



'SMART LASERS', Développement de sources laser intégrant un fonctionnement optimisé via un processus d'apprentissage

Jérémie Girardot

► To cite this version:

Jérémie Girardot. 'SMART LASERS', Développement de sources laser intégrant un fonctionnement optimisé via un processus d'apprentissage. Optique [physics.optics]. Université Bourgogne Franche-Comté, 2022. Français. NNT : 2022UBFCK057 . tel-03857280

HAL Id: tel-03857280

<https://theses.hal.science/tel-03857280>

Submitted on 17 Nov 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



**THÈSE DE DOCTORAT DE L'ÉTABLISSEMENT UNIVERSITÉ DE
BOURGOGNE FRANCHE-COMTÉ**

**PRÉPARÉE AU LABORATOIRE INTERDISCIPLINAIRE CARNOT
DE BOURGOGNE**

École doctorale n° 553
Carnot Pasteur

Doctorat de physique

Par

M. Girardot Jérémie

**“Smart lasers” : développement de sources lasers à
fibre intégrant un fonctionnement optimisé via un
processus d'apprentissage**

Thèse présentée et soutenue à Dijon, le 05/04/2022

Composition du Jury :

M. Vincent COUDERC	Directeur de recherche, Institut de recherche XLIM	Rapporteur
M. Marco ROMANELLI	Maître de conférences HDR, Université de Rennes 1	Rapporteur
Mme Giovanna TISSONI	Professeur, Institut de Physique de Nice - INPHYNI	Examineur
M. Jérôme HAUDEN	Team Manager, iXblue Photonics, besançon	Examineur
M. Philippe GRELU	Professeur, UBFC	Directeur de thèse
M. Edouard HERTZ	Maître de conférences HDR, UBFC	Co-directeur de thèse

Remerciements

Je tiens tout d'abord à remercier Monsieur Philippe Grelu et Monsieur Edouard Hertz qui m'ont encadré tout au long de cette thèse et qui m'ont fait partager leurs brillantes intuitions, ainsi que le temps qu'ils m'ont consacré.

Je remercie également Monsieur Aurélien Coillet, maître de conférences, et Monsieur Franck Billard, ingénieur de recherches tout deux à l'université de Bourgogne. Cette thèse est le fruit d'une collaboration de trois années avec eux. Leur bienveillance et leur disponibilité ont été des facteurs déterminants pour la réussite de ces travaux. Je remercie pour les mêmes raisons Monsieur Malik Nafa, tous ont été co-auteurs des articles que j'ai pu publier durant cette thèse.

J'adresse tous mes remerciements à Monsieur Vincent Couderc ainsi qu'à Monsieur Marco Romanelli pour l'honneur qu'ils m'ont fait en acceptant d'être rapporteurs de cette thèse.

J'exprime ma gratitude à Madame Giovanna Tissoni et à Monsieur Jérôme Hauden qui ont bien voulu être examinateurs.

Enfin, je souhaiterais remercier tous les doctorants, post doctorants et le personnel de l'université de Bourgogne, qui m'ont permis de travailler dans une ambiance sereine et agréable. Merci à vous tous.

Table des matières

1	Introduction	6
2	Dynamique impulsionnelle dans les lasers à fibre	12
2.1	Principes fondamentaux des lasers à fibre	13
2.1.1	Fonctionnement du processus laser	13
2.1.2	Composants essentiels d'optique fibrée	16
2.2	Le blocage de mode	19
2.2.1	La génération d'impulsions ultra-courtes	19
2.2.2	Le soliton dissipatif	25
2.2.3	Dynamiques impulsionnelles au-delà du blocage de mode conven- tionnel	28
2.3	Quel absorbant saturable utiliser ?	34
2.3.1	Choix des composants	34
2.3.2	Calcul déterminant l'orientation des composants	37
2.4	Les autres paramètres déterminant la propagation	39
3	Qu'est ce qu'un algorithme d'évolution ?	42
3.1	L'intelligence artificielle en photonique	43
3.2	Fonctionnement de l'algorithme d'évolution	45
3.3	Simulations du processus de l'algorithme	48
3.3.1	Détermination de la fonction de mérite composite	48
3.3.2	Etude de l'impact des conditions initiales de l'algorithme	49
3.4	Les principaux instruments de mesure	53
3.4.1	L'autocorrélateur	53
3.4.2	L'analyseur de spectre optique (OSA)	55
3.4.3	L'analyseur de spectre radiofréquence (ESA)	55
3.4.4	L'oscilloscope avec photodiode ultrarapide	55
3.5	Famille de fonction de mérite	57
3.6	Développement de l'algorithme pour la cavité	58
3.6.1	Boucle de rétroaction	58
3.6.2	Algorithme développé	62
4	Optimisation du blocage de mode fondamental d'une cavité simple via l'algorithme	66
4.1	Développement du montage laser	67
4.1.1	Détermination de la longueur de fibre dopée	67
4.1.2	Montage laser expérimental	69

4.2	Optimisation selon la largeur spectrale	70
4.2.1	Simulations de l'utilisation de l'algorithme implémenté sur la cavité	70
4.2.2	Fonction de mérite expérimentale	74
4.3	Cartographie des états de blocage de mode	76
4.4	Convergence vers un régime de blocage de mode fondamental à spectre large	77
4.4.1	Optimisation complète depuis une population aléatoire	77
4.4.2	Réoptimisation depuis un individu dégradé	79
4.4.3	Ajustement de la largeur spectrale utilisant des mesures DFT	80
4.4.4	Robustesse du laser et du procédé	82
5	Augmentation de la versatilité : génération de différentes dynamiques impulsionnelles	84
5.1	La ligne à dispersion nulle ou ligne 4f	85
5.1.1	Bref historique du façonnage d'impulsion	85
5.1.2	Principe d'une ligne à dispersion nulle	86
5.2	Propriété du modulateur spatial de lumière à double masque	87
5.3	Développement de la ligne 4f	88
5.3.1	Caractéristiques des composants	88
5.3.2	Orientation des réseaux	92
5.3.3	Mesure de la dispersion dans le plan de Fourier	93
5.4	Calculs de caractéristiques impulsionnelles	94
5.4.1	Calcul du TBP en limite de Fourier	95
5.4.2	Calcul de la constante d'autocorrélation K	96
5.5	Recompression d'impulsions extra-cavité via un algorithme d'évolution	96
5.5.1	Fonction de mérite utilisée	97
5.5.2	Montage développé	97
5.5.3	Processus de recompression extra cavité	98
5.5.4	Recompression avec ajout de fibre	100
6	Optimisations intégrant la modulation de phases spectrales par le SLM	104
6.1	Montage d'une nouvelle cavité laser	105
6.2	Optimisation sur la largeur spectrale	106
6.2.1	Paramètres de l'algorithme	106
6.2.2	Convergence vers un blocage de mode à largeur spectrale contrôlable	107
6.3	Optimisation sur la SHG	109
6.3.1	Paramètres de l'algorithme	109
6.3.2	Convergence vers un blocage de mode à impulsion brève	110
6.3.3	Ajustement de la durée d'impulsion	111
7	Optimisations intégrant la modulation de l'intensité spectrale par le SLM	114
7.1	Optimisation sur l'énergie de l'impulsion et ajustement de la largeur spectrale	115
7.1.1	Paramètres de l'algorithme	115
7.1.2	Convergence vers un régime de blocage de mode	116

7.1.3	Contrôle de la largeur spectrale et de l'énergie de l'impulsion . .	118
7.2	Génération de molécule de soliton à retard temporel pré-déterminé . . .	120
7.2.1	Paramètres de l'algorithme	120
7.2.2	Convergence vers un régime de molécule de soliton	123
7.2.3	Correction des surmodulations	127
7.2.4	Recompression des molécules générées	128
7.2.5	Convergence vers une molécule de soliton avec paramètres corrigés	128
7.2.6	Contrôle du retard temporel entre les solitons	131
8	Conclusion	134

Chapitre 1

Introduction

Historiquement, le développement du laser pose ses bases sur l'optique ondulatoire dans les années 1800, avec une modification du concept même de la lumière. En effet, en 1801, Thomas Young remet en question l'interprétation corpusculaire de la lumière. Grâce à sa célèbre expérience des fentes d'Young il met en évidence le comportement ondulatoire de celle-ci. La théorie ondulatoire de la lumière va ensuite être appuyée par les travaux de nombreux scientifiques, dont Augustin Fresnel, Max Planck ou encore Albert Einstein, menant à terme à la mesure de la vitesse de la lumière, puis à la quantification de l'énergie d'un corpuscule de lumière : le photon. Planck et Einstein recevront, en 1918 et 1921, le prix Nobel pour leurs recherches respectives. De plus, au début du XX^{ème} siècle, un événement clé a permis l'essor de la photonique : en 1917, Einstein découvre un troisième type d'interaction lumière-matière, l'émission stimulée, en plus de l'émission spontanée et de l'absorption. Ses travaux ont ouvert la voie à toute une série d'inventions. Ainsi, la seconde partie du XX^{ème} siècle fut une période de développement intense. Tout d'abord, l'invention du MASER (Microwave Amplification by Stimulated Emission of Radiation) en 1953 par Charles Townes à l'université de Columbia, fut la première grande découverte dans le domaine réunissant l'utilisation de ces trois interactions. Cette invention est le premier pas vers la création du LASER et créa pour la communauté scientifique un nouveau champ d'investigation [1]. Ce dispositif, tout d'abord basé sur l'utilisation de molécules d'ammoniac pour réaliser la transition entre deux états quantiques, a permis l'émission de faisceaux cohérents de micro-ondes. Townes réfléchit à la manière de transposer son appareil utilisant les micro-ondes à des fréquences plus grandes. Gordon Gould conçut théoriquement le laser comme milieu à gain dans une cavité optique résonnante dès 1957. Il fut aussi le premier à introduire l'acronyme "LASER". Mais en 1958, suite à une conversation avec Gould, Townes publia, avec son nouveau collaborateur Arthur Schawlow, un article qui présente une nouvelle idée pour réaliser l'inversion de population : le pompage optique [1]. Cette étape théorique fut la dernière amenant la réalisation expérimentale du LASER.

Le premier LASER (Light Amplification by Stimulated Emission of Radiation) fut construit par Theodore Maiman chez Hughes Research Laboratories à Malibu, en 1960 [2], soit trois ans après la note de Gordon Gould. Ce laser utilise du rubis synthétique comme milieu actif et émet un faisceau lumineux d'un rouge profond dont la longueur d'onde est de 694,3 nm. La première application du laser au rubis concernait des télémètres militaires. L'invention a rapidement montré son intérêt. En 1963, le laser au dioxyde de carbone (CO₂) est mis au point par Kumar Patel [3]. Ce laser est beaucoup moins onéreux et beaucoup plus efficace que le laser au rubis. Puis, l'expansion des lasers se poursuit avec le développement en 1966 des premiers lasers à colorant. Un gros avantage de ces lasers est le contrôle de la longueur d'onde d'émission en variant la concentration du colorant. Le premier laser à blocage de modes passif a d'ailleurs été réalisé par De Maria en 1966 [4], utilisant un absorbant saturable liquide à colorant. Naturellement, ces lasers ont été une révolution dans le domaine de la spectroscopie. Dans le même temps, en 1963, Elias Snitzer développe le premier laser à fibre [5], mais le milieu à gain reste alors un barreau solide de verre de baryum dopé aux ions Nd^{3+} large de quelques dizaines de μm , induisant la propagation de nombreux modes transverses. Dix ans plus tard, un procédé de fabrication de fibres monomode dopées aux terres rares (Erbium, Thulium, Ytterbium...), et à faibles pertes est créé par l'équipe de Poole à l'université de Southampton [6]. Ce procédé est à la base de la création des

premiers laser à fibre commerciaux, qui ont vu le jour au début des années 1990 après le développement de laser à fibre dopée Erbium (Reekie et Mears en 1986 [7,8]). La nature du dopage détermine la longueur d'onde d'émission. Une longueur d'onde de 1550 nm correspond au minimum d'atténuation dans une fibre en silice, expliquant la raison de son utilisation dans le domaine des télécommunications. Ces avancées s'accompagnent du développement de nombreux composants d'optique intégrée comme les coupleurs [9], les multiplexeurs [10] ou encore les séparateurs de polarisation [11]. Ainsi, des lasers plus puissants, plus compacts, et avec une meilleure qualité de faisceau ont pu être fabriqués. En outre, l'échauffement thermique, qui peut être un facteur déterminant pour les lasers standards, est moins critique pour les lasers à fibres où la dissipation de la chaleur est favorisée par le faible diamètre des fibres. Pour terminer, le faisceau de sortie est collimaté, et très maniable. Les lasers à fibres ont donc connu un essor important, mais ne commencent à être utilisés couramment dans l'industrie qu'à partir des années 2010. Aujourd'hui, les lasers sont utilisés de manière courante dans l'industrie : le perçage laser [12], la découpe [13], le soudage [14,15], le marquage [16,17] mais également dans la bio imagerie (microscopie de fluorescence par absorption à deux photons [18]), dans les télécommunications, la métrologie ou encore l'armée (nous citerons le radar optique LIDAR à titre d'exemple [19,20]).

L'avènement du laser, et particulièrement des lasers impulsionnels, est à l'origine de nombreux travaux en régime non linéaire. En effet, les intensités lumineuses générées grâce aux lasers ont permis de produire et étudier divers phénomènes non linéaires. Inversement, la recherche sur la génération de phénomènes non linéaires spécifiques, comme les supercontinua [21,22], nécessitant des énergies très importantes, ont contribué au développement de nouvelles sources laser. Les déphasages non linéaires varient selon le produit de l'intensité et du coefficient non linéaire du milieu [23]. La méthode la plus efficace pour générer ce type d'interaction est l'utilisation de régimes laser impulsionnels, dont la puissance crête est très élevée. Le développement des lasers à fibres a donc suivi cette tendance. Ainsi, en 1991, Duling développe le premier laser à fibre à blocage de mode [24], utilisant un NALM (Boucle Amplificatrice Nonlinéaire) [25]. Ce laser délivrait des impulsions d'une durée de 3.3 ps pour une largeur spectrale de 1 nm autour de 1530 nm, ce qui constituait à cette époque une impulsion relativement brève. La même année, Zirngibl et al [26] démontrent l'utilisation d'un matériau InGaAs/GaAs comme absorbant saturable, permettant la génération d'impulsions de l'ordre de la picoseconde. La génération de tels régimes, que nous considérerons comme impulsionnels ultracourts, repose sur l'utilisation d'absorbants saturables en cavité laser pour induire un mécanisme de blocage de modes [27,28]. Un absorbant saturable est un élément dont la transmission dépend de l'intensité incidente. Il va favoriser les hautes intensités et discriminer les basses, permettant après plusieurs tours de cavité la génération d'impulsions d'une durée de la dizaine de femtoseconde à quelques picosecondes, et dont l'intensité crête est considérable. Ce type de régime est intéressant dans de nombreux domaines, notamment la médecine ou le traitement de surface, pour l'absence d'effets thermiques qu'il provoque sur le substrat. Il constitue également une base pour des études fondamentales des effets non linéaires, la génération de trains d'impulsions ou encore de supercontinua.

Aujourd'hui, la diversité des lasers à fibre laisse présager pour ces sources un avenir radieux dans les domaines industriels et pour la recherche fondamentale. Cependant, certaines limitations restent à surmonter. Dans le domaine industriel, les lasers à blocage de mode proposés à l'heure actuelle sont auto-démarrants : le régime impulsif est généré dès l'allumage du dispositif, sans ajustement de la part de l'utilisateur. Ces lasers sont donc faciles d'utilisation, et permettent à tout utilisateur non initié à la manipulation de cavité laser d'utiliser ce type d'impulsions. La principale critique formulée sur ces lasers est leur manque de versatilité : une fois le régime impulsif établi, il est impossible de changer ses caractéristiques, et encore moins générer d'autres types de régime à blocage de modes sur mesure comme par exemple des régimes multi-impulsifs. Cette limitation en versatilité est principalement due à un manque de degré de liberté accessibles et interfacés au niveau de la cavité laser et se retrouve dans la plupart des lasers développés dans le monde académique [29, 30, 32, 58]. Par ailleurs, les dynamiques non-linéaires complexes et pas encore totalement comprises dans les laser à fibres à blocage de mode rendent encore aujourd'hui impossible la création de relation analytique entre les paramètres d'un système laser et les caractéristiques du régime généré. Les régimes plus complexes ne sont plus forcément auto-démarrants, et pour générer un régime de blocage de mode spécifique, un réglage manuel par tâtonnement, souvent long et minutieux, et exigeant certaines compétences expérimentales doit alors être effectué. Pour un utilisateur non chevronné, cette tâche est ardue. Ainsi est née dans l'équipe l'idée de développer une source laser à fibre à la fois autonome, adaptable, compacte et versatile. Pour créer toutes ces propriétés en même temps, une solution semble se dégager : l'utilisation d'intelligence artificielle, implémentées sur des cavité laser fibrée usuellement développées dans le domaine académique. Cette intelligence doit être capable de piloter les paramètres accessibles des montages lasers, et de les contrôler pour générer le régime désiré par l'utilisateur. Des algorithmes d'évolution seront choisis dans ces travaux, pour leur capacité d'adaptabilité. L'algorithme est utilisé pour déterminer la combinaison de paramètres à appliquer sur le montage laser pour générer un régime désiré optimal. Ce travail de thèse se rapporte au développement d'un "smart laser" et s'inscrit dans la continuité des travaux d'Ugo Andral débutés en 2013 [33] et considérés comme pionniers dans ce domaine.

La génération de régimes variés implique le développement de cavités versatiles, avec de nombreux paramètres contrôlables. Une façon courante de gagner en versatilité est l'utilisation d'absorbants saturables virtuels à fonction de transfert ajustable. Le fonctionnement de ce type d'absorbant saturable est basé sur des interférences non linéaires et n'induit généralement que de faibles pertes, contrairement aux absorbant saturables réels basés sur l'absorption matérielle. Dans la famille des absorbants saturables virtuels, le processus offrant le plus de versatilité reste la rotation (ou évolution) non linéaire de la polarisation. Tamura présenta en premier l'utilisation du phénomène de rotation non linéaire de polarisation pour créer l'effet de blocage de mode dans un laser fibré [34] en 1992. Cette technique utilise un milieu Kerr, comme une fibre optique dont l'indice de réfraction non linéaire permet de changer l'orientation de la polarisation d'un régime en fonction son intensité [35]. Un dispositif de discrimination selon la polarisation (un polariseur) permet alors de sélectionner les hautes puissance et en discriminer les basses, créant le mécanisme d'absorbant saturable. L'utilisation de contrôleurs de polarisation permet alors d'ajuster la polarisation transmise par le

polariseur, et donc la fonction de transfert en puissance [36]. Si les contrôleurs de polarisation peuvent être contrôlés et interfacés par informatique, il devient possible de les piloter par un algorithme d'évolution pour générer un régime optimal.

Ce manuscrit réalise l'auto génération par algorithme d'évolution de régimes lasers plus stables, plus complexes, dans des durées d'optimisation plus rapides et avec un contrôle plus étendu de leurs caractéristiques que dans les travaux préalables. Il s'articule en deux parties : la première partie concerne la génération d'impulsions uniques stables, robustes et à largeur spectrale pré-déterminée. Ce premier travail s'appuie sur l'implémentation d'algorithme sur une cavité laser fibrée simple et courte à dispersion normale. Le processus s'appuie sur l'utilisation d'une technique d'imagerie spectrale en temps réel par oscilloscope, ce qui permet en plus l'utilisation d'un seul instrument de mesure pour obtenir des données spectrales et temporelles. L'algorithme a alors un double objectif, en optimisant des modes de propagation selon deux critères : la stabilité des impulsions et leur largeur spectrale. La seconde partie de manuscrit présente les travaux sur la génération de régimes lasers plus complexes. Ce travail repose sur la complexification de la cavité laser, avec à la fois l'utilisation de deux amplificateurs et l'augmentation du nombre de paramètres accessibles à contrôler par l'algorithme en incluant en cavité un façonneur d'impulsion, ou "pulse shaper". Une modulation de la phase et/ou de l'intensité spectrale de chaque composante permet de modifier finement le régime impulsionnel obtenu. Il a été de fait possible de réaliser entre autre la pré-détermination de 3 caractéristiques impulsionnelles, mais la section clé de cette partie est sans aucun doute la génération de molécules de solitons à deux impulsions, avec retard temporel contrôlable. Ce dernier travail ouvre la perspective d'auto-générer d'autres régimes complexes, repoussant par étape les limites d'utilisation de l'algorithme d'évolution en photonique.

Ces travaux de thèse ont requis deux domaines de compétence : des compétences purement photoniques pour la construction de cavités laser avec la manipulation de fibres, le soudage, les techniques de coupures successives, l'utilisation de divers instruments de mesures mais également l'optique en espace libre. La maîtrise des concepts théoriques sur la propagation de champs optique dans des fibres a dû également être maîtrisée. Le second domaine de compétence est celui de la programmation, avec l'adaptation de l'algorithme à des utilisations différentes, l'interfaçage de montages expérimentaux, le développement de programmes informatiques pour le traitement de données.

Chapitre 2

Dynamique impulsionnelle dans les lasers à fibre

2.1 Principes fondamentaux des lasers à fibre

2.1.1 Fonctionnement du processus laser

Le phénomène laser repose sur les trois types de stimulation des électrons. Prenons comme exemple un système à deux niveaux d'énergies, comme illustré en figure 2.1. Le niveau de plus faible énergie, qui est également le plus stable est appelé niveau "fondamental". S'il interagit avec un photon de longueur d'onde adéquat, il peut passer dans un état excité, et augmente donc son niveau d'énergie. On parle alors de phénomène d'absorption. Une fois sur un niveau excité, l'électron va rapidement retourner vers un état moins énergétique. En se désexcitant, l'électron va émettre un photon d'énergie égale à la différence des énergies des niveaux : c'est le processus d'émission. S'il intervient sans l'action d'un champ incident, on parle d'émission spontanée. En revanche, si il est déclenché par l'effet d'un champ optique, on parle d'émission stimulée. Dans ce dernier cas, les caractéristiques (longueur d'onde et phase) du photon émis sont les mêmes que celui qui le stimule, ce qui conduit à une amplification du champ incident.

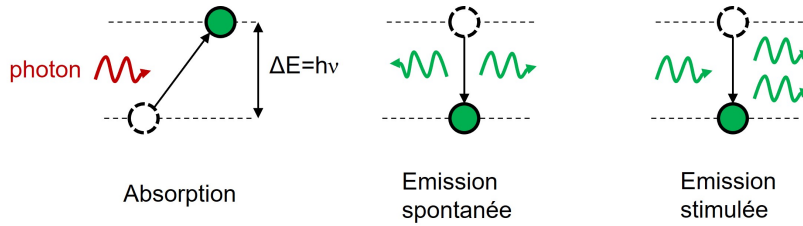


FIGURE 2.1: Illustration des trois interaction lumière matière nécessaire au rayonnement laser

Sans apport d'énergie, la population des niveaux est répartie selon une loi de Boltzmann [39] : $N_2 = N_1 e^{-\frac{E_2 - E_1}{k_b T}}$, avec N_1 et N_2 respectivement la population des niveaux 1 et 2, $E_2 - E_1$ la différence d'énergie entre les niveaux, T la température et k_b la constante de Boltzmann. Dans ce cas, la population est en large majorité répartie sur le niveau d'énergie 1, c'est à dire le fondamental. La réponse du milieu à gain lors de l'interaction avec un photon d'énergie égale à ΔE sera alors une absorption, et non une émission stimulée. Il devient donc impossible d'avoir une amplification de l'émission¹. Pour l'obtenir, il faut donc que la population du niveau excité 2 soit supérieure en nombre à celle du niveau fondamental 1. On parle alors **d'inversion de population**. Ce critère est nécessaire à la génération d'une émission laser, et requiert un apport d'énergie extérieur.

Avec un apport d'énergie par un champ appelé pompe, la population des niveaux N_1 et N_2 suivent l'évolution suivante : $\frac{dN_1}{dt} = Bu(\nu)(N_2 - N_1) + AN_2$ et $\frac{dN_2}{dt} =$

1. dans le cadre de l'approximation des équations de taux

$-Bu(\nu)(N_2 - N_1) - AN_2$, avec A et B respectivement les probabilités d'émission spontanée et d'émission stimulée et d'absorption, et $u(\nu)$ l'intensité du champ de pompe. En régime stationnaire, la dépendance entre N_2 et N_1 s'écrit alors : $N_2 = \frac{Bu(\nu)}{A + Bu(\nu)N_1}$. La probabilité d'absorption est sensiblement la même que la probabilité d'émission stimulée pour un niveau donné. Avec un système à deux niveaux, on aura alors tour à tour des excitations et désexcitation des électrons, sans réaliser d'inversion de population. Au maximum, avec un fort pompage, on aura $N_1 = N_2$. Pour que cette inversion intervienne et donc générer le processus laser, trois niveaux d'énergies minimum sont nécessaires (comme c'est le cas pour le laser à rubis par exemple [2]), dont un fondamental, un excité et un métastable (facilitant l'inversion de population). Dans la plupart des cas, 4 niveaux sont utilisés, avec un niveau métastable pour créer l'inversion de population. Ce processus à quatre niveaux est résumé et illustré en figure 2.2 (a). A noter qu'à température ambiante, les ions Erbium $3+$, utilisés comme milieu à gain dans tous les travaux présentés, représentent un système à quasi 3 niveaux. Comme présenté en figure 2.2 (b), le niveau fondamental de l'Erbium est le $^4I_{15/2}$, le niveau excité grâce à un pompage à 980 nm est le $^4I_{11/2}$ et le niveau métastable est le $^4I_{13/2}$. La largeur du niveau métastable permet la génération d'une bande spectrale large d'environ 30 nm et donc de nombreux modes longitudinaux, permettant le blocage de modes. Ce phénomène sera expliqué en section suivante.

Ces processus physiques, qui permettent à eux seuls de créer une émission laser, ont lieu dans un milieu optiquement actif appelé **milieu à gain**. L'énergie nécessaire pour créer l'absorption est apportée à ce milieu par un dispositif de **pompage**. Une fois pompé, ce milieu se désexcitera par émission spontanée, générant un champ de faible amplitude. Pour amplifier ce champ par émission stimulée, un résonnateur est nécessaire pour le rediriger vers le milieu à gain avec lequel il va de nouveau interagir. Ce résonnateur est créé traditionnellement par deux miroirs, et est appelé **cavité laser**. Enfin, un prélèvement d'une partie de l'émission laser est réalisé pour l'utilisation, créant des **pertes utiles**. Pour résumer, le processus laser requiert 3 ingrédients clés [37] : un milieu à gain, une cavité et un dispositif de pompage. Ce type de procédé permet de générer un rayon lumineux cohérent spatialement et temporellement. En plus de la condition d'inversion de population, deux conditions de fonctionnement sont implicites. Premièrement, le gain total de la cavité (apporté par le milieu à gain) doit être supérieur à toutes les pertes de celle-ci (y compris les pertes utiles). Deuxièmement, les ondes lumineuses ne doivent pas interférer de manière destructive entre elles dans la cavité : il faut donc que les modes longitudinaux autorisés par la longueur de cavité soient à l'intérieur de la courbe de gain.

L'idée d'utiliser des fibres dopées comme milieu à gain et de construire des cavités laser fibrés est apparue à la fin des années 1980. Aujourd'hui, les lasers à fibre sont utilisés dans de nombreux domaines industriels comme le traitement de surface, le forage ou le soudage [15]. Ces derniers ont de nombreux avantages :

- La lumière étant guidée par fibre optique, le faisceau de sortie est facilement dirigeable et manipulable.
- La fibre pouvant être enroulée, l'augmentation de la taille de la cavité est possible sans rendre le laser plus volumineux.

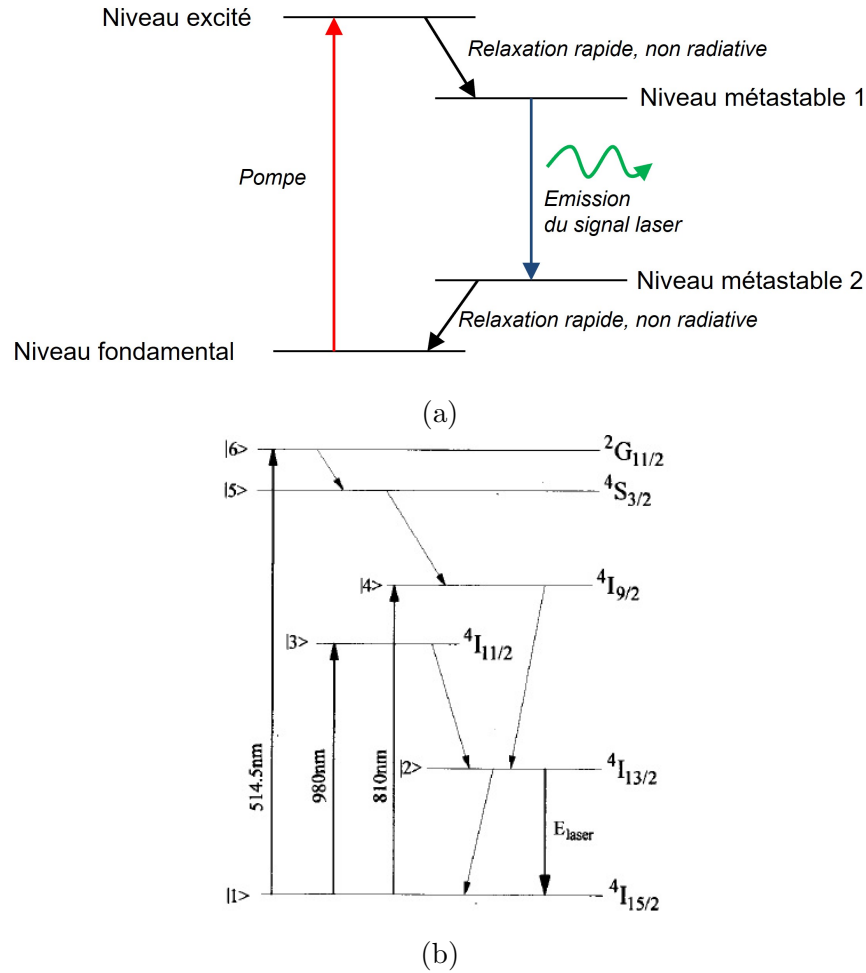


FIGURE 2.2: (a) Illustration du processus laser d'après un système à 4 niveaux, et (b) Niveaux d'énergies quantiques de l'Erbium, tirés de [40]

- La dissipation de la chaleur est rendue aisée par le faible diamètre de fibres (environ $100\mu\text{m}$ pour une fibre passive, comparé à plusieurs centimètres pour des barreaux classiques).
- La qualité du faisceau est excellente ($M^2 \approx 1$).

