IV.5 – Transitoires cohérents dans le Rubidium. La fluorescence (carrés noirs) observée en fonction du délai de pompe-sonde reproduit très fidèlement les transitoires cohérents théoriques attendus (trait noir). L'impulsion de pompe, résonnante, a une phase quadratique de $\phi_p^{(2)} = -8,4 \cdot 10^5 \text{ fs}^2$.

IV.3 Contrôle des transitoires

Nous venons de voir dans la partie précédente que l'évolution des transitoires cohérents atomiques est fortement liée au profil de phase de l'impulsion de pompe. Il s'agit donc d'un processus qui se prête extrêmement bien à des expériences de contrôle cohérent : en jouant sur le profil de phase (et éventuellement d'amplitude) de l'impulsion de pompe, on doit pouvoir contrôler le comportement de la population dans l'état $|e\rangle$. Plusieurs expériences de ce type ont été mises en œuvre, au sein de l'équipe ainsi que dans l'équipe de Yaron Silberberg au Weizmann Institute of Science.

IV.3.1 Quelques expériences.

L'équipe de Silberberg a ainsi maximisé le régime transitoire[16], au détriment de la population asymptotique. Dans cette expérience, l'impulsion de pompe ne présente pas de dérive de fréquence. Par contre, un saut de phase est appliqué dans le domaine spectral, à proximité de la pulsation de résonance ω_{eg} . Ils ont ainsi montré qu'en appliquant des sauts de l'ordre de π , il était possible d'obtenir un régime transitoire intense et très piqué (voir la figure IV.6 page suivante). De façon assez schématique, on peut dire que pour une impulsion de pompe $E_p(t)$ normale (c'est à dire sans saut de phase, limitée par transformée de Fourier), les contributions non résonnantes interfèrent destructivement deux à deux et on obtient un transfert de population en forme de marche (trait noir sur la figure IV.6) lié

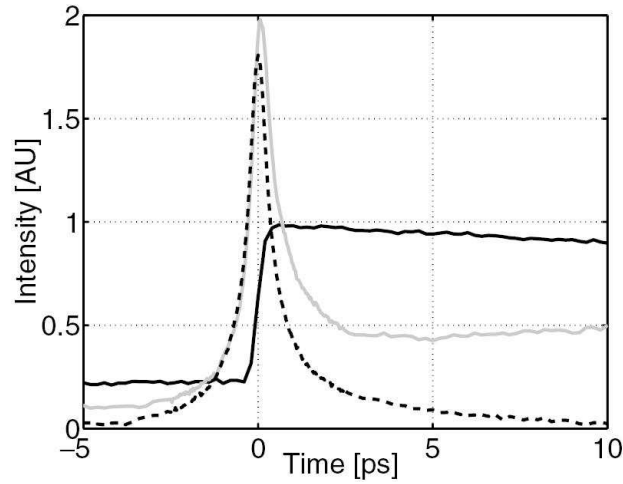


FIG. IV.6 – **Figure tirée de la référence [16]** Transient population measurement as a function of the probe delay for the pulse shapes : transform limited pump pulse (solid black line), on resonant p-step shaped pulse (dashed line), shifted 0.7p step (solid gray line).

à la contribution à résonance. On peut ainsi écrire l'amplitude de probabilité au moment du passage par résonance (c'est-à-dire $a_e(0)$) sous la forme

$$a_e(0) = \frac{a_e(+\infty)}{2} \left[1 + \frac{i}{\pi \tilde{E}_p(\omega_{eg})} \mathcal{P}_P \int_{-\infty}^{\infty} \frac{\tilde{E}_p(\omega)}{\omega_{eg} - \omega} d\omega \right] \quad (\text{IV.20})$$

où $\mathcal{P}_P$ représente la partie principale de Cauchy.

Dans le domaine spectral, en appliquant un saut de phase de π au niveau de la pulsation de résonance, la moitié des composantes non résonnantes change de signe et les interférences deux à deux destructives deviennent constructives. On obtient alors un transitoire très piqué (tirets noirs sur la figure IV.6) avec un transfert asymptotique plus faible. Nous n'entrerons pas plus dans les détails de cette expérience de contrôle cohérent et nous invitons les lecteurs intéressés à consulter l'article d'origine[16] ou la thèse de Jérôme Degert[55].

Au sein de l'équipe, plusieurs expériences ont été réalisées. La première combine phase quadratique et saut de phase pour déphaser les transitoires cohérents observés précédemment. Cette expérience a une importance particulière pour les travaux réalisés au cours de cette thèse et est donc détaillée un peu plus loin (partie IV.3.2 page suivante).

Une deuxième expérience utilise la forte analogie entre les transitoires cohérents et la diffraction dans l'approximation de Fresnel pour réaliser une lentille à zones de Fresnel temporelle[18, 55]. Une lentille à zone de Fresnel est en fait un masque qui bloque certaines zones judicieusement choisies du profil spatial d'un faisceau lumineux afin de le focaliser. Cette approche se transpose parfaitement au contrôle des transitoires cohérents présentés sur la figure IV.2. En effet, si l'on éteint le champ de pompe sur chaque intervalle temporel où la population $P_e(t)$ décroît, on transformera les oscillations en marches d'escalier et on

augmentera la population asymptotique. Cette intuition est confirmée par les simulations théoriques réalisées par Jérôme Degert et présentées sur la figure IV.7.

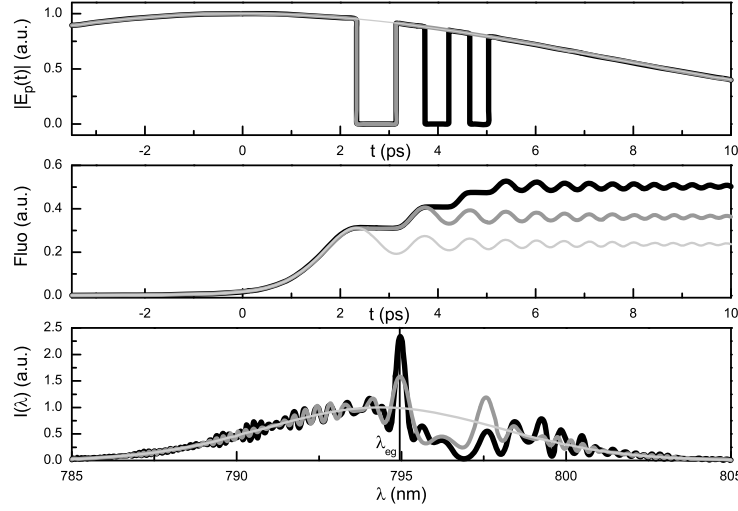


FIG. IV.7 – **Figure tirée de la référence [55]** Application du principe des zones de Fresnels au domaine temporel. L'impulsion pompe ($\lambda_p = 794,4$ nm, $\Delta t_{p0} = 110$ fs, $\phi_p^{(2)} = -4,9 \cdot 10^5$ fs²) est découpée temporellement afin d'accroître la population asymptotique. Les trois graphes représentent respectivement l'amplitude du champ électrique de l'impulsion pompe, la population de l'état excité en fonction du temps et le spectre de l'impulsion. La courbe en gris clair correspond aux transitoires cohérents non modifiés, la courbe en gris au cas où l'on a supprimé la première contribution destructive à $a_e(t)$ et la courbe en noir au cas où l'on a supprimé les trois premières contributions destructives.

L'expérience a été réalisée en collaboration avec l'équipe de Marcus Motzkus du Max-Planck-Institute of Quantum Optics. La figure IV.8 page ci-contre présente le résultat obtenu pour un trou au niveau de la première interférence destructive. On observe bien une focalisation temporelle de la population.

IV.3.2 Déphasage des transitoires

Les deux parties précédentes mettent en avant la grande sensibilité de l'amplitude de probabilité à la phase et à l'amplitude de l'impulsion de pompe. Il n'y a rien d'étonnant à cela : dans l'approximation perturbative, le champ de pompe et l'amplitude de probabilité sont simplement reliés par une relation d'intégration (voir l'équation IV.5 page 97).

Il est ainsi possible de contrôler l'évolution des transitoires cohérents en contrôlant la phase temporelle du champ de pompe[18].

En particulier, les transitoires cohérents observés pour une pompe à forte dérive de fréquence peuvent être inversés en changeant le signe du champ de pompe au cours de l'inter-

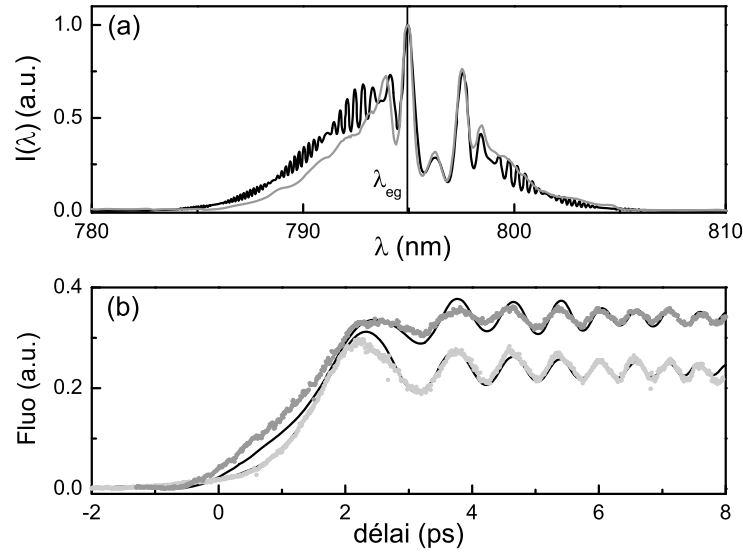


FIG. IV.8 – **Figure tirée de la référence [55]** Augmentation de la population transférée. (a) Spectre calculé (noir) et spectre réalisé expérimentalement (gris). (b) Population de l'état excité pour une impulsion découpée (gris) et pour l'impulsion initiale (gris clair), ainsi que les courbes théoriques obtenues à partir des spectres expérimentaux (noir).

action avec l'atome. Dans la pratique, il est extrêmement difficile, voire même impossible, d'agir directement sur la phase ou l'amplitude temporelle d'une impulsion femtoseconde et on agit dans le domaine spectral.

Nous allons montrer ici qu'il est possible de changer le signe du champ temporel, et donc de déphaser les transitoires cohérents, en appliquant un saut de phase de π dans le domaine spectral.

IV.3.2.1 Homothétie temps fréquence

Pour pouvoir comprendre de façon qualitative ce qui se passe lorsqu'on applique un saut de phase de π à une impulsion à forte dérive de fréquence, nous allons nous appuyer sur l'homothétie temps-fréquence.

Considérons tout d'abord une impulsion de pompe à forte dérive de fréquence. Nous avons vu que la pulsation instantanée d'une telle impulsion varie linéairement avec le temps (équation IV.16 page 99). Dans le cas présent d'une forte dérive ($\phi_p^{(2)} \gg \Delta t_{p0}^2$) on a la relation :

$$\omega(t) = \omega_p + \frac{t}{\phi_p^{(2)}} \quad (\text{IV.21})$$

A chaque instant t_0 au cours de l'impulsion, on retrouve une pulsation apparente $\omega_0 = \omega(t_0)$. Prenons par exemple la pulsation de la transition $|g\rangle \rightarrow |e\rangle$: elle est présente dans l'impulsion à un instant $t_{eg} = \phi_p^{(2)}(\omega_{eg} - \omega_p)$. Autour du temps t_{eg} , la pulsation instantanée

de l'impulsion sera voisine de ω_{eg} . Plus précisément, sur un intervalle Δt centré en t_{eg} , la pulsation de la pompe vaudra ω_{eg} à $\Delta\omega$ près avec :

$$\Delta\omega = \frac{\Delta t}{|\phi_p^{(2)}|} \quad (\text{IV.22})$$

Cependant, comme le temps et la pulsation sont des variables couplées, on doit toujours vérifier :

$$\Delta\omega \cdot \Delta t \geq 4 \ln(2) \quad (\text{IV.23})$$

Donc, même si l'on considère un intervalle spectral $\Delta\omega$ infiniment étroit, on aura toujours un intervalle temporel borné vérifiant :

$$\Delta t \geq \sqrt{4 \ln(2) |\phi_p^{(2)}|} \quad (\text{IV.24})$$

En d'autres termes, la pulsation ω_{eg} sera présente dans l'impulsion sur un intervalle temporel d'au moins $\delta t = \sqrt{4 \ln(2) |\phi_p^{(2)}|}$, centré en t_{eg} . Ce résultat est une façon d'expliquer la durée du passage par résonance des transitoires cohérents, qui est de l'ordre⁶ de δt .

Ce résultat se retrouve graphiquement sur une représentation temps-fréquence de l'impulsion à forte dérive de fréquence (représentation de type Wigner[126] ou Husimi). Une telle représentation est schématisée sur la figure IV.9, en supposant pour la pompe un spectre gaussien. L'intensité de la pompe dans le plan temps-fréquence est donnée par une

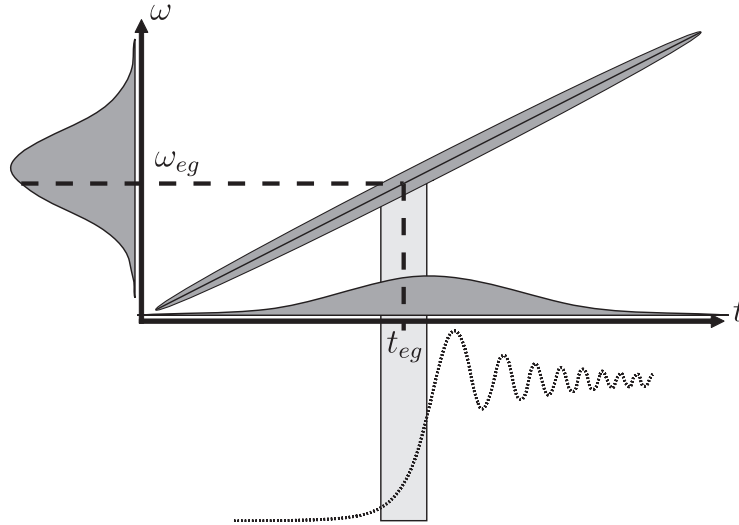


FIG. IV.9 – Représentation temps-fréquence de l'intensité d'une impulsion de pompe à forte dérive de fréquence. Le temps d'arrivée t_{eg} de la pulsation de résonance ω_{eg} correspond au passage par la résonance des transitoires cohérents (pointillés noirs) créés par la pompe.

gaussienne, de section elliptique, schématisée par l'ellipse grise. Cette ellipse est d'autant

⁶La valeur exacte dépendant du critère choisi pour définir la résonance.

plus inclinée par rapport à la verticale que la phase quadratique $\phi_p^{(2)}$ est importante. Les projections sur les deux axes redonnent respectivement l'intensité spectrale (axe vertical) et l'intensité temporelle (axe horizontal). Les transitoires cohérents induits par une telle pompe sont tracés sous l'axe temporel. On retrouve bien sur ce schéma que la pulsation ω_{eg} arrive à l'instant t_{eg} correspondant au passage par la résonance des transitoires cohérents. De plus, ce passage par résonance (intervalle en gris clair) est d'autant plus long que l'ellipse est penchée, c'est à dire que $\phi_p^{(2)}$ est grande.

Considérons maintenant la même impulsion de pompe, à laquelle on ajoute un saut de phase de π , dans le domaine spectral. Ce saut de phase correspond à une phase additionnelle $\phi_\pi(\omega)$ valant :

$$\phi_\pi(\omega) = \begin{cases} 0 & \text{si } \omega < \omega_\pi, \\ \pi & \text{sinon.} \end{cases} \quad (\text{IV.25})$$

Pour comprendre qualitativement l'effet de ce saut de phase spectral dans le domaine temporel, nous reprenons la représentation temps-fréquence précédente où l'on ajoute le saut de phase.

Ceci est schématisé sur la figure IV.10 où l'on a aussi rajouté les transitoires cohérents induits par la pompe avec saut de phase (trait noir). Compte tenu de l'homothétie temps-

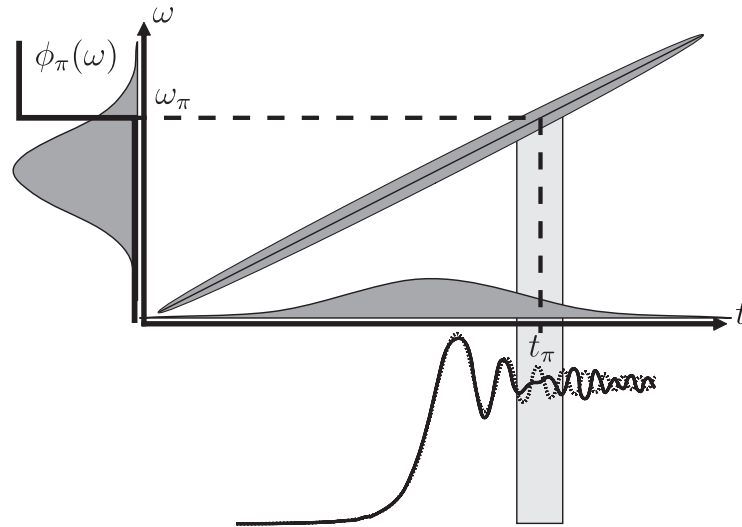


FIG. IV.10 – Effet d'un saut de phase sur une pompe à forte dérive de fréquence. Le saut de phase de π appliqué en ω_π change le signe de l'enveloppe temporelle de la pompe en t_π . Cette inversion dure environ $3,6\sqrt{|\phi_p^{(2)}|}$ (intervalle gris clair). Après l'inversion, les transitoires cohérents avec saut de phase (trait noir) sont en opposition de phase avec ceux observés habituellement (pointillés gris).

fréquence, ce saut de phase de π dans le domaine spectral se traduit par un changement de signe du champ temporel à un instant $t_\pi = \phi_p^{(2)}(\omega_\pi - \omega_{eg})$. Ce changement de signe se

fait sur un intervalle de temps $[t_1, t_2]$ centré en t_π et de durée de l'ordre de δt . La durée exacte de cet intervalle dépend de la modélisation du saut de phase. Pour une fonction de type Heaviside, il a été montré[55] que cette durée est de l'ordre de $3,6\sqrt{|\phi_p^{(2)}|}$.

On s'attend donc à voir des transitoires cohérents habituels avant l'intervalle $[t_1, t_2]$, une période de transition durant cet intervalle, et des transitoires en opposition de phase après.

IV.3.2.2 Réalisation expérimentale

Ce déphasage des transitoires cohérents a été observé expérimentalement lors d'expériences réalisées en collaboration avec l'équipe de Marcus Motzkus. Le même dispositif expérimental que pour l'observation des transitoires cohérents (voir la figure IV.4 page 103) est utilisé mais un façonneur est inséré pour introduire une phase spectrale additionnelle. Le schéma expérimental est présenté sur la figure IV.11. Ce façonneur introduit un saut de phase de π au niveau d'une pulsation ω_π que l'on peut faire varier.

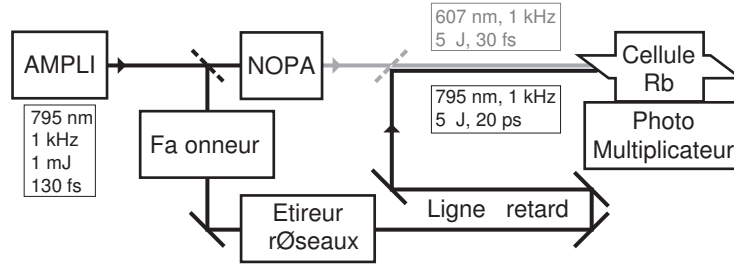


FIG. IV.11 – Dispositif expérimental pour le déphasage de transitoires cohérents. Il s'agit du dispositif présenté figure IV.4 page 103, où l'on a ajouté un façonneur pour contrôler la phase spectrale de l'impulsion de pompe.

La figure IV.12 page suivante présente les résultats obtenus lors de la thèse de Jérôme Degert[55]. Les transitoires cohérents normaux fig.IV.12(a) sont comparés aux transitoires obtenus lorsque l'on ajoute un saut de phase de π centré en $\lambda_\pi = \lambda_{eg} + 0,26$ nm fig.IV.12(b), en $\lambda_\pi = \lambda_{eg} + 0,52$ nm fig.IV.12(c), et en $\lambda_\pi = \lambda_{eg} + 1,04$ nm fig.IV.12(d). Les points noirs correspondent aux signaux expérimentaux et les courbes noires aux ajustements théoriques. Pour les figures IV.12(b),(c) et (d) la courbe en pointillé reproduit l'ajustement théorique en l'absence de saut de phase. Les transitoires expérimentaux en présence de saut de phase de π sont en très bon accord avec le comportement attendu et on observe clairement les trois zones prévues : avant, pendant et après l'intervalle $[t_1, t_2]$. On constate aussi que les transitoires cohérents sont extrêmement sensibles à la position du saut de phase de π . En particulier, entre les courbes IV.12(b) et (c), la position du saut de phase a été décalée de seulement 0,26 nm et on peut facilement distinguer les deux courbes. Cet incrément de 0,26 nm est en fait le plus petit décalage réalisable avec le façonneur utilisé à l'époque : cela correspond à la largeur spectrale d'un pixel du masque de ce façonneur. Des simulations effectuées par Jérôme Degert[55] et présentées sur la figure IV.13 page 112 montrent que

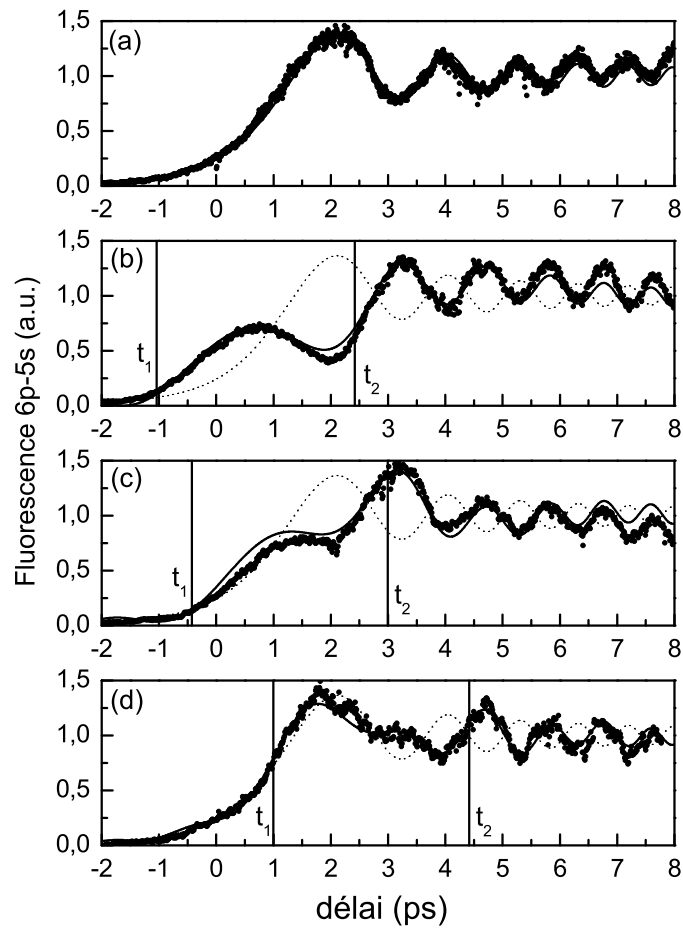


FIG. IV.12 – **Figure tirée de la référence [55].** Transitoires cohérents normaux (a) et en présence de saut de phase de π : en $\lambda_\pi = 795,23$ nm (b), en $\lambda_\pi = 795,49$ nm (c) et en $\lambda_\pi = 796,01$ nm (d). Les signaux expérimentaux sont en ronds noirs et les ajustements théoriques en traits noirs.

les transitoires cohérents doivent permettre de discriminer des décalages bien plus faibles, de l'ordre de 0,09 nm.

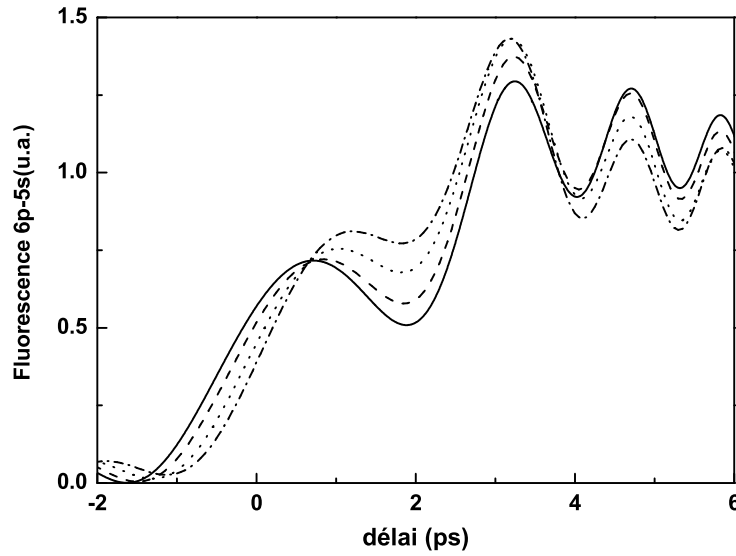


FIG. IV.13 – **Figure tirée de la référence [55].** Etude théorique de l'évolution des transitoires cohérents lorsqu'on déplace le saut de phase de π avec un pas de l'ordre de 0,087 nm). On explore la zone située entre la figure IV.12(b) et la figure IV.12(c). Initialement (courbe en traits pleins) le saut de phase est centré un pixel avant la résonance. Les courbes suivantes correspondent à un décalage de : 1,3 pixels (tirets), 1,6 pixels (pointillés), 1,9 pixels (tirets et pointillés).

IV.4 Sensibilité au champ

Les résultats expérimentaux et théoriques obtenus par Jérôme Degert durant sa thèse dans l'équipe nous ont poussés à explorer la sensibilité des transitoires cohérents à l'aide du nouveau façonneur haute résolution. En effet, celui-ci présente une résolution théorique de 0,06 nm/pixel plus fine que celle utilisée dans les précédentes simulations. Nous souhaitons voir si la sensibilité des transitoires cohérents pouvaient être mise à profit pour la caractérisation et l'alignement de ce nouveau façonneur.

Le dispositif utilisé pour la réalisation expérimentale de ce test est présenté sur la figure IV.14 page suivante. Il est très proche de celui utilisé lors de la réalisation expérimentale du déphasage des transitoires (figure IV.11 page 110). La différence la plus notable est l'absence d'étireur à réseaux. En effet, les capacités de mise en forme de notre façonneur lui permettent d'appliquer à l'impulsion de pompe à la fois la forte phase quadratique (précédemment introduite par les réseaux) et le saut de phase de π .

Nous avons déplacé le saut de phase par pas de 1 pixel, soit 0,06 nm, sur une large plage de longueur d'onde. Dans tous les cas, deux transitoires successifs sont clairement distincts. La figure IV.15 page suivante présente ainsi les résultats obtenus pour des sauts de phase distants de la résonance de 0,30 nm (courbe 1) à 0,54 nm (courbe 5). La partie gauche présente les simulations théoriques et la partie droite les résultats expérimentaux. La

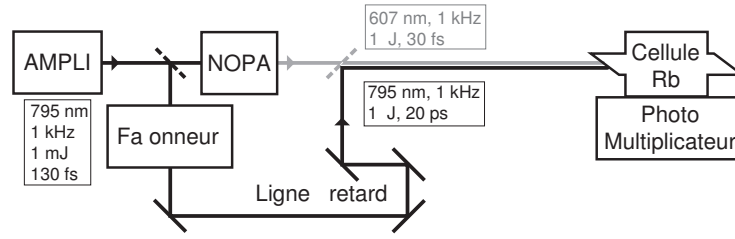


FIG. IV.14 – Dispositif expérimental pour tester la sensibilité des transitoires cohérents. Le façonneur introduit à la fois la forte phase quadratique et le saut de phase de π .

phase quadratique introduite par le façonneur pour cette expérience est de $\phi_p^{(2)} = 9.10^5 \text{ fs}^2$.

On différencie très bien deux transitoires successifs. Le comportement global des transitoires expérimentaux est en bon accord avec les simulations théoriques, notamment la phase relative entre les transitoires avant résonance. La principale différence est le rapport d'amplitude entre l'épaule à 1 ps et le maximum à 3 ps : sur les transitoires expérimentaux, l'épaule est plus intense que prévu. Nous attribuons cette différence à un défaut du façonneur non compensé à l'époque où ces expériences ont été réalisées. En effet, le façonneur génère, en plus de l'impulsion à dérive de fréquence avec saut de phase, une réplique limitée par transformée de Fourier située à $t = 0$. En l'absence de correction, cette réplique vient perturber l'évolution des transitoires cohérents à proximité du passage par la résonance. Malgré ce défaut, les transitoires cohérents permettent une

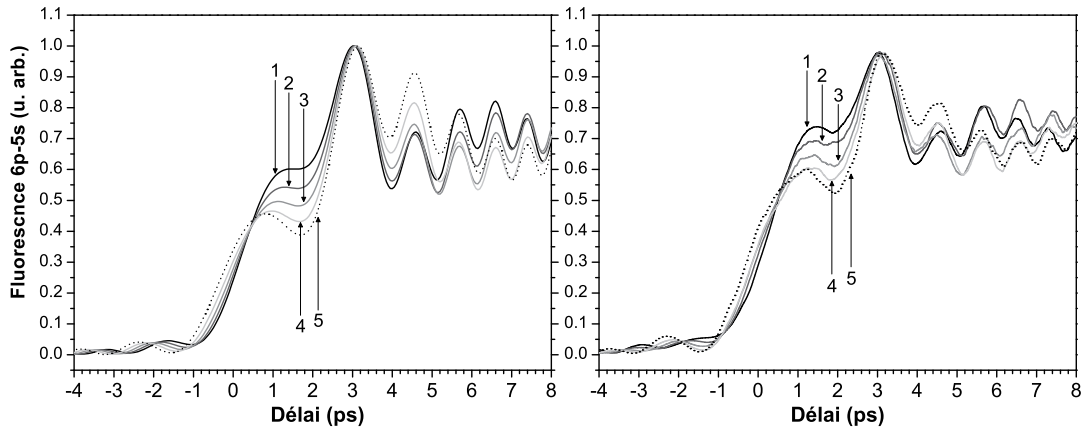


FIG. IV.15 – Transitoires cohérents théoriques (à gauche) et expérimentaux (à droite) obtenus pour différentes positions du saut de phase de π . Les courbes 1 à 5 correspondent respectivement à des sauts placés à 0,30 nm, 0,36 nm, 0,42 nm, 0,48 nm et 0,54 nm de la résonance. Le façonneur introduit en outre une phase quadratique $\phi_p^{(2)} = 9.10^5 \text{ fs}^2$.

mesure extrêmement précise de la position du saut de phase de π . En comparaison, l'autre

possibilité pour mesurer cette position est de regarder le trou en amplitude apparaissant sur le spectre de l'impulsion de pompe (conséquence de l'application du saut de phase) à l'aide d'un spectromètre. Avec le spectromètre dont nous disposons, cette mesure ne se fait qu'à 0,2 nm près. De plus elle ne permet pas de distinguer le saut de phase d'un vrai trou en amplitude.

De façon générale, la plupart des techniques de caractérisation de la phase (ou de l'amplitude et de la phase) d'impulsions ultracourtes ont une résolution bien moins bonne que les 0,06 nm que nous avons observé avec les transitoires cohérents. Il en résulte que les transitoires cohérents sont suffisamment sensibles pour servir à la calibration des façonneurs de haute résolution[47]. De plus cette calibration se fait sans autre matériel (spectromètre, SPIDER ou FROG) que le dispositif expérimental servant à l'étude du système atomique.

Il serait très intéressant de pouvoir utiliser cette sensibilité extrême pour mettre en œuvre une technique de caractérisation complète d'impulsions ultracourtes. Nous allons voir que les transitoires cohérents permettent une telle caractérisation, moyennant quelques modifications du schéma de pompe.

Cette technique de caractérisation passe par une étape intermédiaire où l'on mesure une part de la fonction d'onde atomique. Cette étape est tout à fait intéressante en soi et nous allons la détailler. Nous verrons ensuite comment exploiter cette information sur la fonction d'onde atomique pour reconstruire le champ de pompe. Nous détaillerons enfin les caractéristiques assez inhabituelles de cette technique de mesure complète.

IV.5 Spirographe atomique

Comme nous l'avons vu dans les parties précédentes, les différentes expériences de transitoires cohérents ou de contrôle de ces derniers font interagir les champs de pompe $E_p(t)$ et de sonde $E_s(t)$ avec le système atomique à trois niveaux $\{|g\rangle, |e\rangle, |f\rangle\}$. L'état du système atomique est entièrement déterminé par sa fonction d'onde $|\Psi(t)\rangle$ définie précédemment (équation IV.1 page 96) et dont nous rappelons l'expression :

$$|\Psi(t)\rangle = a_g(t) |g\rangle e^{-i\omega_g t} + a_e(t) |e\rangle e^{-i\omega_e t} + a_f(t) |f\rangle e^{-i\omega_f t} \quad (\text{IV.26})$$

Sous certaines conditions, il est possible de résoudre temporellement, à l'échelle femtoseconde, l'évolution de $a_e(t)$. Nous avons donc accès à l'amplitude de probabilité de $|\Psi(t)\rangle$ sur l'état $|e\rangle$. Ceci constitue une mesure partielle de la fonction d'onde atomique. La mesure de fonction d'onde[101, 127, 128] est un domaine très actif. Dans de nombreux cas, la mesure concerne un paquet d'onde qui évolue librement[101] (comme un paquet d'onde vibrationnel évoluant dans un potentiel électronique donné). Ici, nous proposons une mesure quelque peu différente car nous nous intéressons à un seul niveau propre dont on veut mesurer l'évolution forcée sous l'action du champ de pompe.

Nous allons dans un premier temps montrer comment la combinaison de plusieurs signaux de transitoires cohérents nous permet de remonter à $a_e(t)$; technique semblable à l'holographie d'états quantiques[128]. Nous verrons ensuite quelles difficultés expérimentales nous avons dû résoudre pour mettre en œuvre cette reconstruction. Enfin, les différentes fonctions d'onde reconstruites seront détaillées et nous insisterons sur les paramètres clefs de ces reconstructions.

IV.5.1 Des transitoires cohérents à l'atome

Nous reprenons ici les résultats de la partie IV.1.2. Nous supposons en particulier que la sonde $E_s(t)$ est suffisamment courte pour être assimilée à un delta de Dirac. Sous cette condition, le signal de fluorescence mesuré $S^\delta(\tau)$ est directement proportionnel à la population dans l'état excité (équation IV.14 page 99) :

$$S^\delta(\tau) = S^\delta(\tau) \propto P_e(\tau) = |a_e(\tau)|^2 \quad (\text{IV.27})$$

Dans toute la suite, nous omettrons le facteur de proportionnalité et nous ne mentionnerons pas explicitement la supposition faite sur $E_s(t)$. Nous noterons donc simplement $S(\tau)$ le signal pompe-sonde mesuré, celui-ci valant :

$$S(\tau) = P_e(\tau) = |a_e(\tau)|^2 \quad (\text{IV.28})$$

IV.5.1.1 Séquence d'impulsions de pompe

Pour accéder à l'amplitude de probabilité à partir de la population, nous allons remplacer le champ de pompe $E_p(t)$ par une séquence de deux impulsions $E_{p1}(t)$ et $E_{p2}(t)$ dont nous contrôlerons le délai de façon interférométrique. À l'aide d'au moins deux transitoires cohérents $S(t)$ enregistrés pour des phases relatives différentes entre les deux impulsions, nous pourrions reconstruire l'amplitude de probabilité $a_{e2}(\tau)$ due à la seconde impulsion. L'idée est de tirer profit de la très grande durée de vie du niveau $|e\rangle$ devant les dynamiques étudiées ici pour créer un oscillateur local à l'aide de la première impulsion $E_{p1}(t)$. Le contrôle interférométrique nous permet alors d'ajouter la contribution de la deuxième impulsion $E_{p2}(t)$ avec un déphasage donné. On peut ainsi par exemple faire des mesures homodynes en phase ou en quadrature et donc remonter à la contribution $a_{e2}(\tau)$.

Plus précisément, à partir d'une impulsion de pompe inconnue $E_p(t)$, on génère une séquence d'impulsions $E_{p\Sigma}(t)$. Aucune supposition n'est nécessaire sur cette impulsion. Pour générer la séquence, nous utiliserons notre façonneur pour appliquer à $\tilde{E}_p(\omega)$ une fonction de transfert spectrale définie par

$$H^\theta(\omega) = \frac{1}{2}\mathbb{1} + \frac{1}{2} \exp \left[i\theta + i\phi^{(1)}(\omega - \omega_p) + i\frac{\phi_p^{(2)}}{2}(\omega - \omega_p)^2 \right] \quad (\text{IV.29})$$

où, rappelons-le, ω_p est la pulsation centrale de l'impulsion $E_p(t)$. Cette fonction de transfert génère une première impulsion identique à $E_p(t)$ et une deuxième décalée de $\tau = \phi^{(1)}$, avec une phase quadratique $\phi_p^{(2)}$ et un déphasage de θ . On aura ainsi une séquence d'impulsions $E_{p\Sigma}(t)$ que nous noterons :

$$E_{p\Sigma}(t) = \mathcal{F}^{-1} \left[H^\theta(\omega) \tilde{E}_p(\omega) \right] = E_{p1}(t) + e^{i\theta} E_{p2}(t) \quad (\text{IV.30})$$

Ainsi définies, $E_{p1}(t)$ et $E_{p2}(t)$ sont toujours en phase, la phase relative étant explicitement contenue dans le terme $\exp[i\theta]$. Nous noterons $a_{e1}(t)$ et $a_{e2}(t)$ les contributions respectives de $E_{p1}(t)$ et $E_{p2}(t)$. Compte tenu des notations de l'équation IV.30, et de la dépendance

linéaire en champ, l'amplitude de probabilité totale $a_{e\Sigma}(t)$ due à la séquence d'impulsions se met sous la forme :

$$a_{e\Sigma}(t) = a_{e1}(t) + e^{i\theta} a_{e2}(t) \quad (\text{IV.31})$$

Le signal de transitoire cohérent obtenu, pour une phase relative θ s'écrit quant à lui :

$$S^\theta(\tau) = |a_{e1}(\tau) + e^{i\theta} a_{e2}(\tau)|^2 \quad (\text{IV.32})$$

Nous supposons que le délai $\phi^{(1)}$ séparant les deux impulsions $E_{p1}(t)$ et $E_{p2}(t)$ est grand devant leurs durées respectives. Comme la durée de la sonde est déjà supposée négligeable, cela revient au final à vérifier :

$$\phi^{(1)} \gg \Delta t_{p1}, \Delta t_{p2} \gg \Delta t_s$$

Prendre cette précaution sur $\phi^{(1)}$ permet de bien séparer les contributions de $E_{p1}(t)$ et $E_{p2}(t)$: $a_{e1}(t)$ a atteint sa valeur asymptotique $a_{e1}(\infty)$ avant que $a_{e2}(t)$ ne commence à prendre des valeurs non-nulles.

Si on considère un délai pompe-sonde τ grand devant la durée de la première pompe $E_{p1}(t)$ on peut alors écrire :

$$a_{e\Sigma}(\tau) \stackrel{\tau \gg \Delta t_{p1}}{=} a_{e1}(\infty) + e^{i\theta} a_{e2}(\tau) \quad (\text{IV.33})$$

et donc

$$S^\theta(\tau) \stackrel{\tau \gg \Delta t_{p1}}{=} |a_{e1}(\infty) + e^{i\theta} a_{e2}(\tau)|^2 \quad (\text{IV.34})$$

Dans toute la suite, nous supposons implicitement que les différentes conditions temporelles sont remplies et notamment que τ est suffisamment grand pour utiliser l'équation IV.34. Nous écrirons ainsi :

$$\begin{aligned} S^\theta(\tau) &= |a_{e1}(\infty) + e^{i\theta} a_{e2}(\tau)|^2 \\ &= |a_{e1}(\infty)|^2 + |a_{e2}(\tau)|^2 + 2 \operatorname{Re} [e^{i\theta} a_{e1}^*(\infty) a_{e2}(\tau)] \end{aligned} \quad (\text{IV.35})$$

A l'aide de plusieurs enregistrements $S^\theta(\tau)$, pour des valeurs de θ différentes, il est alors possible de reconstruire $a_{e1}^*(\infty) a_{e2}(\tau)$ pour des délais $\tau \gg \Delta t_{p1}$. A un facteur multiplicatif complexe près, nous avons ainsi accès à $a_{e2}(\tau)$.

IV.5.1.2 Reconstructions possibles

Plusieurs approches sont possibles, nécessitant un nombre d'enregistrements plus ou moins grand. Comme ces enregistrements sont relativement longs à réaliser, nous avons cherché à en minimiser le nombre. Nous présentons ici différentes méthodes envisageables, en terminant par la reconstruction dite géométrique qui nous a semblé la plus pertinente.

Pour chaque reconstruction envisagée, nous montrerons une simulation des transitoires cohérents à enregistrer pour reconstruire $a_e(\tau)$. Sauf mention contraire, les impulsions $E_{p1}(t)$ et $E_{p2}(t)$ correspondront respectivement à une impulsion limitée par transformée de Fourier et à une impulsion décalée de 10 ps possédant une phase quadratique de 2.10^5 fs^2 (durée étirée de l'ordre de 10 ps). Ces deux impulsions ont des durées limitées par

transformée de Fourier de 100 fs, sont de même énergie, mutuellement cohérentes, et sont parfaitement résonnantes : $\omega_p = \omega_{eg}$.

Une première possibilité, donnant une solution unique, est de réaliser quatre enregistrements pour $\theta = \{0, \pi, \pi/2, -\pi/2\}$. Les simulations de ces enregistrements sont données sur la figure IV.16. Au délai $\tau = 0$, la première impulsion arrive et transfère de la population. Comme la première impulsion est limitée par transformée de Fourier, le régime transitoire dure peu et a une forme de marche d'escalier. Pour des délais faiblement positifs (entre 1 et 3 ps sur la figure), le signal n'évolue plus : il vaut $|a_e(\infty)|^2$. Autour de $\tau = 10$ ps, la deuxième impulsion passe par la résonance et sa contribution interfère différemment avec celle de la première impulsion selon leur phase relative θ : (a) 0 en noir et $\pi/2$ en gris et (b) π en noir et $-\pi/2$ en gris.

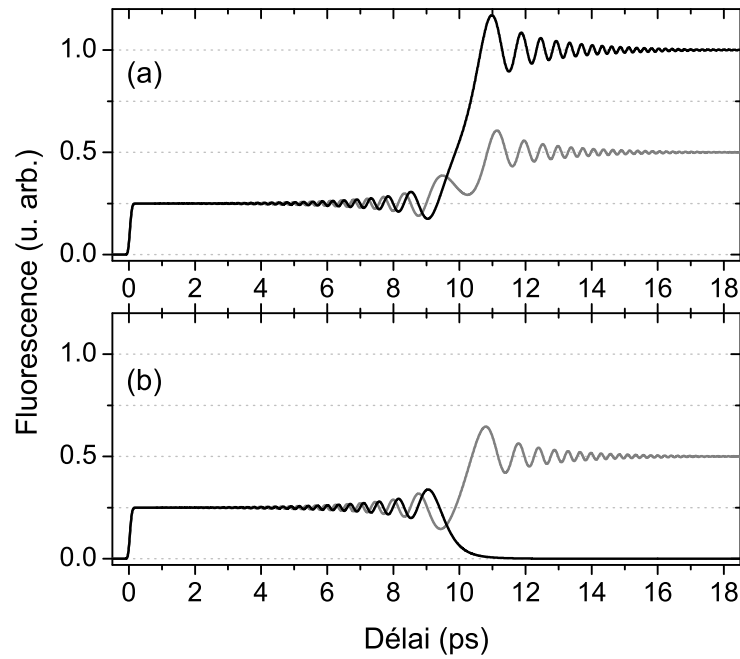


FIG. IV.16 – Signaux théoriques pour une reconstruction à 4 enregistrements : (a) phase relative de 0 en noir, de $\pi/2$ en gris et (b) phase relative de π en noir et de $-\pi/2$ en gris.

Par combinaison linéaire complexe de ces enregistrements on obtient un signal $S^{[4]}(\tau)$

$$S^{[4]}(\tau) = \frac{[S^0(\tau) - S^\pi(\tau)]}{2} + i \frac{[S^{-\pi/2}(\tau) - S^{\pi/2}(\tau)]}{2} = a_{e1}^*(\infty) a_{e2}(\tau) \quad (\text{IV.36})$$

qui correspond bien à $a_{e1}^*(\infty) a_{e2}(\tau)$ de façon univoque.

Il est possible d'économiser un enregistrement en utilisant une deuxième solution qui n'utilise pas que des séquences d'impulsions. L'idée est d'utiliser deux enregistrements avec des séquences d'impulsions pour $\theta = \{0, \pi/2\}$ et un troisième, que nous noterons $S^\emptyset(t)$, où seule l'impulsion $E_{p2}(t)$ est présente. Ces trois enregistrements sont simulés sur

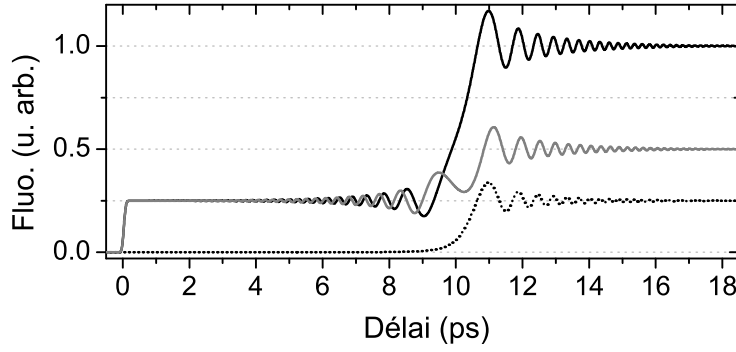


FIG. IV.17 – Signaux théoriques pour une reconstruction à 3 enregistrements : séquence d'impulsions avec phase relative de 0 en noir, de $\pi/2$ en gris et deuxième impulsion seule en pointillés noirs.

la figure IV.17 respectivement en noir, en gris et en pointillés noirs. On peut ainsi construire le signal $S^{[3]}(\tau)$:

$$S^{[3]}(\tau) = \frac{S^0(\tau) + iS^{\pi/2}(\tau) - (1-i)S^\emptyset(\tau)}{2} = a_{e1}^*(\infty)a_{e2}(\tau) + \frac{1-i}{2}|a_{e1}(\infty)|^2 \quad (\text{IV.37})$$

Cette technique nous permet d'accéder à $a_{e1}^*(\infty)a_{e2}(\tau)$ (et donc à $a_{e2}(\tau)$) à une constante près. Cette constante peut être très simplement retirée. En effet, pour un instant τ compris entre la fin de la première impulsion et le début de la seconde (délai entre 1 et 3 ps sur la figure IV.17), on sait que $a_{e2}(\tau)$ doit être nulle et donc le signal reconstruit $S^{[3]}(\tau)$ à cet instant est égal à la constante en trop, à savoir $(1-i)|a_{e1}(\infty)|^2/2$. Il suffit donc de retrancher au signal reconstruit $S^{[3]}(\tau)$ sa valeur "initiale" (c'est à dire celle du plateau entre 1 et 3 ps) pour récupérer $a_{e1}^*(\infty)a_{e2}(\tau)$.

Les dernières solutions, les plus minimalistes, n'utilisent que deux enregistrements pour $\theta = \{0, \pi/2\}$ afin de reconstruire l'amplitude de probabilité $a_{e1}^*(\infty)a_{e2}(\tau)$. Ces signaux sont représentés sur la figure IV.18 page ci-contre, en noir pour $\theta = 0$ et en gris pour $\theta = \pi/2$. Dans le cas général, ces deux signaux ne suffisent pas à reconstruire une solution unique pour $a_{e2}(\tau)$. En effet, ils correspondent aux système d'équation

$$S^0(\tau) = |a_{e1}(\infty)|^2 + |a_{e2}(\tau)|^2 + 2 \operatorname{Re} [a_{e1}^*(\infty)a_{e2}(\tau)] \quad (\text{IV.38a})$$

$$S^{\pi/2}(\tau) = |a_{e1}(\infty)|^2 + |a_{e2}(\tau)|^2 - 2 \operatorname{Im} [a_{e1}^*(\infty)a_{e2}(\tau)] \quad (\text{IV.38b})$$

qui sont toutes les deux quadratiques en $|a_{e2}(\tau)|$. On aura donc a priori pour chaque instant τ quatre solutions possibles. C'est une situation assez inconfortable, d'autant plus que très souvent, ces solutions sont dégénérées en plusieurs instants $\{\tau_i\}$, ce qui rend l'identification de la solution ayant un sens physique extrêmement fastidieuse.

Cependant, dans le cas particulier d'une deuxième impulsion faible par rapport à la première on a $|a_{e2}(\tau)|^2 \ll |a_{e1}(\infty)|^2$ pour tout instant τ . On peut alors dans un premier temps négliger ce terme et le système d'équations IV.38 devient linéaire en $a_{e2}(\tau)$. Par une approche variationnelle on peut ensuite reconstruire $a_{e1}^*(\infty)a_{e2}(\tau)$ en plusieurs itérations. En

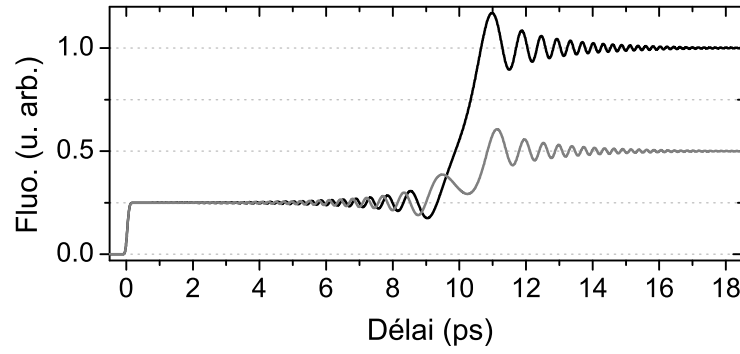


FIG. IV.18 – Signaux théoriques pour une reconstruction à 2 enregistrements : phase relative de 0 en noir, de $\pi/2$ en gris.

pratique, nous n'avons pas appliqué cette approche car utiliser une deuxième impulsion faible limite grandement la quantité de fluorescence détectée. Nous nous sommes donc tournés vers une technique de reconstruction originale qui permet d'obtenir $a_{e1}^*(\infty)a_{e2}(\tau)$ à partir du système IV.38, quel que soit le poids relatif des deux impulsions. Cette méthode dite géométrique donne deux solutions, mais la solution physique est facilement identifiable et séparable de la deuxième solution.

IV.5.1.3 Reconstruction géométrique

L'idée de base de la reconstruction géométrique est d'interpréter le système d'équation IV.38 page précédente en terme de distances dans le plan complexe. Cela permet de ramener la résolution du système à un instant τ au problème d'intersection de deux cercles. Nous allons dans un premier temps faire quelques petits changements par rapport au système d'équations IV.38.

Tout d'abord, nous allons considérer un cas un peu plus général que celui présenté précédemment : le deuxième enregistrement sera fait non pas pour $\theta = \pi/2$, mais pour un angle α non nul quelconque. Les deux enregistrements nous fournissent donc le système d'équations :

$$S^0(\tau) = |a_{e1}(\infty) + a_{e2}(\tau)|^2 \quad (\text{IV.39a})$$

$$S^\alpha(\tau) = |a_{e1}(\infty) + \exp[i\alpha]a_{e2}(\tau)|^2 \quad (\text{IV.39b})$$

Ces deux transitoires cohérents sont représentés sur la figure IV.19(a) : en noir pour $S^0(\tau)$ et en gris pour $S^\alpha(\tau)$ (on a pris $\alpha = \pi/3$ pour ces simulations). Comme la reconstruction géométrique exploite le lien étroit qui existe entre les signaux de transitoires cohérents et les amplitudes de probabilité correspondantes, il nous semble opportun d'explicitier à nouveau ce lien. La partie (b) de la figure IV.19 page suivante représente dans le plan complexe l'évolution de l'amplitude de probabilité dans l'état excité : en noir pour $\theta = 0$ et en gris pour $\theta = \alpha$. On observe bien deux contributions distinctes : celle de la première impulsion qui donne un segment horizontal de l'origine au point $(1/2, 0)$ et qui est indépendante de la phase relative, et celle de la deuxième impulsion qui produit une spirale de Cornu dont

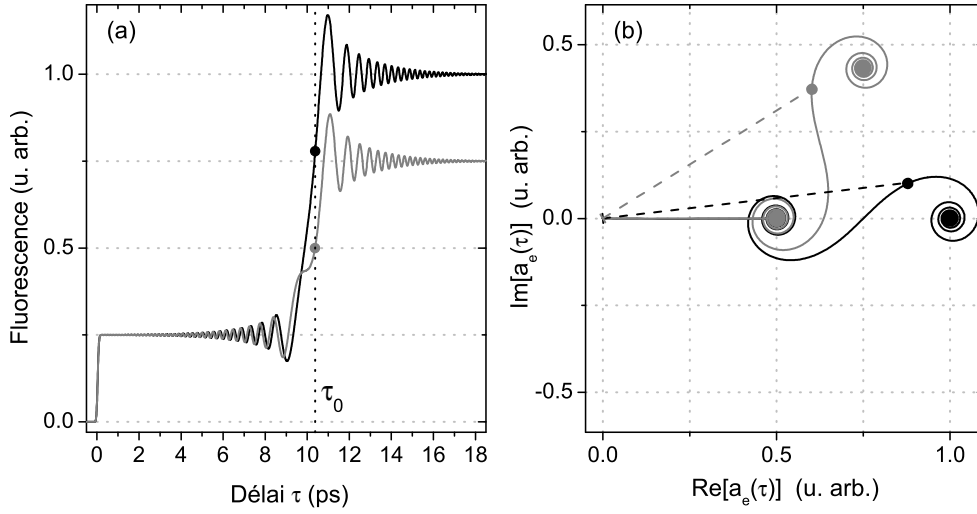


FIG. IV.19 – (a) Signaux théoriques pour une reconstruction géométrique : phase relative de 0 en noir, de α en gris. (b) Amplitudes de probabilité totales correspondantes (phase relative de 0 en noir et de α en gris).

l'orientation dépend de la phase relative. A un instant τ_0 donné, le signal $S^0(\tau_0)$ que l'on mesure (IV.19(a), point noir) correspond dans le plan complexe au carré de la distance entre l'origine du repère et la position courante sur la spirale en noir. Cette distance est matérialisée par le segment en pointillé noir. Au même instant, $S^\alpha(\tau_0)$ donne le carré de la distance entre l'origine du plan complexe et la position sur la spirale grise. Le terme $a_{e2}(\tau)$ que nous voulons reconstruire correspond dans le plan complexe à la spirale noire mais avec un point de départ recentré sur l'origine du repère.