Ils présentent donc des avantages évidents en terme de coûts, compacité, flexibilité et simplicité d'utilisation. Cependant, il est bien connu que les processus non linéaires bornent les cavités laser à fibre impulsives à une énergie modeste comparés aux lasers plus conventionnels [38], ce qui en constitue la principale limitation. Ce type de cavités lasers peuvent se construire selon de nombreux types d'architecture (figures 8, figures 9, cavités σ ...). Les deux architectures les plus simples sont les cavités linéaires, comme les cavités lasers plus classiques et les cavités en anneaux, schématisées en figure 2.3.

— Les cavités linéaires [41]

Elles sont composées d'une fibre dopée avec généralement deux semi miroirs accolés aux extrémités, qui sont le plus souvent des réseaux de Bragg. Le premier miroir, en plus de créer le résonateur pour la cavité, permet de laisser passer la pompe dans un sens, et le second permet la création des pertes utiles du laser,

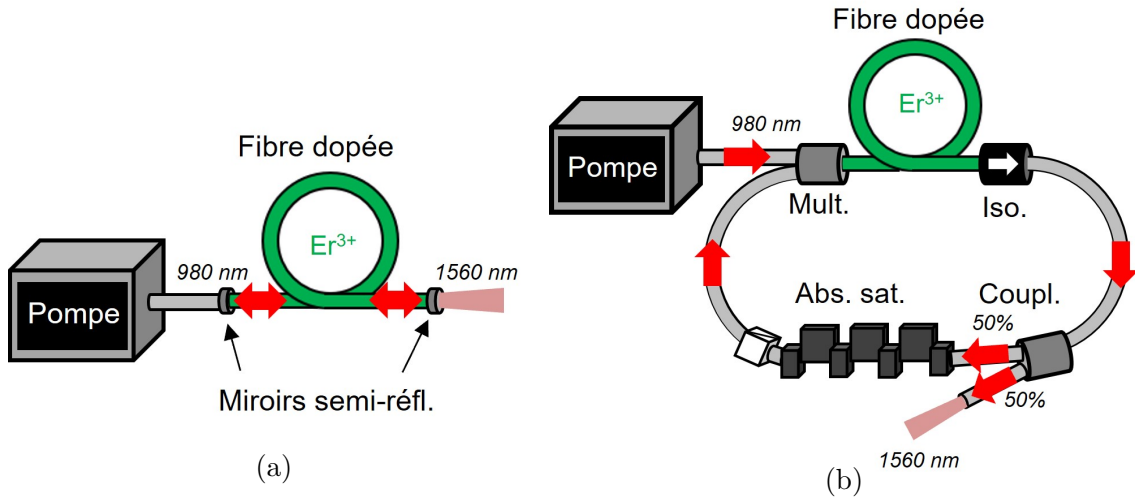


FIGURE 2.3: Schéma d'architecture laser à fibre types linéaires (a), et en anneaux (b)

comme nous pouvons le voir dans la figure 2.3 (a). Ces cavités, peu propices au développement et peu pratiques d'utilisation, sont rarement utilisées pour la génération d'impulsions ultracourtes.

— Les cavités en anneaux [42]

Ce type de cavité, illustré en figure 2.3 (b), se construit sur un circuit dans lequel le champ laser se propage de manière unidirectionnelle, passant donc à chaque tour dans chaque composant. L'avantage de ce type de cavité réside dans la possibilité d'insertion de nouveaux composants, d'où leur utilisation très commune. A titre d'exemple, en figure 2.3 (b), la cavité présentée requiert l'ajout de multiplexeurs (Mult.), d'isolateur (Iso.) ou encore de coupleur (Coupl.), mais permet l'addition d'absorbants saturable (Abs. Sat.) permettant la création de régimes de propagation différents et variés comme le régime à blocage de mode. Nous détaillerons dans les sections suivantes le panel de régimes accessibles. Dans la prochaine section, une liste des composants d'optique fibrée les plus usuels sera présentée.

2.1.2 Composants essentiels d'optique fibrée

Les cavités laser à fibre sont principalement composées de 6 éléments énoncés ci-dessous. Lors du rapprochement des deux fibres, on peut réaliser un couplage par champ évanescent entre des champs circulants dans chaque fibre, permettant de faire passer une partie du champ d'une première fibre vers une seconde. La nature du couplage diffère selon certains critères. Ce couplage nous permet la création des trois éléments suivants clés d'optique fibrée :

— Le démultiplexeur (WDM)

Développés par Digonnet et Shaw en 1983 [10], il permet l'introduction de la pompe dans la cavité en introduisant deux longueurs d'ondes différentes dans une seule fibre (figure 2.4 (a)). On réalise le couplage entre deux fibres sur une certaine longueur L . La longueur de couplage détermine le rapport d'intensité entre les longueurs d'ondes de chaque ports de sortie. Habituellement, on utilise

un port d'entrée qui mélange deux longueurs d'ondes (celle de la pompe et celle du champ laser) et deux ports de sortie qui contiennent chacun une des longueurs d'onde. Ainsi, nous séparons le faisceau selon ses fréquences. Si cet élément est utilisé dans le sens inverse, on parlera plutôt de multiplexeur.

— Le coupleur

Il permet la création de pertes utiles en séparant le champ selon sa puissance (figure 2.4(b)). Dans ce cas, c'est la distance z entre les deux fibres qui détermine le rapport de distribution du champ entre les deux ports de sortie. Le champ est dans ce cas séparé selon sa puissance [9]. En pratique, on peut fabriquer le ratio que l'on veut (50/50, 90/10 par exemple).

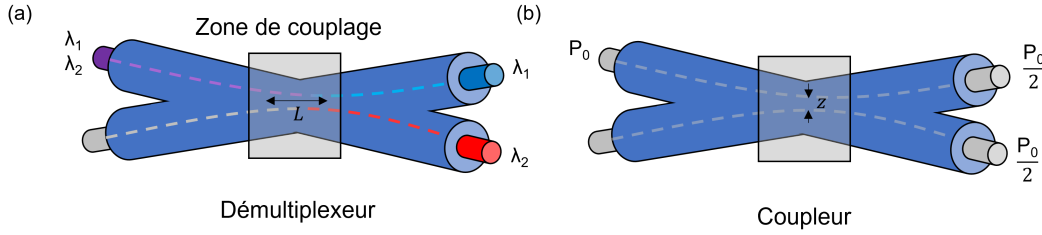


FIGURE 2.4: Schéma de deux éléments clés d'optique intégrée utilisant le couplage d'ondes évanescentes : le coupleur et le démultiplexeur

— Le séparateur de polarisation (Polarization Beam Splitter ou PBS)

Basé sur les mêmes effets de couplage, il permet de séparer un champ en deux en fonction de sa polarisation, typiquement en deux polarisations linéaires verticales et horizontales [11]. Les deux ports de sortie contiennent alors deux états de polarisation orthogonaux. On l'utilise généralement comme polariseur, en créant des pertes sur certaines composantes de polarisation pour créer les phénomènes d'absorbant saturable par rotation non linéaire de polarisation.

— Le contrôleur de polarisation

Il permet le contrôle de l'état de polarisation du champ. Ce contrôle se fait grâce à une biréfringence créée par contrainte mécanique (torsions, pincement, enroulement) sur la fibre. Cet effet a été démontré par Lefevre en 1980 [43], d'où le surnom boucle de Lefevre, souvent utilisé pour dénommer les contrôleurs de polarisation. Premièrement, la fibre est enroulée en 3 boucles sur un faible rayon, créant ainsi une première biréfringence constante (la plupart du temps on aura des déphasages approchés de $\pi/2$, π et $\pi/2$). Puis chaque boucle est montée sur un support permettant la rotation de la boucle par rapport à un axe fixe. Cette torsion permet de réajuster l'axe de la polarisation, créant un dispositif similaire à trois lames biréfringentes dont les axes neutres peuvent être orientés et permettant d'obtenir n'importe quel état de polarisation à partir de n'importe quel état. Dans la figure 2.5, les torsions sont représentées par les flèches noir du haut, tandis que le rayon de courbure de la fibre est représenté par la valeur R et les deux vis représentent la fixation de la fibre.

— L'isolateur

Cet élément, qui permet de fixer un seul sens de propagation au faisceau en créant de fortes pertes sur celui se propageant dans le sens inverse, est illustré en figure

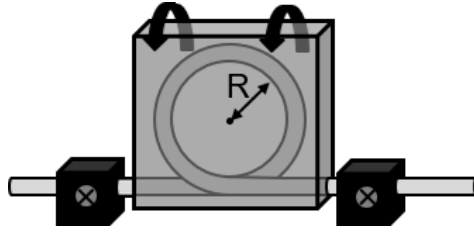


FIGURE 2.5: Schéma d'un contrôleur de polarisation

2.6. Les isolateurs optiques sont basés sur l'effet Faraday non réciproque [44], en utilisant un rotateur de Faraday (Rot. Far.). En pratique, dans la plupart des cas, des isolateurs insensibles à la polarisation sont utilisés [45]. Dans ce cas, deux cristaux biréfringents sont utilisés, séparés par un rotateur de Faraday. La lumière voyageant dans la direction voulue est divisée par le cristal biréfringent d'entrée selon ses composantes de polarisation verticale (0°) et horizontale (90°), appelées respectivement rayon ordinaire et rayon extraordinaire. Le rotateur de Faraday fait tourner le rayon ordinaire et le rayon extraordinaire de 45° . Cela signifie que le rayon ordinaire est maintenant à 45° , et le rayon extraordinaire à -45° . Le cristal biréfringent de sortie recombine ensuite les deux composantes. La lumière se déplaçant dans le sens inverse est séparée en un rayon ordinaire à 45° et un rayon extraordinaire à -45° par le cristal biréfringent. Le rotateur de Faraday fait à nouveau tourner les deux rayons de 45° . Maintenant, le rayon ordinaire est à 90° , et le rayon extraordinaire est à 0° . Au lieu d'être focalisés par le deuxième coin biréfringent, les rayons divergent.

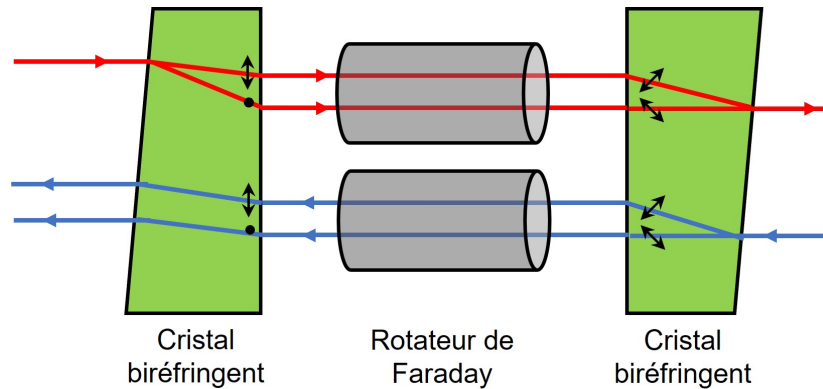


FIGURE 2.6: Schéma du fonctionnement d'un isolateur optique indépendant de la polarisation

— Le circulateur

Cet élément comporte 3 ports dans lequel chaque faisceau entrant par un port ressort par le suivant [46]. Ainsi, le faisceau entre dans le dispositif par le port n°1, ressort par le n° 2, puis rentre à nouveau par le n°2 pour ressortir par le port n°3. Ce composant est particulièrement utilisé dans les télécommunications optiques, et intervient pour mettre en oeuvre des architectures impliquant de séparer l'aller et le retour du faisceau, comme par exemple l'utilisation d'un miroir non linéaire SESAM dans un laser en anneau, ou bien un filtre de Bragg dans une ligne télécom.

Dans ce dernier cas, on utilise un élément pour renvoyer le champ dans le port n°2, comme un réseau de Bragg ou simplement une fibre repliée sur elle même. Son fonctionnement est basé sur les mêmes effets que pour l'isolateur, avec un rotateur de Faraday et deux cristaux biréfringents [47]. En revanche, les pertes créées par le cristal 2 (en admettant que le port 2 soit après celui-ci) sont récupérées pour créer le port 3. L'illustration d'un circulateur optique est donnée en figure 2.7.

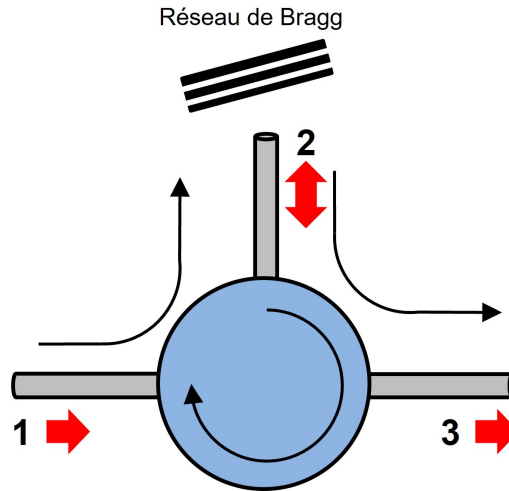


FIGURE 2.7: Schéma d'un circulateur optique

— La fibre monomode SMF 28

Cet élément n'est pas un composant en soit mais est très utilisé pour relier les composants entre eux, et présente de très faibles pertes ($< 0.2 \text{ dB.km}^{-1}$) pour des signaux de la bande C des télécommunications. Elles sont également utilisées pour construire les boucles de Lefevre. Ce genre de fibre ne permet la propagation que d'un seul mode spatial, avec un diamètre de coeur autour du diamètre du mode spatial LP01. Pour un champ se propageant à 1550 nm, sa dispersion est anormale ($\beta_{2,SMF} = -0.0228 \text{ ps}^2.\text{m}^{-1}$). Son utilisation mène donc pour la plupart des cas à des cavités de dispersion anormale. Pour compenser cette dispersion, une fibre à compensation de dispersion (DCF) peut être utilisée. Dans les DCFs, la réduction de la dispersion vers le régime normal se fait par contrainte géométrique, en réduisant le diamètre de coeur. Les pertes d'une DCF sont donc plus élevées que pour une SMF. Comme nous le verrons plus tard, une seconde technique pour compenser la dispersion est d'utiliser une fibre dopée à coeur fin, avec une dispersion normale directement.

2.2 Le blocage de mode

2.2.1 La génération d'impulsions ultra-courtes

Depuis l'avènement du laser à fibre, la diversification des sources a permis la génération de divers régimes de propagation de la lumière et notamment la génération d'impulsions ultracourtes. Les premiers travaux traitant de régimes impulsionsnels générés par cavités laser fibrées autour de 1550 nm sont publiés en 1986 par Mears [7, 8], mais ces

régimes étaient de simples régimes Q-switched à impulsions de durée de 100 ns. Les laser fibrés permettent la génération aisée d'impulsions dites ultracourtes, c'est à dire dont la durée est inférieure à quelques picosecondes. Ce type de régime furent générés pour la première fois quelques années plus tard, avec des impulsions pouvant atteindre une durée de quelques picosecondes à quelques dizaines de femtosecondes [80]. Ces travaux, réalisés par Tamura en 1993, ont également permis un bond conceptuel de la gestion de la dispersion chromatique en cavité laser fibrée, en utilisant une fibre dopée à dispersion normale. La recherche de tels régimes a été motivée par la possibilité de produire et d'étudier des processus très spécifiques. En effet, la lumière, en étant absorbée par la matière, peut déclencher différents types d'interactions avec celle-ci, notamment des interactions photo-chimiques (changement de composition électronique de la matière, isomérisation...), des interactions mécaniques (cassures de liaisons atomiques) ou encore des interactions thermiques (augmentation de la température). Dans de nombreux domaines d'application, les effets thermiques sont indésirables : nous citerons à titre d'exemple la médecine, où lors d'ablations par laser, la préservation des tissus environnants est aussi importante que l'ablation elle même. Pour limiter ces effets, des impulsions à fortes puissances crêtes et très brèves sont nécessaires. Dans le domaine industriel, les impulsions ultracourtes sont très utilisées pour le micro-usinage ou le marquage laser. La durée n'est pas le seul critère à considérer, le taux de répétition des impulsions est aussi un facteur décisif. La génération d'impulsions très courtes, concentrant la puissance dans une très faible durée, permet d'atteindre de telles puissances crêtes. Dans le domaine de la recherche fondamentale, ces impulsions sont très utiles dans l'étude de phénomènes non linéaires et la production de super continua. La spectroscopie est un domaine dans lequel les lasers femtosecondes ont prouvé leur utilité, comme pour les travaux de A. Zewail portant sur les états de transition d'une réaction chimique à l'aide de la spectroscopie à la femtoseconde [48, 49]. Dans notre cas, l'étude des dynamiques impulsionnelles en elle même constitue notre terrain d'investigation.

Le régime d'impulsions ultracourtes en cavité laser est généré par un processus nommé blocage de modes [50]. En effet, la largeur spectrale du gain de la fibre dopée utilisée permet la propagation de nombreux modes longitudinaux dans la cavité. La différence spectrale entre deux modes successifs se propageant dans une cavité est appelée intervalle spectral libre (ISL), et est reliée à la longueur de la cavité et l'indice de réfraction n par $\delta\nu = c/2nL$. Une illustration de l'intervalle spectral libre est donné en figure 2.8. Par exemple, prenons le cas d'une fibre dopée Erbium Er^{3+} (EDF) dans une cavité longue de 5 m entièrement fibrée. L'émission de l'Erbium est autour de 1560 nm, avec une largeur de gain non saturé d'environ 25 nm, soit une bande de gain 3.1 THz. Dans cette bande, tous les modes n'ont pas une intensité suffisante pour être résolus spectralement, ce qui incite à travailler avec des valeurs à mi hauteur de la bande de gain $\Delta\nu$. L'intervalle spectral libre de ce laser est de $\delta\nu = 21$ MHz, ce qui nous donne une propagation de $N = 1,5 \cdot 10^5$ modes longitudinaux. En limite de Fourier, la durée d'impulsion serait de 64 fs. Ces modes longitudinaux, caractérisés par leur pulsation, leur amplitude et leur phase, donnent lieu à des régimes continus bruités chaotiques lorsqu'ils interfèrent entre eux sans aucune synchronisation. En revanche, si le déphasage est constant pour tous les modes longitudinaux, ces modes interfèrent entre eux de manière constructive donnant lieu à un pic d'intensité plus important, comme illustré en figure 2.9. Cette figure présente la somme des champs de 50 modes

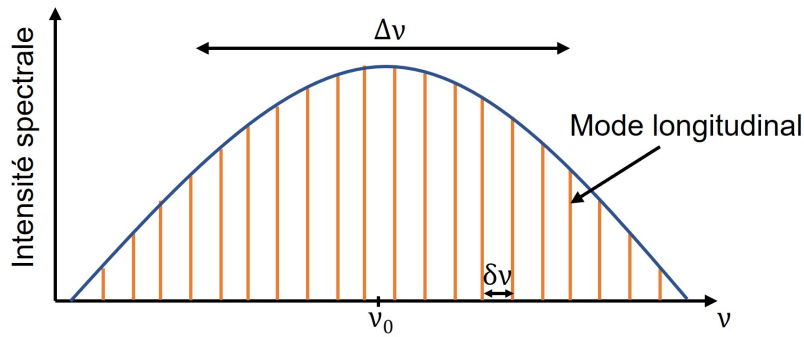


FIGURE 2.8: Schéma présentant la propagation des modes longitudinaux pouvant se propager dans une bande de gain à mi hauteur $\Delta\nu$ avec un intervalle spectral libre de $\delta\nu$

lorsque le déphasage entre les modes n'est pas fixé, menant à un régime continu bruité, tandis que la somme des champs mène à un régime impulsionnel lorsque le déphasage entre les modes est fixé. On dit alors que les modes sont bloqués ou verrouillés en phase.

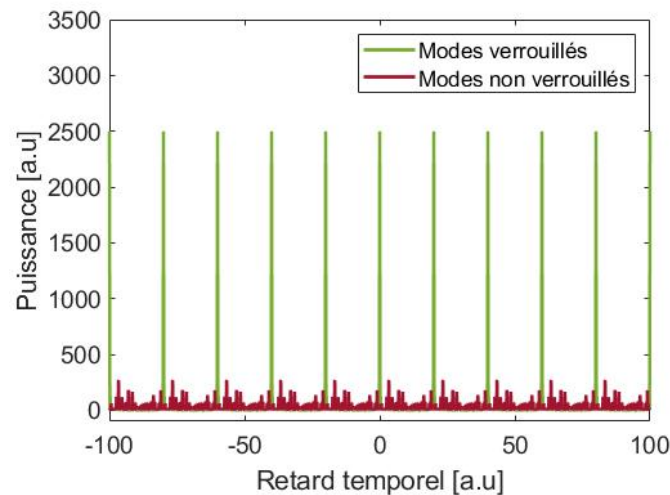


FIGURE 2.9: Schéma illustrant le processus de blocage de mode d'un point de vue temporel, avec la propagation de 50 modes longitudinaux

La durée de l'impulsion à mi hauteur en limite de Fourier τ est liée au nombre de mode bloqués N et son intervalle spectral libre $\delta\nu$ par la relation $\tau = \kappa/\Delta\nu$, avec $\Delta\nu = N\delta\nu$ [33]. Le facteur κ est une constante de valeur différente selon la forme de l'impulsion (gaussienne, sécante hyperbolique...). Plus le nombre de modes circulant est important, plus courte est l'impulsion. D'un point de vue temporel, le blocage de mode peut être vu aussi comme une discrimination des modes de faible intensité par rapport à ceux de forte intensité. Ainsi, on crée de fortes pertes sur les régimes continus, favorisant les régimes impulsionnels. Le blocage de mode peut être obtenu par deux voies différentes :

- **Le blocage de mode actif**

Il requiert un élément extérieur à la cavité, souvent un modulateur électro-optique d'amplitude ou de fréquence. Cette technique utilise certains effets comme l'effet

Pockels, l'effet Kerr ou encore l'effet d'électro-gyration pour créer un couplage entre les modes adjacents. La fréquence de modulation induite aura alors une fréquence correspondante à l'intervalle spectral libre. Le mode le plus intense imposera alors sa phase aux modes adjacents, qui à leur tour l'imposeront à leurs modes adjacents. Le blocage de la phase de chaque mode est ainsi réalisé, ce qui produit les impulsions. Le premier blocage de mode actif utilisant un modulateur électro-optique sur un laser Erbium date de 1989 [51, 52].

— Le blocage de mode passif

C'est le modèle de blocage de mode le plus attractif, car il ne requiert aucune synchronisation extérieure par un élément coûteux. En effet, ce type de blocage de modes inclut directement un élément dans la cavité, que l'on appelle absorbant saturable, pour créer le régime impulsionnel. Cette technique permet de générer des impulsions ultracourtes à bas coût. La prochaine section détaillera ce type spécifique de blocage de mode, le plus couramment utilisé.

Qu'est ce qu'un absorbant saturable ?

Un absorbant saturable (AS) est un élément dont la transmission dépend de la puissance du champ incident. Ainsi les régimes à faible puissance seront discriminés par rapport à ceux à forte puissance. Cet élément est inclus dans des cavité laser pour favoriser les régimes impulsionnels en discriminant les régimes continus. La génération d'impulsion débute par un régime d'intensité relativement élevée, où les modes commencent à être verrouillés. Ce léger pic de puissance sera amplifié par le milieu à gain à chaque tour de cavité si sa puissance est suffisamment importante pour être transmise par l'AS. Après un nombre important de passages dans le complexe AS/fibre dopée, qui peut varier de quelques centaines à quelques milliers, une impulsion de très forte intensité et de durée très brève sera créée, comme illustré dans la figure 2.10. D'un point de vue spectral, de nombreux modes se propageant dans la cavité, le spectre d'une impulsion est très large.

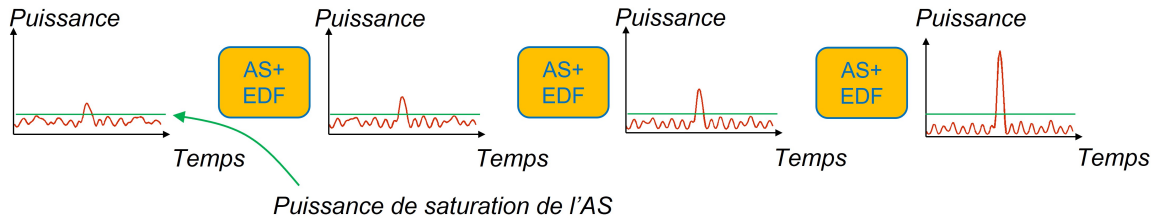


FIGURE 2.10: Schéma décrivant les effets conjugués d'un absorbant saturable (AS) et d'une fibre dopée (EDF) sur le processus de génération d'impulsion

Prenons comme premier exemple celui d'un absorbant saturable réel. Cet élément est un système à deux niveaux, pour lequel l'absorption de l'AS serait maximale tant que le champ incident n'a pas suffisamment de puissance pour exciter les atomes du niveau fondamental. Une fois ce niveau dépeuplé, l'absorption devient presque nulle, et l'AS transmet les puissances supérieures à sa valeur de saturation. Sa transmission en puissance peut être simplement modélisée par l'équation suivante :

$$T = T_{max} - \frac{T_{max} - T_{min}}{1 + \frac{P}{P_{sat}}} \quad (2.1)$$

Avec T_{max} la transmission maximale de l'AS, T_{min} sa transmission minimale, P la puissance du champ incident et P_{sat} la puissance de saturation du dispositif. Ainsi, plus la puissance du champ incident est importante, plus la transmission sera importante. Si le temps de relaxation de l'absorbant saturable est inférieur à la durée de l'impulsion, on parlera d'absorbant saturable rapide [27], tandis qu'on parlera d'absorbant saturable lent si le temps de relaxation du milieu est supérieur à la durée d'impulsion [28]. Il est intéressant de noter qu'un absorbant saturable lent peut permettre la génération d'impulsions ultracourtes dont la durée est inférieure au temps de relaxation du milieu [53], même si cela est plus compliqué. L'avantage d'utiliser un absorbant saturable rapide réside alors dans la stabilité souvent plus importante des régimes générés avec ce type de milieu.

Un absorbant saturable peut être de deux natures différentes, selon les effets physiques mis en jeu. Certains sont basés sur l'absorption matérielle d'un élément. On parlera alors d'absorbant saturable réel, comme pour l'exemple ci-dessus. Nous citerons à titre d'exemple l'absorbant saturable par semi conducteur (semiconductor saturable absorber mirrors, SESAM) [54, 55], les colorants [56] ou encore le graphène [57]. Ces types d'absorbants saturables ne demandent aucun réglage manuel, mais ne permettent aucun ajustement. C'est pourquoi nous préférons utiliser des absorbants saturables virtuels, basés sur des interférences non linéaires, qui permettent un ajustement des caractéristiques du régime généré, mais demandent souvent un long réglage manuel. De plus, ce style d'absorbant saturable comporte généralement de plus faibles pertes comparées aux réels. Nous citerons à titre d'exemple les interféromètres de Sagnac sous leurs formes NOLM (nonlinear loop mirror) [58] ou NALM (nonlinear amplifying loop mirror) [25] ou encore les effets d'absorbant saturable par rotation non linéaire de la polarisation [59, 60], utilisés dans nos expériences, et détaillées dans le prochain paragraphe.

Absorbant saturable par rotation non linéaire de la polarisation

La rotation non linéaire de la polarisation (RNLP), ou évolution non linéaire de la polarisation est un processus induit par la propagation d'un champ électrique à forte amplitude dans un milieu Kerr isotrope. Ce processus est déclenché par deux effets non linéaires : l'automodulation de phase (SPM) et l'intermodulation de phase (XPM). Prenons l'exemple d'une impulsion très intense dans un composant fibré : l'effet Kerr implique que l'indice de réfraction de la fibre devient dépendante de l'intensité du champ :

$$n(I) = n_0 + n_2(I) \quad (2.2)$$

avec n_0 l'indice de réfraction linéaire de la fibre (il vaut 1.44 pour une fibre en silice monomode SMF28), n_2 l'indice de réfraction non linéaire (qui vaut $2.6 \cdot 10^{-20} \text{m}^2/\text{W}$ dans ce cas) et I l'intensité de l'impulsion. Cette variation d'indice va modifier le déphasage $\Delta\phi$ via l'induction d'un déphasage non linéaire $\Delta\phi_{NL} = \gamma PL$ avec γ le facteur de non

linéarité de la fibre tel que

$$\gamma = \frac{2\pi n_2}{\lambda A_{eff}} \quad (2.3)$$

A_{eff} étant l'aire effective de la fibre, P la puissance de l'impulsion et L la longueur de la fibre. Ce phénomène est la SPM.

Continuons notre analyse d'un régime impulsionnel, pour un champ incident $\vec{E}$ ayant des projections E_x et E_y respectivement sur les axes x et y . Pour une simplification des calculs, une représentation circulaire de la polarisation est préférée. Avec cette représentation, $E_+ = (E_x + iE_y)/\sqrt{2}$ et $E_- = (E_x - iE_y)/\sqrt{2}$ correspondent aux états de polarisations circulaire droit et circulaire gauche. Un second phénomène non linéaire, la XPM intervient alors. Couplée à la SPM, ces deux effets entraînent des changements dans les indices de réfraction pour les deux sens de polarisation circulaire, se traduisant par deux déphasages non linéaires pour chaque composantes [33] :

$$\phi_{NL+} = \frac{2\gamma}{3}(P_+ + 2P_-)L \quad (2.4)$$

et

$$\phi_{NL-} = \frac{2\gamma}{3}(P_- + 2P_+)L \quad (2.5)$$

avec P_+ et P_- les puissances de chacune des composantes, menant à un déphasage entre les deux composantes :

$$\Delta\phi_{NL} = \frac{2\gamma}{3}(P_- - P_+)L \quad (2.6)$$

Dans le cas où la polarisation du champ est circulaire gauche ou droite ($P_+ = 0$ ou $P_- = 0$), la XPM n'induit aucun déphasage, seule la SPM intervient. En revanche, si la polarisation du champ est elliptique ($P_+ \neq P_-$), le couple SPM/XPM intervient. Les axes de l'ellipse seront alors réorientés par $\Delta\phi_{NL}$ selon la puissance de chaque composante. Cet effet a été démontré pour la première fois dans des liquides par Maker en 1964 [122].

Si une discrimination sur les axes de polarisation avec un polariseur est effectuée, cette variation de polarisation se transformera en variation d'intensité. Si les axes correspondant aux fortes intensités sont sélectionnés, et donc ceux correspondants aux faibles intensités discriminés, un absorbant saturable ultrarapide et efficace est créé. De plus, si un dispositif permettant d'appliquer une phase variable entre les deux composantes du champ est utilisé, comme des lames à cristaux liquides, la fonction de transfert de l'absorbant saturable créée devient ajustable. Une automatisation de la génération d'impulsions peut alors être réalisée [123], ainsi que le couplage avec un algorithme.