La deuxième modification apportée est de supposer que $a_{e1}(\infty) = 1$. Il s'agit juste d'une commodité, cela n'enlève rien à la généralité de la méthode. Comme les méthodes présentées précédemment, celle-ci ne donne accès qu'à $a_{e1}^*(\infty)a_{e2}(\tau)$, et non pas $a_{e2}(\tau)$. Ce choix correspond à la fois à un choix de phase absolue pour $a_{e2}(\tau)$ et à une normalisation des enregistrements à leur valeur initiale. Nous entendons par valeur initiale, la valeur du plateau qui apparaît pour $\tau > 0$.

Avec ce choix de normalisation notre système d'équations devient :

$$S^0(\tau) = |1 + a_{e2}(\tau)|^2 \quad (\text{IV.40a})$$

$$S^\alpha(\tau) = |1 + \exp[i\alpha]a_{e2}(\tau)|^2 \quad (\text{IV.40b})$$

Nous définissons alors les deux distances suivantes :

$$r_0(\tau) = \sqrt{S^0(\tau)} = |1 + a_{e2}(\tau)| \quad (\text{IV.41a})$$

$$r_1(\tau) = \sqrt{S^\alpha(\tau)} = |1 + \exp[i\alpha]a_{e2}(\tau)| \quad (\text{IV.41b})$$

La figure IV.20 page suivante représente ces distances ainsi que les deux amplitudes de probabilités dans le plan complexe. Ces distances correspondent à celles évoquées précé-

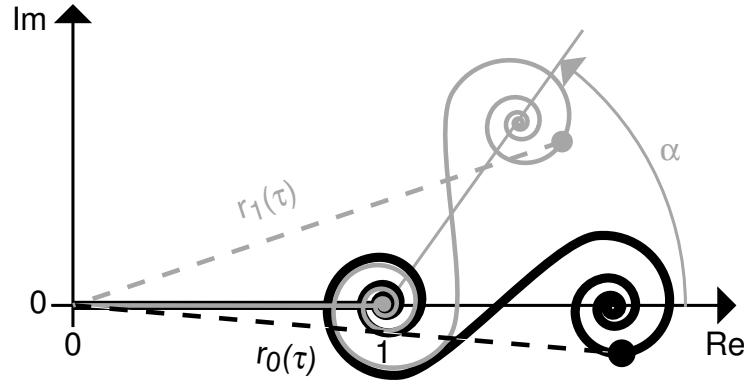


FIG. IV.20 – Evolution de l’amplitude de probabilité dans l’état excité, pour les deux séquences d’impulsions pompe : en phase (noir) et déphasées de α (gris). La spirale de Cornu noire correspond à $a_{e2}(\tau)$, que l’on cherche à reconstruire graphiquement.

demment.

Nous pouvons, sans modifier la distance $r_1(\tau)$, faire tourner mentalement le chemin gris d’un angle $-\alpha$, autour du point $(1,0)$ (point à partir du quel se déroulent les deux spirales). Cela revient en fait à ré-écrire $r_1(\tau)$ sous la forme :

$$r_1(\tau) = |1 + \exp[i\alpha]a_{e2}(\tau)| = |\exp[-i\alpha] + a_{e2}(\tau)| \quad (\text{IV.42})$$

Suite à cette rotation, les spirales grises et noires sont superposées et correspondent toutes deux au terme $a_{e2}(\tau)$ qui nous intéresse.

Nous pouvons ensuite, toujours sans rien changer au système d’équations, choisir le point $(1,0)$ comme origine du repère. On obtient alors la figure IV.21 page suivante. Sur celle-ci, on observe bien que les parties des deux chemins correspondant à $a_{e2}(\tau)$ se superposent. On appelle P_0 et P_1 les points de départ des deux chemins dans ce nouveau repère. Ils ont pour coordonnées $P_0 : (-1,0)$ et $P_1 : (-\cos(\alpha), \sin(\alpha))$. Comme nous connaissons la valeur de α , ces coordonnées sont connues.

On a alors une nouvelle interprétation pour les distances $r_0(\tau)$ et $r_1(\tau)$: elles donnent respectivement la distance entre le chemin $a_{e2}(\tau)$ et deux points distincts P_0 et P_1 (distincts car $\alpha \neq 0$).

Pour reconstruire graphiquement $a_{e2}(\tau)$, il suffit donc de rechercher, pour chaque instant τ , l’intersection de deux cercles (voir la figure IV.22 page suivante) :

- $\mathcal{C}_0$, de centre P_0 et de rayon $r_0(\tau)$.
- $\mathcal{C}_1$, de centre P_1 et de rayon $r_1(\tau)$.

On utilise pour ce faire la reconstruction proposée par Paul Bourke sur la page web : <http://astronomy.swin.edu.au/~pbourke/geometry/2circle/>

On obtient deux solutions $a_{e2}^+(\tau)$ et $a_{e2}^-(\tau)$ qui correspondent aux deux points d’intersection possibles. La solution physique est celle qui prend une valeur initiale nulle, c’est à dire qu’elle débute en $(0,0)$ dans le nouveau repère.

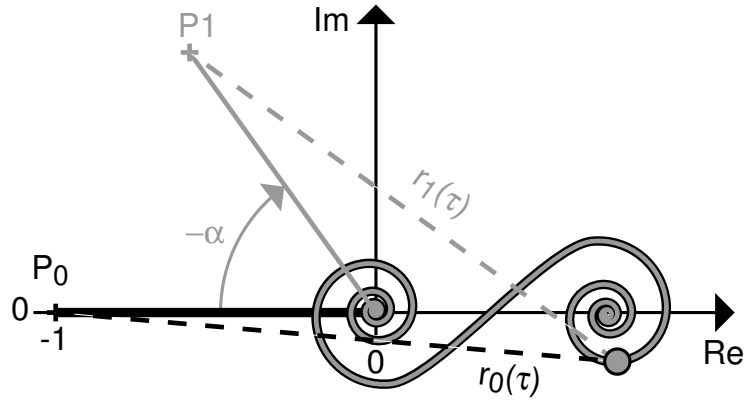


FIG. IV.21 – Réinterprétation du système d'équations IV.41 page 120 : $r_0(\tau)$ et $r_1(\tau)$ donnent la distance entre $a_{e2}(\tau)$ et deux points de référence P_0 et P_1 .

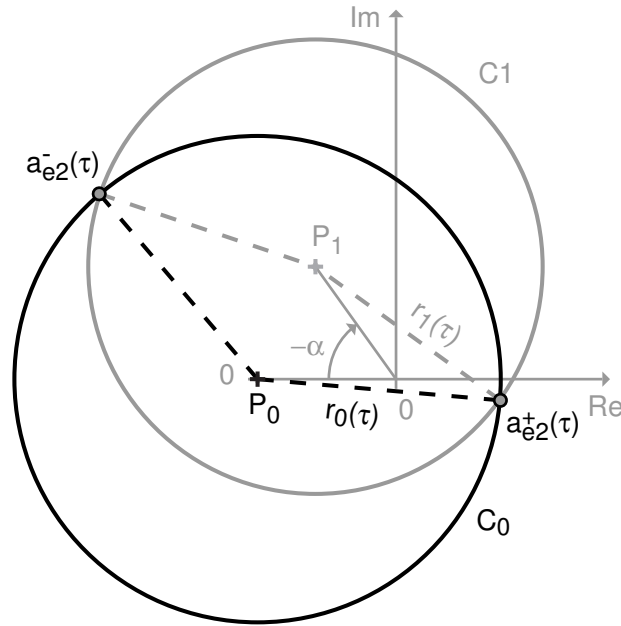


FIG. IV.22 – Reconstruction graphique de $a_{e2}(\tau)$. On obtient deux solutions possibles $a_{e2}^+(\tau)$ et $a_{e2}^-(\tau)$. La solution "physique" est celle commençant en $(0,0)$.

Pour que les deux solutions soient clairement séparées, il suffit que le chemin correspondant à $a_{e2}(\tau)$ ne coupe pas la droite (P_0P_1) . En effet, le long de cette droite les deux solutions sont dégénérées. Pour se tenir à l'écart de cette droite, il faut mettre certaines contraintes à la fois sur l'angle α et sur le chemin $a_{e2}(\tau)$. Ces contraintes ne sont pas simples et nécessitent une connaissance a priori de la forme que prendra le chemin correspondant à

$a_{e2}(\tau)$. On peut toutefois donner des tendances générales dans le cas présent d'une spirale de Cornu (qui correspond à une impulsion $E_{p2}(t)$ gaussienne à forte dérive de fréquence). Tout d'abord, pour un angle α fixé, plus l'excursion de la spirale sera faible devant la distance OP_0 (qui est identique à la distance OP_1) plus on aura de chance de ne pas couper (P_0P_1). Or la distance OP_0 , ici normalisée à un, est en fait liée à la contribution de la première impulsion. Il donc est préférable d'avoir une première impulsion importante par rapport à la deuxième.

La figure IV.23 montre le cas limite d'une première impulsion faible pour laquelle les deux solutions sont encore distinctes. La situation qui nous sert d'exemple depuis le début (points P_0 et P_1 , droite (P_0P_1) symbolisée par le trait noir) y est comparée à la valeur la plus faible de la première impulsion permettant de séparer la solution physique de l'autre (points $P_{0\min}$ et $P_{1\min}$, droite symbolisée par des pointillés noirs). La valeur la plus faible correspond à une droite ($P_{0\min}P_{1\min}$) presque tangente à la spirale.

Dans la pratique, nous avons constaté que deux impulsions d'intensités spectrales comparables donnent deux solutions non dégénérées.

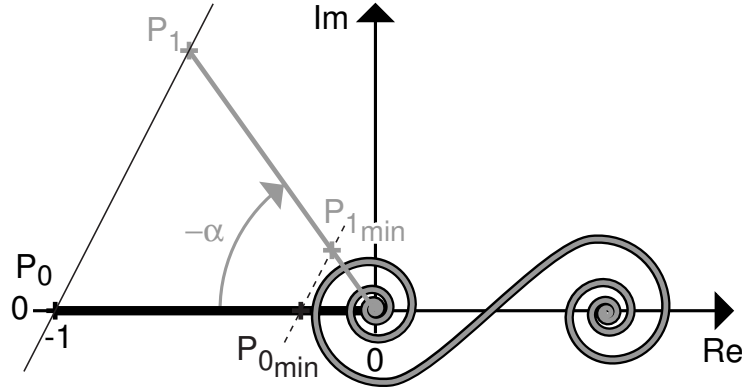


FIG. IV.23 – Comparaison entre la situation prise comme exemple (droite (P_0P_1) en noir) et l'intensité minimale de la première impulsion permettant une reconstruction non dégénérée (droite ($P_{0\min}P_{1\min}$) en pointillés noirs).

Ensuite, pour un rapport d'intensité fixé entre la première et la deuxième impulsion, on a un vaste choix pour l'angle α . Sur la figure IV.24 page suivante cette gamme est représentée par un arc de cercle gris. Pour les valeurs limites, les positions prises par le point P_1 sont représentées par $P_{1\alpha\min}$ et $P_{1\alpha\max}$. Ces valeurs limites correspondent toujours à des droites ($P_0, P_{1\alpha\max}$) et ($P_0, P_{1\alpha\min}$) (représentées en pointillés) quasiment tangentes à la spirale.

La valeur $\alpha = 0$ est bien sûr exclue car elle ne permet aucune reconstruction, les deux signaux expérimentaux étant alors identiques. De façon plus générale, les angles α trop petits sont défavorables. En effet, pour un petit angle, les deux transitoires cohérents sont très semblables ou ce qui revient au même, les points P_0 et P_1 très proches. Dans cette situation, une variation de la partie imaginaire de $a_{e2}(\tau)$ n'entraîne que des changements très faibles des signaux expérimentaux. La partie imaginaire ne sera pas reconstruite avec

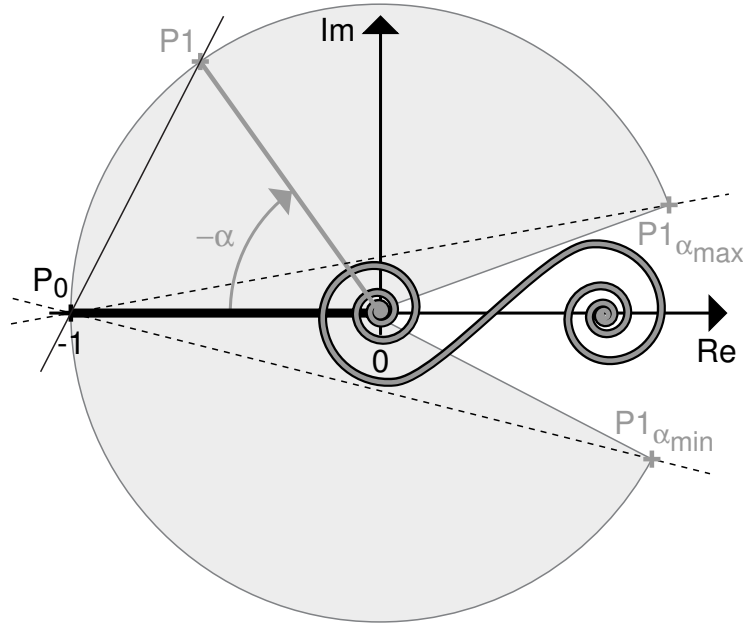


FIG. IV.24 – Gamme d'angles α (arc de cercle gris) permettant une reconstruction non-dégénérée. La valeur $\alpha = 0$ est bien évidemment exclue.

précision.

On peut se demander s'il existe, pour un $a_{e2}(\tau)$ donné, un angle α optimum. En fait, il n'en existe pas. Ou plus exactement, il en existe un différent pour chaque instant τ . Comme les angles trop faibles sont à éviter et que des angles importants augmentent les chances de couper la spirale⁷, nous avons en pratique utilisé le plus souvent des angles de $\pi/2$ ou $-\pi/2$.

IV.5.2 Prérequis expérimentaux

Nous disposons maintenant de plusieurs procédures, et notamment celle dite géométrique, pour reconstruire l'amplitude de probabilité créée par un champ de pompe donné à l'aide de plusieurs signaux expérimentaux. Pour les mettre en œuvre, il faut générer des séquences d'impulsions assez particulières et pas toujours simples à réaliser expérimentalement. En effet, si nous nous intéressons à l'amplitude de probabilité $a_{e2}(t)$ due à un champ $E_{p2}(t)$, nous devons réaliser au moins deux séquences d'impulsions du type

$$E_{p\Sigma}(t) = E_{p1}(t) + e^{i\theta} E_{p2}(t) \quad (\text{IV.43})$$

où les deux impulsions $E_{p1}(t)$ et $E_{p2}(t)$ sont séparées par un délai $\phi^{(1)}$ grand devant leurs durées. Dans la pratique, comme la deuxième impulsion dure plusieurs picosecondes, ce délai $\phi^{(1)}$ doit être de l'ordre de la dizaine de picosecondes. Or il nous faut un contrôle

⁷Pour $\alpha = \pi$ par exemple, les solutions seront dégénérées à chaque changement de signe de la partie imaginaire de $a_{e2}(\tau)$, c'est à dire à chaque demi-tour de la spirale, ainsi qu'au passage par la résonance.

interférométrique (de l'ordre d'une fraction de femtoseconde pour les longueurs d'onde visibles) sur ce délai pour prétendre imposer des phases relatives θ données. Maintenir un tel contrôle pendant les plusieurs minutes que durent chaque enregistrement n'est pas une chose aisée. Cela peut être réalisé à l'aide d'un interféromètre de Michelson stabilisé.

Dans notre cas, nous allons entièrement générer la séquence d'impulsion et la phase quadratique de la deuxième impulsion à l'aide du façonneur. Il nous faut donc nous assurer que nous pouvons générer les délais adéquats avec une précision interférométrique.

IV.5.2.1 Dispositif

Le dispositif expérimental que nous avons utilisé pour la reconstruction d'amplitude de probabilité est présenté sur la figure IV.25. Ce dispositif est très semblable à celui utilisé pour tester la sensibilité des transitoires (figure IV.14 page 113). Nous avons juste rajouté une lame séparatrice devant la cellule de rubidium afin de prélever une partie des faisceaux. Cette partie est utilisée pour réaliser la corrélation entre la pompe et la sonde dans un cristal de BBO type I de $100\ \mu\text{m}$. La détection du signal somme (centré à $344\ \text{nm}$) est assurée par une photodiode. Comme dans les précédentes expériences, la chambre de

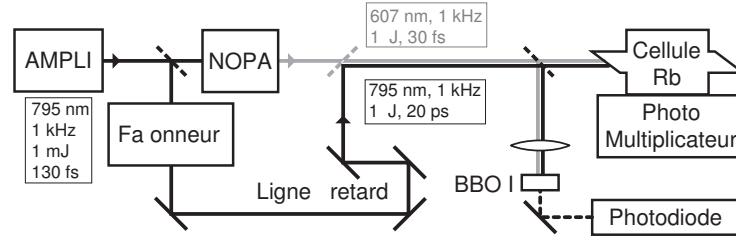


FIG. IV.25 – Dispositif expérimental utilisé pour la reconstruction d'amplitude de probabilité et de champ électrique.

la cellule de rubidium est portée à une température de 83°C et le queusot à 69°C . Cela correspond à une pression de rubidium de l'ordre de $2,7 \cdot 10^{-3}\ \text{Pa}$. Le photomultiplicateur assurant la détection de la fluorescence (RCA IP28) est porté à une tension de $-750\ \text{V}$.

L'impulsion de sonde, générée à l'aide du NOPA, a une durée de l'ordre de 30 fs et une longueur d'onde centrale de 607 nm.

La séquence d'impulsions de pompe est générée par le façonneur à partir de l'impulsion issue de l'amplificateur. Avant mise en forme, elle a une durée de l'ordre de 130 fs et est centrée à 795 nm. Le façonneur crée une première impulsion $E_{p1}(t)$ limitée par transformée de Fourier puis une deuxième $E_{p2}(t)$ à forte dérive de fréquence avec une phase relative de θ . Cela se traduit, dans le domaine spectral, par la multiplication de l'impulsion issue de l'amplificateur par la fonction de transfert :

$$H^\theta(\omega) = \frac{1}{2} \left[a_1 \mathbb{1} + a_2 e^{i[\theta + \phi^{(1)}(\omega - \omega_p) + \phi_p^{(2)}(\omega - \omega_p)^2/2]} \right] \quad (\text{IV.44})$$

où $\phi^{(1)}$ est le délai entre les deux impulsions et $\phi_p^{(2)}$ la phase quadratique de la deuxième. Nous avons aussi introduit deux paramètres supplémentaires a_1 et a_2 (compris entre 0 et 1) qui permettent de changer les poids relatifs des deux impulsions de la séquence.

Les énergies des impulsions de pompe et de sonde sont de l'ordre du μJ en entrée de la cellule de rubidium ; elles sont choisies de façon à rester en régime perturbatif.

IV.5.2.2 Tests préliminaires

Dans un premier temps, nous avons voulu vérifier le contrôle interférométrique de la phase θ pour des impulsions limitées par transformée de Fourier. Nous avons pour cela programmé dans le façonneur deux impulsions, décalées de 3 ps en phase ou en opposition de phase ($\theta = 0$ ou π , $\phi_p^{(1)} = 3$ ps, $\phi_p^{(2)} = 0$). Nous avons choisi deux impulsions de même intensité spectrale, c'est à dire $a_1 = a_2 = 1$. La figure IV.26 présente les intensités de corrélation et les transitoires obtenus avec ces deux séquences d'impulsions. Les traits

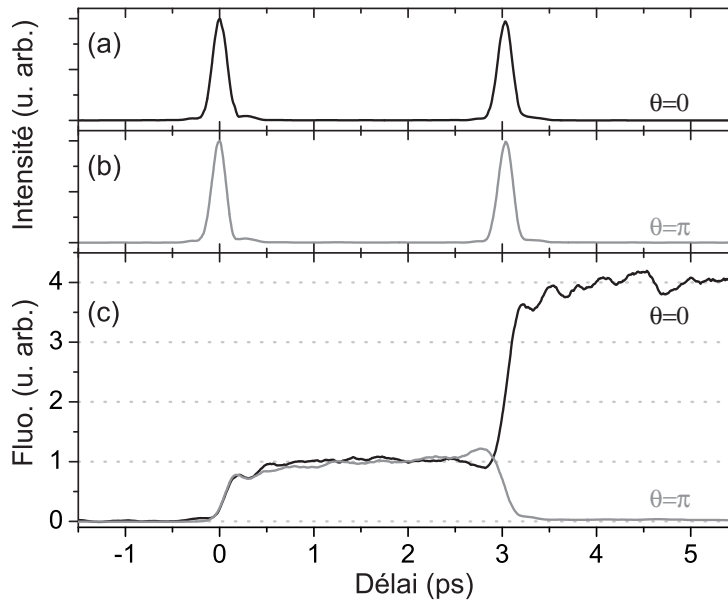


FIG. IV.26 – Séquences de deux impulsions courtes décalées de 3 ps avec une phase relative de 0 (noir) et de π (gris). (a) et (b) intensités de corrélation avec une impulsion courte. (c) signaux de fluorescence observés.

noirs correspondent à la séquence de deux impulsions en phase et les traits gris à celle d'impulsions déphasées de π .

Les corrélations (a) et (b) montrent que les deux impulsions de chaque séquence ont la même intensité temporelle, indépendamment de la phase. Sur les signaux de fluorescence (c), on retrouve le comportement attendu pour la population. Autour de $\tau = 0$, la première impulsion passe par la résonance et crée une contribution à la population dans l'état excité. Nous avons arbitrairement normalisé à 1 le signal de fluorescence correspondant. Le régime transitoire est court et a une forme de marche d'escalier : ceci est dû à l'impulsion limitée par transformée de Fourier. Ensuite, tant que la deuxième impulsion n'est pas arrivée, la population conserve la valeur asymptotique transférée par la première impulsion. Autour de $\tau = 3$ ps, la seconde impulsion apporte sa contribution. Celle-ci est de même importance

que celle de la première impulsion et son effet sur la population transférée dépend de la phase relative entre les deux impulsions. Pour deux impulsions en phase, les contributions interagissent constructivement et après une courte période de transitoire, on atteint un niveau de population 4 fois plus élevé, ce qui correspond bien à une amplitude de probabilité double. Pour deux impulsions en opposition de phase, l'interaction est destructive et à la fin de la période transitoire, la population asymptotique est à zéro.

Ce test simple confirme que le façonneur permet bien de générer des séquences avec une précision interférométrique, et de garder cette précision suffisamment longtemps pour que l'on puisse enregistrer les transitoires.

Pour compléter ce test, nous avons réalisé une deuxième expérience avec une phase variable entre les deux impulsions. Nous nous sommes placés après l'arrivée de la deuxième impulsion, à un délai τ fixe et nous avons enregistré la population asymptotique en fonction de la phase relative θ . Cette population est représentée sur la figure IV.27. Les carrés noirs

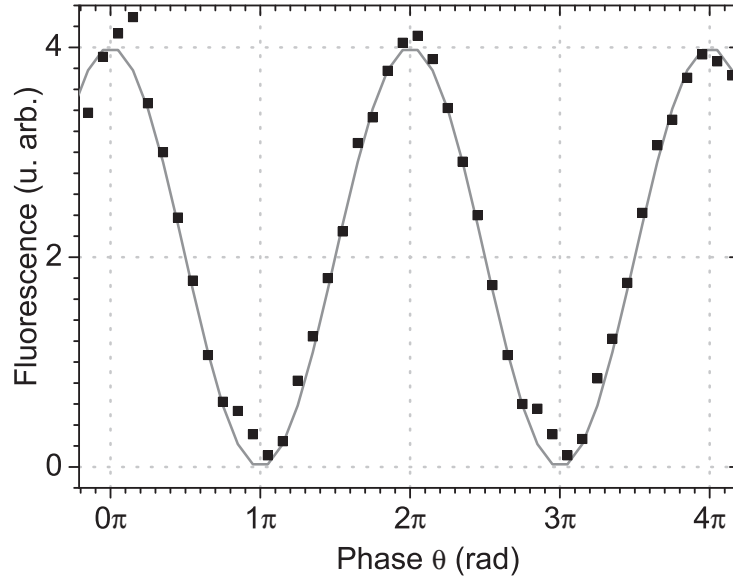


FIG. IV.27 – Population asymptotique transférée par une séquence d'impulsions en fonction de la phase relative θ entre les impulsions : population expérimentale (carrés noirs) et théorique (trait gris).

correspondent au signal de fluorescence expérimental et le trait gris au cosinus de la phase demandée. Aucun ajustement n'est fait sur le signal expérimental à part une normalisation du maximum à la valeur 4. En particulier, le 0 sur la figure correspond au "zéro laser", c'est à dire au signal de fluorescence obtenu en absence de tout champ de pompe. Le contraste est donc extrêmement bon. L'expérience reproduit très fidèlement le comportement attendu, avec toutefois une erreur systématique pour des déphasages de l'ordre de $0,9\pi + 2k\pi$.

Ce deuxième test nous montre que l'on peut faire varier à volonté la phase relative entre les deux impulsions de la séquence.

IV.5.3 Spirographe expérimental

IV.5.4 Premiers résultats

Dans une première expérience, nous avons choisi un délai de $\phi^{(1)} = 4$ ps entre les deux impulsions de pompe et une phase quadratique de $\phi_p^{(2)} = 2.10^5$ fs² pour la deuxième. Les deux impulsions sont de même intensité spectrale : $a_1 = a_2 = 1$. Nous avons enregistré les transitoires cohérents pour deux phases relatives : $\theta = \{0, \pi/2\}$. Ceux-ci sont présentés sur la figure IV.28. Les carrés représentent les transitoires expérimentaux et les traits la théorie. Les tracés noirs correspondent à $\theta = 0$ et les gris à $\theta = \pi/2$.

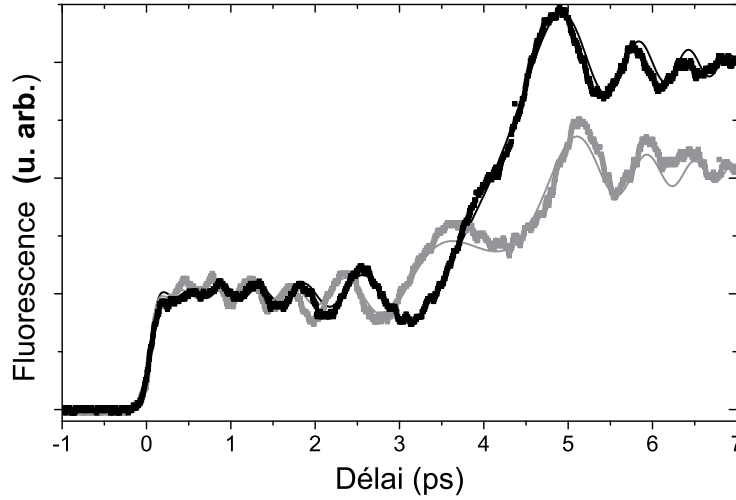


FIG. IV.28 – Transitoires cohérents (trait pour la théorie, carrés pour l'expérience) pour un délai de $\phi^{(1)} = 4$ ps et une phase quadratique de $\phi_p^{(2)} = 2.10^5$ fs² : pour $\theta = 0$ (en noir) et $\theta = \pi/2$ (en gris).

Les paramètres variables des ajustements théoriques sont l'amplitude du signal, l'origine des temps, le poids de la deuxième impulsion, ainsi que la phase relative θ . La phase quadratique n'est pas ajustée, elle est fixée à la valeur demandée au façonneur : 2.10^5 fs².

La phase relative s'avère un paramètre clef des ajustements. Ici les deux phases relatives ajustées sont $-\pi/10$ et $-\pi/10 + \pi/2$ et non pas 0 et $\pi/2$. L'écart de $\pi/2$ est bien respecté, mais il existe un décalage par rapport à 0. Ceci est dû à une erreur sur la pulsation centrale de l'impulsion de pompe ω_p utilisée dans le développement de la phase polynômiale appliquée à la deuxième impulsion (voir l'équation IV.44 page 125). En effet, si la pulsation utilisée pour le développement de la phase diffère de la pulsation centrale, les termes linéaire (délai) et quadratique (dérive de fréquence) vont ajouter une contribution à la phase relative. En appelant $\delta\omega$ l'erreur commise sur la pulsation centrale, on a une erreur sur la phase relative θ qui vaut :

$$\delta\theta = \phi^{(1)}\delta\omega + \frac{\phi_p^{(2)}}{2}\delta\omega^2 \quad (\text{IV.45})$$

Pour une erreur de seulement un pixel au niveau du masque, soit 0,06 nm de décalage par rapport à la longueur d'onde centrale de l'impulsion de pompe (de l'ordre de 795 nm), on

obtient une erreur sur la phase relative de l'ordre de $\delta\theta \simeq \pi/5$. L'écart retrouvé ici, de $\pi/10$ est tout à fait raisonnable, car il correspond à un décalage d'un demi-pixel, entre la position supposée et la position réelle de la pulsation ω_p dans le plan du masque.

Un deuxième paramètre qui diffère légèrement entre la programmation demandée et l'ajustement obtenu est le poids relatif des deux impulsions. L'ajustement donne une deuxième impulsion plus faible que la première ($a_1 = 1, a_2 = 0,75$) alors que nous avons programmé des poids identiques. Cela est dû aux interstices du masque du façonneur : ils ajoutent une contribution à la première impulsion qui se retrouve donc plus importante que prévue. Une compensation des interstices, ou plus simplement un décalage temporel de toute la mise en forme programmée permettrait de s'affranchir de ce problème.

Globalement, malgré ces deux défauts, les transitoires mesurés sont en très bon accord avec la théorie.

En appliquant la reconstruction géométrique à ces deux transitoires cohérents, nous avons reconstruit l'amplitude de probabilité dans l'état $|e\rangle$, présentée sur la figure IV.29. L'amplitude expérimentale (carrés noirs) reproduit bien la spirale de Cornu prévue par

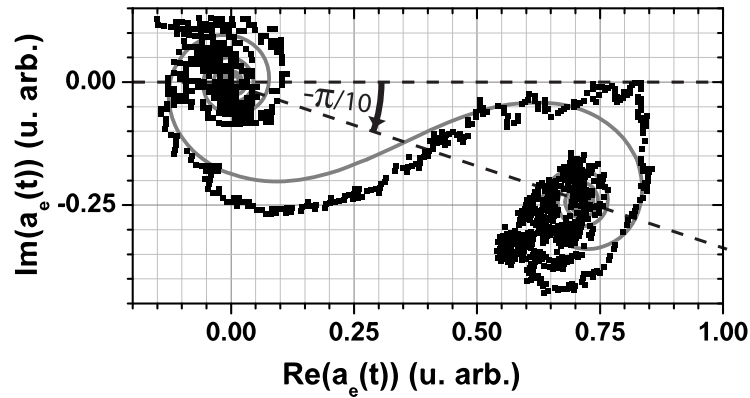


FIG. IV.29 – Amplitude de probabilité (trait gris pour la théorie, carrés noirs pour l'expérience) reconstruite à partir des transitoires de la figure IV.28 page ci-contre (délai de $\phi^{(1)} = 4$ ps et phase quadratique de $\phi_p^{(2)} = 2.10^5$ fs²).

la théorie (trait gris). La valeur asymptotique prise par $a_e(t)$ n'est pas sur l'axe réel, mais légèrement en dessous. Cet écart correspond au décalage de $-\pi/10$ sur les phases programmées. On retrouve cet angle entre l'axe réel et l'axe d'inclinaison de la spirale.

Cette expérience a été répétée pour une phase quadratique négative et des paramètres légèrement différents ($\phi^{(1)} = 6$ ps, $\phi_p^{(2)} = -2.10^5$ fs², $a_1 = 0,3$ et $a_2 = 1$). Les transitoires obtenus sont présentés sur la figure IV.30 page suivante : en gris pour $\theta = 0$ et en noir pour $\theta = \pi/2$. La théorie est représentée par un trait plein, les signaux expérimentaux par des carrés. L'amplitude de probabilité correspondante est tracée sur la figure IV.31 page 131, les carrés noirs correspondant à l'expérience, le trait gris à la théorie. Sur cette reconstruction, on observe un enroulement différent de celui observé précédemment : la spirale de Cornu a une forme de "Z" alors que sur la figure IV.29, elle était en forme de "S". Ceci est dû au

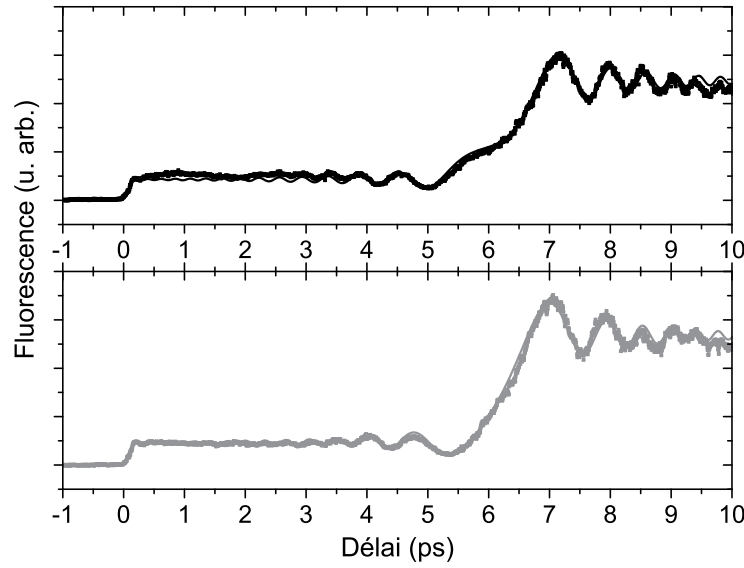


FIG. IV.30 – Transitoires cohérents (trait pour la théorie, carrés pour l'expérience) pour un délai de $\phi^{(1)} = 6$ ps et une phase quadratique de $\phi_p^{(2)} = -2.10^5$ fs² : pour $\theta = 0$ (en gris) et $\theta = \pi/2$ (en noir).

changement de signe de la phase quadratique qui est maintenant négative. Ici encore, les phases relatives sont décalées et la spirale est dirigée selon un axe différent de l'axe réel. Dans cet expérience où nous n'avons pas cherché à minimiser ce décalage, ce dernier est de l'ordre de $\pi/5$.

Afin de mieux percevoir l'évolution temporelle, cette nouvelle reconstruction est présentée sur la figure IV.32 page 132 où l'on a rajouté l'axe temporel : le plan complexe est horizontal et l'axe vertical correspond au temps. On observe l'évolution à l'échelle femtoseconde de la fonction d'onde de l'atome dans l'état excité. La projection dans le plan horizontal redonne la spirale de Cornu habituelle et les projections sur les plans verticaux la partie réelle (à gauche) et la partie imaginaire (à droite). Nous avons donc accès à l'évolution d'une partie de la fonction d'onde atomique, avec une résolution de l'ordre de 30 fs (c'est à dire la durée de l'impulsion de sonde). Cette mesure est évidemment dépendante du choix de jauge que nous avons fait en décrivant l'interaction champ - atome par l'approximation dipolaire électrique.

Cette reconstruction de la fonction d'onde atomique est assez différente de ce qui se fait par ailleurs. En effet, on suit habituellement l'évolution libre d'un paquet d'onde constitué de plusieurs niveaux (vibrationnels et/ou rotationnels)[101]. Ici, nous observons l'évolution d'un seul niveau d'énergie, mais durant son interaction avec le champ de pompe. Il s'agit donc de l'évolution forcée de la fonction d'onde atomique.

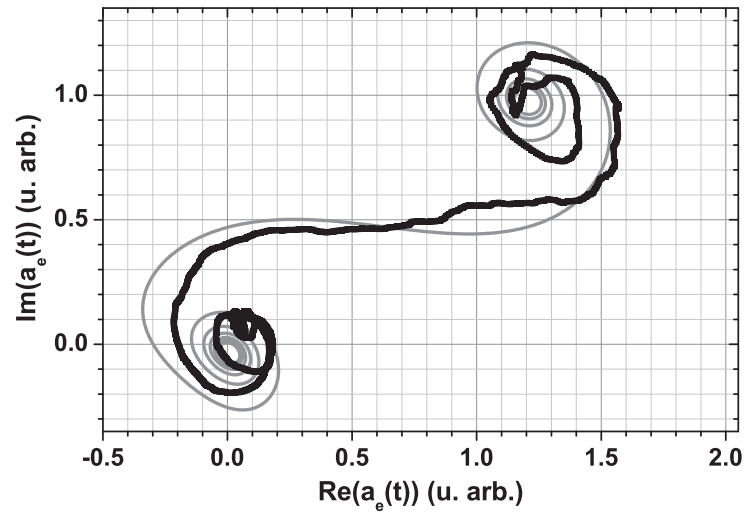


FIG. IV.31 – Amplitude de probabilité (trait gris pour la théorie, carrés noirs pour l'expérience) reconstruite à partir des transitoires de la figure IV.30 page ci-contre (délai de $\phi^{(1)} = 6$ ps et phase quadratique de $\phi_p^{(2)} = -2.10^5$ fs²).

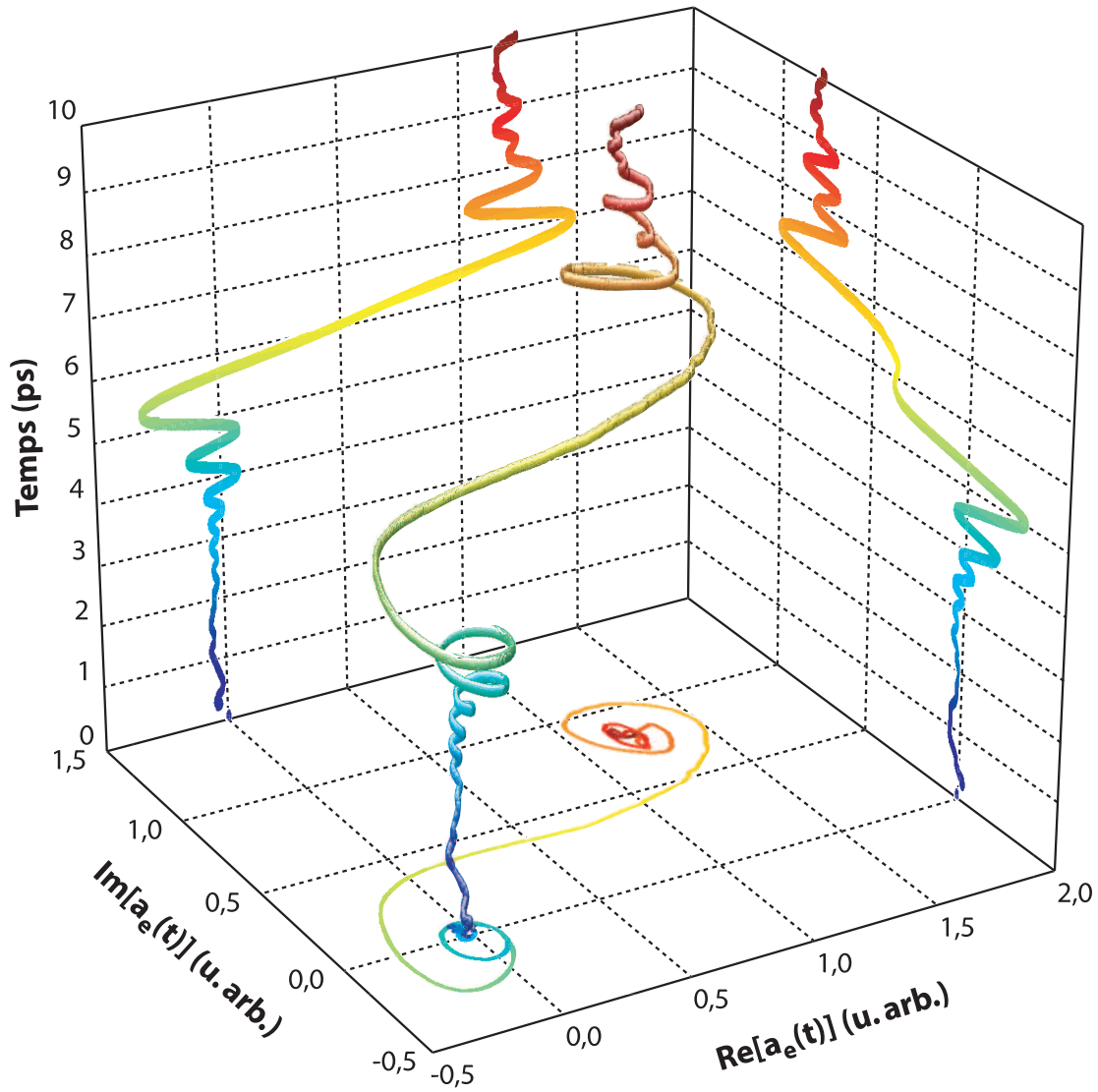


FIG. IV.32 – Reconstruction expérimentale de l'amplitude de probabilité dans l'état excité. On résout l'évolution temporelle à l'échelle femtoseconde de $a_e(\tau)$. La projection dans le plan complexe redonne la spirale de Cornu et les projection sur les axes réel et imaginaire, l'évolution temporelle de la partie réelle et imaginaire de $a_e(\tau)$.

IV.5.5 On peut jouer avec le spirographe

Avec notre méthode de mesure de la fonction d'onde atomique, nous pouvons observer le comportement de l'atome lorsqu'il est soumis à différents types de champ de pompe. Pour cela, il suffit de réaliser les reconstructions présentées précédemment pour des secondes impulsions de pompes $E_{p_2}(t)$ différentes. Nous avons ainsi réalisé trois reconstructions pour un délai $\phi^{(1)}$ de 6 ps, une phase quadratique $\phi_p^{(2)}$ de -2.10^5 fs^2 , deux impulsions de même poids ($a_1 = a_2 = 1$) et avec des phases relatives de la forme $\theta = \{\phi^{(0)}, \phi^{(0)} + \pi/2\}$. D'une reconstruction à l'autre, $\phi^{(0)}$ était augmenté de $2\pi/3$. Les trois amplitudes de probabilités reconstruites sont présentées sur la figure IV.33 en noir, rouge et bleu. Les traits pointillés de

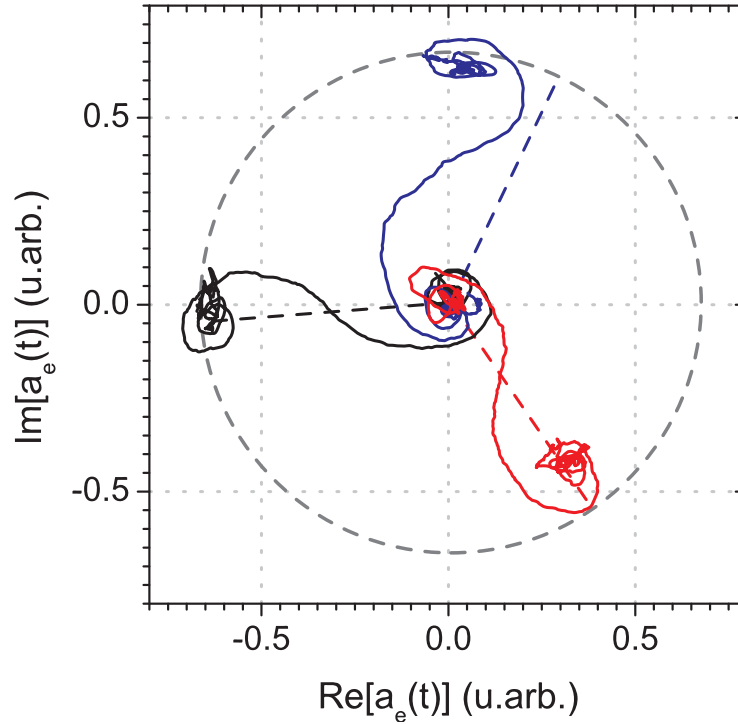


FIG. IV.33 – Amplitudes de probabilité reconstruites pour différentes phases entre la première et la deuxième impulsion de pompe. Les phases sont décalées de $2\pi/3$ d'une mesure à l'autre. Les deux impulsions de la séquence sont de même intensité.

couleur indiquent trois directions faisant deux à deux des angles de $2\pi/3$. Dans l'ensemble, les amplitudes reconstruites montrent le comportement attendu. On retrouve ainsi assez bien les différences de phase de $2\pi/3$ d'une amplitude à l'autre même si l'amplitude en bleu s'écarte quelque peu de la direction théorique.

IV.5.6 Mais attention à la casse

L'expérience précédente n'a pas été réussie du premier coup. Nous avons dans un premier temps tenté l'expérience avec une première impulsion plus faible : ($a_1 = 0,3$ et

$a_2 = 1$). Les autres paramètres restant inchangés. La deuxième impulsion est donc environ 10 fois plus intense que la première. Les amplitudes reconstruites sont présentées sur la figure IV.34. Seule la spirale en rouge semble correctement reconstruite. Les deux

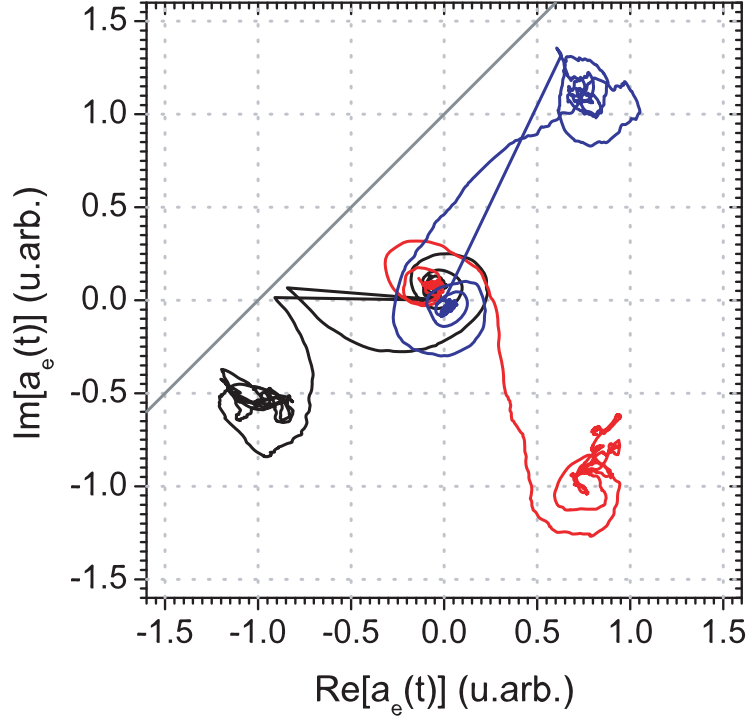


FIG. IV.34 – Amplitudes de probabilité reconstruites pour différents déphasages décalés de $2\pi/3$ d’une reconstruction à l’autre. La deuxième impulsion est environ 10 fois plus intense que la première. Le trait gris indique la droite de dégénérescence des solutions.

autres illustrent les problèmes de dégénérescence des solutions qui apparaît lors de la reconstruction géométrique (voir la partie IV.5.1.3 page 119). En effet, lorsque l’amplitude de probabilité à reconstruire traverse la droite (P_0, P_1) , représentée par le trait gris sur la figure IV.34, la reconstruction passe de la solution physique à la solution non-physique.

Pour visualiser ce phénomène, nous pouvons rajouter sur la même figure la deuxième solution donnée par la reconstruction géométrique. C’est ce que nous avons fait sur la figure IV.35 page ci-contre. Pour chaque déphasage $(0, 2\pi/3 \text{ et } 4\pi/3)$, nous avons représenté les deux solutions données par la reconstruction géométrique dans la même couleur. Les parties n’ayant pas de sens physique sont en trait fin, les parties ayant un sens physique en gras. La droite de dégénérescence (P_0, P_1) est tracée en gris. On observe bien que si l’amplitude de probabilité ne traverse pas (P_0, P_1) , ce qui est le cas de l’amplitude en rouge, la solution physique (trait rouge gras) est clairement distincte de la solution non physique (trait rouge fin). Par contre, dans les deux autres cas où l’amplitude de probabilité traverse (P_0, P_1) , la solution physique est répartie sur les deux solutions que donne la reconstruction géométrique.

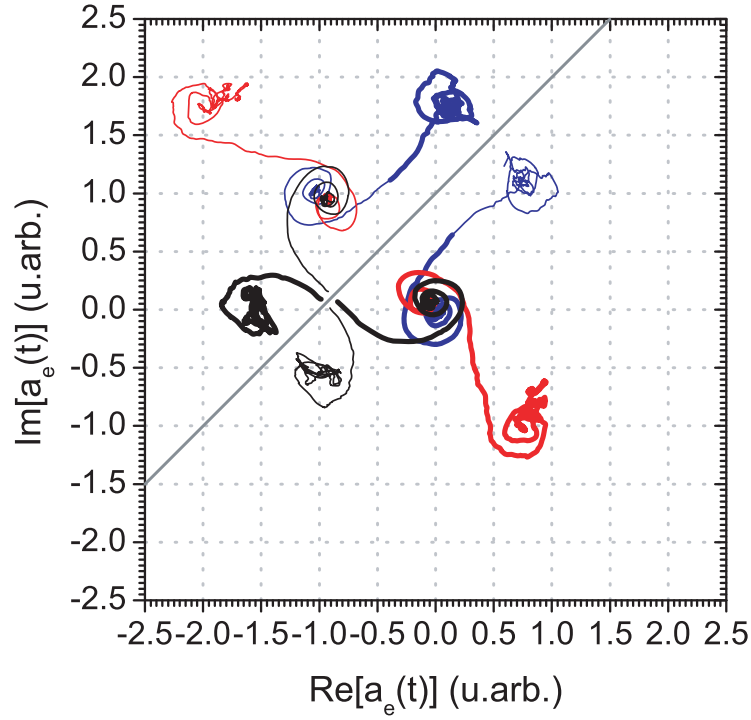


FIG. IV.35 – Les deux solutions données par la reconstruction géométrique pour chacun des trois déphasages sont représentées dans la même couleur. Les traits gras indiquent la part de chaque solutions qui a un sens physique.

Pour pouvoir reconstruire correctement les amplitudes de probabilité représentées en bleu et en noir, il aurait fallu utiliser un angle α différent de $\pi/2$ pour la reconstruction géométrique. En effet, en modifiant cet angle, on change la position du point de référence P_1 et donc la position de la droite de dégénérescence (voir les figures IV.22 page 122 et IV.24 page 124).

IV.6 Caractérisation par transitoires

Il est possible d'utiliser les reconstructions d'amplitude de probabilité, présentées dans la partie précédente, pour caractériser le champ de pompe. Pour l'instant, nous avons assimilé l'impulsion de sonde à un delta de Dirac. Sous cette approximation, les reconstructions nous donnent accès à $a_{e2}(\tau)$ qui est liée à la deuxième impulsion de pompe $E_{p2}(t)$ par une simple intégration (voir l'équation IV.5 page 97) :

$$a_{e2}(\tau) = -\frac{\mu_{eg}}{2i\hbar} \int_{-\infty}^{\tau} E_{p2}(t) e^{i\omega_{eg}t} dt \quad (\text{IV.46})$$

par simple dérivation, nous obtenons :

$$\frac{\partial a_{e2}(\tau)}{\partial \tau} \propto E_{p2}(\tau) e^{i\omega_{eg}\tau} \propto \mathcal{F}^{-1} \left[\tilde{E}_{p2}(\omega + \omega_{eg}) \right] \quad (\text{IV.47})$$

Nous avons donc accès au deuxième champ de pompe, décalé spectralement de la pulsation de transition de pompe ω_{eg} . Cette expression très simple ne tient pas compte des propriétés de l'impulsion de sonde.

Nous allons détailler ici le signal que nous pouvons reconstruire en tenant compte de cette sonde. Pour cela, nous allons reprendre la technique de reconstruction utilisant des séquences d'impulsions de pompe, présentée précédemment (partie IV.5.1.1 page 115), dans le cas d'une sonde de durée finie. Nous étudierons ensuite les particularités du signal obtenu par dérivation. Nous présenterons enfin quelques reconstructions expérimentales et nous étudierons les possibilités d'utilisation de cette méthode, notamment dans des systèmes plus complexes que les trois niveaux utilisés ici.