Dans la plupart des cas, pour ajuster la polarisation, des lames $\lambda/2$ et $\lambda/4$ ou des contrôleurs de polarisation sont utilisés, permettant une modification de la fonction de transfert du dispositif lames+polariseur de manière manuelle. De plus, une bonne combinaison de lames ou contrôleurs de polarisation permet de générer tous les états de polarisation à partir de n'importe quel état. Il devient donc possible d'atteindre toutes les fonctions de transfert possibles. L'ajustement de ces paramètres permet la génération de nombreuses dynamiques impulsionnelles variées, une liste non exhaustive de ces dynamiques est proposée dans une section suivante.

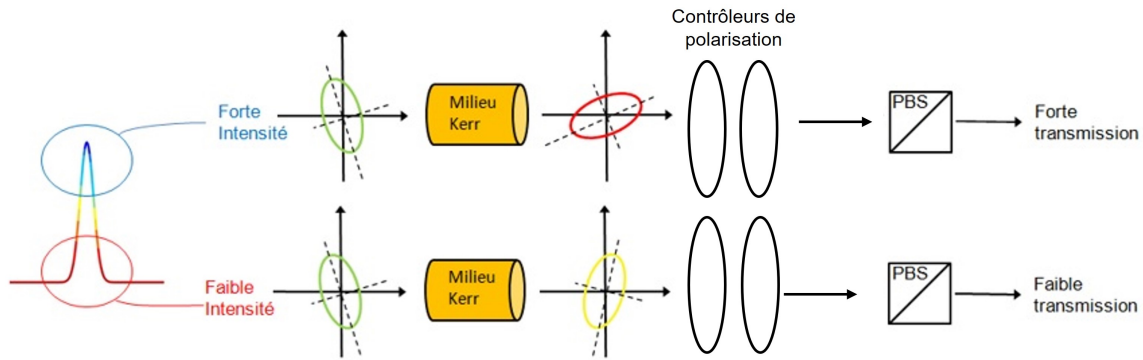


FIGURE 2.11: Schéma décrivant la rotation non linéaire de la polarisation

2.2.2 Le soliton dissipatif

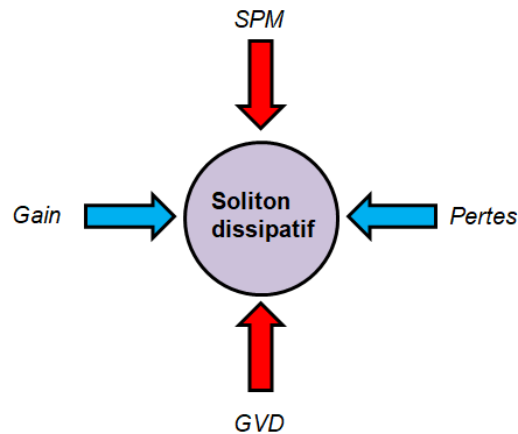


FIGURE 2.12: Schéma montrant l'équilibre entre les phénomènes physiques pour la création de solitons dissipatifs

Un soliton est idéalement une onde qui se propage sans modification ni alteration sur une distance infinie. Il naît grâce à un équilibre entre deux effets physiques : la dispersion de la vitesse de groupe (ou Group Velocity Dispersion, GVD), et l'automodulation de phase (SPM) que nous avons détaillé en section précédente. C'est en 1834 que l'ingénieur écossais John Scott Russel observa pour la première fois la propagation d'une vague dans un canal sans alteration ni de vitesse ni d'amplitude [61]. La GVD, caractérisée par son coefficient β_2 , provient de la dépendance de la vitesse de phase d'une onde à sa longueur d'onde. Si elle est positive, nous sommes en régime de dispersion normal, les grandes longueurs d'onde se propagent plus rapidement que les courtes, ce qui résulte d'un allongement de la durée d'impulsion. Si elle est négative, à l'inverse, les courtes longueurs d'onde se propageront plus rapidement que les grandes, ce qui conduit également à un allongement de la durée d'impulsion. Pour compenser cet élargissement, la SPM entre en jeu. L'équilibre entre les effets de la SPM et ceux de la dispersion mène à la génération de solitons conventionnels, dont la propagation est modélisée par l'équation de Schrödinger non-linéaire prenant en compte seulement ces

deux effets [62]. Ces solitons sont appelés NLS pour Non Linear Schrödinger. Cependant, ce modèle de soliton n'est pas entièrement valide pour les impulsions laser en cavité fibrée. En effet, à chaque tour de cavité, le soliton subira des pertes dues aux éléments optiques constituant la cavité. De même, ces pertes seront compensées par le gain du milieu amplificateur. Ces deux grandeurs ne sont pas prises en compte dans le modèle du soliton NLS, d'où la limitation du modèle pour décrire la propagation d'un champ dans ce type de milieu. Un autre modèle de soliton, appelé soliton dissipatif, incluant un équilibre entre pertes et gain de la cavité en plus de celui SPM/GVD, sera donc développé, plus conforme à décrire la dynamique des impulsions dans les cavités laser fibrées. Dans ce cas, en régime stationnaire, les caractéristiques des solitons créés dépendent entièrement des caractéristiques de la cavité, et non des conditions initiales comme c'est le cas pour les solitons conventionnels. Ces conditions sont illustrées dans la figure 2.12.

Considérons à présent le problème d'un point de vue calculatoire. Au début des années 1980 [63], l'équation de Schrödinger non linéaire était utilisée pour modéliser la propagation d'un champ dans une fibre passive. Cependant, cette équation n'est valable que pour les solitons conventionnels NLS. Haus et al. en 1991 [64] ont donc rajouté des termes de gain et de pertes, afin de décrire la propagation de solitons dissipatifs dans des cavités laser fibrées. L'équation construite, pour un champ électrique E stationnaire, s'écrit :

$$(g - l - i\phi)E + \left(\frac{g}{\Omega_g^2} + iD\right) \frac{\partial^2 E}{\partial t^2} + (\delta - i\gamma)|E|^2 E = 0 \quad (2.7)$$

Avec g et l respectivement le gain et les pertes linéaires de l'élément considéré, ϕ le déphasage linéaire induit par cet élément, Ω_g la largeur de bande du gain, γ le terme des non linéarités et δ le terme de modulation d'amplitude de l'absorbant saturable. D est la dispersion de l'élément, qui est reliée à la GVD par la relation suivante : $D = -\frac{2\pi c}{\lambda^2} \beta_2$, avec D en $\text{ps.nm}^{-1}.\text{km}^{-1}$ et β_2 en $\text{ps}^2.\text{m}^{-1}$. La solution de cette équation est une impulsion de type sécante hyperbolique [64, 65] :

$$E = A \operatorname{sech}\left(\frac{t}{\tau}\right) \exp\left[iB \operatorname{sech}\left(\frac{t}{\tau}\right)\right] \quad (2.8)$$

avec A l'amplitude de l'impulsion, τ sa durée et B le paramètre de "chirp". Le chirp, ou la dérive de fréquence, est un phénomène présent dans toutes les cavité lasers fibrées. Il implique une pulsation qui varie durant l'impulsion. Pour expliquer ce phénomène, considérons la forme générale d'un champ électrique : $E(t) = E_0 e^{i(\omega_0 t + \varphi(t))} = E_0 e^{i\phi(t)}$, avec ω_0 la pulsation propre de l'impulsion et $\omega(t) = \frac{d\phi(t)}{dt} = \omega_0 + \frac{d\varphi(t)}{dt}$ la pulsation instantanée. Le chirp correspond alors à la partie non stationnaire $\frac{d\varphi(t)}{dt}$, et implique alors une phase $\varphi(t)$ variante dans le temps. Ce phénomène est très présent dans les impulsions laser à fibre, particulièrement pour des régimes de dispersion normaux.

L'équation de Haus est cependant instable dans la plupart des cas. Elle est indicative pour les états stationnaires, mais ne permet donc pas d'étudier la dynamique. Ceci en raison de la prise en compte insuffisante des phénomènes de saturation. C'est pourquoi on modélise la dynamique d'une impulsion dans une cavité laser fibrée par l'équation

cubique quintique de Ginzburg-Landau (CGLE) [66], donnée en équation 2.9, et étudiée numériquement par Soto Crespo en 1997 dans le cas d'une cavité à dispersion normale [67].

$$\psi_z - i\frac{D}{2}\psi_{tt} - i|\psi|^2\psi - i\nu|\psi|^4\psi = \delta\psi + \beta\psi_{tt} + \epsilon|\psi|^2\psi + \mu|\psi|^4\psi \quad (2.9)$$

ψ est le champ complexe de l'impulsion qui dépend de deux variables : t et z qui sont respectivement le temps dans le référentiel de l'impulsion et la distance de propagation. ψ_z est la dérivée du premier ordre de ψ selon z et ψ_{tt} est la dérivée du second ordre de ψ selon t . ν est le coefficient de saturation de l'effet Kerr, μ la saturation du gain non linéaire et δ les pertes linéaires. Le terme $\beta\psi_{tt}$ représente le filtrage spectral par le milieu à gain, et le coefficient ϵ représente le gain non linéaire. Les phénomènes de saturations pris en compte dans cette équation assurent le caractère attracteur du soliton dissipatif. Dans l'équation 2.9, plusieurs termes menant à la stabilisation des régimes solitoniques sont présents comparé à 2.7, notamment le terme quintique μ , qui étaient la principale limitation du précédent modèle. La CGLE comprend tous les paramètres des solitons dissipatifs, la rendant conforme à la modélisation de leur propagation.

La CGLE est un modèle distribué. Or, une cavité laser est discontinue car constituée de plusieurs éléments distincts. Pour modéliser plus précisément la dynamique laser, des modèles de propagation à paramètres variables sont donc nécessaires, induisant des variations du soliton. Un phénomène de rayonnement d'ondes dispersives apparaît alors. Ces ondes sont appelées bandes latérales de Gordon-Kelly [38, 68], indiquées par des pics de continus sur le spectre optique et symétriques sur la localisation par rapport à la longueur d'onde centrale du spectre, comme présenté dans le spectre mesuré en figure 2.13.

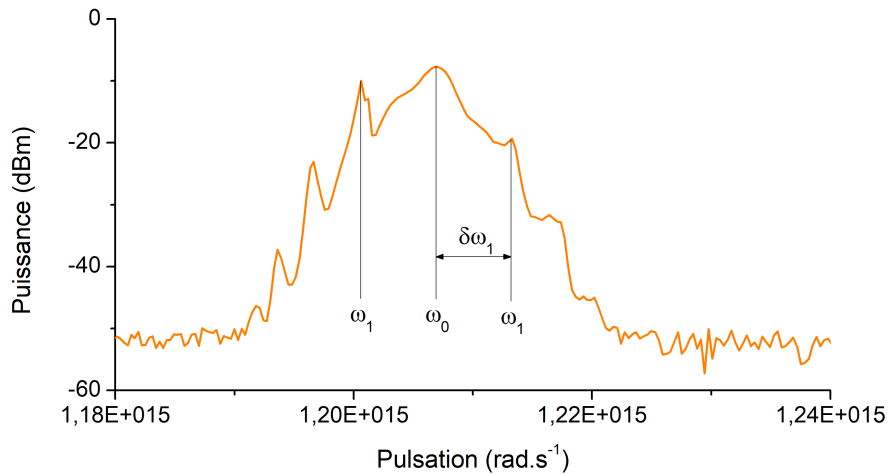


FIGURE 2.13: Spectre optique d'une impulsion usuelle enregistrée dans nos expériences, montrant les bandes latérales de Gordon Kelly

Ces bandes sont un bon indicateur de la dispersion globale de la cavité. Cette dispersion, résultant de la somme de la dispersion de chaque élément, est liée aux bandes

de Gordon Kelly par la relation suivante [69] :

$$\delta\omega_m^2 = -\frac{1}{\tau^2} + \frac{2\pi m}{\beta_{2,net}} \quad (2.10)$$

Avec $\delta\omega_m$ la déviation de la pulsation de la $m^{\text{ième}}$ bande de Gordon Kelly par rapport à la pulsation centrale du champ, τ la durée d'impulsion et $\beta_{2,net}$ la GVD globale. Dans cet exemple, la durée d'impulsion est estimée à $\tau = 250$ fs par autocorrélation, et $\delta\omega_1 = 6.22 \cdot 10^{12}$ rad.s⁻¹, ce qui nous donne une GVD globale de $\beta_{2,net} = -0.115$ ps². A noter que cette relation, basée sur un $\delta\omega$ forcément positif car absolu, donne une valeur absolue de la GVD. Les bandes n'apparaissant uniquement pour des cavités à dispersion anormales, le signe de la GVD est forcément négatif.

2.2.3 Dynamiques impulsionnelles au-delà du blocage de mode conventionnel

Une dynamique impulsionnelle est fortement déterminée par les paramètres de la cavité dans laquelle elle se propage. Ainsi, la gestion de la dispersion, le sens de propagation, la puissance de pompe, la séquence des composants dans la cavité, les non linéarités déterminées par les longueurs de fibre ou encore la fonction de transfert de l'absorbant saturable sont autant de causes déterminantes dans l'obtention d'un régime. En jouant sur ces paramètres, il est alors possible de générer des états de propagation spécifique, à dynamiques complexes.

Le Q-switch mode locking

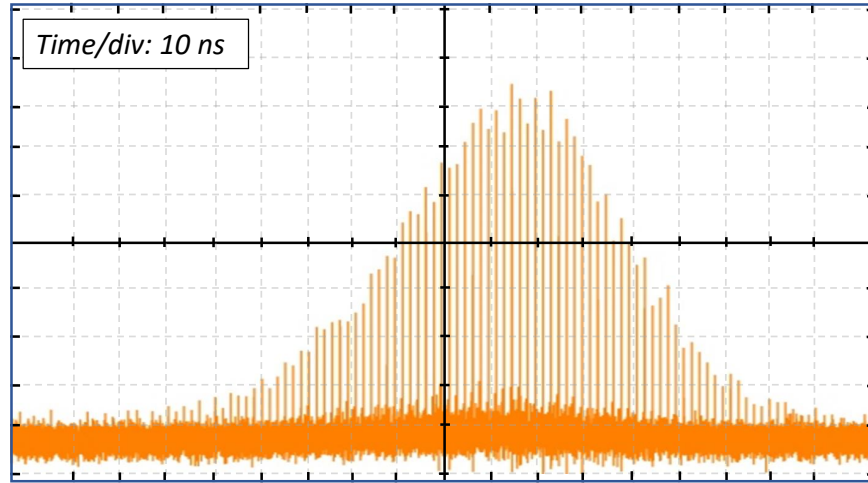


FIGURE 2.14: Trace d'oscilloscope présentant un régime de Q switched mode locking, avec un enveloppe d'une centaine de ns

Le Q switching, ou commutation Q, est un processus de déclenchement permettant la mise en forme de champ en impulsion. Elle est réalisée en modulant le facteur de qualité Q dans une cavité via un dispositif induisant des pertes élevées pendant un court

instant. Ainsi, lorsque les pertes sont très élevées, peu de champ circule dans la cavité et donc l'amplification par émission stimulée reste minimale. La puissance de pompe étant constante, le milieu à gain se dépeuple peu et reste fortement excité, permettant donc une forte inversion de population. A partir d'un certain temps, l'émission spontanée ainsi que d'autres phénomènes pousseront le milieu à entrer dans un régime de saturation, n'augmentant plus l'inversion de population. Les pertes de la cavité sont alors rendues faibles en changeant brutalement le facteur de qualité Q , d'où la terminologie Q switch. La population du niveau métastable se vide brusquement, ce qui mène à une impulsion d'une durée de l'ordre de la nanoseconde. Lors de cette impulsion longue, il est possible de combiner une dynamique d'impulsions ultracourtes issues du blocage de mode, donnant un régime avec une enveloppe due au Q switching, et une modulation en train d'impulsion due au blocage de mode. Cet état est nommé "Q-switching mode locking" (QSML). Les impulsions dues au blocage de mode y sont très intenses, plus qu'avec un simple blocage de mode, mais sont très instables, donnant une concentration d'impulsions ultracourtes pendant un court instant, comme nous pouvons le voir dans la figure 2.14. Une étude complète des régimes de Q-switching mode locking a été présentée en 1999 par Honninger [70]. Cette dynamique, présente quelque soit le régime de dispersion de la cavité, n'est généralement pas recherchée à cause de son caractère très instable. De plus, ce régime peut s'avérer être destructeur pour certains éléments, du fait de ses puissances crêtes extrêmement élevées.

Les solitons incohérents (ou noise like pulses)

Découvert en 1997 [71], le régime de solitons incohérent ou "noise like pulse" dans les lasers est caractérisée par une perte totale de cohérence mutuelle entre les impulsions successives [72], comme illustré en figure 2.15 (a) et (b). Cependant, il reste quasi-stationnaire après son établissement. Ses potentialités semblent intéressantes, avec des impulsions à fortes énergies. Ce régime très chaotique mène à une trace d'autocorrélation multi-coups avec un fort pédestal (long de plusieurs dizaines de picosecondes), comme nous pouvons le remarquer dans la figure 2.15 (c), et un spectre optique très large (jusqu'à une centaine de nm de large avec un laser Erbium en régime de dispersion normal). La trace d'autocorrélation montre également la présence d'un pic de cohérence à délai nul correspondant à l'autocorrélation d'une structure entière avec elle même. En effet, l'organisation des impulsions change complètement au cours de la propagation, donnant des paquets d'ondes non organisés.

Ce genre de régime a tendance à apparaître avec des cavités laser à dispersion normale, et plutôt longue, avec typiquement une forte longueur de fibre DCF. Horowitz et Silberberg [71], précurseurs dans la génération de régimes de noise like pulses, ont réalisés une série d'études permettant d'établir un contrôle de la génération de noise like pulses en cavités laser fibrées [73, 74]. Leur hypothèse était qu'une longueur importante de fibre est nécessaire à la génération de tels régimes, grâce à la présence de fortes non linéarités. Cependant, en 2005, Tang et al démontrent la génération d'impulsions incohérentes dans des cavités à dispersion anormale et courtes [75], pointant certaines limitations dans les études d'Horowitz. Plus récemment, Boucon évoque la possibilité de générer ce type d'impulsion dans un laser Raman [77], puis Soto Crespo démontre numériquement la possibilité de générer de tels régimes sans influence d'effets Raman

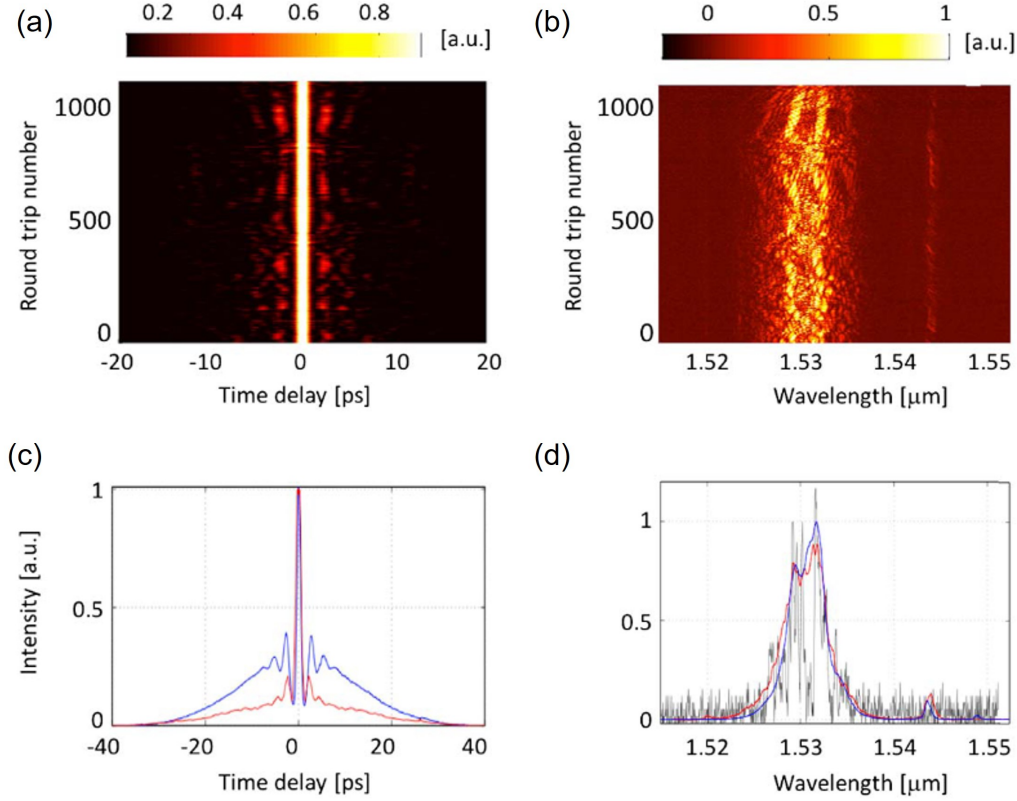


FIGURE 2.15: Illustration du phénomène de solitons incohérents pour un laser à régime de dispersion anormal, avec (a) l'évolution de son intensité temporelle par tour de cavité, (b) l'évolution de son spectre par tour de cavité, (c) comparaison entre la trace d'autocorrélation de second ordre multi-coups normalisée (courbe bleue) et la trace d'autocorrélation de premier ordre moyenne normalisée (courbe rouge) calculé depuis (a), et (d) comparaison entre le spectre moyen OSA enregistré (courbe bleue) et le spectre moyen calculé (courbe rouge) calculé depuis (b), tirés de [79]

ou de grandes biréfringences [78]. En 2017, Krupa démontre l'existence de structures organisées basées sur des composantes de polarisation en interaction dans les régimes incohérents [79], puis Z. Wang enregistre en temps réel la dynamique de construction d'impulsions incohérentes sur plusieurs milliers de tours de cavité, et pour différents régimes de dispersion [76].

Le blocage de mode multi impulsionnel

Dans une cavité laser fibrée, la mise en forme d'une impulsion dépend des paramètres de la cavité non modifiables, comme la dispersion, les non linéarités ou le gain. Le blocage de mode multi impulsionnel prend place lorsque la cavité laser comporte de fortes non linéarités dues à la présence d'importantes longueurs de fibre SFM. En effet, la mise en forme d'une impulsion dans une cavité avec des paramètres fixés est réalisée jusqu'à une puissance crête et une énergie typique. Ainsi, pour une cavité à dispersion anormale et à faibles pertes, la puissance et l'énergie typique d'un soliton sont reliées aux paramètres de celle ci par les relations suivantes :

$$P_0 = \frac{|\beta_2|}{\gamma T_0^2} \quad \text{et} \quad \mathcal{E} = \frac{2|\beta_2|}{\gamma T_0} \quad (2.11)$$

Avec $T_0 = \tau/1.76$ et τ la durée d'impulsion. La puissance de pompe, qui est variable, apporte une certaine quantité d'énergie stockée dans la cavité. Lorsque celle-ci augmente, l'énergie stockée dans la cavité augmente également. Après avoir dépassé le seuil de blocage de mode, une importante puissance intracavité mènera à une impulsion plus intense. En revanche, si l'apport en puissance par la pompe dépasse la valeur caractéristique de la puissance d'un soliton dissipatif, donnée dans 2.11, le régime devient instable, avec tout d'abord une augmentation de l'intensité crête des pics de continus (bandes latérales), puis dans un deuxième temps la génération de nouvelles impulsions. Cela mènera donc à la formation de nouvelles impulsions moins intenses à chaque apport en puissance d'un multiple de P_0 , conduisant à un régime multi impulsif. Ces impulsions multiples, étant parfois proches les unes des autres temporellement, ne sont pas facilement détectable par une photodiode ultrarapide, dont la résolution temporelle est de l'ordre de la dizaine de picoseconde. Ce phénomène est illustré dans la figure 2.16.

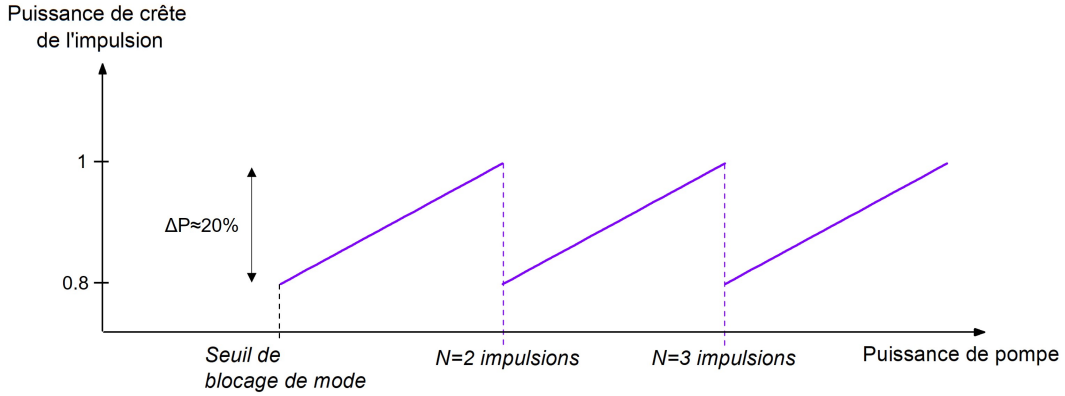


FIGURE 2.16: Illustration du phénomène de génération de plusieurs impulsions dans une cavité laser fibrée

Ce type de dynamique, conduisant à des distributions d'impulsions souvent irrégulières, est dans la plupart des cas préjudiciable pour être utilisé. Cependant, il constitue un banc de test intéressant pour l'étude des interactions entre impulsions. Ces dynamiques peuvent être intéressantes également pour une augmentation du taux de répétition de la cavité lorsque les impulsions sont auto-organisées comme nous le détaillerons par la suite. Pour limiter l'apparition de régimes multi impulsif, des régimes de dispersions normaux et des cavités courtes seront utilisés. On utilisera alors des DCF pour compenser la dispersion anormale des fibres SMF (modèle du "dispersion managed solitons"), ou en utilisant une longue fibre dopée à dispersion normale [80].

Si de nombreux régimes multi-impulsionnels semblent irréguliers, certains régimes présentent des patterns quasi-stationnaires ou stationnaires, et pour différentes longueurs d'ondes. Nous présenterons dans les sous sections suivantes les principaux régimes multi impulsif organisés.

Le blocage de mode harmonique

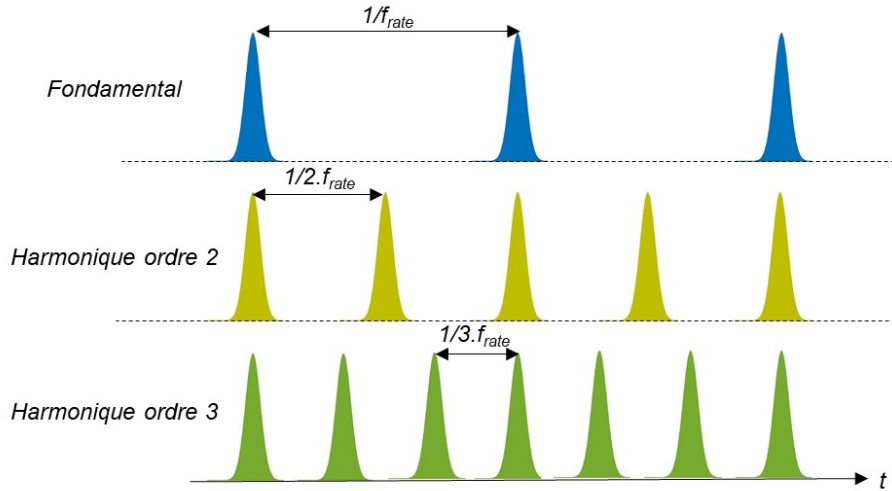


FIGURE 2.17: Illustration du blocage de mode harmonique, allant jusqu'à l'harmonique 3

Dans un régime de blocage de mode harmonique, les solitons sont uniformément répartis dans la cavité, à une distance temporelle identique d'une impulsion à l'autre. Dans ce cas, le taux de répétition du laser est un multiple du taux de répétition fondamental, ce multiple correspondant à l'ordre de l'harmonique généré, comme présenté en figure 2.17. Ce type régime stable a été étudié par Grudinin dans les années 1990 [81,82] dans le but de contrôler le taux de répétition d'un laser à blocage de mode. Plus récemment, des blocages de mode à harmonique supérieures à 10^3 ont pu être observés [83], menant à un taux de répétition de 20 GHz. A travers ces travaux, il a été remarqué que les cavités à dispersion anormale ainsi qu'une forte énergie de pompe favorisent les blocages de modes harmoniques. Les chercheurs s'accordent à penser que cet espacement identique provient d'une force de répulsion. Grudinin [81] prétendait que cette force répulsive provenait d'un effet acousto optique. Lors de la propagation d'une impulsion, un phénomène d'électrostriction apparaît modifiant localement l'indice de réfraction, créant un effet de lentille optique. Cet effet entraîne alors un regroupement en paquet d'ondes des impulsions, comme l'a démontré Pilipetskii [84]. Cependant, l'explication la plus convaincante arrive plus tard, avec les travaux de Kutz [85]. Il explique ce phénomène grâce à l'amplification des impulsions par la fibre dopée. Si plusieurs impulsions se propagent en même temps dans celle-ci, la première va être alors très amplifiée, tandis que les autres potentielles impulsions ne le seront pas, l'effet de la pompe étant moins important pour ces dernières. Ainsi, cela contraint un espacement régulier entre les impulsions. A noter qu'un phénomène de gigue temporelle apparaît généralement dans ces régimes.