IV.6.1 Signal de reconstruction

Comme nous venons de le rappeler, le point de départ des reconstructions de fonctions d'onde atomiques de la partie précédente était de supposer la sonde infiniment courte. On avait alors un signal pompe-sonde donné par l'équation IV.27 page 115, que nous rappelons ici :

$$S^\delta(\tau) \propto P_e(\tau) = |a_e(\tau)|^2 \quad (\text{IV.48})$$

A partir de ce signal, en utilisant des séquences de pompe du type $E_{p_1}(t) + e^{i\theta} E_{p_2}(t)$ (équation IV.43 page 124), et l'une des différentes techniques de reconstruction, nous obtenions la quantité $a_{e_1}^*(\infty)a_{e_2}(\tau)$ (voir par exemple l'équation IV.36 page 117).

Si on tient compte de la sonde, nous avons montré (équation IV.12 page 98) que le signal pompe-sonde s'écrit :

$$S(\tau) = |b_f(\tau)|^2 \quad (\text{IV.49})$$

avec la quantité $b_f(\tau)$ définie comme la corrélation de l'amplitude de probabilité $a_e(t)$ par l'impulsion de sonde, décalée spectralement :

$$b_f(\tau) = -\frac{\mu_{fe}}{2i\hbar} \int_{-\infty}^{\infty} dt' E_s(t') e^{i\omega_{fe}t'} a_e(t' + \tau) \quad (\text{IV.50a})$$

$$b_f(\tau) = -\frac{\mu_{fe}\mu_{eg}}{4\hbar^2} \int_{-\infty}^{\infty} dt' E_s(t') e^{i\omega_{fe}t'} \int_{-\infty}^{t'+\tau} dt E_p(t) e^{i\omega_{eg}t} \quad (\text{IV.50b})$$

Nous pouvons alors suivre le même raisonnement avec $b_f(\tau)$ que celui que nous avons suivi avec $a_e(\tau)$. En effet, le raisonnement utilisé précédemment ne s'appuie que sur la linéarité de $a_e(\tau)$ avec $E_p(t)$ et ici, on voit bien que $b_f(\tau)$ présente lui aussi une dépendance linéaire vis à vis du champ de pompe.

Le point de départ est donc d'écrire la quantité $b_f(\tau)$ associée à une séquence d'impulsions $E_{p_1}(t) + e^{i\theta} E_{p_2}(t)$ sous la forme

$$b_f(\tau) = b_{f_1}(\tau) + e^{i\theta} b_{f_2}(\tau) \quad (\text{IV.51})$$

où b_{f_1} est la contribution de E_{p_1} et b_{f_2} celle de E_{p_2} . Les différentes techniques de reconstruction à 4, 3 ou 2 séquences d'impulsions nous donne le signal $G(\tau)$:

$$G(\tau) = b_{f_1}^*(\infty)b_{f_2}(\tau) \quad (\text{IV.52})$$

où $b_{f_1}(\infty)$ est la valeur asymptotique associée à la première impulsion. Celle-ci vaut :

$$b_{f_1}(\infty) = \frac{\pi\mu_{fe}\mu_{eg}}{-2\hbar^2} \tilde{E}_s(\omega_{fe}) \tilde{E}_{p_1}(\omega_{eg}) \quad (\text{IV.53})$$

$G(\tau)$ s'écrit comme la corrélation entre la sonde décalée spectralement de ω_{fe} et $a_{e_2}(t)$:

$$G(\tau) \propto \int_{-\infty}^{\infty} dt E_s(t) e^{i\omega_{fe}t} a_{e_2}(t + \tau) \quad (\text{IV.54})$$

Pour dériver $G(\tau)$ on utilise la propriété suivante du produit de corrélation :

$$\frac{\partial}{\partial x} \left[\int f_1(y) f_2(x + y) dy \right] = \int f_1(y) f_2'(x + y) dy = - \int f_1'(y) f_2(x + y) dy \quad (\text{IV.55})$$

Pour que cette propriété soit vérifiée, il suffit que toutes les intégrales qui apparaissent soient convergentes.

Ce résultat permet d'écrire la dérivée de $G(\tau)$, que nous noterons $g(\tau)$, comme la corrélation de la sonde par la dérivée de $a_{e_2}(\tau)$, soit :

$$g(\tau) = \frac{\partial G(\tau)}{\partial \tau} \propto \int_{-\infty}^{\infty} dt E_s(t) e^{i\omega_{fe}t} a_{e_2}'(t + \tau) \quad (\text{IV.56a})$$

$$\propto \int_{-\infty}^{\infty} dt E_s(t) e^{i\omega_{fe}t} E_{p_2}(t + \tau) e^{i\omega_{eg}(t + \tau)} \quad (\text{IV.56b})$$

Cette dernière forme montre clairement l'effet de la sonde sur le signal reconstruit : $g(\tau)$ donne le produit de corrélation entre l'impulsion de sonde et l'impulsion de pompe, décalées spectralement de ω_{fe} et ω_{eg} . Dans le domaine spectral, en tenant compte des différents facteurs multiplicatifs laissés de côté, le signal reconstruit s'écrit finalement :

$$\tilde{g}(\omega) = \sqrt{2\pi} \left| \frac{\pi\mu_{fe}\mu_{eg}}{2\hbar^2} \right|^2 \left[\tilde{E}_s(\omega_{fe}) \tilde{E}_{p_1}(\omega_{eg}) \right]^* \tilde{E}_s(\omega_{fe} - \omega) \tilde{E}_{p_2}(\omega_{eg} + \omega) \quad (\text{IV.57a})$$

$$\tilde{g}(\omega) \propto \tilde{E}_s(\omega_{fe} - \omega) \tilde{E}_{p_2}(\omega_{eg} + \omega) \quad (\text{IV.57b})$$

Nous avons donc accès au produit du spectre de la deuxième impulsion de pompe, décalé de ω_{eg} , par le spectre de l'impulsion de sonde, décalé de ω_{fe} et symétrisé (c'est à dire que la partie hautes fréquences de la sonde se multiplie à la partie basses fréquences de la pompe). Cette expression spectrale mérite plusieurs commentaires.

Tout d'abord, cette méthode de caractérisation apparaît clairement comme une mesure à référence : on mesure la pompe par rapport à la sonde (ou inversement). Ces deux champs jouent des rôles symétriques : la connaissance de l'un permet de caractériser l'autre. Ils peuvent par contre être dans des domaines spectraux très différents, grâce aux deux décalages de ω_{fe} et ω_{eg} qui les recentrent vers une fréquence nulle. Cette caractéristique permet de sortir du cercle vicieux des techniques à référence habituelles où les deux champs ont des propriétés spectrales comparables. Dans ces techniques, en effet, il faut disposer d'une référence parfaitement caractérisée qui est spectralement similaire à l'impulsion inconnue : l'obtention de la référence est donc un problème de même type que la caractérisation directe de l'impulsion inconnue.

Ensuite, la reconstruction du champ de pompe ne pourra se faire que si le signal mesuré est non-nul. Pour cela, il faut que $\tilde{E}_s(\omega_{fe} - \omega)$ et $\tilde{E}_{p2}(\omega_{eg} + \omega)$ se recouvrent. Or ces deux champs sont centrés aux fréquences $\omega_{fe} - \omega_s$ et $\omega_p - \omega_{eg}$. Pour assurer leur recouvrement il faut donc que les désaccords de la pompe et de la sonde (par rapport transitions atomiques) soient de signes opposés et de même importance. Cela revient à s'assurer que la transition à deux photons (un de pompe et un de sonde) de $|g\rangle$ vers $|f\rangle$ est résonnante.

La première impulsion courte $E_{p1}(t)$ n'intervient dans le signal reconstruit que par le facteur multiplicatif $\tilde{E}_{p1}(\omega_{eg})$. Elle ne modifie donc pas la phase et l'amplitude spectrale relative des différentes composantes du signal. De plus, ce signal sera d'autant plus intense que l'amplitude spectrale à ω_{eg} de cette première impulsion sera grande. Il peut donc être intéressant d'avoir une première impulsion $E_{p1}(t)$ intense (tout en restant dans le domaine perturbatif), pour augmenter la quantité de signal. On retrouve l'idée que cette première impulsion joue le rôle d'un oscillateur local⁸.

Enfin, pour pouvoir caractériser l'impulsion de pompe $\tilde{E}_{p2}(\omega)$, il n'est pas nécessaire que la sonde soit particulièrement courte devant l'impulsion de pompe. En effet, pour reconstruire $\tilde{E}_{p2}(\omega)$, il suffit de pouvoir déconvoluer par la sonde ; c'est à dire, dans le domaine spectral, diviser $\tilde{g}(\omega)$ par $\tilde{E}_s(\omega_{fe} - \omega)$. Cela suppose donc une connaissance préalable de la sonde (caractérisation à référence). Il faut aussi, pour que le résultat de cette division ait un sens, que le spectre de la sonde soit de largeur comparable ou supérieure à celle du spectre de pompe. La durée de la sonde importe donc peu⁹. Ce qui compte en revanche, c'est sa largeur spectrale¹⁰. Une sonde de même largeur spectrale que la pompe convient parfaitement.

Pour reconstruire $\tilde{E}_{p2}(\omega)$ à partir de $\tilde{g}(\omega)$, il faut procéder de la manière suivante. Tout d'abord, il faut connaître la sonde $\tilde{E}_s(\omega)$ ainsi que les pulsations ω_{eg} et ω_{fe} des transitions atomiques utilisées. On calcule ainsi la sonde décalée et symétrisée $\tilde{E}_s(\omega_{fe} - \omega)$. On divise alors le signal $\tilde{g}(\omega)$ par cette quantité et on décale le tout de ω_{eg} . L'étape de division doit évidemment être réalisée avec prudence et on évitera une division trop brutale au profit de techniques de déconvolution plus adaptées.

IV.6.2 Reconstructions expérimentales

Les reconstructions expérimentales ont été obtenues avec le même dispositif que pour la reconstruction de fonction d'onde atomique (voir la figure IV.25 page 125). Les signaux expérimentaux sont en fait les mêmes, seuls le traitement et l'interprétation diffèrent.

IV.6.2.1 Première expérience

Dans un premier temps, nous avons supposé que la sonde était limitée par transformée de Fourier et nous avons cherché à caractériser la deuxième impulsion $E_{p2}(t)$ de la séquence d'impulsions de pompe. Le délai entre les deux impulsions de pompe est fixé à 6 ps et une

⁸par le biais de sa contribution $b_{f1}(t)$.

⁹du moment que la sonde est parfaitement caractérisée !

¹⁰ou, ce qui revient au même, sa durée équivalente limitée par transformée de Fourier.

phase quadratique de $\phi_p^{(2)} = -2.10^5 \text{ fs}^2$ est appliquée à la seconde impulsion. Comme les impulsions utilisées par le façonneur pour générer la séquence de pompe sont au départ proches de la limite par transformée de Fourier, on s'attend à retrouver sur $E_{p2}(t)$ la phase introduite par le façonneur : la phase spectrale quadratique $\phi_p^{(2)}(\omega - \omega_p)^2/2$ correspond, dans le domaine temporel, à une phase quadratique de signe opposé $-\alpha_p t^2$, le lien entre $\phi_p^{(2)}$ et α_p étant donné par l'équation IV.17b page 99. La reconstruction est présentée sur la figure IV.36, elle se décompose en trois étapes :

- Deux transitoires pour $\theta = 0$ et $\theta = \pi/2$ sont enregistrés.
- Le signal $b_{f2}(\tau)$ est retrouvé par reconstruction géométrique.
- Par dérivation et translation, le champ de pompe est reconstruit en phase (carré gris) et en amplitude (trait gris).

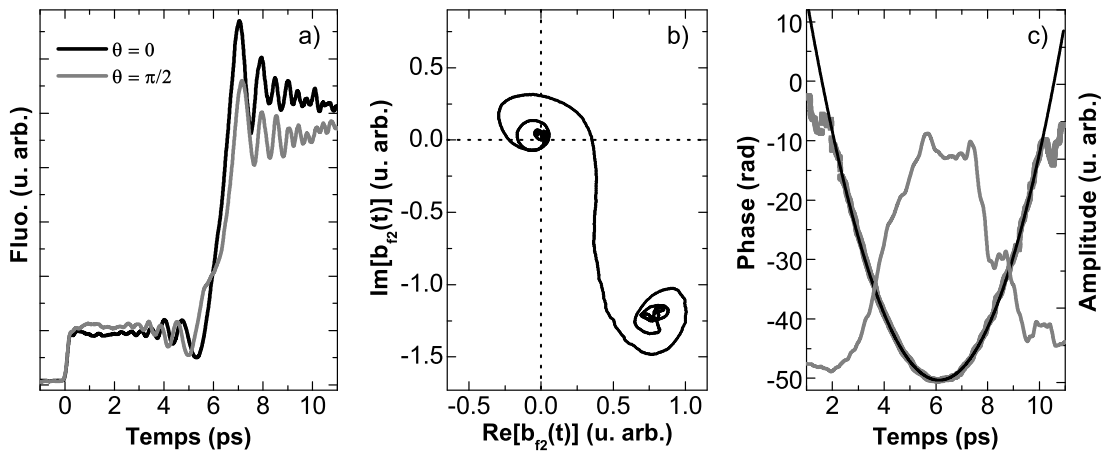


FIG. IV.36 – Reconstruction du champ électrique : a) mesure des transitoires pour des phases relatives de 0 (trait noir) et $\pi/2$ (trait gris) entre les deux impulsions de pompe. b) Amplitude de probabilité reconstruite géométriquement. c) Champ de pompe en phase (carrés gris) et amplitude (trait gris) obtenu par dérivation puis translation. La phase retrouvée est en accord parfait avec celle programmée dans le façonneur (trait noir).

La phase temporelle ainsi mesurée est en parfait accord avec la phase introduite par le façonneur (IV.36 c, trait noir).

On dispose donc d'une méthode, certes quelque peu complexe, permettant de caractériser les impulsions au sein même de la cellule de rubidium que nous utilisons pour nos expériences. Il s'agit là d'une caractéristique importante de cette technique : elle permet une mesure *in situ* du champ électrique, ce qui est un avantage non négligeable pour les impulsions courtes. En effet, ces dernières sont particulièrement sensibles à la dispersion non linéaire au sein des différentes optiques traversées, comme d'éventuelles lentilles ou les fenêtres de la cellule. Une mesure réalisée en amont de la cellule ou sur un autre faisceau peut ainsi donner une mesure déformée étant donné que les optiques traversées ne sont pas les mêmes pour l'impulsion mesurée et l'impulsion utilisée pour l'expérience.

IV.6.2.2 Caractérisation d'un barreau de verre

On peut émettre quelques réserves concernant la mesure précédente. En effet, nous avons retrouvé à l'aide des transitoires cohérents la phase que nous avons introduite pour pourvoir générer ces mêmes transitoires cohérents... Peut-être avons nous tourné en rond.

Afin de s'assurer de la validité de notre technique de caractérisation, nous l'avons utilisée pour mesurer une phase indépendante du dispositif de mesure. Nous avons ainsi caractérisé la phase spectrale $\phi_b(\omega)$ introduite par un barreau de verre. Pour ce faire, nous avons placé un barreau de verre (6 cm de SF10 et 4 cm de SF58) en amont du façonneur et comparé le champ obtenu $E_{p2}^b(t)$ avec celui retrouvé en absence de barreau ($E_{p2}(t)$). Le schéma de principe de cette expérience est présenté sur la figure IV.37. Dans le domaine spectral, la différence entre les deux reconstructions est juste la phase $\phi_b(\omega)$ introduite par le barreau. Cette phase est facilement calculable à l'aide des coefficients de Sellmeier et

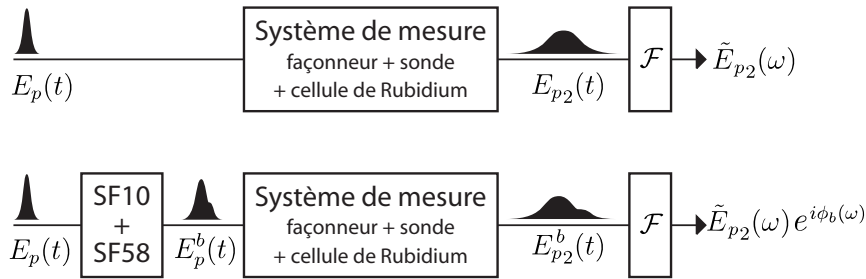


FIG. IV.37 – Mesure de la dispersion introduite par un barreau de verre (SF10+SF58). Par comparaison entre deux mesures, sans et avec barreau, on retrouve la phase spectrale $\phi_b(\omega)$ introduite par le barreau.

représente donc une référence sure.

Les résultats de cette mesure sont rassemblés sur la figure IV.38 page suivante. La partie (a) présente les phases spectrales reconstruites sans barreau (trait noir) et avec (trait gris). Dans les deux cas, nous reconstruisons la phase de la deuxième impulsion de la séquence générée par le façonneur. Celle-ci possède une forte phase quadratique $\phi_p^{(2)}$ introduite par le façonneur, et on observe bien, pour les deux reconstructions, une phase majoritairement quadratique. En présence du barreau, la phase reconstruite contient aussi la contribution ϕ_b de ce dernier. On observe ainsi (trait gris IV.38 (a)) une parabole plus évasée qu'en l'absence de barreau.

Par différence entre les deux mesures, on obtient cette phase $\phi_b(\omega)$. Celle-ci est représentée en carrés gris sur la figure IV.38 (b). La phase théorique déduite des coefficients de Sellmeier du SF10 et du SF58 est tracée en noir. Le trait gris indique quant à lui l'intensité spectrale reconstruite. Dans toutes les phases reconstruites, nous avons retiré d'éventuelles composantes linéaires (qui ne font que traduire l'existence d'un délai) afin d'observer la partie non-linéaire qui seule nous importe ici.

La phase mesurée est en bon accord avec le phase calculée par les coefficients de Sellmeier sur toute la gamme spectrale où l'intensité reconstruite est suffisamment importante.

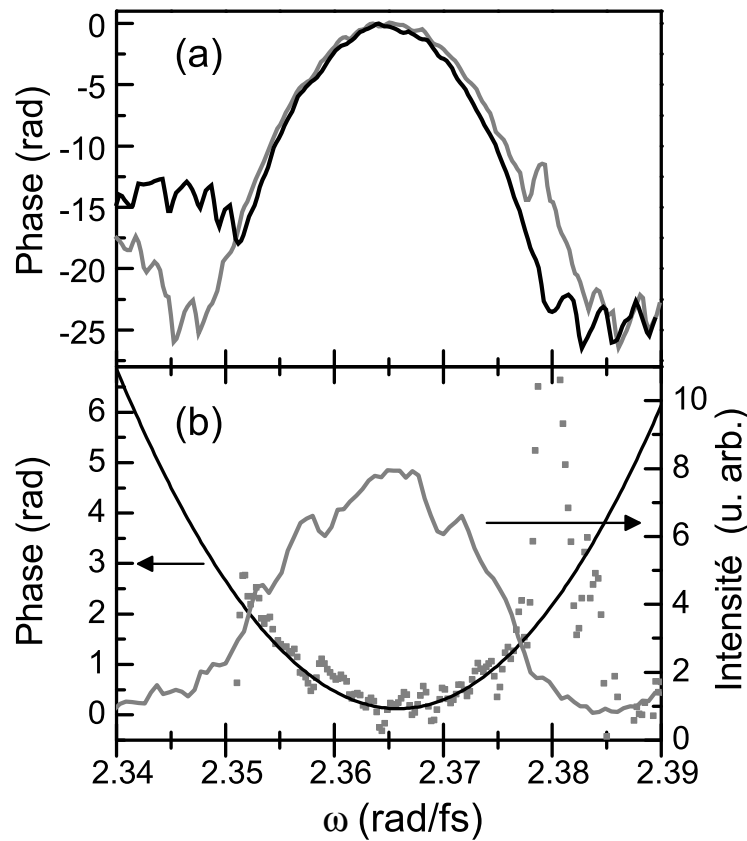


FIG. IV.38 – Mesure de la phase introduite par un barreau de verre. (a) phases reconstruites avec (trait gris) et sans barreau (trait noir). (b) phase due au barreau reconstruite (carrés gris) et théorique (trait noir) et amplitude spectrale reconstruite (trait gris).

Ce test simple montre que la technique de caractérisation par transitoires cohérents permet bien de caractériser des impulsions courtes.

IV.6.2.3 Limites de la mesure

Nous souhaitons ici discuter deux points importants pour cette technique de reconstruction : quelle est la résolution que l'on peut espérer atteindre et quel est l'intérêt de la phase quadratique $\phi_p^{(2)}$ introduite par le façonneur ?

La caractérisation par transitoires mesure avant tout des signaux concernant l'atome à trois niveaux qui interagit avec la séquence d'impulsions de pompe et l'impulsion de sonde. La résolution spectrale de cette mesure va ainsi être limitée principalement par les largeurs spectrales des niveaux utilisés. Dans notre cas, nous devons prendre en compte la durée de vie de l'état excité τ_e et celle de l'état final τ_f (toutes deux supérieures à la vingtaine de ns), ainsi que les largeurs Doppler des transitions $|g\rangle \mapsto |e\rangle$ et $|e\rangle \mapsto |f\rangle$ que nous noterons respectivement Δ_{eg} et Δ_{fe} (chacune de l'ordre de 3 rad/ns). Comme la densité au sein de la cellule de rubidium est faible, nous ne prendrons pas en compte les

effets de collisions qui pourraient limiter la durée de vie des cohérences entre les différents niveaux de notre système.

Dans le domaine spectral, notre technique donne accès au signal $\tilde{g}(\omega)$ défini par l'équation IV.57 page 137, et dont nous rappelons l'expression ici :

$$\tilde{g}(\omega) \propto \tilde{E}_s(\omega_{fe} - \omega) \tilde{E}_{p2}(\omega_{eg} + \omega) \quad (\text{IV.58})$$

Ce signal est donc très clairement dépendant des pulsations de transitions ω_{fe} et ω_{eg} . Les durées de vie finies font que ces pulsations ne sont pas définies à mieux que τ_f^{-1} et τ_e^{-1} respectivement. De plus, de très nombreux atomes contribuent à ce signal, chacun ayant des pulsations de transitions décalées par effet Doppler. Les largeurs Doppler Δ_{eg} et Δ_{fe} contribuent donc elles aussi à l'imprécision sur la valeur de ω_{fe} et ω_{eg} . Sans faire d'étude théorique poussée, nous supposons simplement que ces contributions s'ajoutent. L'imprécision en pulsation sur le signal $\tilde{g}(\omega)$ vaut donc :

$$\delta\omega = \tau_e^{-1} + \tau_f^{-1} + \Delta_{eg} + \Delta_{fe} \propto 7 \text{ rad/ns} \quad (\text{IV.59})$$

C'est cette imprécision qui donne la limite de résolution spectrale que l'on peut obtenir avec cette technique. A titre de comparaison, notre spectromètre a une résolution de 0,1 nm à 800 nm, ce qui correspond à une résolution en pulsation de 300 rad/ns, c'est à dire une résolution 40 fois moins bonne que celle de notre technique. La caractérisation par transitoires cohérents tire donc pleinement profit de la finesse des différents niveaux utilisés. Si nous pouvions éliminer l'élargissement par effet Doppler¹¹, nous pourrions obtenir une résolution améliorée d'un facteur dix environ.

Pour ce qui est du rôle de la phase quadratique, la réponse est plus en demi-teinte. Historiquement, nous nous sommes tout d'abord intéressés à l'évolution de l'amplitude de probabilité dans l'état excité au cours de son interaction avec une impulsion à forte dérive de fréquence. Quand nous avons dérivé de cette mesure une technique de caractérisation de champ, nous l'avons appliquée aux signaux déjà existants. Nous avons donc caractérisé en premier lieu des impulsions présentant une forte phase quadratique. Pour autant, l'intérêt de cette phase quadratique, rajoutée par le façonneur, n'apparaît pas évident. C'est particulièrement vrai pour la mesure de la phase introduite par un barreau de verre présentée dans la partie précédente. Est-il nécessaire d'ajouter cette phase quadratique pour mieux la retrancher ensuite ?

Nous avons pensé dans un premier temps que le $\phi_p^{(2)}$ ajouté par le façonneur permettait une meilleure résistance au bruit. En effet, une impulsion limitée par transformée de Fourier (ou presque) donne des transitoires très brefs en forme de marche abrupte alors qu'en ajoutant une forte phase quadratique, on obtient des transitoires oscillants, s'étalant sur un large intervalle temporel qui semblent plus faciles à observer en présence de bruit.

Cependant, des simulations récentes semblent montrer que l'ajout de $\phi_p^{(2)}$ ne change pas la qualité du champ reconstruit pour un niveau de bruit donné. La figure IV.39 page suivante montre ainsi la reconstruction simulée de l'amplitude de probabilité $a_{e2}(t)$ (IV.39(a) trait noir) et du champ spectral $\tilde{E}_{p2}(\omega)$ (IV.39(b) trait gris pour l'amplitude, carrés gris pour la phase) pour une phase ajoutée par le façonneur de $\phi_p^{(2)} = 2.10^5 \text{ fs}^2$. Pour cette si-

¹¹Avec des atomes refroidis, par exemple.

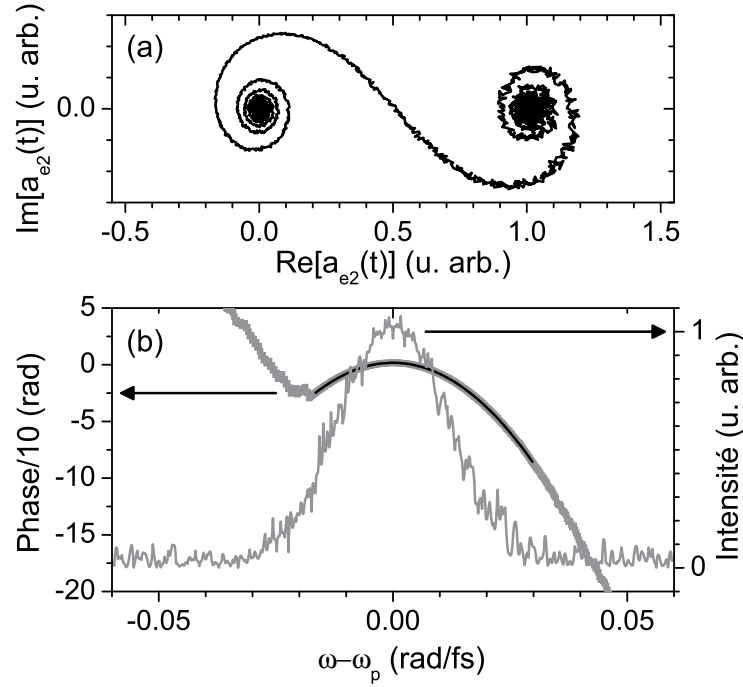


FIG. IV.39 – Simulation de reconstruction pour $\phi_p^{(2)} = 2.10^5 \text{ fs}^2$. (a) amplitude de probabilité reconstruite $a_{e2}(t)$ (trait noir). (b) phase (carrés gris) et intensité (trait gris) reconstruites pour $\tilde{E}_{p2}(\omega)$, le trait noir correspondant à un ajustement polynomial de la phase.

mulation, la sonde est supposée limitée par transformée de Fourier, centrée à 607 nm avec 15 nm de spectre¹². L'impulsion de pompe est centrée à 795 nm avec 10 nm de spectre. A l'entrée du façonneur, celle-ci possède une phase quadratique $\phi_i^{(2)} = 5000 \text{ fs}^2$ et une phase cubique $\phi_i^{(3)} = 2.10^5 \text{ fs}^3$. A partir de cette impulsion de pompe, les deux séquences d'impulsions de pompe $E_{p1}(t) + \exp(i\theta)E_{p2}(t)$ pour $\theta = \{0, \pi/2\}$ sont calculées, ainsi que les transitoires correspondants. Un bruit proportionnel de 1% est ajouté à ces transitoires puis $a_{e2}(t)$ et $\tilde{E}_{p2}(\omega)$ sont reconstruits¹³.

L'ajustement polynômial de la phase reconstruite (IV.39(b) trait noir) donne une phase quadratique totale de $\phi_t^{(2)} = -194909 \text{ fs}^2 \pm 252 \text{ fs}^2$ et une phase cubique de $\phi_t^{(3)} = -230270 \text{ fs}^3 \pm 32700 \text{ fs}^3$. En retranchant la phase quadratique $\phi_p^{(2)}$ due au façonneur, on obtient une mesure assez précise de la phase initiale de l'impulsion de pompe ($\phi_i^{(2)} = -5091 \text{ fs}^2 \pm 252 \text{ fs}^2$ et $\phi_i^{(3)} = -230270 \text{ fs}^3 \pm 32700 \text{ fs}^3$).

Pour comparaison, nous avons refait la même simulation, mais sans ajouter de phase quadratique : $\phi_p^{(2)} = 0 \text{ fs}^2$. La figure IV.40 page suivante montre l'amplitude de probabilité $a_{e2}(t)$ (IV.40(a) trait noir) et le champ spectral $\tilde{E}_{p2}(\omega)$ (IV.40(b) trait gris pour l'amplitude, carrés gris pour la phase) obtenus dans ce cas. La qualité de la mesure est très semblable à celle de la simulation précédente : l'ajustement polynômial (IV.40(b) trait noir) donne

¹²Largeur à mi-hauteur en intensité.

¹³A l'aide de la technique à trois transitoires cohérents.

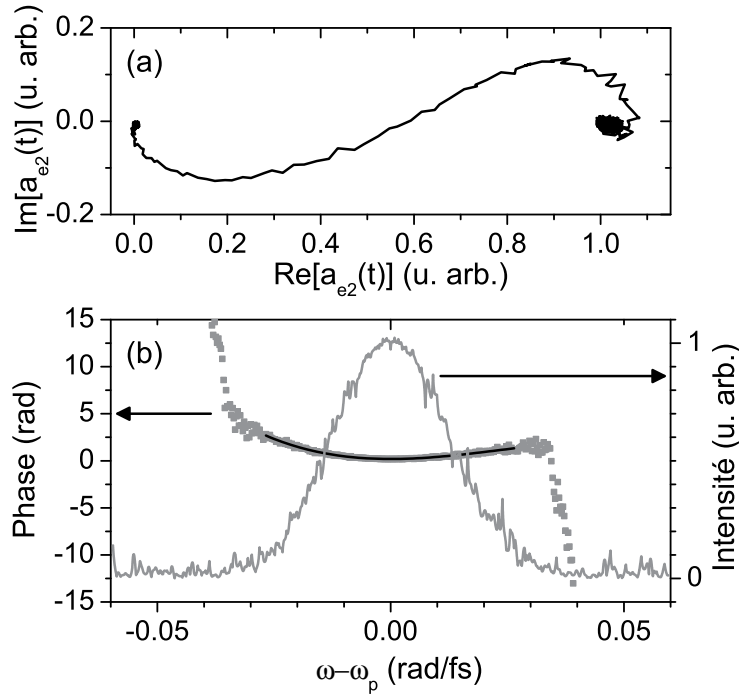


FIG. IV.40 – Simulation de reconstruction pour $\phi_p^{(2)} = 0 \text{ fs}^2$. (a) amplitude de probabilité reconstruite $a_{e2}(t)$ (trait noir). (b) phase (carrés gris) et intensité (trait gris) reconstruites pour $\tilde{E}_{p2}(\omega)$, le trait noir correspondant à un ajustement polynomial de la phase.

une phase quadratique totale de $\phi_t^{(2)} = 5093 \text{ fs}^2 \pm 65 \text{ fs}^2$ et une phase cubique de $\phi_t^{(3)} = -209270 \text{ fs}^3 \pm 14500 \text{ fs}^3$.

D'autres simulations ont été menées, pour différents niveaux de bruit (entre 10^{-6} et 10^{-2}) et différentes phases initiales pour l'impulsion de pompe. Sans surprise, le signal reconstruit se dégrade au fur et à mesure que le bruit augmente. Il reste cependant tout à fait exploitable pour des niveaux de bruit atteignables expérimentalement. Pour un niveau de bruit raisonnable, et des paramètres de reconstruction fixés, nous avons reconstruit avec ces simulations toute sorte d'impulsions : limitées par transformée de Fourier, fortement étirées, ou présentant un trou spectral très étroit (0,1 nm), ...

En conclusion sur ce point particulier, on peut dire que la phase quadratique ajoutée par le façonneur ne permet en rien de modifier la pertinence de la mesure. Notre technique semble très bien fonctionner pour des impulsions limitées par transformée de Fourier et des expériences sont à l'étude. Le seul intérêt de cette forte phase quadratique a été de montrer expérimentalement que notre technique fonctionne très bien pour des impulsions fortement étirées qui sont habituellement difficiles à caractériser. De plus, il n'y a aucun paramètre à changer lorsque l'on passe d'une impulsion courte à une impulsion étirée : à intensité spectrale fixée, le niveau de signal est le même et les mêmes paramètres peuvent être utilisés. En particulier, l'intervalle temporel ΔT sur lequel on reconstruit le champ ne doit pas être dicté uniquement par la durée de l'impulsion à caractériser mais aussi par la

résolution spectrale $\delta\omega = 2\pi/\Delta T$ que l'on souhaite atteindre.

IV.6.3 Extensions et perspectives

Les résultats présentés ici montre qu'avec un système à trois niveaux, il est possible de faire de la caractérisation d'impulsions ultracourtes. Cependant, l'atome de rubidium utilisé présente plusieurs états finaux qui vont pouvoir contribuer à la fluorescence mesurée (voir la figure IV.3 page 103). Notre technique reste-t-elle valable lorsqu'il y a plusieurs états finaux ?

Un autre cas très intéressant est le cas d'un continuum d'états finaux. Cela correspond par exemple à une impulsion de sonde qui ionise le système étudié. Est-ce que dans cette situation la caractérisation est toujours possible et a-t-on toujours un rôle symétrique de la pompe et de la sonde ? Nous nous proposons ici de répondre en partie à ces questions.

IV.6.3.1 Plusieurs niveaux finaux

On reprend ici l'approximation perturbative et on utilise toujours une séquence d'impulsions de pompe mais on suppose que plusieurs états finaux $|n\rangle$ peuvent contribuer à la fluorescence. On note $P_n(\infty, \tau)$ la population asymptotique dans ces états pour un délai pompe sonde τ donné. Comme nous avons défini $b_f(\tau)$ pour l'état $|f\rangle$ (voir l'équation IV.12 page 98), nous introduisons :

$$b_n(\tau) = -\frac{\mu_{ne}}{2i\hbar} \int_{-\infty}^{\infty} dt' E_s(t') e^{i\omega_{ne}t'} a_e(t' + \tau) \quad (\text{IV.60a})$$

$$b_n(\tau) = -\frac{\mu_{ne}\mu_{eg}}{4\hbar^2} \int_{-\infty}^{\infty} dt' E_s(t') e^{i\omega_{ne}t'} \int_{-\infty}^{t'+\tau} dt E_p(t) e^{i\omega_{eg}t} \quad (\text{IV.60b})$$

où μ_{ne} est le moment dipolaire de la transition $|e\rangle \mapsto |n\rangle$. On note aussi ω_{ne} les pulsations des transitions $|e\rangle \mapsto |n\rangle$. La fluorescence s'écrit maintenant comme la somme incohérente des contributions des différents états. En supposant qu'ils contribuent tous de manière équivalente¹⁴, on peut écrire pour un délai τ suffisamment grand :

$$S^\theta(\tau) \propto \sum_n P_n(\tau) \propto \sum_n |b_{n1}(\infty) + b_{n2}(\tau) \exp[i\theta]|^2 \quad (\text{IV.61})$$

où l'on a introduit les notations $b_{n1}(\tau)$ et $b_{n2}(\tau)$ qui sont les contributions respectives de $E_{p1}(t)$ et $E_{p2}(t)$.

La reconstruction géométrique n'apparaît pas clairement réalisable sur cette expression car on ne peut plus identifier simplement cette somme au carré d'une distance. On peut par contre toujours utiliser la reconstruction utilisant quatre transitoires, donnée dans l'équation IV.36 page 117, pour reconstruire le signal $G(\tau)$:

$$G(\tau) = \frac{[S^0(\tau) - S^\pi(\tau)]}{2} + i \frac{[S^{-\pi/2}(\tau) - S^{\pi/2}(\tau)]}{2} \propto \sum_n b_{n1}^*(\infty) b_{n2}(\tau) \quad (\text{IV.62})$$

¹⁴C'est à peu près le cas dans le rubidium : le produit du moment dipolaire μ_{ne} par la probabilité de fluorescence donne le même ordre de grandeur pour les états impliqués.

On peut alors dériver cette quantité par rapport à τ :

$$g(\tau) = \frac{\partial G(\tau)}{\partial \tau} \propto \sum_n b_{n1}^*(\infty) \int_{-\infty}^{\infty} dt E_s(t) e^{i\omega_{ne}t} E_{p2}(t+\tau) e^{i\omega_{eg}(t+\tau)} \quad (\text{IV.63a})$$

Ce résultat est très proche de celui obtenu pour un seul état final (donné par l'équation IV.56 page 137) : on obtient une somme de convolutions du champ de pompe avec le champ de sonde avec des décalages spectraux dépendant des pulsations ω_{ne} . Quoi qu'il en soit, ce résultat peut toujours se mettre sous la forme de la convolution d'une fonction réponse $h(t)$ par le deuxième champ de pompe $E_{p2}(t)$, décalé spectralement de ω_{eg} .

$$g(\tau) = \int_{-\infty}^{\infty} h(t) E_{p2}(t+\tau) e^{i\omega_{eg}(t+\tau)} dt \quad (\text{IV.64})$$

Cette fonction réponse dépend de façon complexe de la sonde $E_s(t)$ et des différents états finaux :

$$h(t) \propto \sum_n b_{n1}^*(\infty) E_s(t) e^{i\omega_{ne}t} \quad (\text{IV.65})$$

Dans le domaine spectral, on a ainsi accès à la quantité :

$$\tilde{g}(\omega) \propto \tilde{h}(-\omega) \tilde{E}_{p2}(\omega_{eg} + \omega) \quad (\text{IV.66})$$

où $\tilde{h}(\omega)$ est la transformée de Fourier de $h(t)$. Les rôles joués par la sonde et la pompe sont donc maintenant clairement dissymétriques. Il est toujours possible de mesurer la deuxième impulsion de pompe $\tilde{E}_{p2}(\omega)$ à condition de connaître parfaitement $\tilde{h}(\omega)$. Pour cela, il faut connaître $\tilde{E}_s(\omega)$ ainsi que tous les moments dipolaires μ_{ne} et toutes les pulsations de transitions ω_{ne} . Les contraintes sur $\tilde{h}(-\omega)$ pour que la reconstruction soit valable sont celles énoncées pour $\tilde{E}_s(\omega_{fe} - \omega)$ dans le cas d'un seul niveau final (voir la partie IV.6.1 page 137). Par contre, il ne semble pas possible de reconstruire la sonde connaissant $\tilde{E}_{p2}(\omega)$.

IV.6.3.2 Continuum d'états finaux

On reprend exactement la même démarche que pour quelques états finaux, en remplaçant la somme discrète par une intégrale sur toutes les pulsations possibles pour les états finaux. Cela revient à considérer une ionisation loin du seuil dans un continuum plat. Des calculs simples permettent d'obtenir la quantité :

$$g(\tau) \propto \int_{-\infty}^{\infty} dt |E_s(t)|^2 E_{p2}(t+\tau) e^{i\omega_{eg}(t+\tau)} \quad (\text{IV.67})$$

On a donc accès à la convolution du profil d'intensité de l'impulsion de sonde par le champ de pompe. Une fois encore, les rôles sont nettement dissymétriques. Cette fois par contre, la situation est plutôt favorable. En effet, pour connaître la pompe, il n'est plus nécessaire d'avoir une caractérisation complète de la sonde : son profil d'intensité suffit. Cela peut s'avérer intéressant pour une sonde dans l'XUV où il existe quelques méthodes de corrélations intensimétriques mais pas de méthode de caractérisation complète.

Si on dispose d'un moyen de résoudre en énergie les éléments issus de l'ionisation (spectromètre de photo-électrons ou de photo-ions), on peut se ramener à un système à trois niveaux. En effet, l'étape d'ionisation résolue en énergie joue le rôle d'un troisième niveau lié ayant une largeur spectrale donnée par la résolution du spectromètre utilisé.

Il pourrait donc être possible d'étendre cette technique dans l'UV et le XUV, ce qui est particulièrement intéressant étant donné le peu de techniques de caractérisation dans ces gammes spectrales. Des expériences sont à l'étude.

IV.6.4 Conclusion

La mesure complète d'impulsions par transitoires cohérents présente des caractéristiques plutôt atypiques. Il s'agit tout d'abord d'une mesure à référence, mais avec une référence qui peut être dans un domaine spectral totalement différent de celui de l'impulsion inconnue. Ceci est très intéressant car la caractérisation d'impulsions femtoseconde ne présente pas la même difficulté dans différents domaines spectraux. On peut ainsi imaginer utiliser une référence dans un domaine où la caractérisation se fait aisément (le proche infrarouge par exemple) et un champ à caractériser dans un domaine plus problématique comme l' XUV[129, 63, 130, 131].

Cette technique se prête très bien aux impulsions fortement mises en forme, comme nous l'avons vu avec la reconstruction d'impulsions à forte dérive de fréquence. Bien sûr, elle fonctionne aussi pour des impulsions proches de la limite par transformée de Fourier. Mais très souvent, il existe d'autres techniques performantes pour ce genre d'impulsions qui ont l'avantage d'être auto-référencées : FROG[65, 132, 133], SPIDER[68], et leurs dérivées[66, 134, 64], [135, 112, 74, 73, 72, 136, 137]. Par contre ces techniques ne se prêtent pas très bien à la caractérisation d'impulsions fortement façonnées. En effet, pour ce genre d'impulsions, les mesures FROG nécessitent des traces 2D avec un grand nombre de points temporels et spectraux. Cela rend la reconstruction itérative très lente et peu aisée. Dans le cas du SPIDER, il peut être difficile de trouver un compromis acceptable entre quantité de signal et résolution spectrale.

Une autre difficulté pour des techniques habituelles est la caractérisation d'impulsions constituées de plusieurs "morceaux" nettement séparés (dans le domaine temporel ou spectral) comme par exemple une séquence de plusieurs impulsions. Il est alors difficile de relier entre elles les phases reconstruites sur chacun des morceaux. Nous avons pu constater que la reconstruction par transitoires cohérents est particulièrement efficace dans ces situations¹⁵.

Dernier intérêt de cette technique, elle permet une caractérisation *in situ* des impulsions femtosecondes : nous avons caractérisé nos impulsions au sein même de la cellule de rubidium que nous utilisons pour nos expériences de contrôle cohérent.

¹⁵Ceci semble être dû au signal expérimental qui est à peu de chose près l'intégrale du champ à caractériser : même si ce dernier s'annule, le signal expérimental ne s'annule pas et la phase reconstruite garde un sens.

CONCLUSION

Au cours de cette thèse, nous avons tout d'abord développé différents types de façonneurs pour impulsions femtoseconde. Nous les avons ensuite utilisés pour contrôler de manière cohérente un système atomique simple. Ces expériences de contrôle cohérent ont permis d'observer l'évolution d'une fonction d'onde atomique à l'échelle femtoseconde. Une technique de caractérisation complète d'impulsions ultra-courtes a été dérivée de ces observations.

Le premier dispositif que j'ai conçu utilise un masque spatial à cristaux liquides (LC-SLM) placé au sein d'une ligne à dispersion nulle. Il permet la mise en forme d'impulsions à 800 nm, de 130 fs de durée initiale pour une énergie de plusieurs centaines de μJ . L'étalement du spectre de l'impulsion incidente au niveau du masque spatial permet d'utiliser ce dernier comme un filtre spectral. Afin d'obtenir un contrôle complet sur la forme des impulsions en sortie, la phase et l'amplitude des composantes spectrales transmises peuvent être ajustées indépendamment. Pour cela, deux masques haute résolution indépendants, ne contrôlant que la phase, ont été associés. Cette solution originale assure une résolution en phase et en amplitude jusqu'ici inégalée. Elle permet la réalisation d'impulsions aux profils temporels très complexes s'étalant sur plusieurs dizaines de picosecondes. Ces capacités de mise en forme ont été exploitées dans les différentes expériences de contrôle cohérent de cette thèse. Une étude des limitations de ce façonneur, notamment les effets de pixellisation et de couplage spatiotemporel, a été menée. Différentes solutions pour compenser une partie de ces défauts ont été proposées et mises en œuvre.

Un deuxième dispositif plus simple a été construit au STEACIE Institute dans le but de réaliser des expériences de contrôle optimal. Pour ces mêmes expériences, deux compresseurs à prismes et à miroirs déformables ont été développés pour pouvoir optimiser la durée d'impulsions visibles accordables. Divers soucis techniques ont retardé ces expériences, elles sont finalement prévues pour cet été.

Enfin, un AOPDF fonctionnant dans le visible a été étudié et optimisé pour la compression et la mise en forme des impulsions issues de NOPA. La possibilité de recomprimer la sortie d'un NOPA sur toute la gamme visible sans aucun réalignement a été démontrée. Nous avons aussi prouvé qu'en plus de la compression, l'AOPDF disposait de capacités de mise en forme tout à fait intéressantes, en conservant un rendement élevé. Ce type de compresseur-façonneur devrait prochainement permettre de simplifier grandement les expériences de contrôle cohérent ou de contrôle optimal dans la gamme visible.

Lors de la fabrication, de l'étude ou de l'utilisation de ces différents façonneurs, nous nous sommes heurtés au problème de la caractérisation des impulsions fortement mises en forme. En effet, les nombreuses techniques de caractérisation complète ou partielle qui existent ne sont pas tout à fait adaptées. La plupart de ces techniques trouvent leurs origines dans la course aux impulsions courtes qui a fait rage durant ces trois dernières décades. L'attention s'est donc portée sur la mesure précise de la phase spectrale résiduelle qui allonge la durée des impulsions. Dans le cas d'impulsions fortement mises en forme, le problème est tout autre : on s'intéresse à des phases très grandes, présentant éventuellement

des variations très brusques. De plus, l'amplitude spectrale est elle aussi très torturée et peut par exemple présenter des trous importants.

Une première technique simple à laquelle nous avons pensé est l'interférométrie spectrale. En envoyant dans un spectromètre deux impulsions décalées temporellement, l'une prélevée en entrée du façonneur, l'autre en sortie, il est possible de caractériser complètement l'action du façonneur (c'est-à-dire d'obtenir sa fonction de transfert). L'idée est d'utiliser l'impulsion non-mise en forme comme référence pour mesurer les changements subis par l'impulsion mise en forme. Malheureusement, cette technique réclame une précision interférométrique sur le délai séparant les deux impulsions, ce qui est très difficile à obtenir en pratique. En effet, le façonneur haute résolution introduit un chemin optique de plus de 5 mètres. Pour pouvoir recombinaison les deux impulsions dans le spectromètre, un chemin équivalent doit être introduit sur le parcours de l'impulsion prélevée en entrée. On se retrouve alors avec un interféromètre ayant deux bras extrêmement grands et il s'avère très compliqué d'obtenir la précision sur le délai à cause notamment des courants d'air et des vibrations. L'AOPDF quant à lui introduit des fluctuations sur le temps de sortie de l'impulsion mise en forme, ce qui rend la précision interférométrique impossible. Il existe toutefois des AOPDF sans gigue qui permettent de limiter très fortement ces fluctuations et d'obtenir la précision voulue. Malheureusement, les prototypes que nous avons utilisés ne proposaient pas cette stabilisation.

Devant l'impossibilité d'utiliser une référence, nous avons essayé des techniques auto-référencées comme le SPIDER ou le FROG. Cependant, ces deux techniques présentent très souvent des résolutions spectrales faibles (de l'ordre du nm) devant celle des façonneurs (0,06 nm pour le façonneur haute résolution). Bien que rien de fondamental dans ces techniques n'empêche une amélioration de leur résolution, la mise en œuvre pratique est complexe et coûteuse. En effet, dans les deux cas, il faut disposer de spectromètres de haute résolution, relativement onéreux. Pour le FROG, la quantité de données à traiter pour caractériser convenablement une impulsion de forte complexité peut s'avérer rédhibitoire. Pour le SPIDER, trouver un compromis entre résolution et quantité de signal est problématique. Différentes variantes du SPIDER permettent de résoudre ce problème mais leur mise en œuvre n'est pas simple.

Enfin, ces deux techniques s'accommodent très mal d'un trou spectral : au niveau d'un trou, on perd (partiellement ou totalement) l'information sur la phase relative entre les deux morceaux se trouvant de part et d'autre de celui-ci. Cela peut poser problème pour la caractérisation de séquences d'impulsions dont la génération requiert de fortes modulations dans le domaine spectral.

Finalement, nous avons eu le plus souvent recours à des techniques de caractérisation incomplètes, comme les cross-corrélations, ou indirectes, par le biais de la réponse d'un système simple aux impulsions à caractériser.

Nous avons ainsi utilisé l'atome de rubidium comme une sonde pour observer les modifications introduites par notre façonneur haute résolution sur les impulsions à 800 nm. Pour cela, nous avons excité de manière résonnante une transition à deux niveaux à l'aide d'une impulsion mise en forme de faible intensité et observé le comportement transitoire de la population transférée. Cette observation se fait à l'aide d'une deuxième impulsion

(visible) décalée temporellement (technique de pompe-sonde). Ces transitoires cohérents se sont révélés extrêmement sensibles aux changements de l'impulsion mise en forme. Ils ont permis de développer une technique de calibration très précise pour les façonneurs haute résolution.

Par la suite, nous avons exploité les capacités de notre façonneur pour mettre au point une technique de mesure de fonction d'onde atomique. Celle-ci permet de suivre l'évolution transitoire de la fonction d'onde atomique sur le niveau excité au cours de l'interaction avec une impulsion donnée. Cette technique tire profit de la linéarité en régime perturbatif de la réponse atomique avec l'impulsion excitatrice. Pour reconstruire l'effet sur la fonction d'onde atomique d'une impulsion donnée, on la fait précéder d'une autre impulsion. On enregistre ainsi plusieurs transitoires cohérents, obtenus pour différentes phases relatives entre les deux impulsions. En combinant ces enregistrements, il est possible de reconstruire l'amplitude de probabilité dans l'état excité due à la première impulsion. Nous avons ainsi pu observer l'évolution de la fonction d'onde atomique lorsque l'atome est soumis à différentes formes d'impulsions. Bien que cette mesure ait quelques ressemblances et partage des techniques de mesure avec les expériences de tomographie quantique ou d'holographie quantique, elle donne accès à une information quelque peu différente. En effet, ces deux techniques permettent de mesurer l'évolution libre d'un paquet d'onde constitué de plusieurs niveaux. Ici, nous observons l'évolution forcée, sous l'influence de l'impulsion mise en forme, de l'amplitude de probabilité sur un seul niveau atomique. Il semble toutefois possible d'étendre cette technique à une superposition de niveaux en s'inspirant de ce qui se fait en holographie quantique. On pourrait ainsi réaliser une reconstruction plus poussée de l'évolution des degrés de liberté interne de l'atome au cours de son interaction avec une impulsion mise en forme.

Enfin, nous avons dérivé de cette mesure de fonction d'onde atomique une technique de caractérisation d'impulsion. Pour cela on exploite le simple lien d'intégration qui existe en régime perturbatif entre l'impulsion excitatrice et l'amplitude de probabilité transférée. En dérivant l'amplitude de probabilité mesurée, il est possible de retrouver l'impulsion. Cette technique de mesure est en fait une méthode à référence : l'impulsion visible utilisée pour sonder l'atome apporte elle aussi sa contribution au signal atomique. On dispose ainsi d'une technique permettant la caractérisation d'une impulsion infrarouge à l'aide d'une référence visible (ou le contraire). L'appartenance des deux impulsions à des domaines spectraux différents est une réelle originalité pour une technique à référence. Autre particularité de la méthode, elle semble facilement transférable à d'autres longueurs d'onde. Il suffit de trouver un système à trois niveaux dans lequel les deux transitions successives sont dans les domaines spectraux voulus pour l'impulsion de pompe et celle de sonde. On pourrait ainsi tirer profit de ces deux particularités pour réaliser la caractérisation d'une impulsion dans une gamme spectrale où les techniques de caractérisations complètes sont difficiles (par exemple dans l'XUV ou l'infrarouge moyen) à l'aide d'une référence appartenant à une gamme spectrale moins problématique (comme le proche infrarouge, autour de 800 nm).

De plus, cette méthode de caractérisation semble tout à fait adaptée à la caractéri-

sation d'impulsions fortement mise en forme et ce pour deux raisons. Tout d'abord, elle est capable de reconstruire correctement la phase de part et d'autre d'un trou en amplitude, que ce soit dans le domaine spectral ou le domaine temporel. Ceci est sûrement lié à l'utilisation d'un signal atomique s'exprimant comme l'intégrale du champ électrique. Ensuite, les mêmes paramètres d'acquisition et de reconstruction peuvent être utilisés pour des impulsions très différentes : pour une intensité spectrale donnée, la quantité de signal est la même pour une impulsion limitée par transformée de Fourier ou une impulsions très fortement étirée. Cette robustesse vis-à-vis de la forme de l'impulsion est assez inhabituelle et tout à fait intéressante pour la caractérisation de façonneurs.