La molécule de solitons

La molécule de soliton, ou les cristaux de solitons sont des complexes organisés de solitons. Ces structures organisées reposent sur des interactions au sein d'un paquet d'ondes

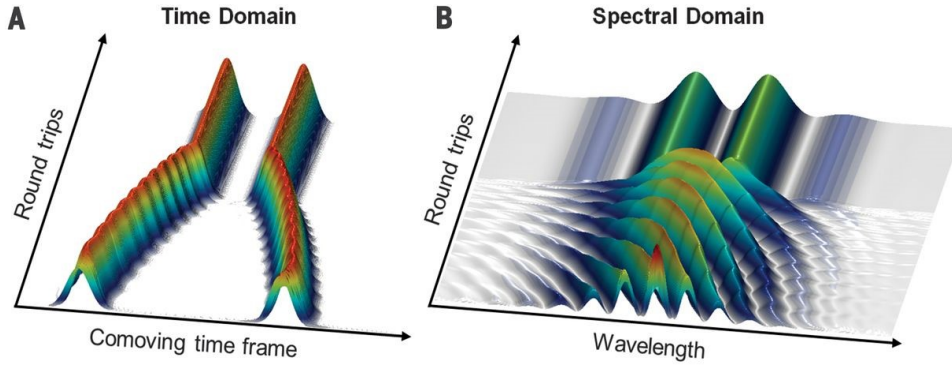


FIGURE 2.18: Illustration de la formation et la propagation de molécules de solitons pour un déphasage constant, tirés de [86]

composé de plusieurs solitons décrites par Malomed [87], mais dont le retard temporel les séparant est très inférieur à la taille de la cavité. Ainsi, avec une faible distance temporelle, les interactions régissant ces paquets sont très fortes, générant des attracteurs [29], et donnant lieu à des structures organisées, stables et robustes, comme illustré en figure 2.18. Les forces liant deux solitons entre eux sont décrites mathématiquement par Gordon en 1983 [88]. Plus la puissance de pompe est importante, plus le nombre de solitons dans la molécule ou le cristal sera élevé. La nature de ces interactions (répulsives ou attractives) varie selon le déphasage entre les solitons : si celui-ci est constant, la distance entre les solitons sera également constante. En revanche, si le déphasage varie, l'interaction sera tantôt répulsive tantôt attractive, donnant lieu à un mouvement régulé des solitons les uns par rapport aux autres. Si la phase est constante, le complexe sera évidemment plus stable. Des travaux expérimentaux ont été réalisés par la suite par Mitschke en 1987 [89] et Grelu en 2002 [90] mettant en évidence expérimentalement l'existence des molécules de solitons, avec différentes relations de phases, variables ou fixes. Dans le dernier chapitre de ce manuscrit, nous présenterons un nouveau moyen d'autogénérer des molécules de soliton à retard temporel prédéterminé. La dynamique des molécules de solitons sera détaillée dans ce point.

La pluie de solitons

La pluie de solitons, illustrée en figure 2.19, a été découverte en 2009 par Chouli [91,92]. Ce type de régime est composé de nombreuses impulsions créées spontanément à partir d'un fond continu bruité, dérivant à une vitesse constante jusqu'à l'atteinte d'une phase condensée de solitons. Cet état est donc composé de trois composantes distinctes : un paquet d'ondes composé de plusieurs dizaines de solitons liés, concentrés en phase condensée, un fond continu de modes non bloqués et donc très bruité, et des solitons créés de manière continue à partir du fond continu dérivant rapidement vers la phase condensée ("drifting solitons"). Ces derniers apparaissent à des positions aléatoires, de manière spontanée si l'énergie de la cavité est suffisante, puis sont attirés vers la phase condensée. La dérive de ces solitons est attribuée aux effets des ondes dispersives du fond continu, induisant un déplacement asymétrique dans la propagation de ces ondes. Une fois en contact avec la phase condensée, les interactions sont si importantes avec celle-ci que les solitons sont piégés à l'intérieur.

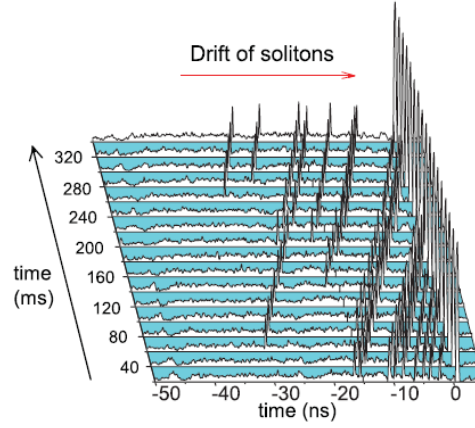


FIGURE 2.19: Illustration du phénomène de pluie de soliton, tiré de [93]

Les dynamiques de pluies de solitons clôturent le catalogue des régimes impulsionnels couramment rencontrés dans les lasers à fibres ultrarapides. Les caractéristiques des régimes solitoniques dépendent principalement des paramètres fixés de la cavité, comme la dispersion, mais également de la fonction de transfert de l'absorbant saturable. Dans la prochaine section, le choix de cet absorbant saturable sera détaillé pour nos travaux.

2.3 Quel absorbant saturable utiliser ?

2.3.1 Choix des composants

Ce choix est dans notre cas dicté par 3 problématiques. Premièrement, une des caractéristique désirée du laser concerne la grande versatilité des régimes générés. Un absorbant saturable réel, avec paramètres fixés, ne permet la génération que d'un seul régime, et ne conviendrait donc pas à nos attentes. Un absorbant saturable virtuel, avec plusieurs paramètres ajustables, est plus adapté pour notre utilisation. Le type d'absorbant saturable permettant la plus grande versatilité des régimes générés est celui obtenu par évolution non linéaire de polarisation, décrit en section 2.2.1. Nous choisirons donc ce type d'absorbant saturable pour nos expériences. Une seconde caractéristique intrinsèque à notre étude, implémentant un algorithme sur un montage laser, est le contrôle informatique de ces degrés de liberté. L'utilisation de contrôleurs de polarisation électroniques basés sur des effets thermiques par exemple (EPC) est une solution, comme montré dans les travaux préalables [95–97]. Cependant, la dernière problématique à résoudre est la génération d'impulsions stables. Pour augmenter la stabilité des régimes, il est préférable de réduire la longueur de fibre utilisée dans le montage, responsable des instabilités. De plus, de tels dispositifs, basés sur des effets thermiques, présentent une durée de stabilisation longue (de l'ordre de la seconde), ce qui conduit à des durées d'optimisations importantes. Un dispositif en espace libre regroupant toutes ces prérequis sera donc choisi : une combinaison de lames à retard de phase variable à cristaux liquides (LC), précédemment utilisées dans les travaux de Winters [98], à temps de réponse rapide (≈ 20 ms comparé à ≈ 1 s pour des EPCs). Ces dernières permettront un pilotage informatique de la fonction de transfert de l'absorbant saturable, et constitueront les degrés de liberté de notre montage laser.

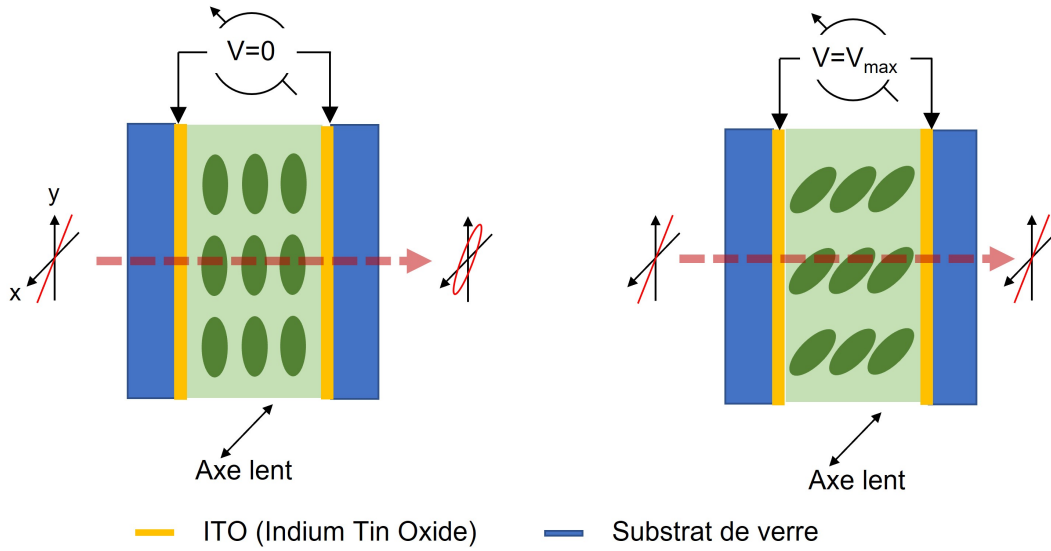


FIGURE 2.20: Illustration du fonctionnement d'une lame à retard à cristaux liquides

Le fonctionnement d'une lame à retardement à cristaux liquides est schématisé dans la figure 2.20. Concrètement, une LC est une cavité contenant un élément nématique, c'est à dire un élément dont l'état combine des propriétés d'un liquide conventionnel et celles d'un solide cristallisé [99]. Nous pouvons le voir comme un liquide visqueux dans lequel baignent des cristaux non sphériques. L'application d'un champ électrique sur la lame modifie l'orientation de ces cristaux, et donc modifie le trajet optique d'un champ se propageant à l'intérieur. Ainsi, un retard entre les composantes du faisceau x et y de la polarisation sera induit, se traduisant par un déphasage. En l'absence de champ électrique, les cristaux sont orientés parallèlement aux bord du support, donc perpendiculairement à la propagation du faisceau, ce qui conduit à une différence de trajet optique maximum (à un déphasage maximum selon x ou y selon leur orientation). En revanche, avec un champ maximal (dans notre cas 10 V), les cristaux sont orientés selon la propagation du faisceau, induisant un déphasage quasi nul entre les deux composantes. La génération du champs électrique étant pilotée par ordinateur, l'utilisations de LCs permet également l'interfaçage de notre cavité, permettant l'implémentation de l'algorithme d'évolution, comme nous le décrirons par la suite.

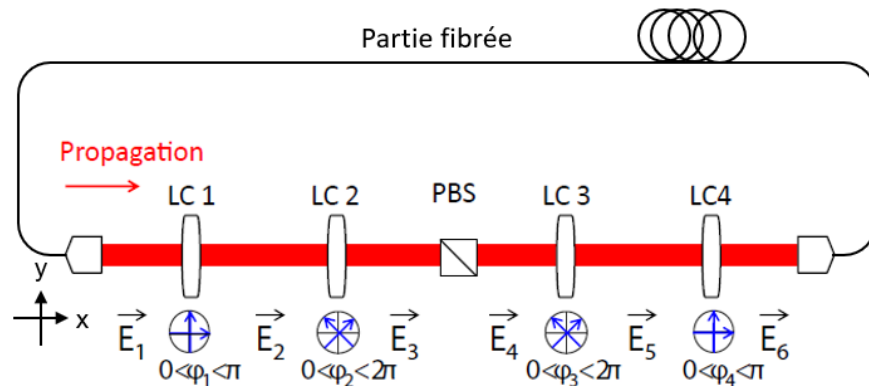


FIGURE 2.21: Illustration de la partie en espace libre de notre montage expérimental

Ainsi, il est possible d'ajuster la polarisation d'un champ incident [99, 100], comme pour un EPC. Dans nos expériences, en cavité, une combinaison astucieuse de LCs est requise pour permettre un contrôle complet des états de polarisations. Cette combinaison est illustrée en figure 2.21. La partie fibrée, étant un milieu Kerr, est responsable de la rotation non linéaire de la polarisation, et modifiera la polarisation des signaux selon leur puissance. Une combinaison de 2 LCs (LC1 et LC2) permettra d'ajuster cette polarisation, ajustant de fait la transmission par un cube séparateur de polarisation (PBS). Cet ajustement permettra, dans un cas souhaité, de discriminer les régimes à faible puissance et de transmettre uniquement les fortes puissances dans la cavité laser. Un couple de 2 nouvelles LCs (LC3 et LC4) permet d'ajuster l'état de polarisation du champ réinjecté dans la partie fibrée après le polariseur, où la polarisation incidente est connue. Ainsi, la fonction de transfert de l'absorbant saturable sera ajustée de la manière la plus versatile qui soit, c'est à dire la possibilité d'induire n'importe quel état de polarisation à partir d'un état quelconque. Pour résumer, un dispositif composé de 4 lames à retard variable à cristaux liquides (LC1, LC2, LC3 et LC4, induisant respectivement des déphasages de φ_1 , φ_2 , φ_3 et φ_4) et un PBS est monté en espace libre. Les pertes créées par le PBS, c'est à dire les composantes verticales selon y de la polarisation, sont utilisées comme faisceau de sortie pour le diagnostic du régime. Les composantes horizontales selon x sont quant à elles transmises par le PBS et réinjectées dans la cavité.

Les lames utilisées permettent l'induction d'un déphasage compris entre 0 et π avec LC1 et LC4 (lames de la marque Meadowlark), et entre 0 et 2π avec LC2 et LC3 (lames de la marque Thorlabs). Chaque lame possède, comme pour une lame d'onde classique, un axe lent et un axe rapide. Faisons un parallèle avec le cas des lames d'ondes : avec ces lames de phases classiques, les déphasages sont fixés, mais c'est l'orientation des axes qui peut être ajustée, permettant la modification de la polarisation du champ. Dans notre cas, avec des LCs, c'est l'inverse : les axes sont fixés tandis que le déphasage est ajustable. Il est alors évident que l'orientation de ces axes d'une lame par rapport à l'autre est une donnée critique pour la génération d'une plage d'état de polarisation générable maximale. Après calcul et réflexion, les LCs sont placés à 45° les unes des autres, ce qui couplé aux différents déphasages accessibles, permet d'atteindre toutes les fonctions de transfert possibles par le PBS.

Nous distinguons dans cette zone deux sous parties : la partie après le PBS, incluant LC3 et LC4, que nous appellerons entrée par la suite, le faisceau entrant dans la partie fibrée après cette LC4. La partie située avant le PBS comportant LC1 et LC2 sera par analogie appelée sortie. Chaque sous partie est composée de deux lames à retard à cristaux liquide orientées à 45° l'une de l'autre. Nous détaillerons dans la prochaine sous partie les calculs qui ont permis de déterminer l'orientation des LCs, permettant d'avoir la fenêtre d'opération la plus large possible [101]. Dans ces deux sous sections, les composantes du champ sont écrites d'après le référentiel classique, avec comme base les axes horizontaux et verticaux x et y .

2.3.2 Calcul déterminant l'orientation des composants

L'entrée

Juste après le PBS, le champ électrique est connu, il est polarisé linéairement selon x :

$$\vec{E}_4 = E_0 \begin{pmatrix} 1 \\ 0 \end{pmatrix} \begin{matrix} \vec{x} \\ \vec{y} \end{matrix} \quad (2.12)$$

Le passage du faisceau dans LC3 modifie sa polarisation de la manière suivante :

$$\vec{E}_5 = E_0 \begin{pmatrix} \cos\left(\frac{\varphi_3}{2}\right) \\ -i \sin\left(\frac{\varphi_3}{2}\right) \end{pmatrix}, \quad (2.13)$$

et donc après LC4, il devient :

$$\vec{E}_6 = E_0 \begin{pmatrix} \cos\left(\frac{\varphi_3}{2}\right) \\ \sin\left(\frac{\varphi_3}{2}\right) e^{i(\varphi_4 - \frac{\pi}{2})} \end{pmatrix}. \quad (2.14)$$

Dans ces conditions, avec $\varphi_3 \in [0, 2\pi]$ et $\varphi_4 \in [0, \pi]$, tous les points de la sphère de Poincaré, représentant l'espace des états de polarisation, peuvent être induits, ce qui signifie que tous les états de polarisation peuvent être injectés dans la partie fibrée de la cavité, avec une combinaison de φ_3 et φ_4 spécifique. La figure 2.22, obtenue en scannant tous les états possibles avec les différentes combinaisons de φ_3 et φ_4 possibles grâce à un programme Matlab, illustre ce résultat. D'après cette figure, nous remarquons que la densité de point est équitablement répartie autour de la sphère, sauf pour les points $x = \pm 1, y = 0$ et $z = 0$, qui présentent une densité bien plus importante. Cela signifie que pour une combinaison de φ_3 et φ_4 donnée, la chance de créer un état de polarisation correspondant à $x = \pm 1, y = 0$ et $z = 0$ est plus importante (5 fois plus) malgré le fait que la sphère de Poincaré soit recouverte. Un travail directement en état de polarisation plutôt qu'en phase amènerait donc un travail plus homogène.

La sortie

Dans cette partie, la polarisation incidente est inconnue, du fait qu'elle dépend non seulement du jeu de paramètres de φ_3 et φ_4 , mais aussi et surtout de la modification du déphasage induit par la cavité, impossible à mesurer en régime non linéaire, ce qui rend le travail calculatoire plus compliqué. Nous allons démontrer qu'à partir de n'importe quel état de polarisation incident, il est possible d'atteindre toutes les transmissions possibles par le PBS, seule la projection de l'état sur l'axe x étant importante dans notre cas. Comme pour l'entrée, LC1 et LC2 ont leurs axes rapides orientés à 45° l'un de l'autre. Dans cette configuration, le champ incident inconnu s'écrit de la manière la plus générale :

$$\vec{E}_1 = E_0' \begin{pmatrix} \cos \theta \\ \sin \theta e^{i\varphi} \end{pmatrix} \quad (2.15)$$

avec $0 < \theta < \frac{\pi}{2}$ et $0 < \varphi < 2\pi$, ce qui permet de couvrir toute la sphère de Poincaré (le champ d'entrée est supposé totalement polarisé). Après LC1, dont l'axe lent se situe sur l'axe y , le champ s'écrit :

$$\vec{E}_2 = E_0' \begin{pmatrix} \cos \theta \\ \sin \theta e^{i(\varphi + \varphi_1)} \end{pmatrix}, \quad (2.16)$$

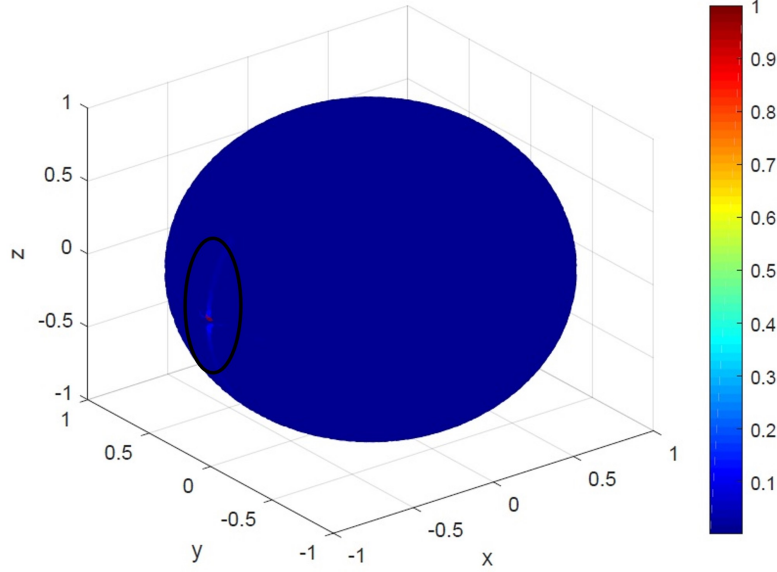


FIGURE 2.22: Illustration de la couverture de la sphère de Poincaré par les différentes combinaisons possibles de φ_3 et φ_4

Après LC2, qui a ses axes neutres orientés à 45° du plan $[x, y]$, le champ peut s'écrire de la forme suivante :

$$\vec{E}_3 = E_0' \begin{pmatrix} \cos \theta \cos \left(\frac{\varphi_2}{2} \right) - i \sin \theta \sin \left(\frac{\varphi_2}{2} \right) e^{i(\varphi + \varphi_1)} \\ \sin \theta e^{i(\varphi + \varphi_1)} \cos \left(\frac{\varphi_2}{2} \right) - i \cos \theta \sin \left(\frac{\varphi_2}{2} \right) \end{pmatrix}. \quad (2.17)$$

Le déphasage de φ est inconnu mais comme la plage de déphasage induit par LC1 est supérieur à π , il devient alors possible d'atteindre le critère suivant :

$$\forall \varphi, \exists \varphi_1 \in [0, \pi] \quad \text{and} \quad \exists k \in \mathbb{N} : \varphi + \varphi_1 = (2k + 1) \frac{\pi}{2} \quad (2.18)$$

Quand ce critère est atteint la transmission est un extremum, soit maximale soit minimale. Avec les bornes de déphasage, une seule valeur de k est possible et la transmission en intensité par le PBS (le carré du module du champ divisé par E_0') peut être simplifié par :

$$\begin{aligned} T &= \left| \cos(\theta) \cos \left(\frac{\varphi_2}{2} \right) \pm \sin(\theta) \sin \left(\frac{\varphi_2}{2} \right) \right|^2 \\ &= \cos^2 \left(\theta \mp \frac{\varphi_2}{2} \right). \end{aligned} \quad (2.19)$$

La période d'un $\cos^2$ étant égal à π , une plage de déphasage de $\varphi_2 \in [0, 2\pi]$, comme c'est le cas, permet d'ajuster la transmission entre 0 et 1. Ainsi, toutes les transmissions par le PBS peuvent être atteintes à partir de n'importe quel champ incident en sortie, juste avec l'application judicieuse d'une combinaison de φ_1 et φ_2 .

Ces sous sections ont démontrées que l'induction de toutes les fonctions de transfert de l'absorbant saturable est possible. Comme énoncé dans l'introduction, nos travaux

ont comme objectif la génération de manière automatique d'un régime laser spécifique, et l'optimisation de celui-ci selon un ou plusieurs critères désirés. Pour comprendre l'intérêt d'utiliser un algorithme d'évolution comme celui que nous présenterons ensuite, il convient tout d'abord de comprendre les nombreuses limitations des montages expérimentaux. Ceux-ci proviennent des très nombreux paramètres à considérer pour la génération d'impulsion, en plus de ceux contrôlables facilement (la puissance de pompe, les paramètres de l'absorbant saturable, la longueur de la cavité, le régime global de dispersion ou encore la séquence de éléments dans la cavité). Ces autres paramètres ont également un fort impact sur les propriétés de propagation d'un champ dans la cavité, amenant des limitations rendant impossible la prédiction par le calcul d'un résultat spécifique précis avec une configuration expérimentale et dans un environnement donné. Les trois principaux autres problèmes que nous rencontrerons sont détaillés dans la section suivante.

2.4 Les autres paramètres déterminant la propagation

— Les phénomènes d'hysteresis

Le premier problème que nous remarquons lorsqu'on tente de générer un régime spécifique avec une cavité fibrée, en agissant sur des lames d'onde par exemple, est la présence d'un fort phénomène d'hysteresis [102]. Ce phénomène est d'autant plus présent pour les états impulsions : lors de l'ajustement d'une lame d'onde en cavité, on peut sentir une optimisation de l'impulsion vers un état optimal. Cependant, si la lame est trop tournée dans un sens, le régime impulsions est perdu. Et si on tente de revenir à une combinaison de paramètres précédente en tournant la lame dans le sens inverse, l'impulsion n'est souvent plus générée, ou du moins plus avec des caractéristiques optimales. Ce phénomène est également très présent sur la puissance de pompe : pour générer un régime impulsions, un seuil minimum de puissance de pompe doit être dépassé. En revanche, une fois l'état établi, il est possible de diminuer cette puissance en dessous du seuil tout en conservant le régime de blocage de mode. Le nombre d'impulsions générées avec une puissance de pompe donnée dépendra également du chemin utilisé pour atteindre cette puissance (croissante ou décroissante) [103]. Le problème est donc facilement identifiable : une fois un état perdu lorsqu'on essaie de l'optimiser, il est impossible de le retrouver juste en revenant en arrière. Toutefois, le phénomène d'hysteresis est aussi un problème pour un algorithme d'évolution. Pour l'éviter, comme il sera décrit par la suite, un état neutre qui servira de point de départ à chaque mesures sera utilisé, garantissant une répétabilité de celles-ci.

— L'environnement

Les effets de l'environnement sur nos cavités sont les facteurs limitants les plus importants. Ce problème provient principalement de la voie utilisée pour générer des régimes de blocage de mode : la rotation non linéaire de la polarisation. Les conditions environnementales (pression, température...) influent significativement sur la réponse du milieu en terme de polarisation. Ainsi, la combinaison de paramètres à

appliquer pour l'obtention d'un état de polarisation spécifique, générant un régime de propagation spécifique, diffère selon les conditions environnementales. Cette propriété rend impossible l'utilisation d'une combinaison de paramètres type à appliquer pour générer un régime spécifique, et la création d'une bibliothèque de combinaisons n'aurait alors que peu de pertinence ni d'intérêt. Un moyen de limiter l'effet de l'environnement variable est l'isolement thermique du montage laser [33], ce qui demande un travail long et souvent inutile pour des conditions extrêmes. Toutefois, l'expérience a montré qu'un régime est capable de s'adapter dans le temps si celui-ci n'est pas perturbé, rejoignant également l'idée d'hysteresis. Par exemple, si la combinaison des paramètres reste inchangée pendant un certain temps, chaque allumage de la pompe mènera à la génération de régimes similaires, légèrement dégradés par rapport au premier régime optimisé. En revanche, avec une légère perturbation, le régime impulsionnel est généralement perdu. L'algorithme doit alors être réutilisé.

Les effets de l'environnement sont également visibles lors de perturbations extérieures. Parfois, un léger choc sur la table, une fenêtre qui s'ouvre ou une porte qui claque peut être à l'origine d'une perte d'un régime impulsionnel. L'énergie apportée par cette perturbation déstabilise le régime de propagation, qui peut alors être perdu. Un défi à relever consistera alors à développer une cavité robuste, en fixant toutes les fibres et en raccourcissant aux maximum la longueur de la cavité.

— Les contraintes mécaniques

L'impact des contraintes mécaniques sur une fibre a été démontré par Winful en 1985 [104]. Ce sont d'ailleurs ces propriétés qui permettent la construction de contrôleurs de polarisation à base de fibre optique. Chaque contrainte mécanique appliquée sur une fibre induit un déphasage par biréfringence. On peut distinguer trois types de contrainte mécanique : l'écrasement de la fibre (à l'aide d'une vis par exemple), la torsions sur elle même, et la courbure de la fibre. Ce dernier type est plutôt difficile à maîtriser, rendant impossible la création de deux cavités exactement identiques. Le rayon de courbure des fibres, qui a probablement le plus d'impact parmi les trois contraintes mécaniques, modifie donc l'état de polarisation du champ dans la cavité de manière imprévisible, augmentant encore les incertitudes quant à l'état de polarisation du régime généré d'une mesure sur une autre. Si une partie de fibre bouge, le régime généré avec une même combinaison de paramètre accessible variera. Il est impossible d'éliminer complètement ces effets, mais une fixation de ceux-ci permet, au même titre que les phénomènes d'hysteresis, de les éliminer de la liste des paramètres à investiguer pour la génération de régimes optimaux. La solution choisie sera donc de fixer toutes les fibres afin qu'elles bougent le moins possible d'une manipulation à l'autre, limitant l'impact des contraintes mécaniques.

Le paragraphe sur les paramètres non contrôlables ayant une incidence sur la génération de régimes conclue la section et le chapitre, démontrant le grand nombre de paramètres gouvernant la propagation d'un champ dans une cavité laser fibrée à blocage de mode. Ce nombre est tellement important qu'il rend impossible, en tous cas pour le mo-

ment, de construire une relation analytique entre les paramètres d'une cavité et les caractéristiques du champ qu'elle va générer. L'absence de relation ne nous permet pas de prédire avec précision quel régime va se propager avec une combinaison de paramètres accessibles à appliquer. Pour générer un régime précis, optimisé selon nos besoins, nous devons accepter ces limitations sans chercher à les éviter, en se concentrant uniquement sur les paramètres contrôlables du montage. Les algorithmes d'évolution, dont le principe sera présenté en section suivante, répondent particulièrement bien à ces problématiques, ce qui nous a orienté vers l'utilisation de ce type d'algorithme.

Chapitre 3

Qu'est ce qu'un algorithme
d'évolution ?

Ce nouveau chapitre, qui s'articule en 6 sections, présentera tout d'abord l'historique de l'utilisation d'intelligence artificielle dans le domaine de la photonique, puis détaillera le fonctionnement d'un algorithme d'évolution, en introduisant toutes les notions nécessaires à sa compréhension. Dans une troisième section, des simulations déterminant les conditions initiales de l'algorithme à utiliser pour une performance optimale seront réalisées. Ensuite, une liste des principaux instruments de mesure couramment utilisés dans le domaine de la photonique ultra rapide pour le diagnostic de régimes sera détaillée, permettant la création d'une famille de fonctions d'objectifs présentée en section 5. Enfin, la construction de l'algorithme en lui même, par une combinaison de programmes Labview/ Matlab, sera présentée en dernière section.

3.1 L'intelligence artificielle en photonique

L'intelligence artificielle a déjà une histoire relativement ancienne puisque les premiers travaux de John Holland portant sur les systèmes adaptatifs remontent à 1962 [105], et de nombreuses variantes de cet algorithme ont été développées par la suite pour tendre aux processus utilisés depuis une dizaine d'années. L'utilisation d'intelligence artificielles dans le domaine élargi de la photonique est plus récent, datant des années 2000. Aujourd'hui, l'intelligence artificielle dans le domaine de la photonique est utilisée pour de multiples tâches, comme la génération d'impulsions ultracourtes en cavité laser, la génération de supercontinua [106, 107] mais également le design d'éléments optiques [108, 109], la caractérisation de dynamiques lasers [110, 111] ou le façonnage d'impulsions (pulse shaping) [112]. Ces tâches, dépendantes d'une multitude de paramètres à considérer, ne permettent généralement pas d'établir un protocole analytique strict entre la combinaison des paramètres appliquée et le résultat obtenu. La modification d'un résultat demande alors un ajustement simultané de plusieurs paramètres, avec une combinaison à appliquer non prédictible par le calcul. L'intelligence artificielle est une solution bien appropriée pour ce type de problème. En photonique, deux types d'intelligence artificielle sont utilisés couramment : les algorithmes d'évolution (ou algorithmes évolutionnistes), et les processus d'apprentissage automatique (le "machine learning"). Nous ne détaillerons pas dans cette section le fonctionnement précis de chaque type (celui de l'algorithme d'évolution sera présenté en section suivante), nous nous concentrerons ici sur les différences entre ces deux approches, et leurs avantages respectifs.

Le deep learning, basé sur l'utilisation de réseaux de neurones [113], est un système informatique dont le fonctionnement est inspiré de celui des neurones biologiques. Les neurones sont conçus comme des automates dotés d'une fonction de transfert qui transforme ses entrées en sortie selon des règles précises. Par exemple, un neurone somme ses entrées, compare la somme pondérée résultante à une valeur seuil, et répond en émettant un signal si cette somme est supérieure ou égale à ce seuil. L'utilisation de procédés de machine learning implémentés sur des cavités laser est récente, les premiers travaux traitant de ce sujet datant de 2015 par Kutz [114], l'efficacité de l'utilisation de procédés de machine learning pour le pilotage de cavités lasers ayant été démontrée en 2014 par Brunton [115]. L'utilisation de réseaux de neurones pour contrôler la génération d'impulsion expérimentalement été démontrée par Baumeister en 2018 [116]. Plusieurs

types de réseaux de neurones existent, comme les réseaux à propagation, les réseaux convolutifs ou encore les réseaux récurrents. Plusieurs couches de neurones peuvent être utilisées, répétant ce processus plusieurs fois. De plus, l'efficacité de la transmission des signaux d'un neurone à l'autre peut varier : on parle de « poids synaptique », et ces poids sont modulés par la phase d'apprentissage, voir la figure 3.1. Ainsi, le réseau va créer des connections mathématiques (comme des synapses) entre des paramètres d'entrée et un résultat, sans que celles-ci soient véritablement identifiées. Pour créer ces connections, une première phase d'entraînement de l'algorithme est indispensable, déterminant les valeurs des seuils de chaque synapses et le poids synaptique associé à chaque neurone. Cette phase est chronophage et requiert l'enregistrement de plusieurs centaines à plusieurs milliers de données expérimentales. Une fois cette phase réalisée, le réseau est capable de retrouver, avec une bonne précision (moins de 5% d'erreur est rapporté par les travaux), une combinaison de paramètres à appliquer pour obtenir un résultat désiré, même si celui-ci n'a pas fait partie des données d'entraînement. De plus, le processus est très rapide, présentant un intérêt industriel important. Cependant, ce processus obtient de mauvais résultats lors d'extrapolations de combinaisons à appliquer en dehors des données d'entraînement. En outre, la phase d'entraînement doit être réalisée pour chaque configuration environnementale, ce qui rend le laser non adaptable. Le réseau de neurones se rapproche alors d'une cartographie des états générés, avec une interpolation des données. Ce type d'algorithme ne convient donc pas pour notre projet, celui de la création d'un laser autonome, adaptable et versatile, avec des conditions environnementales qui peuvent évoluer, ce qui explique pourquoi nous avons choisi d'utiliser un algorithme d'évolution pour nos travaux.