Un autre intérêt de cette technique est le fait que la caractérisation est faite *in situ* et que l'on mesure donc bien l'impulsion qui interagit avec le système étudié. Enfin, il semble possible d'étendre cette technique de caractérisation au cas où l'étape de pompe ionise le système, c'est-à-dire le cas où le niveau final est un continuum au lieu d'un état lié. Ceci pourrait peut-être permettre d'utiliser cette technique pour la caractérisation *in situ* d'impulsions XUV dans des systèmes moléculaires.

Annexe A

NOTATIONS POUR EXPÉRIMENTATEUR

A.1 Transformée de Fourier

A.1.1 Définitions

$\mathcal{F}$ désigne la transformée de Fourier directe. On notera $\tilde{E}(\omega)$ la transformée de Fourier directe de $E(t)$. Celle-ci est définie par :

$$\tilde{E}(\omega) = \frac{1}{\sqrt{2\pi}} \int e^{i\omega t} E(t) dt = \mathcal{F}[E(t)](\omega) \quad (\text{A.1})$$

La transformée de Fourier inverse est notée $\mathcal{F}^{-1}$ et définie par :

$$E(t) = \frac{1}{\sqrt{2\pi}} \int e^{-i\omega t} \tilde{E}(\omega) d\omega = \mathcal{F}^{-1}[\tilde{E}(\omega)](t) \quad (\text{A.2})$$

A.1.2 Propriétés et transformées de Fourier usuelles

Un pic de Dirac temporel, centré en t_0 a pour transformée de Fourier :

$$\mathcal{F}[\delta(t - t_0)](\omega) = \frac{1}{\sqrt{2\pi}} e^{i\omega t_0} \quad (\text{A.3})$$

Une impulsion monochromatique, centrée en ω_0 s'écrit dans le domaine temporel comme :

$$\mathcal{F}^{-1}[\delta(\omega - \omega_0)](t) = \frac{1}{\sqrt{2\pi}} e^{-i\omega_0 t} \quad (\text{A.4})$$

De plus, on a les relations suivantes entre produit et produit de convolution (noté $*$) :

$$\mathcal{F}[f.g] = \frac{1}{\sqrt{2\pi}} \mathcal{F}[f] * \mathcal{F}[g] \quad (\text{A.5})$$

$$\mathcal{F}[f*g] = \sqrt{2\pi} \mathcal{F}[f] . \mathcal{F}[g] \quad (\text{A.6})$$

En combinant ces différents résultats, on obtient la relation suivante dans le domaine spectral pour un décalage temporel de t_0 :

$$\mathcal{F}[E(t - t_0)](\omega) = \mathcal{F}[E(t)](\omega) e^{i\omega t_0} \quad (\text{A.7})$$

A.2 Champ électrique

A.2.1 Définitions - Notations

Pour faire plaisir à Schrödinger sans pour autant froisser l'intuition, nous avons pris les notations suivantes...

Le champ électrique est une quantité réelle, que nous noterons $\mathcal{E}(t)$ et que nous décomposerons comme :

$$\mathcal{E}(t) = \frac{1}{2} [E(t) + E^*(t)] = \text{Re} [E(t)] \quad (\text{A.8})$$

$E(t)$ est le champ complexe correspondant à sa partie de fréquence positive, il se met sous la forme :

$$E(t) = A(t)e^{-i\omega_0 t + i\varphi(t)} \quad (\text{A.9})$$

où ω_0 est la pulsation centrale du champ électrique, $A(t)$ son enveloppe et $\varphi(t)$ sa phase temporelle.

On note $\tilde{E}(\omega)$ la transformée de Fourier directe de $E(t)$, que l'on décompose en amplitude et phase :

$$\tilde{E}(\omega) = A(\omega)e^{i\phi(\omega)} \quad (\text{A.10})$$

Attention, il ne faut pas se laisser abuser par les notations prises dans les équations A.9 et A.10 : $A(t)$ et $A(\omega)$ ne sont pas liées par transformée de Fourier. De même, $\varphi(t)$ et $\phi(\omega)$ n'ont pas de lien direct. Ces notations traduisent juste la séparation habituelle en enveloppe réelle lentement variable et phase.

A.2.2 Durée - largeur spectrale

Les expérimentateurs le savent bien : on ne mesure que des largeurs à mi-hauteur (ou presque). Le plus souvent même, on mesure la largeur à mi-hauteur d'une intensité. Il est en effet très rare que l'on arrive à mesurer une demi-largeur à $1/e$ d'un simple coup d'oeil sur un signal expérimental... C'est pourquoi j'ai pris le parti d'écrire les formules de cette thèse aussi souvent que possible à l'aide de largeur à mi-hauteur en intensité. Je présente ici les notations que j'ai employé ainsi que le lien avec d'autres quantités souvent utilisées.

Nous prenons l'exemple d'un champ temporel de pulsation centrale ω_0 , ayant une phase temporelle $\varphi(t)$ nulle et possédant une enveloppe gaussienne.

Nous appellerons durée et nous noterons Δt sa largeur à mi-hauteur en intensité :

$$E(t) = A(t)e^{-i\omega_0 t} \quad \text{avec} \quad A(t) = 4^{-\frac{t^2}{\Delta t^2}} = e^{-2\ln(2)\frac{t^2}{\Delta t^2}} \quad (\text{A.11})$$

Celle-ci est liée à la demi-largeur à $1/e$, que nous noterons T , par :

$$\Delta t = T\sqrt{2\ln(2)}$$

Dans le domaine spectral ce champ est aussi représenté par une enveloppe gaussienne et une phase spectrale plate.

On appellera largeur spectrale de ce champ, que l'on notera $\Delta\omega$ la largeur à mi-hauteur en intensité de cette gaussienne :

$$\tilde{E}(\omega) = A(\omega) = 4^{-\frac{(\omega-\omega_0)^2}{\Delta\omega^2}} = e^{-2\ln(2)\frac{(\omega-\omega_0)^2}{\Delta\omega^2}} \quad (\text{A.12})$$

où $\Delta\omega$ est la largeur spectrale à mi-hauteur en intensité. Celle-ci est liée à la demi-largeur spectrale à $1/e$, notée Ω , par :

$$\Delta\omega = \Omega\sqrt{2\ln(2)}$$

Ainsi définies, la largeur spectrale et la durée de ce champ limité par transformée de Fourier sont liées par la relation :

$$\Delta t.\Delta\omega = 4\ln(2)$$

Les demi-largeurs à $1/e$ vérifient quant à elles :

$$T.\Omega = 2$$

Annexe B

COUPLAGE SPATIOTEMPOREL

Cette annexe est un complément à la partie II.9.3 page 54. Nous ne détaillerons pas ici complètement les calculs permettant d'obtenir le champ de sortie en fonction de la mise en forme programmée. Ceux-ci ont été détaillés dans plusieurs articles, notamment les références [86, 95]. Nous nous contenterons de donner les éléments essentiels du calcul, l'expression du champ au plan de Fourier ainsi qu'à la sortie du façonneur. Nous donnerons enfin l'expression analytique du champ de sortie dans le cas d'une phase linéaire et d'une phase quadratique.

B.1 Conventions et notations

Le couplage spatiotemporel est dû à l'étalement des pulsations dans le plan du masque (plan de Fourier de la ligne 4f). L'expression complète du champs électrique nous importe peu. Du point de vue spatial, seule nous importe la coordonnée transverse horizontale x selon laquelle se fait l'étalement des pulsations au sein de la ligne. Selon la verticale y , il n'y a pas de changement et nous omettons cette coordonnée. Nous nous intéresserons donc à l'expression du champs en fonction de la coordonnée transverse horizontale x et du temps ou de la pulsation. Dans tous les calculs présentés ci-après, nous omettons les différents facteurs multiplicatifs pouvant apparaître, afin d'alléger les équations.

B.1.1 Champs en entrée

Nous supposons que le champs en entrée possède un profil transverse gaussien de largeur à mi-hauteur en intensité Δx . Nous supposons que les dépendances spatiales et spectrales sont factorisables¹. Nous l'écrivons donc sous la forme :

$$E_E(\omega, x) = E_E(\omega) e^{-2\ln(2)\left(\frac{x}{\Delta x}\right)^2} = E_E(\omega) 4^{-\left(\frac{x}{\Delta x}\right)^2} \quad (\text{B.1})$$

La pulsation ω est définie par rapport à la pulsation centrale ω_0 .

B.2 Champs au plan de Fourier

Pour comprendre l'action du masque spatial $M_X(X)$, il faut connaître l'expression du champs électrique au niveau du plan de Fourier de notre façonneur. Nous allons pour

¹C'est à dire qu'il n'existe pas de couplage entre les deux : ni dérive spatiale de fréquence, ni aberration chromatique...

cela détailler l'action d'un réseau et celle d'un miroir cylindrique. En combinant les deux, nous pourrions déterminer le champ électrique avant et après le masque programmable.

B.2.0.1 Action d'un réseau

Nous nous intéressons à l'action du réseau d'entrée sur un champs $E_0(x_0, \omega)$ de pulsation centrale ω_0 . La figure B.1 résume les différentes notations. Juste après le réseau le

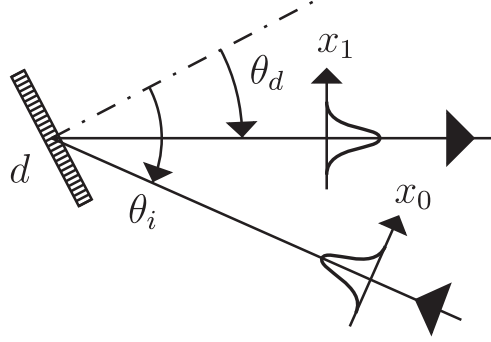


FIG. B.1 – Dispersion d'un faisceau par un réseau.

champs devient[138, 87]

$$E_1(\omega, x_1) = E_0(x_1/\beta, \omega) e^{ik\gamma\omega x_1} \quad (\text{B.2})$$

où β est le facteur d'échelle spatial du réseau et γ sa dispersion angulaire et où l'on a posé $k = \omega_0/c$. c est la vitesse de la lumière. Les deux paramètres du réseau s'écrivent

$$\beta = \frac{\cos(\theta_d)}{\cos(\theta_i)} \quad (\text{B.3a})$$

$$\gamma = \frac{2\pi c}{d\omega_0^2 \cos(\theta_d)} \quad (\text{B.3b})$$

θ_i et θ_d étant les angles d'incidence et de diffraction sur le réseau et d le pas du réseau.

On peut donc voir l'action du réseau comme l'application d'un facteur d'échelle $1/\beta$ sur le profil spatial et l'introduction d'une dispersion angulaire définie par γ . Dans la ligne 4f, γ est liée à la dispersion linéaire dans le plan de Fourier α et à la focale f des miroirs par : $\gamma = \alpha/f$ (voir équation II.21 page 44).

L'action du réseau de sortie est similaire. Toutefois, afin de compenser l'effet du réseau d'entrée, ses paramètres β' et γ' valent :

$$\beta' = 1/\beta \quad (\text{B.4a})$$

$$\gamma' = -\beta\gamma \quad (\text{B.4b})$$

B.2.0.2 Action d'un miroir cylindrique

Nous considérons maintenant l'action d'un miroir cylindrique sur un champ $E_1(x_1, \omega)$. La figure B.2 page suivante présente les notations dans le cas, plus simple à représenter,

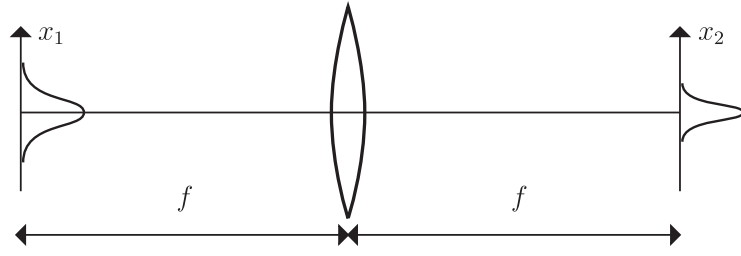


FIG. B.2 – Action d'une lentille sur un faisceau.

d'une lentille. $E_1(x_1, \omega)$ étant l'expression du champs dans le plan focal objet du miroir de focale f , on obtient dans son plan focal image un champ $E_2(x_2, \omega)$ donné par :

$$E_2(\omega, x_2) = \int e^{i\frac{2\pi x_1 x_2}{f\lambda_0}} E_1(\omega, x_1) dx_1 \quad (\text{B.5})$$

L'action du miroir se traduit donc comme une transformée de Fourier selon la coordonnée transverse x : x_1 et x_2 sont conjuguées par transformée de Fourier, à un facteur d'échelle près.

B.2.0.3 Champ au plan de Fourier

En utilisant les équations B.2 page ci-contre et B.5, on trouve pour le champ entrant dans le plan de Fourier une expression simple :

$$E_{\text{pf},E}(\omega, X) = E_E(\omega) e^{-2\ln(2)\left(\frac{X-f\gamma\omega}{\Delta X}\right)^2} = E_E(\omega) 4^{-\left(\frac{X-f\gamma\omega}{\Delta X}\right)^2} \quad (\text{B.6})$$

où X est la position transverse dans le plan de Fourier et ΔX est donnée par :

$$\Delta X = \frac{2\ln(2)\lambda_0 f}{\pi\beta\Delta x} = \frac{4\ln(2)cf}{\beta\omega_0\Delta x} \quad (\text{B.7})$$

On retrouve ainsi le champ spectral entrant dans la ligne $E_E(\omega)$, multiplié par une gaussienne. Cette gaussienne traduit l'étalement du spectre dans le plan de Fourier par le biais du terme $(X - f\gamma\omega)$ qui couple position et pulsation. De plus, elle correspond au profil spatial que donnerait chaque composante spectrale prise individuellement.

Après la traversée du masque $M_X(X)$, ce champ devient :

$$E_{\text{pf},S}(\omega, X) = E_{\text{pf},E}M(X) = E_E(\omega) M_X(X) 4^{-\left(\frac{X-f\gamma\omega}{\Delta X}\right)^2} \quad (\text{B.8})$$

B.3 Champ en sortie

On obtient le champ en sortie en propageant le champ donné dans l'équation B.8 au travers de la deuxième demi-ligne. On réutilise pour cela les équations B.2 et B.5, en prenant garde que le réseau de sortie a un effet opposé à celui de l'entrée (voir l'équation B.4 page ci-contre).

Les calculs permettent ainsi de mettre le champ de sortie sous la forme :

$$E_S(\omega, x) = E_E(\omega) \int e^{iKX} M_X(X + X_0) 4^{-\left(\frac{x}{\Delta X}\right)^2} dX \quad (\text{B.9})$$

où l'on a introduit

$$K(x) = \frac{\beta\omega_0 x}{cf} \quad (\text{B.10})$$

et

$$X_0(\omega) = f\gamma\omega \quad (\text{B.11})$$

K apparaît comme la variable conjuguée de la position dans le plan de Fourier et X_0 comme un décalage dans le plan de Fourier. L'équation B.9 donne le champ de sortie en fonction du spectre d'entrée et de la transformée de Fourier du masque décalé et multiplié par une gaussienne.

Pour des raisons qui apparaîtront évidentes par la suite, nous n'allons pas utiliser l'équation B.9 telle quelle. Nous allons utiliser la relation de dispersion $X_0(\omega) = f\gamma\omega$ pour faire les changements de variable :

$$\omega' = \frac{X}{f\gamma} \quad (\text{B.12a})$$

$$\Delta\omega' = \frac{\Delta X}{f\gamma} \quad (\text{B.12b})$$

Nous définissons aussi v , vitesse de couplage spatiotemporel :

$$v = \frac{c}{\gamma\beta\omega_0} = \frac{-dc \cos(\theta_i)}{\lambda_0} \quad (\text{B.13})$$

Fort de ces changements, on obtient :

$$E_S(\omega, x) = E_E(\omega) \int e^{i\frac{x}{v}\omega'} M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' \quad (\text{B.14})$$

où l'on a utilisé la notation $M_w(\omega) = M_X(f\gamma\omega)$. Bien qu'implicite, le couplage spatiotemporel apparaît très clairement dans cette équation : la position transverse est, à un facteur $1/v$ près, la variable conjuguée de la pulsation. De plus, cette équation permet de calculer analytiquement le champ en sortie pour certains masques simples.

B.4 Filtrage spatial

Dans de très nombreux cas, il est impossible d'obtenir une forme simple pour le champ de sortie où les variables spatiales et spectrales sont séparées. On utilise alors une forme approchée en développant l'équation II.33 page 55 sur les différents modes hermito-gaussiens[29] et en n'en gardant que le mode gaussien. Cette approximation est tout à fait justifiée lorsque l'on filtre spatialement la sortie du façonneur. Le filtrage le plus simple se résume à un diaphragme de faible diamètre, comme cela est fait sur notre façonneur (diaphragme D0 sur la figure II.16 page 31).

Sous ces conditions, on peut écrire pour le champ de sortie, filtré[86, 48, 95, 28] :

$$E_S(\omega, x) = E_E(\omega, x) \int M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' \quad (\text{B.15})$$

où l'on voit apparaître la fonction de transfert $H(\omega)$ du façonneur qui vaut :

$$H(\omega) = \int M_w(\omega' + \omega) 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' = \int M_w(\omega') 4^{-\left(\frac{\omega' - \omega}{\Delta\omega'}\right)^2} d\omega' \quad (\text{B.16})$$

Le masque programmé est ainsi convolué par une gaussienne de largeur à mi-hauteur en intensité $\Delta\omega'$. Cela revient à dire que la mise en forme temporelle correspondant à ce masque sera multipliée par une enveloppe gaussienne de largeur à mi-hauteur :

$$\Delta T = \frac{4 \ln(2)}{\Delta\omega'} = \frac{\Delta x}{|v|} \quad (\text{B.17})$$

B.5 Exemples

B.5.1 Cas d'un délai

Pour retarder d'un délai τ l'impulsion traversant la façonneur, le masque doit introduire une phase spectrale linéaire de la forme $\phi(\omega) = -\omega\tau$. Cela correspond à un masque :

$$M_w(\omega) = e^{-i\omega\tau} \quad (\text{B.18})$$

Nous reprenons les notations de la partie précédente. Le champ en sortie s'écrit comme :

$$E_S(\omega, x) = e^{-i\tau\omega} E_E(\omega) 4^{-\left(\frac{x/v - \tau}{\Delta T}\right)^2} \quad (\text{B.19})$$

ce qui, compte tenu de la définition de $E_E(\omega, x)$ peut se réécrire sous la forme.

$$E_S(\omega, x) = e^{-i\tau\omega} E_E(\omega) 4^{-\left(\frac{x - v\tau}{\Delta x}\right)^2} = e^{-i\tau\omega} E_E(\omega, x - v\tau) \quad (\text{B.20})$$

En définissant $E_S(t, x)$ comme la transformé de Fourier inverse de $E_S(\omega, x)$ on peut écrire ce même résultat sous la forme :

$$E_S(t, x) = E_E(t - \tau, x - v\tau) \quad (\text{B.21})$$

On obtient ainsi bien un champ décalé temporellement de τ , mais aussi décalé spatialement de $v\tau$.

Si on applique un filtrage spatial, on obtiendra

$$E_S(\omega, x) = 4^{-\left(\frac{\tau}{\Delta T}\right)^2} e^{-i\tau\omega} E_E(\omega, x) \quad (\text{B.22})$$

c'est à dire la mise en forme souhaitée, mais avec une atténuation gaussienne.

B.5.2 Cas d'une phase quadratique

Pour réaliser une dérive de fréquence, le masque doit introduire une phase quadratique en pulsation soit

$$M_w(\omega) = e^{-i\frac{1}{2}\phi^{(2)}(\omega)^2} \quad (\text{B.23})$$

Le champ de sortie, donné par l'équation B.14 page 160, peut se ré-écrire sous la forme

$$E_S(\omega, x) = E_E(\omega) \int e^{i\frac{x}{v}\omega'} e^{-i[\frac{1}{2}\phi^{(2)}(\omega')^2 + \phi^{(1)}\omega' + \phi^{(0)}]} 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' \quad (\text{B.24})$$

en posant :

$$\phi^{(1)} = \phi^{(2)}\omega \quad (\text{B.25a})$$

$$\phi^{(0)} = \frac{1}{2}\phi^{(2)}(\omega)^2 \quad (\text{B.25b})$$

Pour la partie intégrale de l'équation B.24, on a donc :

$$I(\omega, x) = \int e^{i\frac{x}{v}\omega'} e^{-i[\frac{1}{2}\phi^{(2)}(\omega')^2 + \phi^{(1)}\omega' + \phi^{(0)}]} 4^{-\left(\frac{\omega'}{\Delta\omega'}\right)^2} d\omega' \quad (\text{B.26})$$

On voit apparaître la transformée de Fourier d'une gaussienne spectrale à laquelle on a appliqué une phase polynômiale d'ordre deux.

Cette intégrale est connue[55] : elle est semblable à l'expression temporelle d'une impulsion gaussienne ayant subi une phase polynômiale. Elle vaut :

$$I(\omega, x) = e^{i[\alpha(x/v - \phi^{(1)})^2 - \phi^{(0)} + \theta]} 4^{-\left(\frac{x/v - \phi^{(1)}}{\Delta x_p/v}\right)^2} \quad (\text{B.27})$$

avec

$$\Delta x_p^2 = \left(\frac{4 \ln(2) \phi^{(2)} v}{\Delta x}\right)^2 + \Delta x^2 \quad (\text{B.28a})$$

$$\alpha = \frac{2\phi^{(2)}}{(2\phi^{(2)})^2 + \left(\frac{\Delta x^2}{2 \ln(2) v^2}\right)^2} \quad (\text{B.28b})$$

$$\theta = -\arctan\left(\frac{4 \ln(2) \phi^{(2)} v^2}{\Delta x^2}\right) \quad (\text{B.28c})$$

$$(\text{B.28d})$$

Compte tenu des définitions de l'équation B.25, on peut écrire le champ de sortie :

$$E_S(\omega, x) = e^{-i\frac{1}{2}\phi^{(2)}\omega^2} E_E(\omega) e^{i\left[\frac{\alpha}{v^2}(x - v\phi^{(2)}\omega)^2 + \theta\right]} 4^{-\left(\frac{x - v\phi^{(2)}\omega}{\Delta x_p}\right)^2} \quad (\text{B.29})$$

Cette forme permet de bien mettre en avant les caractéristiques du champ de sortie. On retrouve tout d'abord bien la phase spectrale quadratique que l'on souhaitait appliquer à l'impulsion. On voit ensuite clairement sur le profil spatial une dérive de fréquence spatiale :

le profil est élargi (Δx_p) et la direction de propagation apparente varie avec la position transverse.

Enfin, la position du maximum d'intensité varie avec la pulsation. Cette dernière caractéristique se comprend mieux en se rappelant que, dans le domaine temporel, tout se passe comme si chaque pulsation ω subissait un délai $\tau = \phi^{(2)}\omega$. On retrouve donc un couplage en $x - v\tau$ identique à celui rencontré pour une phase spectrale linéaire (voir l'équation B.21 page 161).

Contrairement au cas de la phase linéaire, il est ici impossible de séparer la dépendance spatiale de la dépendance spectrale.

Bibliographie

- [1] D. J. Tannor and S. A. Rice, "Control of selectivity of chemical reaction via control of wave packet evolution," *Journal of Chemical Physics* **83**(10), 5013–5018 (1985).
- [2] D. J. Tannor and S. A. Rice, "Coherent pulse sequence induced control of selectivity of reactions : Exact quantum mechanical calculations," *Journal of Chemical Physics* **85**(10), 5805–5820 (1986).
- [3] M. Shapiro and P. Brumer, "Coherent and incoherent laser control of photochemical reactions." *International Reviews in Physical Chemistry* **13**(2), 187–229 (1994).
- [4] C. Chen, Y.-Y. Yin, and D. Elliott, "Interference between Optical Transitions," *Physical Review Letters* **64**, 507–10 (1990).
- [5] R. J. Gordon and S. A. Rice, "Active control of the dynamics of atoms and molecules," *Annual Review of Physical Chemistry* **48**, 601–641 (1997).
- [6] R. S. Judson and H. Rabitz, "Teaching lasers to control molecules," *Physical Review Letters* **68**(10), 1500–3 (1992).
- [7] A. Assion, T. Baumert, M. Bergt, T. Brixner, B. Kiefer, V. Seyfried, M. Strehle, and G. Gerber, "Control of Chemical Reactions by Feedback-optimized Phase-shaped femtosecond laser pulses," *Science* **282**, 919 (1998).
- [8] C. Daniel, J. Full, L. González, C. Kaposta, M. Krenz, C. Lupulescu, J. Manz, S. Minemoto, M. Oppel, P. Rosendo-Francisco, S. Vajda, and L. Wöste, "Analysis and control of laser induced fragmentation processes in $\text{CpMn}(\text{CO})_3$," *Chemical Physics* **267**(1-3), 247–260 (2001).
- [9] R. J. Levis, G. M. Menkir, and H. Rabitz, "Selective Bond Dissociation and Rearrangement with Optimally Tailored, Strong-Field Laser Pulses," *Science* **292**(5517), 709–713 (2001).
- [10] S. Zamith, T. Martchenko, Y. Ni, S. A. Aseyev, H. G. Muller, and M. J. J. Vrakking, "Control of the production of highly charged ions in femtosecond-laser cluster fragmentation," *Physical Review A (Atomic, Molecular, and Optical Physics)* **70**(1), 011,201 (2004).
- [11] C. Rangan and P. H. Bucksbaum, "Optimally shaped terahertz pulses for phase retrieval in a Rydberg-atom data register," *Physical Review A* **64**(3), 033,417/1–5 (2001).
- [12] J. L. Herek, W. Wohlleben, R. J. Cogdell, D. Zeidler, and M. Motzkus, "Quantum control of energy flow in light harvesting," *Nature* **417**(6888), 533 (2002).
- [13] R. Bartels, S. Backus, E. Zeek, L. Misoguti, G. Vdovin, I. Christov, M. M. Murnane, and C. Kapteyn, "Shaped-pulse optimization of coherent emission of high-harmonic soft X-rays," *Nature* **406**, 164–166 (2000).
- [14] D. H. Reitze, S. Kazamias, F. Weihe, G. Mullot, D. Douillet, F. Audebert, O. Albert, V. Rannanathan, L. P. Chambaret, D. Hulin, and P. Balcou, "Enhancement of high-order harmonic generation at tuned wavelengths through adaptive control," *Optics Letters* **29**(1), 86–8 (2004).
- [15] D. Meshulach and Y. Silberberg, "Coherent Quantum Control of 2-Photon Transitions by a Femtosecond Laser-Pulse," *Nature* **396**(6708), 239–242 (1998).