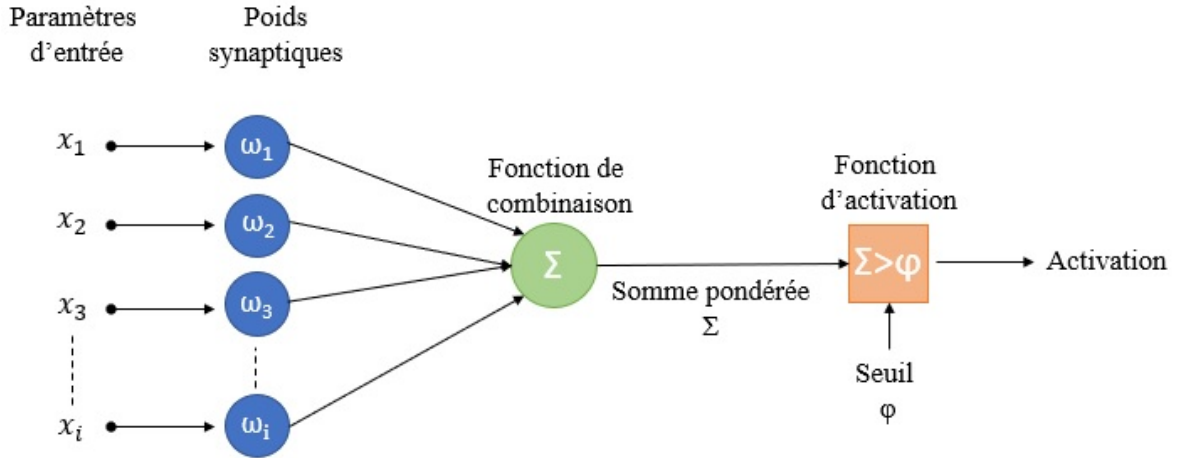


FIGURE 3.1: Illustration du fonctionnement d'un réseau de neurones

Les algorithmes génétiques, ou algorithmes d'évolutions, ont commencé à être utilisés dans le domaine de la photonique dans les années 90, avec les travaux de Mahlab [117], pour la reconnaissance de motifs écrits. Par la suite, Kihm publia ses travaux sur l'utilisation d'un algorithme génétique pour de la tomographie [118], mais c'est seulement depuis l'année 2015 que ce type d'algorithme est utilisé pour la génération d'impulsions ultracourtes directement en cavité laser fibrés [95]. Les travaux d'Andral présentent alors deux principales limitations : la première était l'importante durée d'optimisation,

ce qui poussa Winters à utiliser des dispositifs plus rapides [98], et la seconde était le manque de versatilité du système, permettant un façonnage très limité de l'impulsion, ce qui poussa Woodward à considérer plusieurs objectifs d'optimisation [97]. Ce type d'intelligence artificielle ne demande pas de phase d'entraînement, contrairement aux réseaux de neurones, mais demande en revanche la retranscription mathématique d'une caractéristique précise désirée par l'utilisateur (dans notre cas le régime impulsif est souhaité par exemple). Cette caractéristique sera optimisée par l'algorithme à travers une fonction de mérite, ce qui n'est pas le cas pour les réseaux de neurones. Une bibliothèque des différentes fonctions de mérite (ou d'objectif) utilisées au préalable est résumée dans les travaux de Meng parus en 2020 [119]. Le processus de l'algorithme d'évolution débute sur la création de combinaisons de paramètres aléatoires, et l'utilisation de boucles de rétroaction pour diagnostiquer la qualité des régimes générés. Ainsi, un algorithme d'évolution est beaucoup plus lent que les réseaux de neurones (de quelques minutes à dizaines de minutes comparé à quelques secondes), mais ne débute sur aucune données pré-enregistrées, ce qui confère au laser une grande adaptabilité. Ce type d'algorithme sied complètement à nos attentes, à condition que les éléments utilisés permettent une optimisation dans une durée convenable, c'est à dire quelques minutes. Ce défi sera relevé dans cette thèse, comme nous le verrons dans les prochains chapitres.

3.2 Fonctionnement de l'algorithme d'évolution

Cette section explique le fonctionnement de l'algorithme d'évolution de manière générale, en introduisant ses notions clés et ses étapes essentielles. Cet algorithme sera testé de manière expérimentale dans un chapitre suivant. Un algorithme d'évolution (AE), ou algorithme évolutionniste, est une catégorie d'algorithme dont le fonctionnement s'inspire de la théorie de l'évolution. Il s'inspire donc de l'évolution naturelle d'une population dans un environnement donné, sélectionnant pour chaque génération d'individus les plus robustes, qui se reproduisent pour donner lieu à une génération future. Il s'agit donc de méthodes de calculs bioinspirés. Ces algorithmes sont dits stochastiques, car ils reposent sur des processus aléatoires utilisés de manière itérative. C'est la raison pour laquelle les AE sont la solution à nos limitations précédemment énoncées, l'aspect chaotique du fonctionnement laser pouvant être assimilé à un fonctionnement aléatoire. Ce type d'algorithme est couramment utilisé pour des problèmes d'optimisation lorsque la méthode des gradients ne peut pas être utilisée, et lorsqu'un problème ne comporte pas de solution analytique. Le principe de tels algorithmes est illustré dans la figure 3.2.

La première étape d'un AE consiste à créer, le plus souvent de manière aléatoire, une population d'individus, dans notre cas des régimes de propagation, composés eux même de gènes. Les gènes sont dans notre cas les paramètres ajustables à considérer, c'est à dire les 4 phases appliquées par les LCs dont l'application génère l'individu en question. Cette population va constituer la génération 0 de notre algorithme. Sur le schéma, la génération 0 est composée de 4 individus situés en haut à gauche de couleurs différentes. Les nombres correspondent aux phases φ_1 , φ_2 , φ_3 et φ_4 appliquées par les LCs. En d'autres termes, un grand nombre d'état de polarisation induisant chacun un état de propagation est généré en appliquant successivement des combinaisons de phases

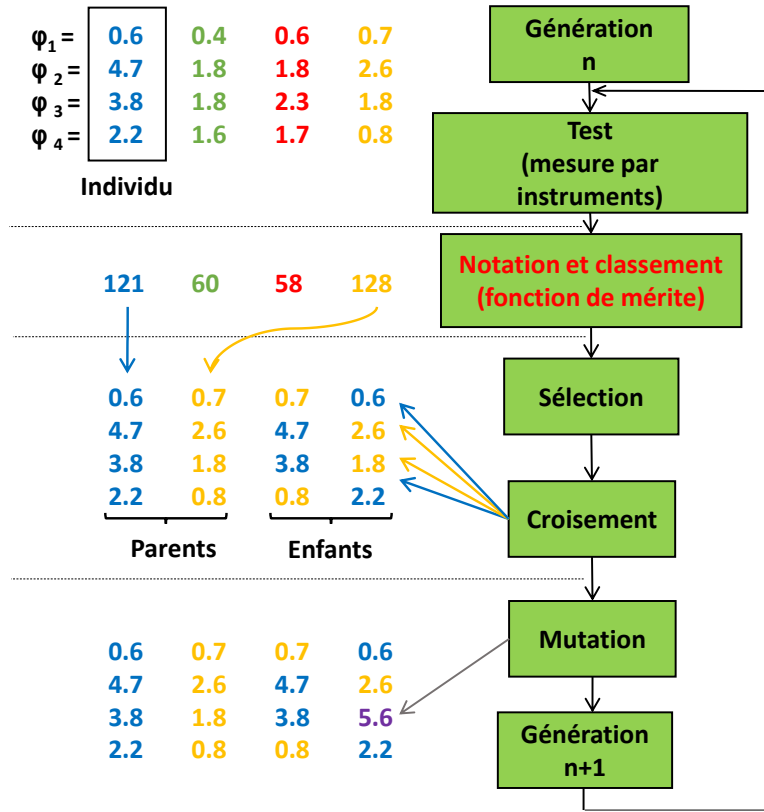


FIGURE 3.2: Illustration du fonctionnement d'un algorithme d'évolution

choisies aléatoirement ou non. Dans notre cas, considérant nos objectifs d'atteindre un état de propagation optimal en toutes circonstances, et n'ayant aucune donnée sur les zones de polarisation à forte probabilité de succès, la génération 0 est créée de manière aléatoire. Pour avoir une probabilité plus importante de sélectionner une zone à forte probabilité de succès, une génération 0 comportant un nombre important d'individus est nécessaire. Cette étape qui constitue l'étape initiale de notre AE n'est pas répétée dans la suite de l'algorithme.

La deuxième étape d'un AE est le test expérimental de chaque individu par les instruments de détection choisis au préalable. Le but de cette étape est d'attribuer un score à chaque individu, permettant de classer les individus entre eux dans la population. Le score de chaque individu est accordé selon une fonction cruciale dans l'algorithme : c'est la fonction de mérite ou fonction d'objectif. Cette fonction est l'ingrédient clé dans l'algorithme car elle détermine vers quel état va converger celui-ci, en d'autres termes quel individu sera considéré comme optimal. Elle est déterminée par l'utilisateur selon ses objectifs d'optimisation, et doit être modifiée si d'autres types d'optimisations sont voulues. Pour nos travaux, le principal problème a donc été de construire une fonction de mérite qui tende le plus vers notre objectif, c'est à dire une fonction qui sous-évalue les régimes continus et Q switch mode locking, et qui alloue les meilleures notes aux régimes de blocage de mode. Dans nos expériences, avec une analyse visuelle, il est facile de le déterminer : un train d'impulsion sur l'oscilloscope et un spectre élargi sur l'analyseur de spectre. En revanche, retranscrire ce raisonne-

ment en une fonction mathématique est plus compliqué. Il convient donc de faire un minutieux travail préalable pour sélectionner les instruments de détection capable de construire une fonction de mérite menant à notre objectif, demandant une phase de test de chaque type de régime de propagation. Dans ces travaux, une famille fonction de mérite multi-objectif permettant l'optimisation d'un régime impulsif incluant une seconde caractéristique sera présentée dans une section suivante.

Une fois le score de mérite évaluée pour chaque individu, ces derniers sont classés selon leur score. Cette étape permet de sélectionner des individus avec les meilleurs scores dans la population constituant la génération n . Les individus sélectionnés sont appelés parents. Les parents sélectionnés sont inclus dans la génération suivante, permettant la sauvegarde de la séquence génétique des meilleurs individus d'une génération sur une autre. En revanche, les individus les moins bien notés sont éliminés. Une nouvelle population d'enfants est alors créée tout d'abord par croisement des gènes des parents, comme illustré dans la figure 3.2. Les gènes des parents sont sélectionnés aléatoirement. L'hypothèse derrière le croisement des gènes parentaux pour l'optimisation est la suivante : un gène peut être responsable de la bonne notation du parent, déterminant l'aire des états de polarisation à investiguer. Il faut donc l'identifier et le préserver pour les générations futures. Nous entendons par "génération future" toute génération qui n'est pas la génération 0. Ainsi, si un enfant hérite de ce gène, il aura également un score élevé lors des futures notations, parfois même plus que ses parents. Ces derniers seront alors sélectionnés, permettant une optimisation. Cependant, si seul le croisement composait la seule source de génération d'enfant, l'optimisation resterait très rudimentaire, tournant autour de quelques individus seulement. Il convient donc de rajouter une dernière étape cruciale à la génération d'enfants.

Cette dernière étape est celle de la mutation. Elle laisse une part aléatoire dans l'optimisation et éventuellement un changement de zone des gènes à explorer si la zone initiale n'est pas optimale. En effet, dans la pratique, de nombreux maxima locaux sont présents. Lorsque le point d'optimisation des parents est déterminé d'après un maximum local, sans mutation, l'AE va converger vers ce maximum qui n'est pas le maximum global. Une certaine proportion d'enfants, déterminée au préalable est mutée : c'est à dire que certains de leurs gènes vont être modifiés. Ces gènes sont soit modifiés de manière totalement aléatoire, soit pour une variation aléatoire autour du gène parental avec une faible amplitude. En pratique, on privilégiera une mutation totalement aléatoire lorsque l'algorithme cherche encore grossièrement son point de convergence dans les premières générations, puis la seconde option est utilisée en fin de procédure d'optimisation pour affiner légèrement le score autour d'un point optimal. Sur la figure, le gène violet en bas est muté, provenant d'une mutation totalement aléatoire. Parents et enfants composent alors la génération $n+1$ dont la taille reste constante, et le processus de test/sélection/reproduction se reproduit pendant un nombre fixé de générations. Après un certain nombre de génération, l'AE converge vers un optimum et à la fin du processus, une population optimale est créée comportant l'individu optimal de l'algorithme.

Les paramètres initiaux de l'AE (nombre maximum de génération, taux d'enfants mutés, nombre de parents pour chaque génération, taille de la population initiale, créée aléatoirement ou non...) sont autant de facteurs déterminants quant à la création rapide et répétable d'un individu optimal. Il convient donc de faire des travaux préliminaires en cavité pour déterminer les conditions initiales de l'algorithme induisant les meilleures performances. Une partie de simulations a été également réalisée pour les déterminer, présentée dans la section suivante.

3.3 Simulations du processus de l'algorithme

Pour avoir une première indication sur les conditions initiales de l'algorithme à utiliser, un travail simulateur est réalisé. Cependant, compte tenu des limitations énoncées précédemment, il est difficile de prédire précisément la propagation des impulsions dans des cavités laser à fibre par simulation. En effet, les absorbants saturables modélisés dans les simulations sont généralement des fonctions de transfert simplifiées. De plus, les capacités de calculs dont nous disposons ne sont pas suffisantes pour calculer des équations (équation 2.7) qui se rapprocheraient le plus d'une propagation réelle dans une durée convenable. Pour modéliser un AE dans un temps de simulation convenable, dans cette sous section, la propagation du champ dans la cavité sera très simplifiée, avec une simple phase induite par la partie fibrée. La transmission du PBS sera calculée facilement d'après l'équation 2.13. La fonction de mérite reste la dernière caractéristique à déterminer pour cette partie simulateur. Pour cette première investigation, une fonction de mérite composite sera construite. Celle-ci est liée au champ transmis par le PBS, c'est à dire la composante horizontale de la polarisation E_{3x} dans la figure 2.21. Cette fonction est détaillée dans la sous section suivante.

3.3.1 Détermination de la fonction de mérite composite

Pour cette section, nous souhaitons une fonction de mérite composite liée à la transmission en champ du faisceau par le PBS. Sur un modèle de calcul similaire à celui de la section 2.3.2, celle-ci est calculée, mais cette fois en assumant que la partie fibrée de la cavité induit un déphasage choisi aléatoirement pour chaque optimisation φ_{cav} (déphasage entre les axes x et y). Ainsi, les calculs se simplifient et $\vec{E}_1$ dépend directement de φ_3 , φ_4 et φ_{cav} . Il devient donc possible de calculer facilement $\vec{E}_3$ en fonction de φ_1 , φ_2 , φ_3 , φ_4 , et φ_{cav} .

Nous débutons donc le calcul de la transmission par $\vec{E}_4$ connue, puis nous procédons comme dans la section 2.3.2 pour calculer les champs successifs. Nous obtenons donc après calcul :

$$\vec{E}_6 = E_0 \begin{pmatrix} \cos\left(\frac{\varphi_3}{2}\right) \\ \sin\left(\frac{\varphi_3}{2}\right) e^{i(\varphi_4 - \frac{\pi}{2})} \end{pmatrix}. \quad (3.1)$$

La partie fibrée induit son déphasage, donnant :

$$\vec{E}_1 = E_0 \begin{pmatrix} \cos\left(\frac{\varphi_3}{2}\right) \\ \sin\left(\frac{\varphi_3}{2}\right) e^{i(\varphi_4 + \varphi_{cav} - \frac{\pi}{2})} \end{pmatrix}. \quad (3.2)$$

ce qui nous permet de calculer :

$$\vec{E}_3 = E_0 \begin{pmatrix} \cos\left(\frac{\varphi_3}{2}\right) \cos\left(\frac{\varphi_2}{2}\right) - i \sin\left(\frac{\varphi_3}{2}\right) \sin\left(\frac{\varphi_2}{2}\right) e^{i\Delta\varphi} \\ \sin\left(\frac{\varphi_3}{2}\right) \cos\left(\frac{\varphi_2}{2}\right) e^{i\Delta\varphi} - i \cos\left(\frac{\varphi_3}{2}\right) \sin\left(\frac{\varphi_2}{2}\right) \end{pmatrix} = E_0 \begin{pmatrix} E_{3x} \\ E_{3y} \end{pmatrix}. \quad (3.3)$$

Avec $\Delta\varphi = \varphi_1 + \varphi_4 + \varphi_{cav} - \pi/2$. La transmission en champ correspond alors à $|E_{3x}|$. Cependant, utiliser une fonction de mérite similaire à celle-ci, c'est à dire dont le comportement est linéaire, serait trop éloigné de la réalité, car elle ne présenterait qu'un seul maximum (quand $|E_x| = 1$). En réalité, les problématiques en lasers fibrés comportent toujours des maxima locaux, justifiant l'utilisation d'un algorithme d'évolution pour être piloté. En conséquence, une modulation de la transmission, donnée en figure 3.3, a été construite pour représenter ces propriétés.

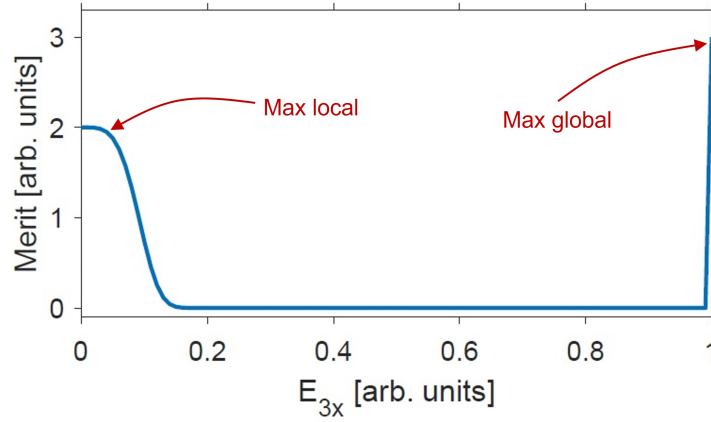


FIGURE 3.3: Graphique présentant la fonction de mérite composite utilisée pour nos premières simulations, construite dans le but de tester l'influence des conditions initiales sur les performances de l'algorithme et présentant un maximum local et un maximum global

Avec cette fonction modulée, la chance de converger vers un maximum local (pour $|E_{3x}| = 0$) est plus importante que pour le maximum global (à $|E_{3x}| = 1$). Elle permet donc de tester les performances de l'algorithme dans des cas extrêmes. Le score maximal à atteindre est de 3 pour $|E_{3x}| \approx 1$, tandis qu'un maximum local moins localisé à 2 se situe aux alentours de $|E_{3x}| \approx 0$. Avec cette fonction de mérite composite, l'impact des conditions initiales sur la performance de l'algorithme peut être étudié. Ce travail sera présenté en sous section suivante.

3.3.2 Etude de l'impact des conditions initiales de l'algorithme

Cette étude considère 4 paramètres initiaux : la proportion d'individus mutés par génération, la taille de la population initiale, la proportion de parent pour chaque génération et la taille des populations futures. La figure 3.4 présente l'impact de ces 4 paramètres, tous les autres paramètres étant fixes. Pour commencer, les conditions initiales sont fixées à une population initiale de 30 individus, des populations futures de 20 individus, un taux de mutation de 50% et une proportion de 25% de parents pour chaque population future (soit 5 parents). Le taux de mutation correspond ici à la proportion d'enfants ayant un gène muté, et cette mutation est complètement

aléatoire. A chaque étude, un seul paramètre varie. La figure 3.4 (a) présente l'impact de la proportion d'enfants mutés par génération sur l'optimum atteint, la figure 3.4 (b) présente celui de la proportion de parents dans la population d'une génération, la figure 3.4 (c) présente l'impact de la taille des populations futures et la figure 3.4 (d) présente celui de la taille de la population initiale. Le score optimal présenté sur les figures est obtenu d'après une moyenne de 500 optimisations sur 20 générations.

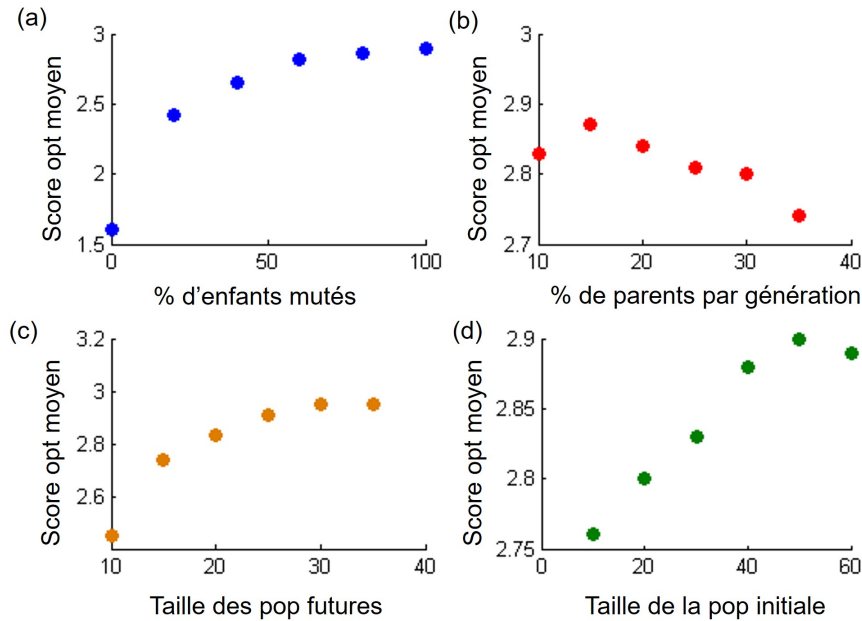


FIGURE 3.4: Graphiques présentant l'impact des conditions initiales sur la performance de l'algorithme

— Proportion d'enfants mutés

La figure 3.4(a) indique que plus la proportion aléatoire d'enfant mutés est importante, plus le score de mérite optimal atteint est important. L'évolution est logarithmique, avec un plateau de score optimal atteint à un taux d'enfants mutés de 70%. Une donnée non prise en compte dans ces simulation est la répétabilité du procédé, plus la proportion d'enfants mutés est importante, plus l'erreur standard l'est. La proportion optimale semble donc être la plus petite après atteinte du plateau, soit 70%. Cette valeur sera utilisée pour les simulations suivantes et pour l'algorithme en cavité.

— La proportion de parents dans la génération future

La figure 3.4(b) indique que le score de mérite optimal décroît avec l'augmentation de la proportion de parents dans la génération. Une proportion de 15% de parents dans chaque génération semble être donc la valeur optimale.

— La taille des population futures

Logiquement, plus la taille des populations futures est importante, plus le nombre d'enfants testés est important et donc plus le score optimal est important, comme nous le montre la figure 3.4(c). Toutefois, une donnée importante à considérer en conditions expérimentales est la durée d'optimisation. Plus une population est importante plus elle prendra de temps à être testée. Il semble donc évident que

si les populations sont très importantes, l'algorithme tendra vers une solution optimale plus importante, mais dans un temps plus long, avec le risque de dérives expérimentales. La meilleure performance de l'algorithme dans notre cas, est l'atteinte d'un résultat optimal dans un temps minimal. Il convient donc de trouver le meilleur compromis optimisation/temps avec ces simulations. La valeur la plus intéressante pour nous sera alors la valeur la plus faible après atteinte du plateau, soit de 20 individus.

— **La taille de la population initiale**

Sans surprise, la figure 3.4(d) indique que plus la population initiale est importante, plus le score optimal sera important. L'impact de la population initiale dans l'optimum atteint semble minime (2% de gain pour une population multipliée par 5), mais présente un optimum pour une population initiale de 50 individus. Dans notre étude, l'algorithme atteint dans la très large majorité des cas une valeur optimale globale avec une population initiale supérieure à 40 individus. Cependant, dans un cas plus réaliste, avec de nombreux optima locaux, cette valeur sera plus importante.

Ces premiers travaux ont permis de fixer certaines conditions initiales pour optimiser au maximum les performances de l'algorithme. Une dernière investigation sur le pourcentage de réussite de l'algorithme, c'est à dire la convergence vers un score maximal en fonction du nombre d'individus composant la génération 0, est réalisé dans le paragraphe suivant. La figure 3.5 présente les résultats de cette étude, avec la figure 3.5(a) et (b) présentant respectivement la proportion d'atteinte de maxima globaux et locaux avec un taux de mutation de 0%, tandis que la figure 3.5(c) et (d) présentent les mêmes résultats pour un taux de mutation de 70%. Les résultats sont détaillés dans les points suivants.

- La comparaison entre les figures 3.5 (a) et (c) confirme tout d'abord l'importance du taux de mutation élevé : dans le premier cas, le score optimal est atteint dans 60% des cas (avec une population initiale de 50 individus) tandis que le second cas converge vers l'optimum dans 97% des cas. Un fort taux de mutation est donc essentiel à une bonne performance de l'algorithme.
- Dans les deux cas, plus la population initiale est importante, plus la proportion de convergence vers un optimum l'est. Sans mutation la proportion augmente de manière quasi linéaire avec la taille de la population initiale, tandis qu'avec un fort taux de mutation un plateau à partir de 30 individus dans la génération 0 semble être atteint.
- La comparaison entre (b) et (d) indique que sans mutation la proportion de convergence vers le maximum local augmente avec l'augmentation de la taille de la génération 0, tandis que la proportion décroît avec mutation. Cela s'explique par un taux de convergence vers le maximum global plus important dans le second cas.
- Le pourcentage d'échec de l'algorithme, c'est à dire la non atteinte d'un maximum, qu'il soit global ou local, est bien plus important pour l'étude sans mutation (30% d'échec sans mutations pour 2% avec pour une population initiale de 50 individus).

Compte tenu de ces deux études, une population initiale de 30 individus et des populations futures de 20 individus composés de 3 parents et 17 enfants, dont 70 % ont un gène

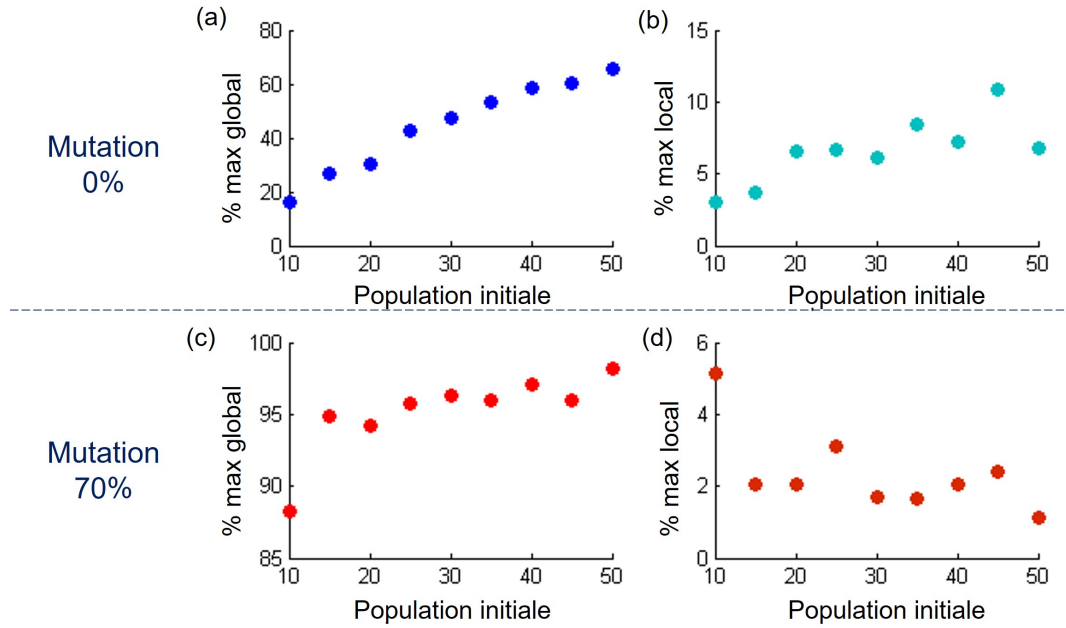


FIGURE 3.5: Graphiques présentant l'impact des conditions initiales sur la performance de l'algorithme, avec (a) la proportion de maximum global atteint par l'algorithme avec un taux de mutation de 0%, (b) sa proportion de maximum local, (c) la proportion de maximum global atteint par l'algorithme avec un taux de mutation de 70%, et (d) sa proportion de maximum local

muté seront utilisées pour des performances optimales. Nous rappelons que ces valeurs dépendent aussi et surtout de la fonction de mérite utilisée. Pour une utilisation de l'algorithme en cavité laser, ces paramètres devront être ajustés. Avec ces paramètres, des simulations sont réalisées. Une courbe d'optimisation type avec une telle fonction de mérite et de telles conditions initiales est donnée en figure 3.6.

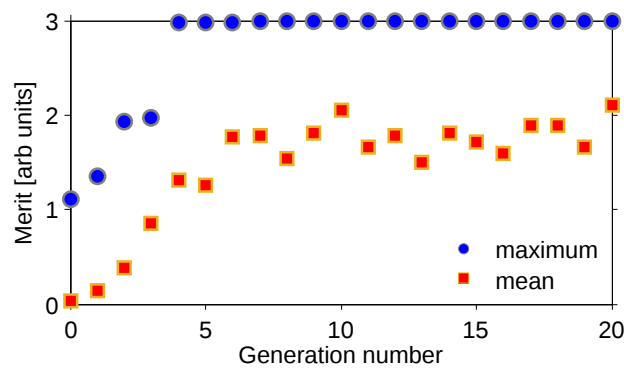


FIGURE 3.6: Courbe d'optimisation simulée, présentant l'évolution du meilleur individu (en bleu) et la moyenne des scores (en rouge) de chaque génération

Nous remarquons une optimisation, avec un individu optimal obtenu à la génération 5. L'optimisation est alors de 300% par rapport au meilleur individu de la génération 0. Sachant que les meilleurs individus sont conservés comme parent, il est logique de voir le maximum de chaque génération augmenter de génération en génération. Le score

moyen est optimisé également, mais fluctue d’une génération sur une autre. Ce comportement est dû au fort taux de mutation utilisé.

Cette étude a permis de déterminer les conditions initiales d’utilisation optimale de l’algorithme. Nous les adapterons pour une utilisation expérimentale, ceux-ci pouvant varier selon l’objectif ciblé et la cavité utilisée. L’objectif désiré, traduit par la fonction de mérite, s’appuie sur le diagnostic réalisé de chaque individu testé. Pour les caractériser, certains instruments de mesure sont couramment utilisés dans le domaine de l’optique ultrarapide. Ceux-ci sont décrits dans la section suivante.

3.4 Les principaux instruments de mesure

En travail expérimental, la fonction de mérite se construit sur des mesures expérimentales de certaines caractéristiques, avec pour objectif de maximiser celles désirées par l’utilisateur. Pour réaliser des mesures variées et caractériser au mieux un régime, plusieurs instruments de mesures sont usuellement utilisés dans les travaux traitant des impulsions ultracourtes en cavité laser fibrée. Ces instruments, qui permettront de construire une fonction de mérite capable d’optimiser un régime selon nos objectifs, sont décrits dans la sous section suivante.

3.4.1 L’autocorrélateur

Encore aujourd’hui, les photodiodes les plus rapides ont un temps de réponse de l’ordre d’une dizaine de picosecondes. Par exemple, dans notre cas, la photodiode utilisée à une résolution de 22 ps (45 GHz). Il est donc impossible de caractériser des impulsions de quelques centaines de femtosecondes avec des instruments directs comme ceux-ci. Pour obtenir une mesure temporelle sur des impulsions ultrarapides, on utilise alors un autocorrélateur optique non linéaire du second ordre multicoup, schématisé en figure 3.7.

Ce dispositif, fonctionnant comme un interféromètre de Mach-Zehnder, est basé sur la corrélation croisée d’un faisceau par lui même : il est tout d’abord séparé en deux faisceaux d’égales puissances, réparties sur deux bras. Un des bras (bras 2) possède une ligne à retard variable, induisant une différence de marche par rapport au premier (bras 1). La polarisation sur chaque bras est ajustée pour optimiser un doublage de fréquence dans un cristal doubleur Bêta-Borate de Baryum (BBO) de type 2. Grâce à celui-ci, l’autocorrélation intensimétrique est mesurée. La mesure par le tube photomultiplicateur PM enregistre l’intensité de ce doublage de fréquence $S(\tau)$, proportionnelle au produit de corrélation des intensités : $S(\tau) \propto \int I(t)I(t - \tau)dt$. La différence de marche induite par le second bras permet le balayage temporel de l’impulsion du bras 1 par rapport à celle du bras 2, en d’autres termes l’autocorrélation. La durée d’impulsion à mi-hauteur pourra alors être estimée d’après la durée de la trace d’autocorrélation (sur le schéma 3.7 T_0), en la divisant par un facteur K dépendant du profil de l’impulsion. Pour une gaussienne, ce coefficient est de 1.414 et pour une sécante hyperbolique, il est de 1.55. Cette différence apporte également une des limites de l’autocorrélation : elle donne des informations sur la durée d’impulsion mais en faisant une hypothèse sur sa forme. Cette technique n’apporte également aucune information sur la phase temporelle

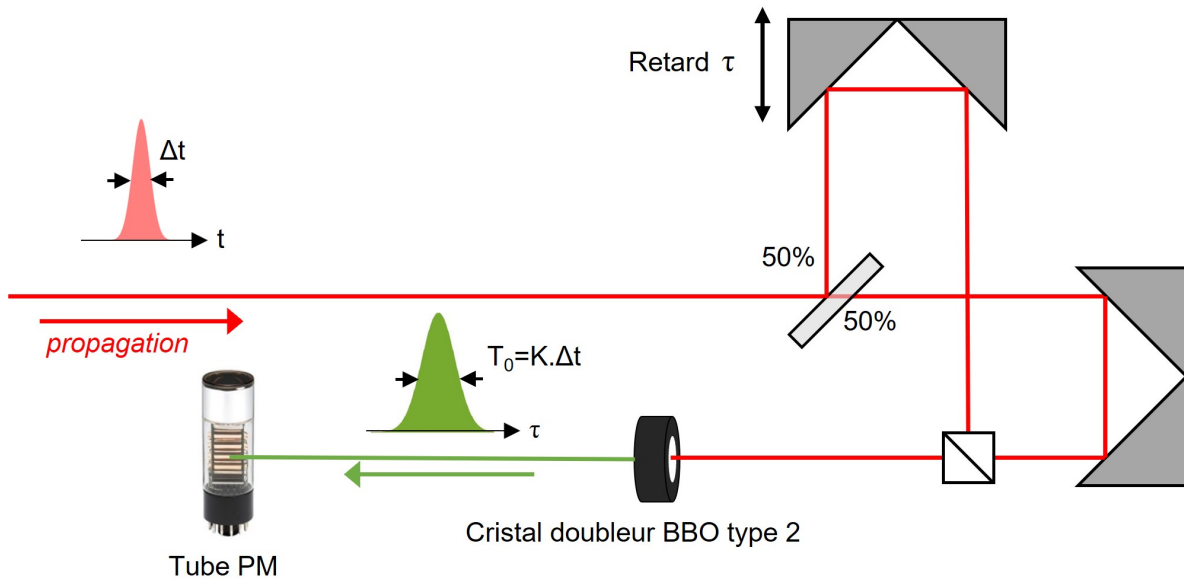


FIGURE 3.7: Montage d'un autocorrélateur intensimétrique multicoup

ou spectrale de l'impulsion.

Dans le cas où le régime mesuré est multi impulsionnel, l'autocorrélation non linéaire permet la détermination du nombre d'impulsions ainsi que la visualisation de leur organisation entre elles. Pour un champ comportant n impulsions, la trace d'autocorrélation comportera alors $2n - 1$ pics. Prenons l'exemple d'un faisceau à 2 impulsions, l'impulsion n°1 et n°2. Une illustration de l'autocorrélation intensimétrique de deux impulsions est donnée en figure 3.8.

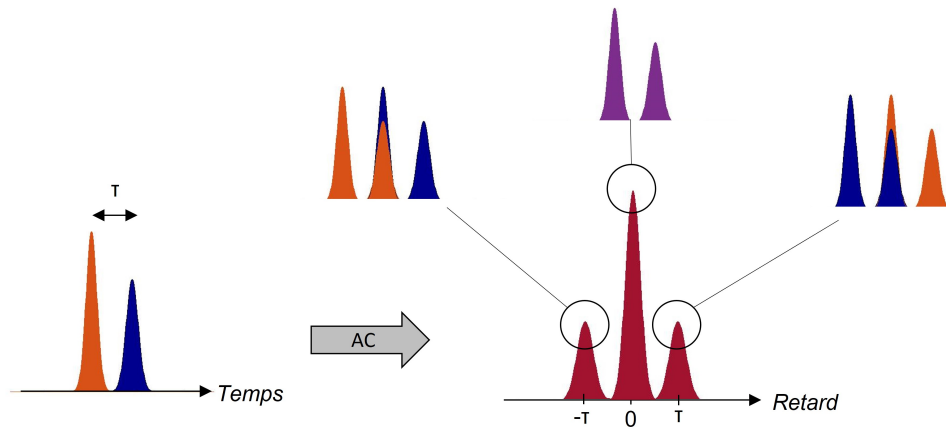


FIGURE 3.8: Schéma représentant la trace d'autocorrélation intensimétrique d'un champ à 2 impulsions

Le premier pic d'autocorrélation correspond au recouvrement de l'impulsion n°1 (en bleu sur la figure) avec l'impulsion n°2 (en orange). Puis le second pic d'autocorrélation, ayant une amplitude plus importante que le premier, correspond au recouvrement des deux impulsions entre elles, et enfin le troisième pic d'autocorrélation correspond au

recouvrement de l'impulsion n°2 avec l'impulsion n°1, donnant la même amplitude que le premier. La trace d'autocorrélation comporte alors 3 pics, un pic intense au délai 0 et deux pics plus faible autour de ce premier à délai correspondant au retard entre les impulsions, et symétriques par rapport au pic central.

3.4.2 L'analyseur de spectre optique (OSA)

Une seconde manière de caractériser une impulsion est d'enregistrer son spectre optique, dans le domaine des longueurs d'onde ou de fréquence. On utilise pour cela un analyseur de spectre optique (OSA) ou spectromètre . Dans nos expériences, l'OSA utilisé est un Anritsu 0.6-1.75 μm avec une résolution de 0.07 nm. Le fonctionnement de cet OSA repose sur la diffraction par un réseau du faisceau, chaque longueur d'onde étant déviée selon un angle différent. Le réseau est monté sur un moteur permettant la rotation de celui-ci et donc le balayage de toutes les longueurs d'ondes possibles. Un miroir concave est utilisé pour la collecte du faisceau vers la détection de l'intensité. Une acquisition prend environ une seconde. Sachant que dans un train d'impulsion, les impulsions sont espacées seulement de quelques dizaines de nanosecondes, l'OSA ne peut en réalité que détecter une valeur moyennée sur 1 seconde du spectre optique. Dans ce cas, il est impossible de détecter des variations spectrales d'une impulsion à une autre (sur deux tours de cavités consécutifs), et donc de collecter des informations précises sur la stabilité du faisceau. Il est également impossible d'obtenir des informations sur la phase spectrale avec cet instrument, le signal mesuré étant l'intensité spectrale $I(\omega) = |E(\omega)|^2$.

3.4.3 L'analyseur de spectre radiofréquence (ESA)

Le fonctionnement d'un analyseur de spectre électrique (ESA) repose sur le même fonctionnement qu'un OSA, mais en analysant le faisceau dans un domaine fréquentiel. Les mesures réalisées peuvent aller de quelques dizaines de Hertz à plusieurs centaines de GHz. L'ESA possédé par notre laboratoire est Agilent, allant de 9 kHz à 3 GHz. Un ESA permet la mesure du spectre radiofréquence (RF) d'un train d'impulsion, donnant des informations sur le taux de répétition des impulsions, et sur la stabilité de celui-ci. Pour un régime de blocage de mode harmonique, la mesure du spectre RF comprend toutes ses composantes harmoniques. L'intensité d'un pic RF à la fréquence de répétition fondamentale d'une cavité est essentielle, car une intensité élevée de celle ci est caractéristique du régime de blocage de mode fondamental. La maximisation de cette caractéristique revient à maximiser la puissance d'une impulsion d'un blocage de mode fondamental.

3.4.4 L'oscilloscope avec photodiode ultrarapide

Cet élément est indispensable pour détecter et caractériser un train d'impulsions, avec une durée d'affichage très rapide. Un oscilloscope permet la visualisation d'un champ électrique créé par une photodiode, convertissant le faisceau optique incident en courant électrique. Comme énoncé dans les sous sections précédentes, ce dispositif n'est pas suffisant pour mesurer le profil d'intensité d'une impulsion femtoseconde. En effet, dans le cas d'impulsions beaucoup plus courte que le temps de réponse le signal mesuré

vaut $S(t) \propto R(t) \int I(t')dt'$, avec $S(t)$ le signal mesuré et $R(t)$ la réponse de la photodiode. Il permet de visualiser certaines de ses caractéristiques, comme sa stabilité. Pour nos expériences, un oscilloscope Lecroy à bande passante de 6 GHz. Il permet également de réaliser des calculs directs du faisceau mesuré, comme par exemple la transformée de Fourier du train d'impulsion, correspondant à son spectre radiofréquence. De plus, une technique appelée Dispersive Fourier Transformation technique (DFT) permet d'accéder à des informations sur les composantes spectrales grâce à un oscilloscope.

En 1973, Desbois et al utilisèrent deux réseaux de diffraction pour créer un élément dispersif et enregistrer le spectre d'une impulsion picoseconde sur un oscilloscope [120]. Cette technique, popularisée en 2013 par Goda et Jalali [121], utilise une fibre dispersive et s'inspire de la diffraction d'onde en champ lointain : une longue pièce de fibre très dispersive est utilisée pour étaler le spectre de l'impulsion temporellement permettant d'obtenir sa transformée de Fourier, qui en champ lointain correspond à son spectre. Une illustration du montage pour cette technique est donnée en figure 3.9.

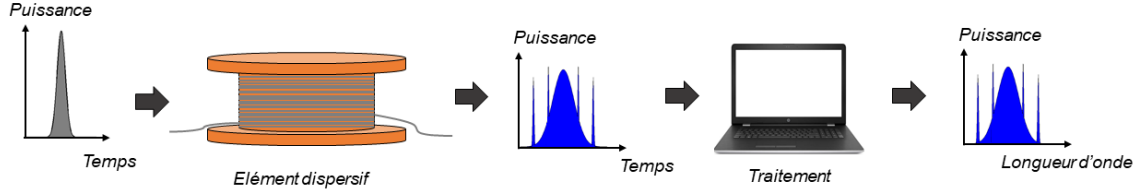


FIGURE 3.9: Illustration de la technique de la DFT

Il est ainsi possible de mesurer le spectre de chaque impulsion du train mais dans un domaine temporel, qui peut se convertir en spectre optique après traitement de ces données. Pour la conversion, une transformée de Fourier inverse, débouchant sur la relation simple 3.4 est utilisée :

$$\lambda = \frac{t}{DL} + \lambda_0 \quad (3.4)$$

Avec D la dispersion de la fibre dispersive (en $\text{ps} \cdot \text{nm}^{-1} \cdot \text{km}^{-1}$) et L sa longueur (en km). λ_0 correspond à la longueur d'onde centrale du spectre (en nm) et t le temps (en ps). La résolution spectrale du dispositif est de $1/D \cdot L \cdot BW_{osc}$, ce qui correspond dans notre cas à 1.5 nm. L'utilisation de cette technique permet donc de mesurer le spectre optique de l'impulsion sans moyennage, contrairement à l'OSA, et dans temps plus court mais avec une résolution moins bonne. Pour un champ autour de $2\mu\text{m}$, les pertes des fibres dispersives en silice usuellement utilisées deviennent trop importantes, l'utilisation de réseaux de diffraction sera alors privilégié. Ainsi, l'utilisation d'un oscilloscope ultrarapide avec traitement DFT des données permet d'obtenir des informations à la fois temporelles et spectrales, permettant l'utilisation d'un seul instrument de détection plutôt que trois (oscilloscope, ESA et OSA).

Ces 4 instruments couramment utilisés dans les travaux traitant de l'optique ultrarapide permettent de caractériser au mieux un régime laser. En piochant dans ces différents diagnostics, il est alors possible d'optimiser un régime selon n'importe quelle caractéristique désirée. Mais pour ce faire, il convient de déterminer une fonction de

mérite (ou d'objectif) capable de favoriser ces caractéristiques au détriments de celles non désirées. La prochaine section présente la réponse à cette problématique.

3.5 Famille de fonction de mérite

La première et principale caractéristique des régimes souhaités pour nos travaux est le blocage de mode stable. Cette caractéristique est le fondement même de nos travaux et de la plupart des travaux sur l'optique ultrarapide : comment générer des régimes de blocage de mode à la demande, sans réglage manuel ? De plus, nous ne souhaitons pas uniquement un critère binaire sur la génération de régimes impulsionnels, nous souhaitons que le régime soit le plus stable et le plus intense possible. Le diagnostic utilisé doit donc apporter en premier lieu un meilleur score aux régimes impulsionnels par rapport aux régimes continus ou aux régimes de Q switched mode locking, puis accorder un meilleur score encore aux régimes impulsionnels les plus stables et intenses. L'idéal serait une caractéristique accordant un score intermédiaire aux régimes de Q switch mode locking, souvent assez proches des régimes de blocage de mode conventionnels. Une caractéristique connue des régimes de blocage de mode et utilisée couramment dans les travaux [94, 95, 123], permettant de répondre à ces attentes est la puissance du spectre radio fréquence ($P_{RF}(FSR)$, en dBm) au taux de répétition fondamental de la cavité laser. Elle peut être obtenue directement par un ESA, ou calculée d'après la trace directe de l'oscilloscope en faisant la transformée de Fourier (FFT). Cette seconde manière de procéder sera utilisée pour la compacité de la boucle de rétroaction à considérer (un seul instrument de mesure), ainsi que pour sa rapidité de réponse (quelques dizaines de ms pour calculer le spectre RF par l'oscilloscope contre environ une seconde pour mesurer le spectre RF par un ESA). Nous notons cependant que les valeurs des puissances calculées en absolue ne sont pas comparables avec celles mesurées par ESA, seulement les variations de ces scores seront à considérer. Pour plus de précision, à cette puissance $P_{RF}(FSR)$, nous soustrairons la puissance maximale du fond continu et bruité $P_{RF}(Bck)$ (en dBm également, Bck pour Background). Ainsi, les régimes impulsionnels uniques et stables seront favorisés. Cette première composante de nos fonctions de mérite est illustrée en figure 3.10.

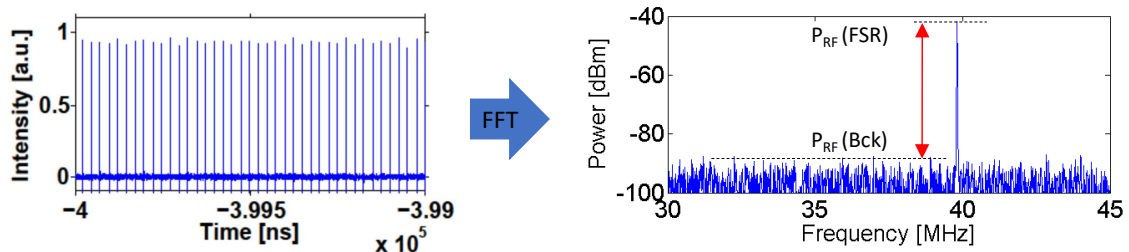


FIGURE 3.10: Illustration de la méthode de calcul de la puissance radiofréquence d'après une trace d'oscilloscope

Ensuite, notre objectif n'est pas uniquement de générer des impulsions stables. Il est aussi d'optimiser une caractéristique de ce régime, désirée par l'utilisateur. Une

fonction de mérite avec une seule contribution (la puissance RF) ne permettrait d'autogénérer que des impulsions stables, mais sur lesquelles aucune autre caractéristique intéressante peut être optimisée. Ainsi, il convient de construire une famille de fonctions de mérite multi-objectives, composée de deux contributions. Cette seconde contribution changerait selon les désirs de l'utilisateur, et permettrait de générer une caractéristique précise d'un régime impulsif. Pour résumer, la première contribution de la fonction de mérite permet de sélectionner les régimes de blocage de mode stables, et la seconde contribution permet d'optimiser une caractéristique désirée (puissance moyenne, largeur spectrale, intensité de doublage de fréquence...) de ces régimes impulsifs. La famille de fonction de mérite utilisée tout au long de cette thèse est présentée en équation 3.5 .

$$F_{Merit} = P_{RF}(FSR) - P_{RF}(Bck) + \alpha U \quad (3.5)$$

Dans cette fonction, U représente la caractéristique du régime impulsif à optimiser. U provient de différents diagnostics, et peut être n'importe quelle caractéristique désirée. Il est alors possible d'adapter cette fonction de mérite pour toutes les utilisations. Dans nos travaux, U sera la largeur spectrale provenant de la trace DFT, l'intensité de seconde harmonique (SHG) ou encore la puissance moyenne mesurée d'après un puissance-mètre. Entre les deux contributions, un poids relatif α est utilisé pour favoriser une contribution par rapport à l'autre. La variation de ce poids relatif a une importance capitale, et peut être utilisée pour contrôler U plus finement, comme nous le verrons dans les prochaines sections. Dans certains cas, les deux contributions ne peuvent pas être optimisées complètement. Le score de mérite optimal sera alors un compromis entre les deux scores provenant de chaque contribution. En d'autres termes, il faudra accepter d'avoir un P_{RF} moins important pour un U plus grand. La valeur de ce paramètre dépend également de la nature de U , et peut être négative. Un problème peut alors se poser pour des valeurs trop importantes en positif ou en négatif. Il faut en effet que la proportion de la seconde contribution reste suffisamment faible, pour que la sélection des régimes de blocages de mode reste prépondérante. Prenons l'exemple d'une optimisation sur la puissance moyenne : celle-ci est plus importante pour des régimes continus que pour des régimes impulsifs. Il faut donc que le poids de la seconde contribution reste suffisamment faible pour que la fonction de mérite continue à sélectionner des régimes impulsifs. Mais il faut également que la seconde contribution garde un poids suffisamment important pour garder un rôle dans l'optimisation. Des tests expérimentaux se doivent donc d'être réalisés pour chaque caractéristique (pour chaque " U ") et déterminer la fenêtre de valeurs du paramètre α qui convient. Idéalement, il faut que la répartition des deux contributions dans le score soit entre 60/40 et 80/20 % du score total.

3.6 Développement de l'algorithme pour la cavité

3.6.1 Boucle de rétroaction

Pour implémenter un AE sur notre cavité, le montage d'une boucle de rétroaction extra cavité est requis. Cette boucle doit être capable de réaliser les 4 actions suivantes : (i) donner l'ordre aux paramètres accessibles de la cavité (ici les LCs) d'appliquer une combinaison, (ii) appliquer cet ordre, (iii) mesurer le régime généré suite à celui-ci et (iv)

traiter les données mesurées (dans le but d'attribuer le score de mérite à l'individu en question grâce à la fonction de mérite). Les étapes de sélection, classement, croisement et mutation sont réalisées qu'au début de chaque génération, et sont détaillées dans les sections précédentes. Les autres actions requises de la boucle sont schématisées dans la figure 3.11.

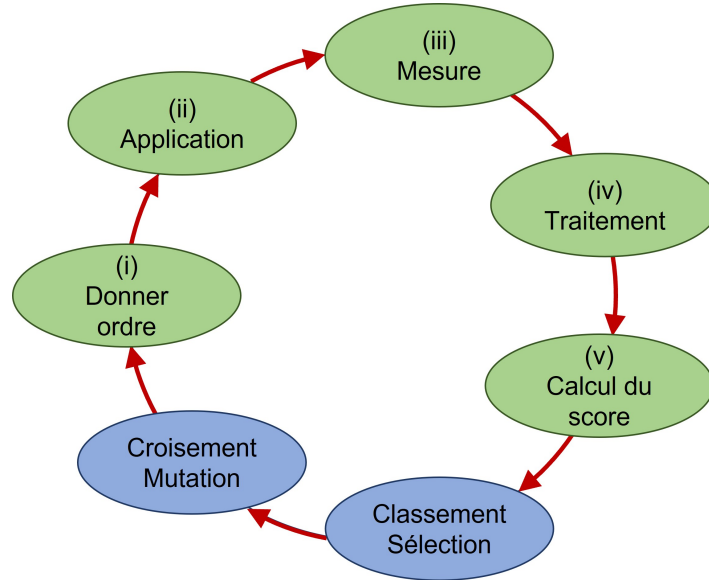


FIGURE 3.11: Schéma présentant les différentes actions requises pour notre boucle de rétroaction

- (i) L'ordre d'appliquer un individu est donné par informatique, ainsi que les étapes propres à l'algorithme en soi (création de la population initiale, notation, tri, croisement, mutation) (voir section 3.2). L'algorithme est construit sous le logiciel Labview, et comporte un appel à un script Matlab pour les étapes énoncées ci-dessus. La connexion entre l'oscilloscope et l'ordinateur se fait par câble ethernet en PC2I, permettant une interaction rapide entre ces deux éléments. L'ordinateur permet donc de donner l'ordre d'appliquer un individu aux contrôleurs des LCs, mais aussi de donner l'ordre aux instruments de mesure de lancer l'acquisition d'une trace après établissement et stabilisation de cet individu. L'ordinateur permet également de recevoir les caractéristiques de l'individu et de lui attribuer son score de mérite après calcul de celui-ci. Cette étape marque la fin d'une boucle, et l'ordre d'appliquer un autre individu est donné, jusqu'à la fin de l'algorithme. L'ordinateur permet donc de construire cette boucle en entrée et sortie.
- (ii) L'application de l'ordre est réalisée par les contrôleurs des LCs. Dans notre cas, on utilise deux contrôleurs de la marque Meadowlark, délivrant un courant alternatif compris entre 0V et 10V. Chaque contrôleur délivre une tension à deux LCs, soit celles en entrée (LC3 et LC4 en figure 2.21) soit en sortie (LC1 et LC2). Deux lames Thorlabs (LC1 et LC4), chacune placée directement avant et après le PBS induisent une phase entre 0 et 2π , tandis que deux lames Meadowlark (LC2 et LC3) permettent d'induire un déphasage entre 0 et π . Comme démontré en section 2.3.2, cette combinaison des lames de phase permet de couvrir entièrement l'espace des polarisations. Dans les deux cas, la phase minimale est induite par une tension

maximale (ici 10V) et la relation entre la phase induite et la tension délivrée n'est pas linéaire. Cette relation est donnée en figure 3.12 . Cette relation non linéaire

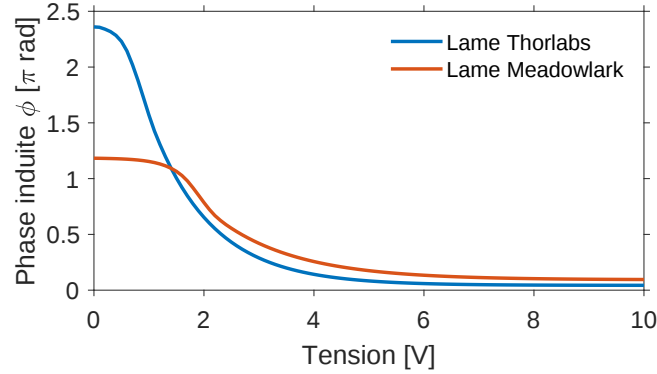


FIGURE 3.12: Graphique présentant l'évolution de la phase induite par les différentes LCs en fonction de la tension appliquée dessus, donné par les constructeurs

nous incite à travailler en phase, voire en état de polarisation pour l'entrée, plutôt qu'en tension pour créer une population initiale aléatoire plus représentative des états de polarisation. Cette population en phase est ensuite convertie en tension pour donner les ordres aux contrôleurs, permettant une représentation de l'espace des états de polarisation plus homogène.

- (iii) La construction d'une fonction de mérite requiert une détermination des instruments de détection à utiliser. Dans cette section, les travaux du chapitre 4 seront considérés pour illustrer les propos, mais la boucle de rétroaction sera légèrement modifiée au fur et à mesure des expériences selon les besoins. La première idée venue lors de cette thèse était l'optimisation d'un régime impulsionnel en fonction de sa largeur spectrale. Un dispositif capable de mesurer la largeur spectrale d'un régime est alors requis. Mais, comme présenté en section précédente, un second critère de mérite, la puissance radiofréquence, doit être utilisé pour sélectionner les régimes de blocage de mode. Ces deux objectifs seront directement mesurés par l'oscilloscope, comme détaillé en section 3.4. La mesure de l'individu et son traitement se fait donc directement par un oscilloscope ultrarapide, après établissement et stabilisation de celui-ci. Pour caractériser au mieux une dynamique impulsionnelle ultrarapide, une photodiode rapide est requise. Une photodiode ayant une réponse de l'ordre de la vingtaine de picoseconde est alors utilisée (avec une bande passante de 45 GHz), branchée sur un oscilloscope LeCroy (40 GS.s^{-1} et une bande passante de 6 GHz). Deux photodiodes sont utilisées pour la détection, sur deux chaînes de l'oscilloscope. La chaîne n°1 mesure la trace directe du faisceau, tandis que la chaîne n°2 mesure la trace DFT de celui-ci, après son passage dans une fibre de dispersion $D_{DFT} = 100 \text{ ps.nm}^{-1}.\text{km}^{-1}$ de $L_{DFT} = 1.1 \text{ km}$ de longueur.
- (iv) Le traitement du signal se fait directement grâce à l'oscilloscope, qui permet le calcul de la puissance RF du signal en utilisant la transformée de Fourier de la voie n°1. Une mesure de la largeur spectral se fait directement sur la voie n°2. D'après ces deux grandeurs, renvoyées à l'ordinateur, il est possible de construire la fonction de mérite capable de sélectionner un régime de blocage de mode et de l'optimiser selon sa largeur spectrale avec un unique instrument de mesure. Cette

boucle permet donc de simplifier le montage expérimental au maximum. Le calcul de la fonction de mérite (équation 3.5) se fait par l'ordinateur, ce qui permet de boucler le circuit d'information.

Ainsi, la boucle de rétroaction expérimentale pour les premiers travaux, dont le schéma est donné en figure 3.13, est montée.