- [16] N. Dudovich, D. Oron, and Y. Silberberg, "Coherent transient enhancement of optically induced resonant transitions," *Physical Review Letters* **88**(12), 123,004–1–4 (2002).
- [17] S. Zamith, J. Degert, S. Stock, B. de Beauvoir, V. Blanchet, M. A. Bouchene, and B. Girard, "Observation of Coherent Transients in Ultrashort Chirped Excitation of an undamped Two-Level System," *Physical Review Letters* **87**(3), 033,001 (2001).
- [18] J. Degert, W. Wohlleben, B. Chatel, M. Motzkus, and B. Girard, "Realization of a Time-Domain Fresnel Lens with Coherent Control," *Physical Review Letters* **89**(20), 203,003 (2002).
- [19] B. Chatel, J. Degert, S. Stock, and B. Girard, "Competition between sequential and direct paths in a two-photon transition," *Physical Review A (Atomic, Molecular, and Optical Physics)* **68**(4), 041,402R (2003).
- [20] V. V. Lozovoy, I. Pastirk, K. A. Walowicz, and M. Dantus, "Multiphoton intrapulse interference. II. Control of two- and three-photon laser induced fluorescence with shaped pulses," *Journal of Chemical Physics* **118**(7), 3187–96 (2003).
- [21] M. Renard, E. Hertz, B. Lavorel, and O. Faucher, "Controlling ground-state rotational dynamics of molecules by shaped femtosecond laser pulses," *Physical Review A (Atomic, Molecular, and Optical Physics)* **69**(4), 043,401 (2004).
- [22] R. J. Levis and H. A. Rabitz, "Closing the loop on bond selective chemistry using tailored strong field laser pulses," *Journal of Physical Chemistry A* **106**(27), 6427–44 (2002).
- [23] W. Wohlleben, T. Buckup, J. L. Herek, and M. Motzkus, "Coherent Control for Spectroscopy and Manipulation of Biological Dynamics," *Chemphyschem* **6**(5), 850–857 (2004).
- [24] M. Dantus and V. V. Lozovoy, "Experimental Coherent Laser Control of Physicochemical Processes," *Chemical Reviews* **104**(4), 1813–1859 (2004).
- [25] B. Chatel and B. Girard, "Coherent Control of Atomic Dynamics with Chirped and Shaped Pulses," in *Femtosecond Laser Spectroscopy*, P. Hannaforde, ed., pp. 267–304 (Springer, 2005).
- [26] R. Ueberna, M. Khalil, R. Williams, J. Papanikolas, and S. Leone, "Phase and amplitude control in the formation and detection of rotational wave packets in the E 1S+g state of Li₂," *Journal of Chemical Physics* **108**(22), 9259–9274 (1998).
- [27] F. Verluise, V. Laude, Z. Cheng, C. Spielmann, and P. Tourniois, "Amplitude and phase control of ultrashort pulses by use of an acousto-optic programmable dispersive filter : pulse compression and shaping," *Optics Letters* **25**(8), 575–7 (2000).
- [28] A. M. Weiner, "Femtosecond pulse shaping using spatial light modulators," *Review of Scientific Instruments* **71**(5), 1929–60 (2000).
- [29] R. N. Thurston, J. P. Heritage, A. M. Weiner, and W. J. Tomlinson, "Analysis of picosecond pulse shape synthesis by spectral masking in a grating pulse compressor," *IEEE Journal of Quantum Electronics* **QE-22**(5), 682–96 (1986).
- [30] A. M. Weiner, D. E. Leaird, J. S. Patel, and J. R. Wullert, "Programmable femto-

- second pulse shaping by use of a multielement liquid-crystal phase modulator," *Optics Letters* **15**(6), 326–8 (1990).
- [31] C. W. Hillegas, J. X. Tull, D. Goswami, D. Strickland, and W. S. Warren, "Femtosecond laser pulse shaping by use of microsecond radio-frequency pulses," *Optics Letters* **19**(10), 737–9 (1994).
- [32] E. Zeek, K. Maginnis, S. Backus, U. Russek, M. Murnane, G. Mourou, and H. Kapteyn, "Pulse compression by use of deformable mirrors," *Optics Letters* **24**(7), 493–495 (1999).
- [33] A. M. Weiner, "Femtosecond optical pulse shaping and processing," *Progress in Quantum Electronics* **19**(3), 161 (1995).
- [34] C. Froehly, B. Colombeau, and M. Vampouille, "Shaping and analysis of picosecond light pulses," *Progress in optics* **20**, 65–153 (1983).
- [35] A. Monmayrant and B. Chatel, "A new phase and amplitude high resolution pulse shaper," *Review of Scientific Instruments* **75**(8), 2668–71 (2004).
- [36] T. Brixner and G. Gerber, "Femtosecond polarization pulse shaping," *Optics Letters* **26**(8), 557–9 (2001).
- [37] P. Tournois, "Acousto-optic programmable dispersive filter for adaptive compensation of group delay time dispersion in laser systems," *Optics Communications* **140**(4–6), 245–9 (1997).
- [38] A. Efimov, C. Schaffer, and D. H. Reitze, "Programmable shaping of ultrabroadbandwidth pulses from a Ti:sapphire laser," *Journal of the Optical Society of America B (Optical Physics)* **12**(10), 1968–80 (1995).
- [39] M. Pittman, S. Ferré, J. P. Rousseau, L. Notebaert, J. P. Chambaret, and G. Chériaux, "Design and characterization of a near-diffraction-limited femtosecond 100-TW 10-Hz high-intensity laser system," *Applied Physics B : Lasers and Optics* **74**(6), 529 (2002).
- [40] D. Zeidler, T. Hornung, D. Proch, and M. Motzkus, "Adaptive compression of tunable pulses from a non-collinear-type OPA to below 16 fs by feedback-controlled pulse shaping," *Applied Physics B* **70**(June), S125–31 (2000).
- [41] A. Monmayrant, A. Arbouet, B. Girard, B. Chatel, A. Barman, B. J. Whitaker, and D. Kaplan, "AOPDF-shaped optical parametric amplifier output in the visible," *Applied Physics B : Lasers and Optics* **81**(2–3), 177–180 (2005).
- [42] T. Witte, K. Kompa, and M. Motzkus, "Femtosecond pulse shaping in the mid infrared by difference-frequency mixing," *Applied Physics B* **76**(4), 467–471 (2003).
- [43] M. Hacker, T. Feurer, R. Sauerbrey, T. Lucza, and G. Szabo, "Programmable femtosecond laser pulses in the ultraviolet," *Journal of the Optical Society of America B (Optical Physics)* **18**(6), 866–71 (2001).
- [44] N. Belabas, J. P. Likforman, L. Canioni, B. Bousquet, and M. Joffre, "Coherent broadband pulse shaping in the mid infrared," *Optics Letters* **26**(10), 743–5 (2001).
- [45] M. Hacker, G. Stobrawa, R. Sauerbrey, T. Buckup, M. Motzkus, M. Wildenhain, and A. Gehner, "Micromirror SLM for femtosecond pulse shaping in the ultraviolet," *Applied Physics B-Lasers And Optics* **76**(6), 711–714 (2003).

- [46] B. Chatel, J. Degert, and B. Girard, "Role of quadratic and cubic spectral phases in ladder climbing with ultrashort pulses," *Physical Review A (Atomic, Molecular, and Optical Physics)* **70**(5), 053,414 (2004).
- [47] W. Wohlleben, J. Degert, A. Monmayrant, B. Chatel, M. Motzkus, and B. Girard, "Coherent transients as a highly sensitive probe for femtosecond pulse shaper," *Applied Physics B (Lasers and Optics)* **79**(4), 435 – 439 (2004).
- [48] M. M. Wefers and K. A. Nelson, "Analysis of programmable ultrashort waveform generation using liquid-crystal spatial light-modulators," *Journal of the Optical Society of America B (Optical Physics)* **12**(7), 1343–62 (1995).
- [49] G. Stobrawa, M. Hacker, T. Feurer, D. Zeidler, M. Motzkus, and F. Reichel, "A new high-resolution femtosecond pulse shaper," *Applied Physics B* **72**(5), 627–30 (2001).
- [50] A. Monmayrant, B. Chatel, and B. Girard, "Femtosecond pulse shape reconstruction from coherent transient measurements," *Optics Letters in preparation* (2005).
- [51] A. Monmayrant, B. Chatel, and B. Girard, "Atomic spirograph : measurement of the excited state wave function using coherent transients," *Physical Review Letter in preparation* (2005).
- [52] A. Monmayrant, B. Chatel, and B. Girard, "Towards electric field reconstruction using coherent transients in a two-level system," in *Ultrafast Phenomena XIV*, T. Kobayashi, ed., pp. 91–93 (Springer, Niigata, 2004).
- [53] N. V. Vitanov, "Transition times in the Landau-Zener model," *Physical Review A* **59**(2), 988–94 (1999).
- [54] C. Nicole, "Dynamique de paquets d'ondes et contrôle cohérent temporel dans les atomes et les molécules d'alcalins," Thèse, Paul Sabatier (1998).
- [55] J. Degert, "Manipulation cohérente d'atomes et de molécules diatomiques avec des impulsions mises en forme," Thèse, Univ. P. Sabatier (2002).
- [56] C. Le Blanc, "Principes et réalisation d'une source laser terawatt femtoseconde basée sur le saphir dopé au titane. Caractérisation des impulsions produites et démonstration du régime d'intensité au niveau de 10^{18} W/cm²," Thèse, Ecole Polytechnique (1993).
- [57] G. Cerullo and S. De Silvestri, "Ultrafast optical parametric amplifiers," *Review of Scientific Instruments* **74**(1), 1–18 (2003).
- [58] T. Wilhelm, J. Piel, and E. Riedle, "Sub-20-fs pulses tunable across the visible from a blue-pumped single-pass noncollinear parametric converter," *Optics Letters* **22**(19), 1494–1496 (1997).
- [59] A. Shirakawa and T. Kobayashi, "Noncollinearly phase-matched femtosecond optical parametric amplification with a 2000 cm⁻¹ bandwidth," *Applied Physics Letters* **72**(2), 147–9 (1998).
- [60] G. Cerullo, M. Nisoli, and S. De Silvestri, "Generation of 11 fs pulses tunable across the visible by optical parametric amplification," *Applied Physics Letters* **71**(25), 3616–18 (1997).
- [61] M. Cavallari, "Interactions paramétriques intracavité en régime femtoseconde. Opti-

- misation d'un oscillateur paramétrique optique visible," Thèse, Ecole Polytechnique (1997).
- [62] G. M. Gale, F. Hache, and M. Cavallari, "Broad-bandwidth parametric amplification in the visible : femtosecond experiments and simulations," *IEEE Journal of Selected Topics in Quantum Electronics* **4**(2), 224–9 (1998).
 - [63] E. Goulielmakis, M. Uiberacker, R. Kienberger, A. Baltuska, V. Yakovlev, A. Scrinzi, T. Westerwalbesloh, U. Kleineberg, U. Heinzmann, M. Drescher, and F. Krausz, "Direct Measurement of Light Waves," *Science* **305**(5688), 1267–1269 (2004).
 - [64] S. Linden, J. Kuhl, and H. Giessen, "Amplitude and phase characterization of weak blue ultrashort pulses by downconversion," *Optics Letters* **24**, 569 (1999).
 - [65] R. Trebino, *Frequency-Resolved Optical Gating : The Measurement of Ultrashort Laser Pulses* (Kluwer Academic, 2000).
 - [66] T. Y. F. Tsang, M. A. Krumbugel, K. W. DeLong, D. N. Fittinghoff, and R. Trebino, "Frequency-resolved optical-gating measurements of ultrashort pulses using surface third-harmonic generation," *Optics Letters* **21**(17), 1381–1383 (1996).
 - [67] D. J. Kane, "Recent progress toward real-time measurement of ultrashort laser pulses," *IEEE Journal of Quantum Electronics* **35**(4), 421 (1999).
 - [68] C. Iaconis and I. A. Walmsley, "Spectral phase interferometry for direct electric-field reconstruction of ultrashort optical pulses," *Optics Letters* **23**(10), 792–794 (1998).
 - [69] M. Françon, *Optical interferometry* (Academic Press, 1966).
 - [70] C. Dorrer, "Implementation of spectral phase interferometry for direct electric-field reconstruction with a simultaneously recorded reference interferogram," *Optics Letters* **24**(21), 1532–4 (1999).
 - [71] I. Kang, C. Dorrer, and F. Quochi, "Implementation of electro-optic spectral shearing interferometry for ultrashort pulse characterization," *Optics Letters* **28**(22), 2264–2266 (2003).
 - [72] C. Dorrer, P. Londero, and I. A. Walmsley, "Homodyne detection in spectral phase interferometry for direct electric field reconstruction," *Optics Letters* **26**(19), 1510 (2001).
 - [73] M. Hirasawa, N. Nakagawa, K. Yamamoto, R. Morita, H. Shigekawa, and M. Yamashita, "Sensitivity improvement of spectral phase interferometry for direct electric-field reconstruction for the characterization of low-intensity femtosecond pulses," *Applied Physics B : Lasers and Optics* **74**(0), s225 (2002).
 - [74] V. Messager, F. Louradour, C. Froehly, and A. Barthelemy, "Coherent measurement of short laser pulses based on spectral interferometry resolved in time," *Optics Letters* **28**(9), 743–745 (2003).
 - [75] J. C. Vaughan, T. Hornung, T. Feurer, and K. A. Nelson, "Diffraction-based femtosecond pulse shaping with a two-dimensional spatial light modulator," *Optics Letters* **30**(3), 323–325 (2005).
 - [76] T. Feurer, J. C. Vaughan, R. M. Koehl, and K. A. Nelson, "Multidimensional control of femtosecond pulses by use of a programmable liquid-crystal matrix," *Optics Letters* **27**(8), 652–654 (2002).

- [77] C. Dorrer, F. Salin, F. Verluise, and J. P. Huignard, "Programmable phase control of femtosecond pulses by use of a nonpixelated spatial light modulator," *Optics Letters* **23**(9), 709–11 (1998).
- [78] A. M. Weiner, J. P. Heritage, and E. M. Kirschner, "High-resolution femtosecond pulse shaping," *Journal of the Optical Society of America B (Optical Physics)* **5**(8), 1563–72 (1988).
- [79] F. Verluise, "Filtre acousto-optique programmable en amplitude et en phase : application aux impulsions laser ultrabrèves," Ph.D. thesis, Ecole Polytechnique (1999).
- [80] K. Takayoshi and B. Andrius, "Sub-5 fs pulse generation from a noncollinear optical parametric amplifier," *Measurement Science and Technology* **13**(11), 1671 (2002).
- [81] M. Roth, M. Mehendale, A. Bartelt, and H. Rabitz, "Acousto-optical shaping of ultraviolet femtosecond pulses," *Applied Physics B : Lasers and Optics* **80**(4 - 5), 441 (2005).
- [82] N. Dudovich, D. Oron, and Y. Silberberg, "Quantum Control of the Angular Momentum Distribution in Multiphoton Absorption Processes," *Physical Review Letters* **92**(10), 103,003–1 103,003–4 (2004).
- [83] D. Zeidler, "Coherent Control of Molecular Dynamics with Shaped Femtosecond Pulses," Ph.D. thesis, Ludwig-Maximilians-Universität München (2001).
- [84] G. Chériaux, "Influences des distortions de phase sur le profil d'impulsions femtosecondes dans l'amplification à dérive de fréquence. Application à la génération d'impulsions de 30 TW à 10 Hz dans le saphir dopé au titane." Thèse, Orsay (1997).
- [85] H. Wang, Z. Zheng, D. E. Leaird, A. M. Weiner, T. A. Dorschner, J. J. Fijol, L. J. Friedman, H. Q. Nguyen, and L. A. Palmaccio, "20-fs pulse shaping with a 512-element phase-only liquid crystal modulator," *IEEE Journal of Selected Topics in Quantum Electronics* **7**(4), 718–27 (2001).
- [86] M. Danailov and I. Christov, "Time-space shaping of light pulses by Fourier optical processing," *Journal of modern optics* **36**(6), 725–731 (1989).
- [87] O. E. Martinez, "3000 times grating compressor with positive group velocity dispersion : application to fiber compensation in 1.3-1.6 μ m region," *IEEE Journal of Quantum Electronics* **QE-23**(1), 59–64 (1987).
- [88] L. Wang and A. M. Weiner, "Programmable spectral phase coding of an amplified spontaneous emission light source," *Optics Communications* **167**(1-6), 211–24 (1999).
- [89] C. Dorrer, "Caractérisation et modulation d'impulsions laser femtoseconde," Ph.D. thesis, Ecole Polytechnique (1999).
- [90] C. Dorrer, N. Belabas, J. P. Likforman, and M. Joffre, "Spectral resolution and sampling issues in Fourier-transform spectral interferometry," *Journal of the Optical Society of America B Optical Physics* **17**(10), 1795–802 (2000).
- [91] A. H. Zewail, *Femtochemistry : Ultrafast Dynamics of the Chemical Bond*, vol. I and II (World Scientific, Singapore, 1994).
- [92] A. Stolow, "Femtosecond time-resolved photoelectron spectroscopy of polyatomic molecules." *Annual Review of Physical Chemistry* **54**(1), 89–119 (2003).

- [93] J. Cao and K. R. Wilson, "A simple physical picture for quantum control of wave packet localization," *Journal of Chemical Physics* **107**(5), 1441–1450 (1997).
- [94] B. Kohler, V. V. Yakovlev, J. Che, J. L. Krause, M. Messina, K. R. Wilson, N. Schwentner, R. M. Whithell, and Y. Yan, "Quantum Control of Wave Packet Evolution with Tailored Femtosecond Pulses," *Physical Review Letters* **74**(17), 3360 (1995).
- [95] M. M. Wefers and K. A. Nelson, "Space-time profiles of shaped ultrafast optical waveforms," *IEEE Journal of Quantum Electronics* **32**(1), 161–72 (1996).
- [96] C. Dorrer and I. A. Walmsley, "Simple linear technique for the measurement of space-time coupling in ultrashort optical pulses," *Optics Letters* **27**(21), 1947–9 (2002).
- [97] J. Delagnes, "Contrôle de la propagation d'impulsions femtoseconde dans une vapeur atomique dense et optiquement résonante," Thèse, Univ. P. Sabatier (2005).
- [98] J. Delagnes, A. Monmayrant, P. Zahariev, B. Chatel, B. Girard, and M. A. Bouchene, "Compensation of resonant dispersion in atomic medium using a pulse shaper," *Optics Letters* (submitted).
- [99] R. L. Fork, O. E. Martinez, and J. P. Gordon, "Negative dispersion using pairs of prisms," *Optics Letters* **9**(5), 150–2 (1984).
- [100] E. Treacy, "Optical pulse compression with diffraction gratings," *I.E.E.E. Journal of Quantum Electronics* **QE-5**, 454–458 (1969).
- [101] I. A. Walmsley and L. Waxer, "Emission tomography for quantum state measurement in matter," *Journal of Physics B* **31**, 1825–1863 (1998).
- [102] M. Armstrong, P. Plachta, E. Ponomarev, and R. Miller, "Versatile 7-fs optical parametric pulse generation and compression by use of adaptive optics," *Optics Letters* **26**, 1152 (2001).
- [103] A. Baltuska, T. Fuji, and T. Kobayashi, "Visible pulse compression to 4 fs by optical parametric amplification and programmable dispersion control," *Optics Letters* **27**, 306 (2002).
- [104] G. Cerullo, M. Nisoli, S. Stagira, and S. De Silvestri, "Sub-8-fs pulses from an ultrabroadband optical parametric amplifier in the visible," *Optics Letters* **23**(16), 1283 (1998).
- [105] M. Zavelani-Rossi, G. Cerullo, S. De Silvestri, L. Gallmann, N. Matuschek, G. Steinmeyer, U. Keller, G. Angelow, V. Scheuer, and T. Tschudi, "Pulse compression over a 170-THz bandwidth in the visible by use of only chirped mirrors," *Optics Letters* **26**(15), 1155–1157 (2001).
- [106] S. Lochbrunner, T. Wilhelm, J. Piel, S. Sporlein, and E. Riedle, "Sub-20-fs tunable pulses in the visible and NIR by noncollinear optical parametric amplification (NOPA)," *Advanced Solid State Lasers. OSA Trend in Optics and Photonics Series* **26**, 366–71 (1999).
- [107] A. Shirakawa, I. Sakane, and T. Kobayashi, "Pulse-front-matched optical parametric amplification for sub-10-fs pulse generation tunable in the visible and near infrared," *Optics Letters* **23**(16), 1292–4 (1998).
- [108] P. Baum, S. Lochbrunner, L. Gallmann, G. Steinmeyer, U. Keller, and E. Riedle,

- “Real-time characterization and optimal phase control of tunable visible pulses with a flexible compressor,” *Applied Physics B* **74**(9), S219 (2002).
- [109] T. Oksenhendler, “Mesures et contrôles temporels dans le domaine des lasers ultra-brefs,” Ph.D. thesis, Ecole Polytechnique (2004).
- [110] D. Kaplan and P. Tournois, “Theory and performance of the acousto optic programmable dispersive filter used for femtosecond laser pulse shaping,” *Journal de Physique IV* **12**, Pr5–69 (2002).
- [111] V. Voloshinov, “Acousto-optic phenomena in crystals with strong anisotropy of optic and elastic properties,” in *World Congress on Ultrasonics* (Paris, 2003).
- [112] A. Monmayrant, M. Joffre, T. Oksenhendler, R. Herzog, D. Kaplan, and P. Tournois, “Time-domain interferometry for direct electric-field reconstruction by use of an acousto-optic programmable filter and a two-photon detector,” *Optics Letters* **28**(4), 278 (2003).
- [113] J. Seres, A. Müller, E. Seres, K. O’Keeffe, M. Lenner, R. F. Herzog, D. Kaplan, C. Spielmann, and F. Krausz, “Sub-10-fs, terawatt-scale Ti :sapphire laser system,” *Optics Letters* **28**(19), 1832–1834 (2003).
- [114] T. Tanabe, M. Yamanaka, T. Okamoto, and F. Kannari, “Compensation for a Transfer Function of a Regenerative Amplifier to Generate Accurately Shaped Ultrashort Pulses in Both the Amplitude and Phase,” *I.E.E.E. Journal of Quantum Electronics* **10**(1), 221–228 (2004).
- [115] Z. Bor, “Distortion of femtosecond laser pulses in lenses,” *Optics Letters* **14**(2), 119–121 (1988).
- [116] I. Walmsley, L. Waxer, and C. Dorrer, “The role of dispersion in ultrafast optics,” *Review of Scientific Instruments* **72**(1), 1–29 (2001).
- [117] D. Zeidler, T. Witte, D. Proch, and M. Motzkus, “Optical parametric amplification of a shaped white-light continuum,” *Optics Letters* **26**(23), 1921–3 (2001).
- [118] M. Shapiro and P. Brumer, “Laser control of product quantum state populations in unimolecular reactions,” *J. Chem. Phys.* **84**(7), 4103 (1986).
- [119] E. D. Potter, J. L. Herek, S. Pedersen, Q. Liu, and A. H. Zewail, “Femtosecond Laser Control of a Chemical Reaction,” *Nature* **355**(6355), 66–68 (1992).
- [120] M. Shapiro and P. Brumer, “Coherent control of molecular dynamics,” *Reports on Progress in Physics* **66**(6), 859–942 (2003).
- [121] T. Amand, V. Blanchet, B. Girard, and X. Marie, “Coherent control in atoms, molecules and solids,” in *Femtosecond laser pulses. Principles and Experiments*, C. Rullière, ed., pp. 333–394 (Springer, 2004).
- [122] S. Zamith, “Dynamique femtoseconde dans des atomes et molécules : précession de spin et dynamique de photoélectrons, transitoires cohérents, dynamique des états excités de l’acétylène,” Thèse d’université, Paul Sabatier (2001).
- [123] G. Grynberg, A. Aspect, and C. Fabre, *Introduction aux lasers et à l’optique quantique* (Ellipses, 1997).
- [124] D. Maas, C. W. Rella, P. Antoine, E. S. Toma, and L. Noordam, “Population transfer via adiabatic passage in the rubidium quantum ladder system,” *Physical Review A* **59**(2), 1374–81 (1999).

- [125] J. Pérez, *Optique géométrique et ondulatoire* (Masson, Paris, 1994).
- [126] J. Paye, “The Chronocyclic Representation of Ultrashort Light Pulses,” *I.E.E.E. Journal of Quantum Electronics* **28**(10), 2262–2273 (1992).
- [127] M. Morinaga, I. Bouchoule, J.-C. Karam, and C. Salomon, “Manipulation of Motional Quantum States of Neutral Atoms,” *Physical Review Letters* **83**(20), 4037–4040 (1999).
- [128] C. Leichtle, W. P. Schleich, I. S. Averbukh, and M. Shapiro, “Quantum State Holography,” *Physical Review Letters* **80**(7), 1418–21 (1998).
- [129] Y. Mairesse and F. Quere, “Frequency-resolved optical gating for complete reconstruction of attosecond bursts,” *Physical Review A (Atomic, Molecular, and Optical Physics)* **71**(1), 011,401 (2005).
- [130] E. Cormier, I. A. Walmsley, E. M. Kosik, A. S. Wyatt, L. Corner, and L. F. DiMauro, “Self-Referencing, Spectrally, or Spatially Encoded Spectral Interferometry for the Complete Characterization of Attosecond Electromagnetic Pulses,” *Physical Review Letters* **94**(3), 033,905 (2005).
- [131] E. Constant, D. Taranukhin, A. Stolow, and P. Corkum, “Methods for the measurement of the duration of high-harmonic pulses,” *Physical Review A* **56**(5), 3870–3878 (1997).
- [132] D. J. Kane and R. Trebino, “Characterization of Arbitrary Femtosecond Pulses Using Frequency-Resolved Optical Gating,” *I.E.E.E. Journal of Quantum Electronics* **29**(2), 571–579 (1993).
- [133] R. Trebino, K. W. DeLong, D. N. Fittinghoff, J. N. Sweetser, M. A. Krumbügel, and B. A. Richman, “Measuring ultrashort laser pulses in the time-frequency domain using frequency-resolved optical gating,” *Review of Scientific Instruments* **68**, 3277–3295 (1997).
- [134] P. O’Shea, M. Kimmel, X. Gu, and R. Trebino, “Highly simplified device for ultrashort-pulse measurement,” *Optics Letters* **26**(12), 932–934 (2001).
- [135] C. Ventalon, J. M. Fraser, and M. Joffre, “Time-domain interferometry for direct electric field reconstruction of mid-infrared femtosecond pulses,” *Optics Letters* **28**(19), 1826–8 (2003).
- [136] C. Dorrer, E. M. Kosik, and I. A. Walmsley, “Direct space-time characterization of the electric fields of ultrashort optical pulses,” *Optics Letters* **27**(7), 548–50 (2002).
- [137] P. Baum, S. Lochbrunner, and E. Riedle, “Zero-additional-phase SPIDER : full characterization of visible and sub-20-fs ultraviolet pulses,” *Optics Letters* **29**(2), 210–212 (2004).
- [138] O. Martinez, “Grating and prism compressors in the case of finite beam size,” *J Opt Soc Am B Opt Physics* **3**(7), 929–934 (1986).