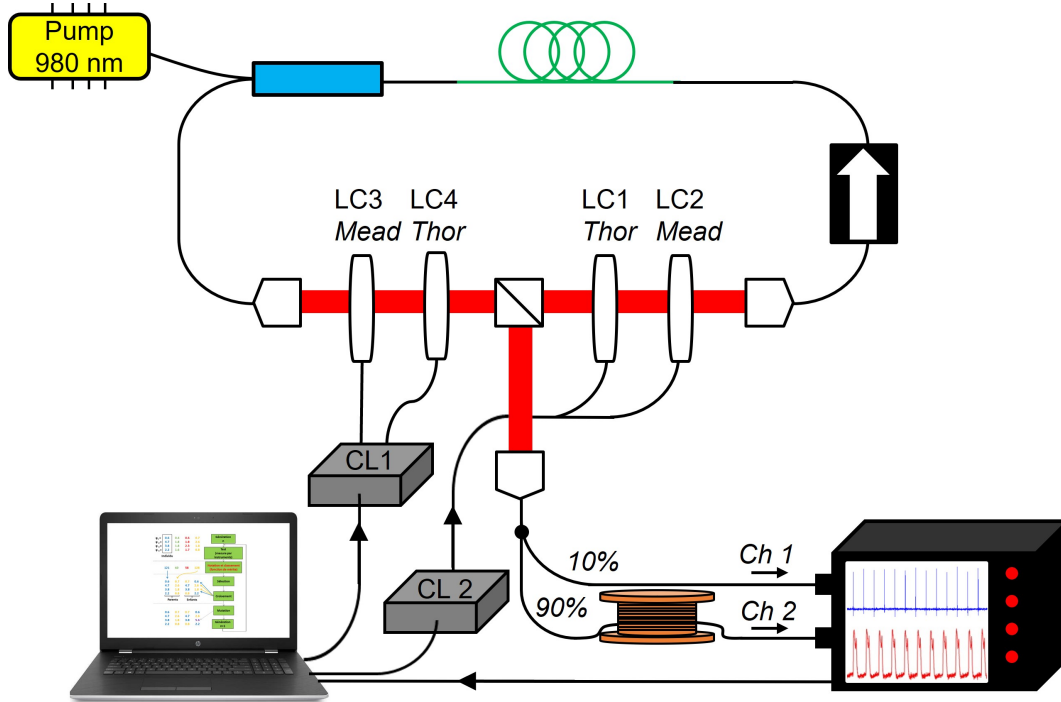


FIGURE 3.13: Illustration de la boucle de rétroaction montée, comportant un oscilloscope ultrarapide, deux contrôleurs de lames à cristaux liquides contrôlant 4 lames, et un dispositif DFT permettant la mesure de données spectrales

La réalisation de toutes ces actions n'est pas instantanée. Les limitations électroniques de chaque instrument, couplées au temps nécessaire à la stabilisation d'un individu, rend la boucle chronophage. Dans ces travaux, chaque test prend au minimum 1.8 sec. Un chronogramme de l'opération est donné en figure 3.14. La limitation électronique des contrôleurs des LCs (CL1 et CL2 sur la figure) implique que chaque ordre donné par ceux-ci doit être séparé d'au minimum 300 ms. Les ordres envoyés aux 2 LC ne peuvent donc pas être simultanés. Ainsi, un individu demandant deux ordres (individu neutre et individu testé) par CL (pour chaque LC), l'application d'un individu se fait en 600 ms sans temps de stabilisation. On appelle temps de stabilisation le temps pour un état impulsionnel de se construire et de se stabiliser dans la cavité. Ce temps se couple avec le temps de régénération des CL. Une étude préliminaire a donné un temps de construction d'un régime de blocage de mode de 500 ms maximum. Nous décidons donc de fixer ce temps de stabilisation à 600 ms, laissant le temps pour le régime de se stabiliser. L'état neutre, étant un état continu très bruité, ne nécessite pas de temps de stabilisation. L'acquisition et le traitement (Acq) durent environ 300 ms, mettant fin à la boucle.

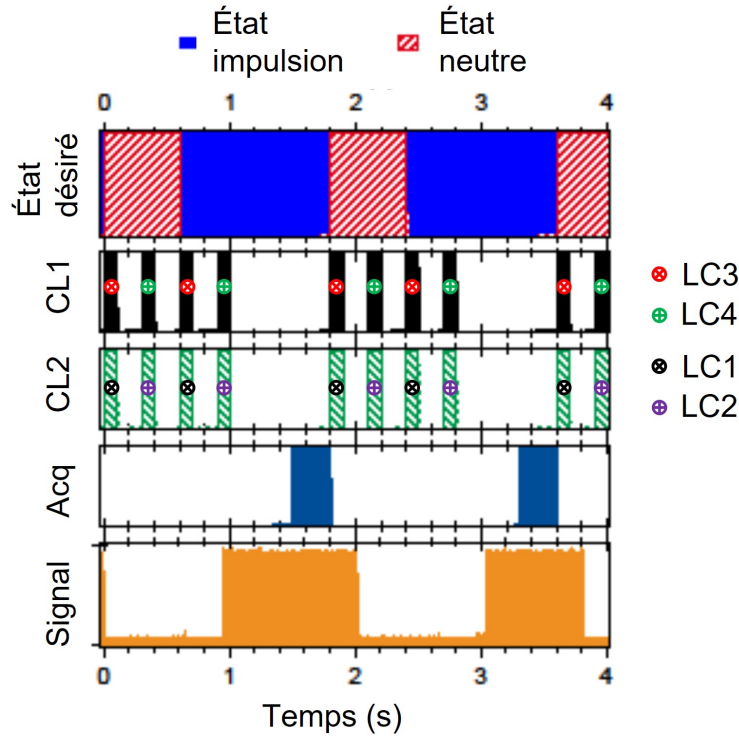


FIGURE 3.14: Schéma d'un chronogramme de notre boucle de rétroaction, présentant la séquence des différentes actions réalisées

Le temps de réponse des LCs (d'environ 20 ms) est négligeable dans notre cas. Dans ces conditions, une optimisation de 10 générations serait longue de 9 min. Ainsi, compte tenu de toutes les données citées, un algorithme d'évolution a été développé sous le logiciel Labview. Un algorithme a déjà été développé dans les précédents travaux [94,95]. Cependant, les composants de la boucle de rétroaction étant différents dans notre cas, un grand nombre d'ajustements a dû être réalisé pour rendre l'algorithme conforme à nos utilisations. Certaines fonctionnalités ont également été ajoutées. Nous les expliquerons et détaillerons dans la sous section suivante.

3.6.2 Algorithme développé

La face avant de cet algorithme est donnée en figure 3.15. Encore une fois, cet algorithme correspond à celui utilisé pour les premiers travaux, donné pour illustration. Pour les derniers travaux (chapitres 6 et 7), l'algorithme a dû être complété et adapté, mais le principe et la majeure partie des variables restent identiques pour tous les travaux. Nous énumérerons et commenterons rapidement les différentes variables de cet algorithme dans la liste ci après.

- **L'initialisation** : cette partie permet d'ouvrir une connexion entre l'ordinateur et les autres composants de la boucle de rétroaction. "Device Count" et "Located USB Devices" sont deux indicateurs de la bonne connexion entre le PC et les contrôleurs de LC. Deux LED virtuelles s'allument lorsque les deux CLs sont opérationnels. "Add Scope" est une commande permettant d'ouvrir une liaison entre le PC et l'oscilloscope, et éventuellement de choisir celui qui sera utilisé

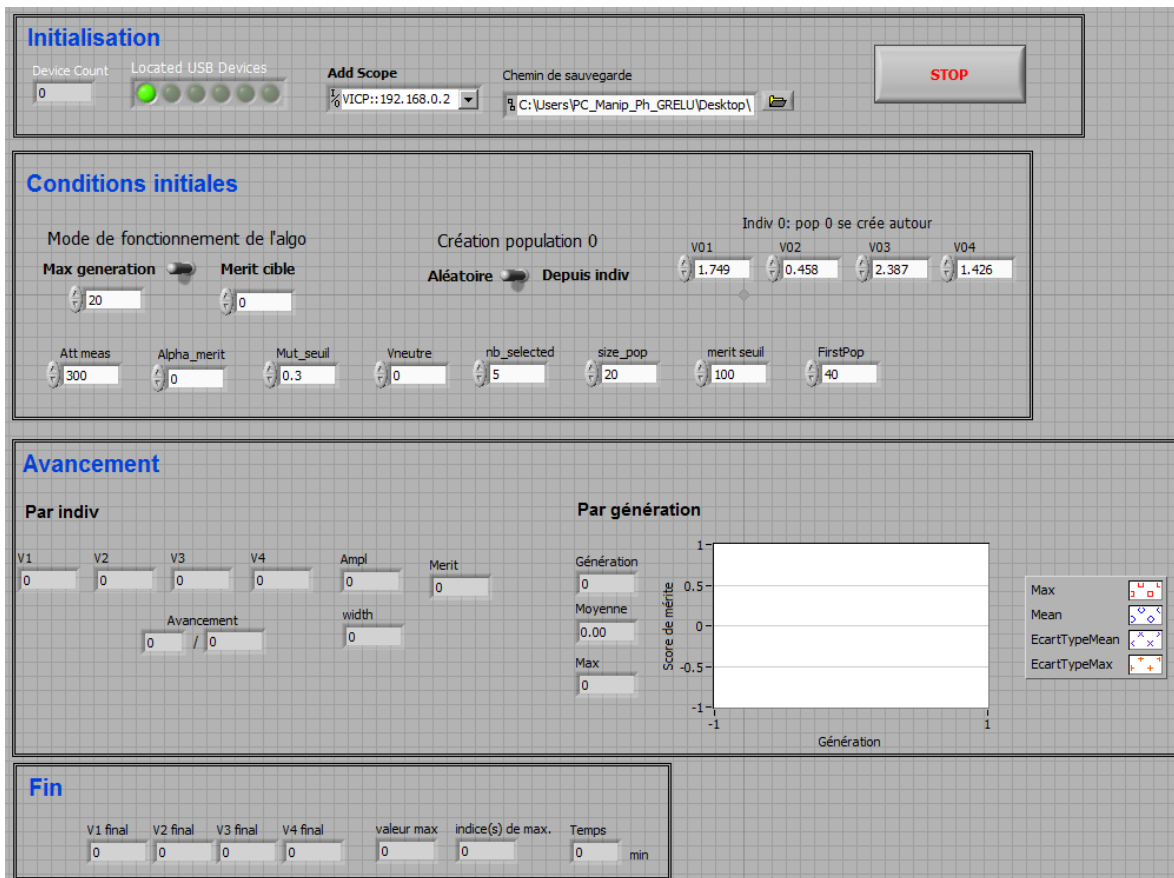


FIGURE 3.15: Schéma présentant la face avant de notre algorithme sous Labview

pour la mesure dans le cas où plusieurs oscilloscopes seraient connectés. “Chemin de sauvegarde” permet de choisir dans quel dossier s’enregistreront les fichiers de l’avancement de l’algorithme et ses résultats.

- **Les conditions initiales** : cette partie présente les différentes fonctionnalités de démarrage et de fonctionnement de l’algorithme. Toutes ces variables sont des commandes. “Att meas” induit un temps d’attente entre l’ordre et la mesure, ce qui correspond au temps de stabilisation décrit dans la section précédente. “Alpha_merit” est le poids relatif du second objectif de la fonction de mérite, “Mut_seuil” correspond au taux de gènes mutés dans une population future, “FirstPop” et “size_pop” déterminent le nombre d’individu pour la population initiale et les populations futures, et “nb_selected” détermine le nombre d’enfants par génération. “merit seuil” correspond à un seuil de la fonction de mérite après lequel l’amplitude de mutation des gènes n’est plus totalement aléatoire, mais restreint autour d’un point de fonctionnement intéressant. Ce seuil correspond par exemple au score minimum d’un état de blocage de mode. Une fois le score atteint, l’algorithme va investiguer autour des individus avec une amplitude faible (5 %). Deux boutons poussoirs permettent de choisir entre deux modes de fonctionnement de l’algorithme et deux manières de créer la population initiale. Le fonctionnement “Max generation” permet de réaliser une optimisation sur un nombre de génération donné. Ce type de fonctionnement permet d’optimiser au maximum un régime, mais est coûteuse en temps. C’est pourquoi parfois une op-

timisation par fonctionnement “Merit cible” sera réalisée, permettant de stopper l’optimisation une fois qu’un score de mérite prédéterminé est atteint. On réalise alors un gain de temps. Une autre manière de la faire est de créer une population initiale autour d’un point de fonctionnement déjà connu “Depuis indiv”, qui permet de construire une population initiale moins importante. Dans le cas d’une création “Aléatoire”, un important nombre d’individus dans la population initiale est recommandé, pour balayer au maximum les différents états de polarisation. Nous détaillerons ces différents fonctionnement dans des sections suivantes.

- **L’avancement** : cette partie est constituée de tous les indicateurs sur l’avancement de l’optimisation, par individu et par génération. Les gènes en tension de chaque individu sont donnés par “V1”, “V2”, “V3” et “V4”, avec leur puissance radiofréquence “Ampl”, leur largeur spectrale “width” et leur score de mérite “Merit”. Un graphique tracé en temps réel permet de suivre visuellement la progression de l’algorithme, donnant pour chaque génération le score de mérite moyen “Moyenne” et le score maximum “Max” de chaque génération.
- **La fin** : indique simplement les caractéristiques de l’individu optimal, avec son score de mérite “valeur max”, sa position dans la dernière génération “indice(s) max”, et la durée d’optimisation “Temps” en minute.

Cette section clôture le dernier chapitre de cette thèse présentant les informations théoriques nécessaires à la compréhension des travaux pratiques. Le chapitre suivant présentera les résultats des premières optimisations réalisées en cavité avec l’algorithme précédemment présenté.

Chapitre 4

Optimisation du blocage de mode fondamental d'une cavité simple via l'algorithme

L'objectif des travaux de ce chapitre sera d'utiliser l'algorithme avec une boucle de rétroaction pour optimiser un régime impulsif selon sa stabilité et d'en contrôler sa largeur spectrale. Pour cela, la première étape est de construire un montage laser correspondant à nos objectifs. Ce développement est décrit en section suivante.

4.1 Développement du montage laser

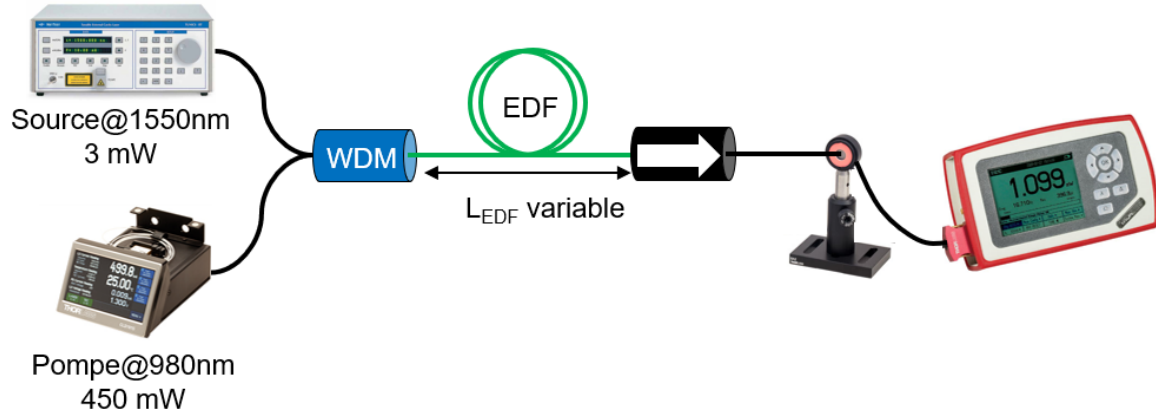
Pour reprendre et résumer les problématiques théoriques, la génération d'un régime impulsif dépend de la combinaison de 4 paramètres : la dispersion, les non linéarités, le gain et les pertes. La gestion de la dispersion et le gain sont apportés par une fibre fortement dopée Erbium (EDF), tandis que le dispositif induisant le mécanisme d'absorbant saturable apporte les pertes essentielles à la sélection des blocages de mode. Si ce dernier est fibré, comme pour les contrôleurs de polarisation, il modifiera l'équilibre dispersion/non linéarités. Si on considère que la longueur de fibre monomode reliant les composants entre eux est fixée à son minimum, que la séquence des éléments utilisés est déterminée, et que la nature de la fibre dopée est sélectionnée, une problématique pratique reste encore à élucider : quelle longueur de fibre dopée doit être utilisée pour générer les régimes souhaités ? La réponse à cette problématique est détaillée dans la sous section suivante.

4.1.1 Détermination de la longueur de fibre dopée

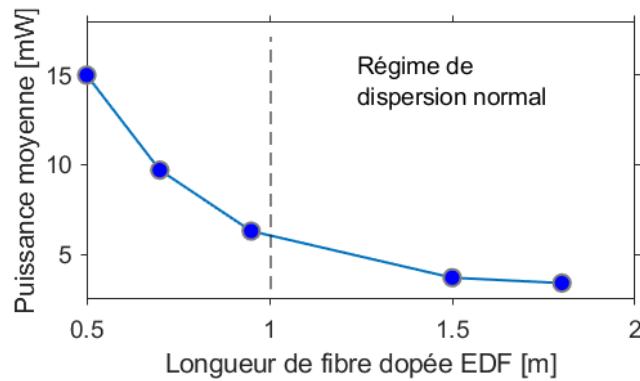
En ce qui concerne le gain, une idée reçue consiste à penser que plus la longueur de fibre dopée est importante, plus l'amplification sera importante. Cette proposition est fautive. En effet, la puissance de pompe étant fixée, elle sera donc en grande partie absorbée par le début de la fibre dopée, créant un fort gain dans cette partie de fibre, mais seulement très faible pour la fin de la fibre si l'apport d'énergie par la pompe n'est plus suffisant. Cependant, des phénomènes de réabsorption du champ peuvent apparaître si la fibre dopée est trop longue, induisant des pertes non négligeables et non compensées par le gain de la fibre. En d'autres termes, la balance gain/perte de la fibre dopée n'est pas constante sur sa longueur. Si la longueur de fibre est trop importante pour une puissance de pompe fixe, l'amplification du faisceau diminue.

Il est très difficile d'estimer la longueur optimisant l'amplification par calcul, en raison du peu de données constructeur relatives au gain dépendant de la puissance pompe et de la longueur de fibre. Dans la pratique, on va donc utiliser une technique de coupure par étapes successives pour la déterminer. Un montage composé d'une source à 1560 nm, d'une pompe à 980 nm, d'un multiplexeur et du morceau de fibre dopée est utilisé et présenté en figure 4.1 (a). La source Tunic d'Anritsu apporte un régime continu d'une puissance de 3 mW à 1550 nm et une diode laser Thorlabs apporte une puissance de pompe de 450 mW à 980 nm. Un puissancemètre est utilisé pour mesurer la puissance moyenne du faisceau en sortie de montage, permettant de calculer l'amplification de la source, apportée par une fibre dopée coupée par étapes. Dans notre montage, le puissancemètre utilisé ne permet pas la discrimination des longueurs d'onde. C'est la raison pour laquelle un isolateur est utilisé, permettant le filtrage des résidus

de pompe non absorbés par la fibre dopée grâce à sa bande passante réduite autour de la fréquence du champ laser. Ainsi, il est possible de mesurer sa puissance directement.



(a)



(b)

FIGURE 4.1: Schéma du processus de coupures successives utilisé pour déterminer la longueur optimale de gain, avec (a) son montage expérimental et (b) son résultat

Le graphique présentant la puissance moyenne du faisceau amplifié en fonction de la longueur de fibre dopée est donné en figure 4.1 (b). Nous notons que la fibre dopée n'est pas soudée avec le WDM et l'isolateur, ce qui induit des pertes non négligeables à ce niveau, et explique les puissances moyennes mesurées inférieures à celles attendues (seulement de 500% à $L_{EDF} = 0.5$ m). Cela ne change pas la conclusion de l'étude sur la longueur de fibre optimale. Le graphique montre que la longueur de fibre optimisant l'amplification se situe en dessous de 0.5 m. L'optimum n'est donc pas atteint dans cette étude.

Cependant, comme énoncé dans les sections précédentes, l'amplification n'est pas la seule caractéristique à prendre en compte. Un régime de dispersion global normal est souhaité, favorisant la génération de régimes mono-impulsionnels. Compte tenu des longueurs de fibres monomodes utilisées (environ 2 m), à dispersion anormale, et comme présenté dans la figure, un régime de dispersion normale requiert une longueur minimale de 1 m d'EDF. Avec ces deux problématiques, et pour obtenir un bon équilibre entre régime de dispersion normale et gain important, une EDF de 1.2 m sera utilisée pour construire la cavité laser.

4.1.2 Montage laser expérimental

Dans nos travaux, nous souhaitons simplifier au maximum la cavité développée dans les précédents travaux [33, 94, 95], notre but étant de générer des dynamiques plus simples, avec des solitons uniques, stables et robustes. Un soliton est stable s'il est identique à chaque tour de cavité pour un endroit de la cavité donné, tandis que la robustesse se définit comme la capacité d'un soliton à faire face aux perturbations extérieures. Ces deux problèmes sont souvent liés, ont les mêmes origines mais n'apportent pas les mêmes conséquences. Dans ce souci de stabilité et robustesse, une cavité courte en diminuant les non linéarités par le raccourcissement de la fibre monomode (SMF) pour connecter les éléments au maximum est souhaitée. Dans notre cas, la longueur de fibre monomode sera réduite au minimum à 2 m. La fibre EDF 80 de ofs optics de 1.2 m, fortement dopée erbium 3^+ est donc utilisée. Cette fibre a une GVD de $\beta_{2,EDF} = 0.06\text{ps}^2.\text{m}^{-1}$, soit environ 3 fois la dispersion de la fibre SMF en valeur absolue.

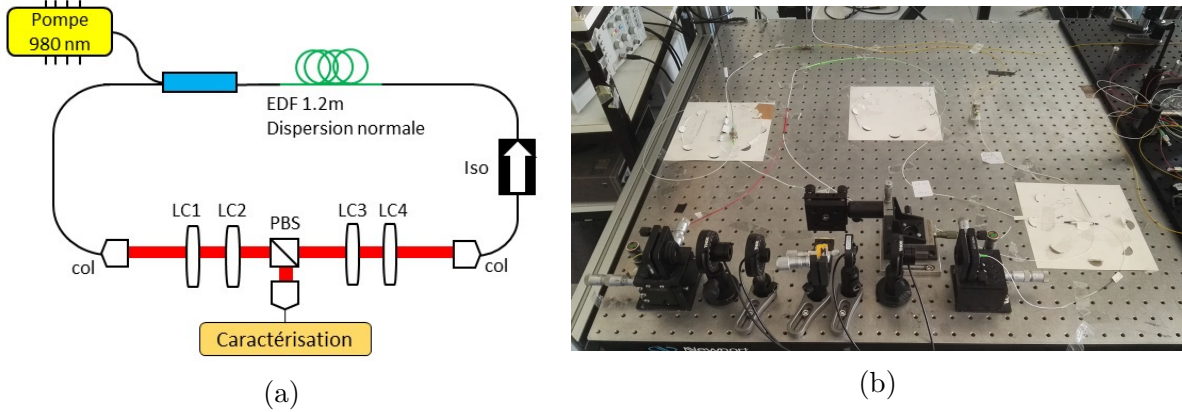


FIGURE 4.2: Figure présentant (a) le schéma du montage laser développé avec (b) sa photo

Ce montage est une cavité en anneau à gestion de dispersion classique [80] mais où la dispersion moyenne est délibérément décalée dans le domaine normal afin de favoriser les régimes monoimpulsionnels sur une plus grande plage de paramètres : ces derniers seront donc “chirpés” et de plus forte énergie [124]. En plus de l’EDF, un isolateur est utilisé pour sélectionner une seule direction de propagation au faisceau. Celui-ci est placé en régime de contre propagation, c’est à dire qu’il sélectionne une direction de propagation du faisceau laser opposée à celle de la pompe, afin d’optimiser l’absorption de la pompe par l’EDF. Notre montage expérimental est présenté dans la figure 4.2 (a). La longueur totale de fibre de cette cavité est de 3.2 m et comporte un espace en espace libre de 42 cm, conduisant à un taux de répétition de 58 MHz. La pompe est réalisée par la diode laser à 980 nm, et est introduite dans la cavité grâce à un démultiplexeur WDM, puis est filtrée par l’isolateur avec sa bande passante réduite autour de la longueur d’onde du champ (1560 nm). Considérant la longueur de SMF utilisé, sa GVD ($\beta_{2,SMF} = -0.0222\text{ps}^2.\text{m}^{-1}$) et la dispersion de l’espace en espace libre nulle, la GVD globale de cette cavité est bien normale, autour de $\beta_{2,net} = 0.03 \text{ ps}^2$. Pour réduire au maximum les pertes induite par les éléments composant la cavité, ils sont soudés entre eux, aucun connecteur n’est utilisé. Pour la comparaison, une

soudure induit moins de 0.1 dB de pertes si elle est correctement réalisée, contre 0.2 dB pour un connecteur classique. Avec les soudures et les pertes liées au dispositif en espace libre (on peut considérer 3 dB de pertes), le dispositif complet mène à une efficacité $P_{\text{faisceau}}/P_{\text{pompe}}$ maximal de 15%. La puissance de pompe utilisée est entre 500 et 600 mW, conduisant donc à une puissance intracavité d'environ 50 mW de puissance moyenne en continu. Cette cavité permet de générer des impulsions uniques avec des dynamiques simples. Tous les composants sont fixés au support avec du scotch, conduisant à des régimes robustes et à une forte répétabilité des états.

4.2 Optimisation selon la largeur spectrale

4.2.1 Simulations de l'utilisation de l'algorithme implémenté sur la cavité

Paramètres de propagation du champ laser

Un second travail de simulations numérique, décrivant grossièrement la propagation du champ dans la cavité expérimentale, plus complet mais également plus long que le précédent travail de simulation (présenté en section 3.3.2) est présenté dans cette section. Les dynamiques laser pourront de fait être visualisées, ainsi que l'utilité de l'algorithme. Certains phénomènes pourront alors être compris et mieux interprétés qu'avec le seul travail pratique. La cavité simulée, correspondant à celle expérimentale est donnée en figure 4.3.

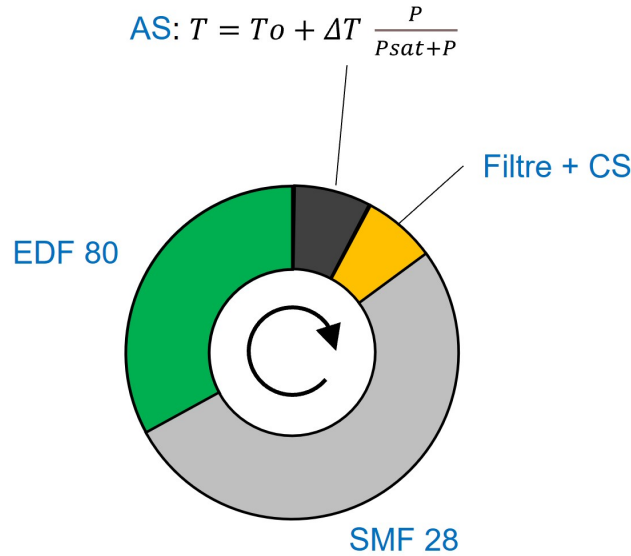


FIGURE 4.3: Schéma du montage laser simulé, comprenant une fibre dopée EDF, une absorbant saturable AS, un filtre porte, un coupleur de sortie CS et une fibre SMF

Cette cavité comporte 4 sections : une fibre dopée EDF, dont les caractéristiques sont les mêmes que celle du montage expérimental, une fibre passive SMF 28 de 2 m, un coupleur de sortie avec filtre et un absorbant saturable. Ces deux derniers éléments sont simplifiés par rapport au montage expérimental : l'absorbant saturable impose

simplement une fonction de transfert dont l'équation est donnée sur la figure, le filtre est une fonction porte centrée à 1550 nm avec une bande passante variable, et le coupleur de sortie induit simplement 3 dB de pertes. Ces deux éléments constituent, dans la réalité, tout le montage en espace libre de la cavité réelle, avec les 4 LCs et le PBS. L'isolateur n'est pas modélisé, le sens de propagation du champ dans la cavité étant déterminé par la séquence des composants. Dans les deux parties fibrées, la propagation du champ est calculée grâce à l'équation de Schrödinger non linéaire généralisée, plus adéquate pour modéliser localement la dynamique dans une fibre que l'équation de Ginzburg-Landau présentée en équation 2.6. Cette relation, limitée à l'ordre 2 de la dispersion, est donnée en équation 4.1.

$$\frac{\partial E}{\partial z} = (\alpha + g)E + i\frac{\beta_2}{2}\frac{\partial^2 E}{\partial t^2} - i\gamma|E|^2 E \quad (4.1)$$

Avec α le terme d'atténuation des fibres, g le gain et z la direction de propagation du champs E . Le terme α est fixé à -0.2 dB.km⁻¹ quelle que soit la fibre considérée. Entre chaque élément, des pertes dûes aux connexions sont fixées à -0.5 dB. Le gain est nul pour la fibre SMF, et se calcule pour la fibre dopée uniquement d'après une puissance de saturation P_{edf} liée à la puissance de pompe. Cette façon de modéliser l'amplificateur est extrêmement simple, car elle ne prend pas en compte la courbe de gain de la fibre dopée. Dans notre cas, la puissance de saturation est fixée à $P_{edf} = 60$ mW, ce qui induit un gain de $g \approx 6$, proche de celui du montage expérimental. Les caractéristiques des deux fibres modélisées sont données en tableau 4.1, correspondantes aux caractéristiques expérimentales.

	EDF	SMF
gain [a.u]	6	0
L [m]	1.2	2.1
n_2 [10^{-20} m ² .W ⁻¹]	2.1	2.5
D [ps.nm ⁻¹ .km ⁻¹]	-48	17
β_2 [ps ² .m ⁻¹]	0.06	-0.022
β_3 [ps.nm ⁻² .km ⁻¹]	0	0.07
MFD [μ m]	6.5	10

TABLE 4.1: Tableau présentant les caractéristiques des fibres dopées EDF et passives SMF modélisées dans les simulations

Pour résoudre l'équation 4.1, une technique de simplification est aujourd'hui couramment utilisée : la méthode dite de "Split-Step Fourier" (SSFM). Cette technique a d'abord été utilisée par Weideman et Herbst en 1984 pour la résolution de la GNLSE [125], puis optimisée pour la propagation de champ optique dans les éléments fibrés par Sinkin en 2003 [126]. Elle repose sur la dualité des deux composantes de la GNLSE : si chaque composante (linéaire et non linéaire) de l'équation a une solution analytique, la GNLSE regroupant les deux parties n'en a pas. Cette technique assume donc que pour une faible distance de propagation, les effets linéaires et non linéaires agissent indépendamment l'un de l'autre. $\hat{L}$ correspondra à l'opérateur linéaire, prenant en compte les effets de dispersion et de gain, tandis que $\tilde{N}$ correspondra à l'opérateur

non linéaire gouvernant les effets non linéaires. Mathématiquement, ces opérateurs s'écrivent de la manière suivante :

$$\tilde{L} = \left(\alpha + g - i \frac{\beta_2}{2} \frac{d^2}{dt^2} \right) \text{ et } \tilde{N} = -i\gamma|E|^2 \quad (4.2)$$

La résolution de la partie non linéaire $\tilde{N}$ est simple. En effet l'équation devient une équation différentielle du premier ordre ($-i\gamma|E|^2 E(z, t) = \frac{\partial E}{\partial z}$, dont la solution est triviale : $E(z, t) = E_0 e^{-i\gamma|E|^2 z}$. Il est alors facile de calculer la propagation sur z : $E(z + h, t) = E_0 e^{-i\gamma|E|^2 z} e^{-i\gamma|E|^2 h} = e^{h\tilde{N}} E(z, t)$.

La résolution de la partie linéaire $\tilde{L}$ est plus compliquée, avec la présence de la dérivée seconde en t . On a dans ce cas : $\frac{\partial E}{\partial z} = (\alpha + g)E + i \frac{\beta_2}{2} \frac{\partial^2 E}{\partial t^2}$. Pour trouver une solution à cette équation, il faut utiliser les propriétés de la transformée de Fourier. En effet, sous condition d'existence, la TF échange dérivation et multiplication par i fois la variable d'intérêt. Ainsi, on aura $\tilde{L}(\omega) = \alpha + g - i \frac{\beta_2}{2} \omega^2$, permettant d'obtenir une simple équation différentielle du premier ordre avec solution $E(z + h, \omega) = e^{h\tilde{L}} E(z, \omega)$. Une transformée de Fourier inverse permet alors de calculer le champ dans le domaine temporel. Ainsi, la propagation dans la partie fibrée du champ électrique E peut être calculée.

En ce qui concerne l'absorbant saturable, la fonction de transfert expérimentale ne peut être modélisée. En effet, celle ci, reposant sur des effets de polarisation non linéaires, et étant dépendante de nombreux paramètres, ne peut être mesurée. Dans ce travail de simulation numérique, l'absorbant saturable aura une fonction de transfert simplifiée, couramment utilisée pour ce type de travail, et donnée dans la figure 4.3. Cette fonction sera la suivante : $T = T_0 + \Delta T \frac{P}{P_{sat} + P}$, avec T la transmission de l'AS, T_0 la transmission minimale de l'AS, ΔT la différence de transmission et P_{sat} la puissance de saturation, correspondant à une transmission de $T = T_0 + \frac{\Delta T}{2}$. Dans le travail de simulation, les trois paramètres fixes de l'absorbant saturable (T_0 , ΔT et P_{sat}) seront utilisés comme gènes dans la simulation de l'algorithme.

Le filtre est un filtre spectral en fonction porte, centré en 0 et de bande passante variable BP_{Filtre} . La bande passante du filtre sera utilisée également comme gène dans la simulation d'algorithme. Le coupleur de sortie induira simplement des pertes de 3 dB, modélisant l'effet du PBS dans la réalité.

Simulation du processus d'algorithme d'évolution

Comme énoncé dans la sous section précédente, dans ces travaux, 4 gènes seront utilisés : T_0 , ΔT , P_{sat} et BP_{Filtre} . Dans ce cas, le champ est visualisé en temps réel, dans chaque élément de propagation. Les régimes de blocage de mode à soliton unique seront donc facilement déterminés par un comptage du nombre de maxima locaux de la trace temporelle, si plusieurs maxima sont détectés, l'état aura un score de 0. Pour répondre à l'objectif de ces travaux, la fonction de mérite sera directement la largeur spectrale du régime. L'étude se porte sur 10 générations, avec des populations futures de 20 individus et une population initiale de 60 individus. Le taux d'enfants mutés est

fixé à 70%. Chaque individu débute sur du bruit, et sa propagation est simulée sur plusieurs milliers de tours de cavité, après stabilisation. Cette étude testera donc 260 individus, correspondant à un temps de calcul de 6h. Avec ces conditions de travail, l'évolution du meilleur score et du score moyen pour chaque génération est donné en figure 4.4.

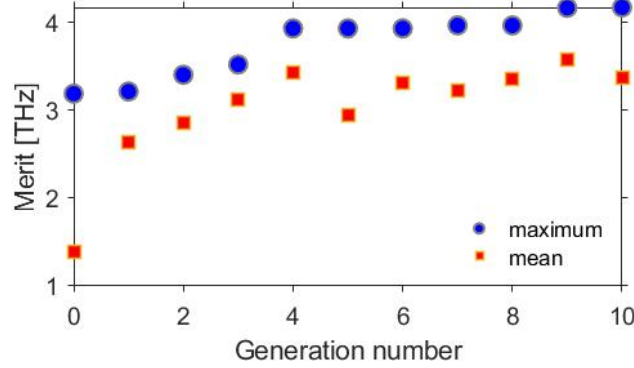


FIGURE 4.4: Graphique présentant l'évolution du score maximale de chaque génération (en bleu) et le score moyen de la génération (en rouge), en fonction de la génération

L'évolution du meilleur score de chaque génération présente une augmentation importante jusqu'à la génération 5, puis une stabilisation à 4 THz, démontrant une optimisation de 30% par rapport au meilleur individu de la génération 0. Comme les meilleurs individus sont sélectionnés à chaque génération, le meilleur score ne fluctue plus une fois le score optimal atteint. Ce comportement est purement théorique, des instabilités apparaissent lors des travaux expérimentaux donnant lieu à des légères variations. Le score moyen est également optimisé de 300%, avec un score optimal atteint à la génération 5 également, mais fluctuant ensuite. Cela s'explique par la diversité des enfants générés par croisement et mutation. L'individu optimal est de $T_0 = 48$, $\Delta T = 0.52$, $P_{sat} = 46.3$ mW et $BP_{Filtre} = 31.3$ nm. Pour toutes les simulations réalisées (une dizaine), les fonctions de transferts générant les régimes les plus larges spectralement ont une transmission maximale proche de 1 ($T_0 + \Delta T \approx 1$), avec une transmission minimale entre 0.5 et 0.7. La bande passante du filtre est généralement autour de $BP_{Filtre} \approx 35$ nm, proches des largeurs spectrales générées. En revanche, les valeurs de P_{sat} varient beaucoup d'une simulation sur l'autre ($8 < P_{sat} < 47$ W), sans tendance générale, ne permettant pas d'apporter une conclusion. En revanche, quelles que soient les valeurs de P_{sat} , les largeurs spectrales optimales restent dans la même plage autour de 4 THz.

Le spectre de l'individu optimal, ainsi que sa trace temporelle, sont donnés en figure 4.5. Son spectre présente une forme caractéristique du régime de dispersion normal de notre cavité, avec un plateau constant au centre et des pentes très abruptes sur les côtés. Sa largeur à mi hauteur est de 4.2 THz, correspondant à une largeur de 35.6 nm. Il est intéressant de remarquer que la largeur spectrale du régime mesurée après la propagation dans la fibre SMF est plus large que la bande passante du filtre, démontrant un phénomène de respiration spectrale. L'impulsion a une durée de 2 ps, et l'énergie de l'impulsion est de 0.6 nJ. Comme attendu avec un régime de dispersion normale, les régimes générés sont loin d'être en limite de Fourier, ayant un fort chirp.

Les formes spectrales et temporelles semblent bien correspondre à celles que nous nous attendons à générer avec ce type de cavité.

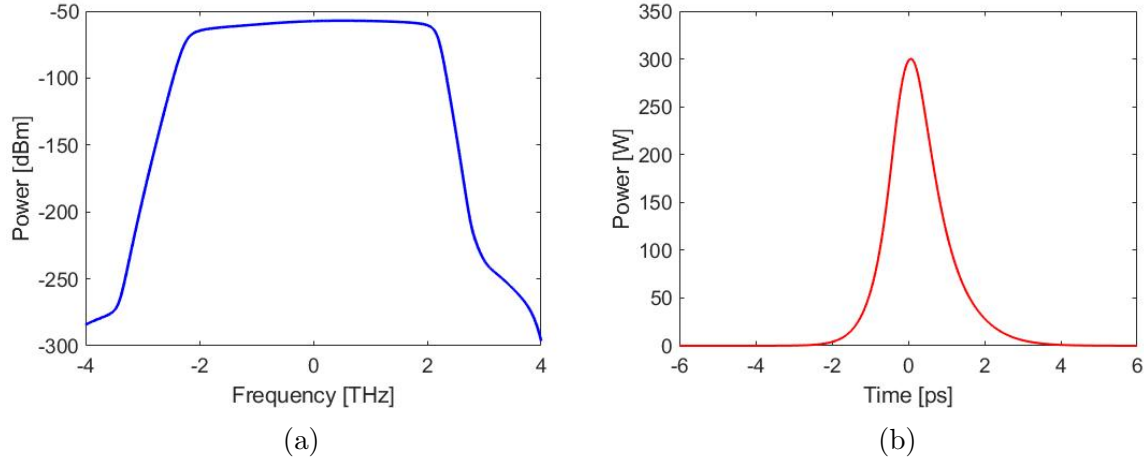


FIGURE 4.5: Graphiques présentant (a) le spectre de l'individu optimal généré par simulations, et (b) sa trace temporelle

Dans la partie simulatoire, la fonction de mérite pouvait se limiter à la largeur spectrale. Dans l'expérience, celle-ci se doit d'être plus complexe et complète. La prochaine sous-section présentera la fonction de mérite utilisée pour l'expérience d'optimisation de la largeur spectrale par l'algorithme.

4.2.2 Fonction de mérite expérimentale

L'objectif de ce premier chapitre est l'optimisation par algorithme d'une impulsion selon sa largeur spectrale. Pour cela, une fonction de mérite appartenant à la famille présentée en section 3.5 sera utilisée, avec comme seconde contribution la largeur spectrale de l'impulsion, calculée directement d'après la trace DFT. Ainsi, la fonction de mérite utilisée pour ce premier chapitre sera celle présentée en équation 4.3.

$$F_{merit} = P_{RF@58MHz} - P_{RF,background} + \alpha < \Delta\lambda_{DFT} > \quad (4.3)$$

Comme énoncé en section précédente, la première contribution est la puissance radio fréquence au taux de répétition fondamental de la cavité, c'est à dire à 58 MHz. Celle-ci est calculée directement d'après la voie n°1 de l'oscilloscope, mesurant la trace directe du train d'impulsion. Dans ce cas, cette grandeur se situe entre 40 dBm (pour les régimes continus) et 70 dBm (pour les blocages de modes les plus stables). Pour un régime impulsif, sa valeur varie entre 55 dBm et 70 dBm. Une étude préliminaire, mesurant la puissance RF pour 40 régimes de natures différentes (continus CW, Q switched et mode-locked) est réalisée et présenté en figure 4.6 (les régimes sont classés par ordre croissant de puissance RF). Nous remarquons que les régimes de blocage de mode ont une puissance plus élevée que les autres régimes, ce qui était attendu et désiré. Cependant, les régimes continus ou Q switched ont sensiblement la même puissance RF, avec toutefois une moyenne légèrement en faveur des régimes Q switched. Ce score est donc bien conforme à nos attentes pour la cavité développée.

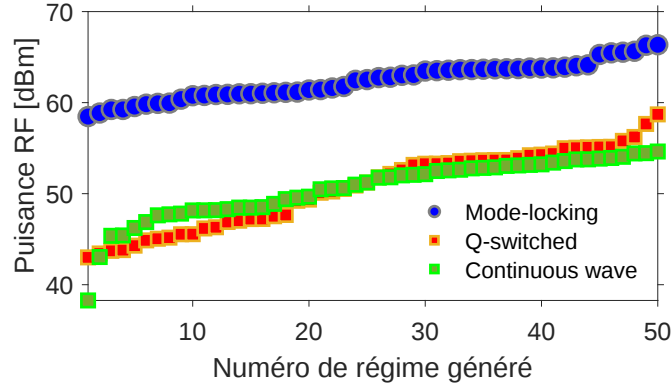


FIGURE 4.6: Graphique présentant l'évolution de la puissance RF au taux de répétition fondamental de nombreux régimes de nature différentes, classés par ordre croissant en puissance

La seconde contribution de la fonction de mérite (appelée U en section 3.5) est la largeur spectrale à mi hauteur de la trace DFT (voie 2 de l'oscilloscope), moyennée sur 9 valeurs ($< \Delta \lambda_{DFT} >_9$). Cette contribution permet d'optimiser un état de blocage de mode selon sa largeur spectrale. Pour des régimes continus, cette valeur est totalement aléatoire, puisque la trace DFT d'un tel régime est très bruitée. Elle peut alors être artificiellement élevée, et donc surpasser les scores attribués aux régimes impulsionnels. Nous ajoutons donc à cette valeur un seuil maximal à ne pas dépasser, qui est déterminé selon les valeurs maximales de largeurs spectrales pour des impulsions obtenues dans nos expériences. Dans notre cas, la valeur de cette contribution varie entre 2.5 ns et 4.5 ns (correspondant à des largeurs spectrales de 25 à 40 nm), le seuil est donc fixé à 5 ns. De plus, si la première contribution reste globalement stable d'une mesure sur l'autre pour un même individu, la mesure DFT peut fluctuer, d'où le moyennage sur 9 valeurs.

Entre ces deux contributions, le poids relatif α permet de favoriser plus ou moins la contribution sur la largeur spectrale. Un problème se pose alors pour des valeurs du paramètre α trop négatives dans ce cas de figure, utilisées pour trouver des régimes dont la largeur spectrale est faible (équation 4.3). Pour qu'un régime de blocage de mode reste sélectionné par notre fonction de mérite, il faut que la variation continu/blocage de mode de la première contribution reste suffisamment importante pour discriminer les régimes continus. En d'autres termes, il faut que la seconde contribution ne vienne pas réduire le score de mérite en dessous de la puissance radiofréquence d'un régime impulsionnel, et donc que sa valeur ne soit pas plus importante que la variation de $P_{RF@58MHz}$. Compte tenu des différentes valeurs (variation de P_{RF} pour rester en impulsionnel $\Delta P_{RF} \approx 20$ dBm et $< \Delta \lambda_{DFT} > \approx 4.5$ ns), la valeur minimale de α doit être de -4 . En revanche, des valeurs de α très importantes ne posent pas de problèmes majeurs grâce aux seuils utilisés. Cependant, plus le poids relatif est élevé, plus les instabilités de mesures le seront. Utiliser un poids relatif α compris entre -4 et 10 semble donc être une bonne fenêtre d'étude pour ce premier chapitre.

4.3 Cartographie des états de blocage de mode

Tout d'abord, avant de tester l'algorithme, nous réalisons une cartographie grossière des états de blocage de mode dans le but de tester les performances de notre cavité, ainsi que de tester la capacité de notre fonction de mérite à les sélectionner au détriment des états continus. Un programme de scan des phases potentielles des 4 lames à cristaux liquides de notre cavité est développé sous Labview. Un score de mérite supérieur à une valeur seuil (dans ce cas le seuil sera de 55 avec $\alpha = 0$, correspondant à une puissance radiofréquence de 55 dBm) détermine un état de blocage de mode, chaque spectre étant enregistré pour vérifier a posteriori cette proposition. Chaque individu est appliqué après un état neutre, permettant également de tester l'influence de cet état sur la génération de régime impulsionnels.

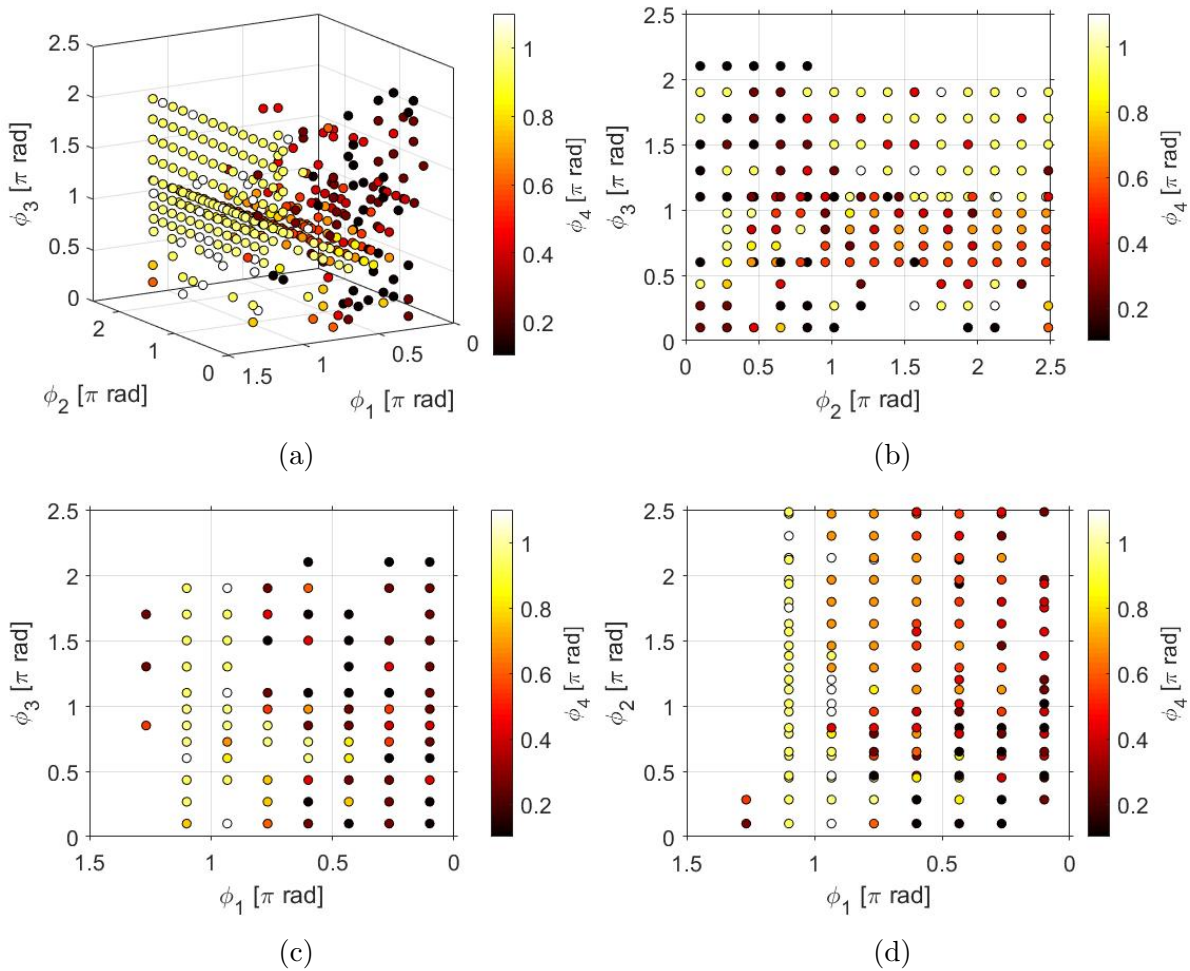


FIGURE 4.7: Graphique présentant (a) une cartographie 3D des états de blocage de mode générés par notre cavité, pour chaque coordonnée en phase ϕ_1 , ϕ_2 , ϕ_3 , ϕ_4 , (b) sa coupe 2D ϕ_3 fonction de ϕ_2 indépendamment de ϕ_1 , (c) sa coupe 2D ϕ_3 fonction de ϕ_1 indépendamment de ϕ_2 et (d) sa coupe 2D ϕ_2 fonction de ϕ_1 indépendamment de ϕ_3

Cet état neutre est défini par les phases maximales de chaque LC, soit des tensions appliquées nulles. Le pas de scan est défini à 0.2π pour toutes les lames. Les résultats de cette cartographie sont présentés en figure 4.7(a,b,c,d). Chaque point de cette car-

tographie correspond à un régime impulsif, dont les coordonnées en phase sont indiquées. Ce graphique montre tout d’abord la quantité de régime de blocage de mode générable par notre cavité (520 sur 4356 individus testés), indiquant une bonne performance de notre cavité. Les régimes sont répartis équitablement sur tout l’espace des phases. Chaque individu est testé en 1.8 secondes (figure 3.14), menant à une cartographie d’une durée de 2h30 min. Nous verrons par la suite que la durée d’optimisation par notre algorithme très inférieure à cette valeur, justifiant par elle même l’utilisation d’algorithme plutôt que d’avoir recours à des scans systématiques pour explorer l’espace des paramètres en entier. La puissance RF fluctue beaucoup lors du scan, ne permettant pas de visualiser des zones précises à forte puissance.

Quelques tests sur l’influence de l’état neutre nous a démontré qu’un régime continu bruyé (comme c’est le cas dans notre étude précédente) est l’état neutre qui favorise le plus la génération de blocage de mode (plus qu’un régime de blocage de mode lui même). Nous choisirons donc cet état neutre par la suite. Cependant, en raison des propriétés variables de la cavité en fonction des conditions extérieures, nous ne pouvons pas certifier la stabilité des conditions expérimentales.

4.4 Convergence vers un régime de blocage de mode fondamental à spectre large

4.4.1 Optimisation complète depuis une population aléatoire

Pour débiter cette étude, nous choisissons un poids relatif $\alpha = 8$ de la fonction de mérite, cette valeur permettant un équilibre plus important entre les deux contributions. De plus, l’expérience nous a montré que cette valeur permettait une optimisation notable de la largeur spectrale. Dans ce cas, environ 70% du score de mérite provient de la première contribution de la fonction, contre 30% pour la seconde contribution.

L’algorithme démarre donc sur une population initiale de 100 individus aléatoires, puis se poursuit sur des populations futures de 20 individus composés de 5 parents et 15 enfants. Une proportion de 70% des enfants ont un gène muté, dont l’amplitude de mutation est variable. En effet, le gène muté est aléatoire au début de l’optimisation, puis varie selon une amplitude faible autour d’un gène parent si la fonction de mérite moyenne des parents est supérieure à 80 (“merit seuil” dans la section 3.6). On considère cette valeur comme suffisante pour caractériser les blocages de modes, et une mutation limitée permet de chercher finement des états optimaux autour d’un parent. Les tensions appliquées pour les individus neutres sont de 0V, correspondant à la phase maximale, et qui permet d’utiliser un cheminement favorable à la génération d’états de blocage de mode (selon notre expérience sur notre cavité). Les optimisations se déroulent sur 20 générations. Pour plus de clarté, un tableau présentant les caractéristiques initiales du fonctionnement de l’algorithme est présenté en tableau 4.2. Ces paramètres étant modifiés au cours de nos travaux, ils seront rappelés avant chaque optimisation.

Avec ces conditions initiales, une optimisation typique par l’algorithme est présentée en figure 4.8. Cette figure indique une optimisation de 10% du score de mérite maxi-

4.4. CONVERGENCE VERS UN RÉGIME DE BLOCAGE DE MODE FONDAMENTAL À SPECTRE LARGE

Gènes	$\varphi_{1,2,3,4}$
Fonction de mérite	$P_{RF} - P_{bck} + \alpha\Delta\lambda_{DFT}$
Population initiale	100
Population finale	20
P_{pompe} (W)	1.0
Nombre génération	20
Taux d'enfants mutés	70 %

TABLE 4.2: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail sur l'optimisation de la largeur spectrale avec 4 LCs comme gène

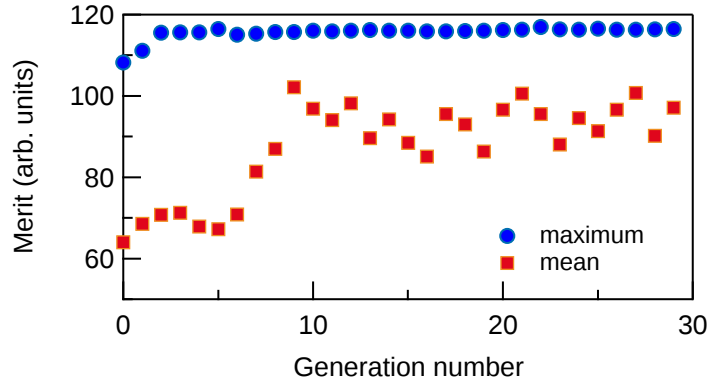


FIGURE 4.8: Graphique présentant l'évolution du score maximale de chaque génération (en bleu) et le score moyen de la génération (en rouge), en fonction de la génération

mum, réalisé en 6 générations, ce qui correspond à une optimisation de 5 minutes. Une fois cet optimum atteint, le score de mérite maximum reste constant, indiquant une stabilité des mesures. Le score moyen, en revanche, est optimisé mais varie beaucoup d'une génération sur l'autre. Cette variation est due au fort taux d'enfants mutés utilisé. L'individu optimal est présenté en figure 4.9. Le spectre de l'individu optimal, présenté en figure 4.9 (a), présente une largeur de 38.4 nm centrée à 1555 nm. Cela correspond à une optimisation de 2.5 nm par rapport au meilleur individu de la génération 0, soit d'environ 6%. Sa forme en fonction porte est caractéristique du régime de dispersion de la cavité. L'évolution de ce spectre sur quelques dizaines de tours de cavité, enregistré grâce la technique de la DFT, est présenté dans la partie basse de cette figure. Le spectre évolue peu pendant cette période, indiquant une grande stabilité de l'état. Cette caractéristique est confirmée par les figures 4.9 (b) et (c), où l'on peut observer respectivement la puissance radiofréquence au taux de répétition fondamental (d'environ 65 dBm) et l'histogramme des largeurs spectrales à mi hauteur, indiquant la majorité des mesures entre 38 et 39 nm. Ces données, et plus particulièrement celle de la DFT, sont limitées par le bruit électronique de l'oscilloscope.

Ce régime optimal est suffisamment stable et robuste pour pouvoir être réappliqué pendant plusieurs jours dans les conditions de notre laboratoire, et ce avec les mêmes caractéristiques. Cependant, dans ces conditions, une dégradation du régime apparaît après quelques jours. Typiquement, après 3 jours, un pic de quasi continu apparaît

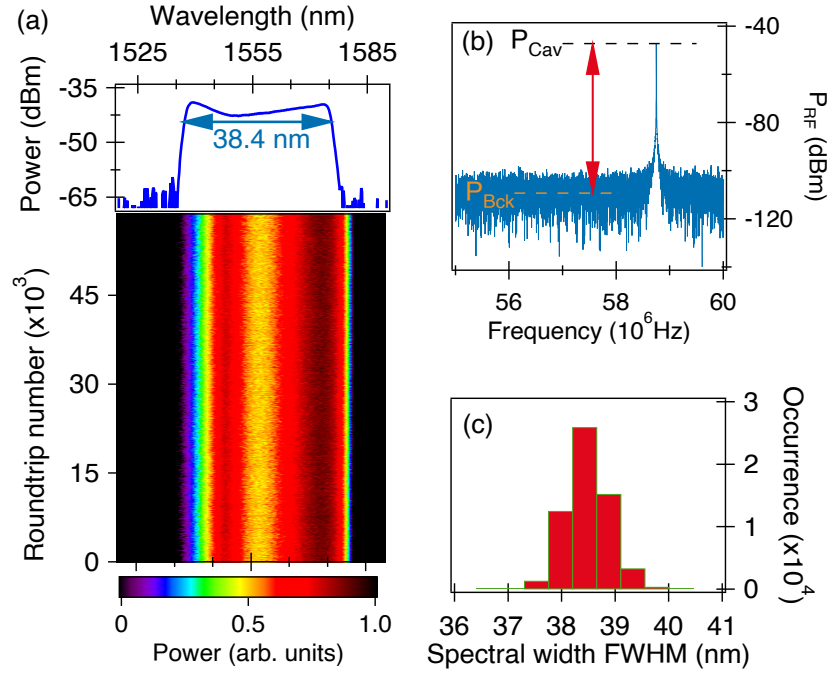


FIGURE 4.9: Graphique présentant (a) l'évolution du spectre du régime obtenu après optimisation avec un poids relatif $\alpha = 8$, (b) son spectre RF au taux de répétition fondamental de la cavité, et (c) l'histogramme de ses largeurs spectrales sur 50000 tours de cavité.

au centre du spectre optique. Un processus de réoptimisation rapide a été testé pour résoudre ce problème, présenté dans la prochaine sous section.

4.4.2 Réoptimisation depuis un individu dégradé

Pour une réoptimisation, la population initiale est créée autour des gènes de l'individu dégradé, en ajoutant des variations aléatoires de faible amplitude ($\Delta\varphi \approx 0.1\pi$) à ses différents gènes. Dans ce cas, une population initiale composée de 40 individus est suffisante, et 3 générations sont suffisantes pour retrouver un optimum, similaire à celui d'une optimisation complète. La durée d'une telle réoptimisation est alors de 3 min. Les spectres de ces 3 états sont présentés dans la figure 4.10.

Cette figure montre des résultats similaires pour une réoptimisation et une optimisation totale, avec des spectres larges de 38 nm et sans pic d'onde quasi continue. Lorsque l'on débute par un état de blocage de mode, il est alors possible de réoptimiser un état plus rapidement avec le même résultat qu'une optimisation complète. De plus, notre algorithme développé possède une fonction permettant de stopper l'optimisation lorsqu'une valeur de fonction de mérite, que l'on juge satisfaisante, est atteinte ("Merit cible" dans la figure 3.15). Ces arrangements permettent une réduction des temps d'optimisations et rendent le procédé très pratique et potentiellement utilisable dans des domaines industriels.

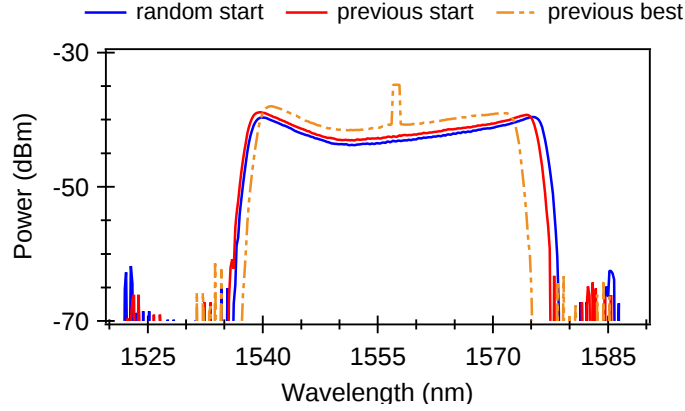


FIGURE 4.10: Graphique présentant les spectres d'un état dégradé à réoptimisé (en orange), le spectre optimal issu d'une réoptimisation (en rouge) et celui issu d'une optimisation complète (en bleu)

4.4.3 Ajustement de la largeur spectrale utilisant des mesures DFT

Dans le but de tester l'influence du poids relatif α sur les optimisations, une optimisation référence où $\alpha = 0$ est réalisée. Dans ce cas, la fonction de mérite redevient une fonction plus conventionnelle, où un seul objectif usuel est présent : la puissance radiofréquence au taux de répétition fondamental de la cavité. L'évolution des scores au cours de l'optimisation est présentée en figure 4.11.

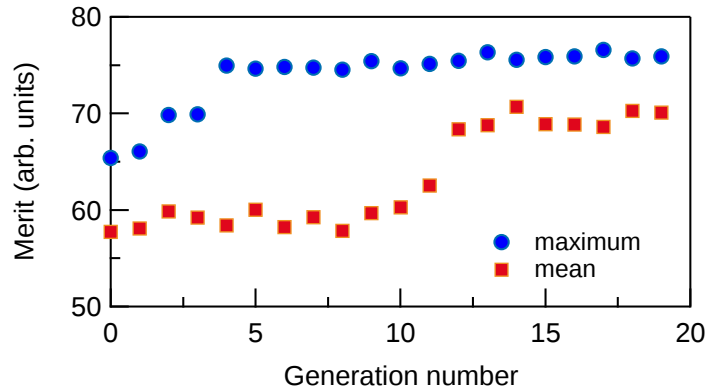


FIGURE 4.11: Graphique présentant l'évolution du score maximale de chaque génération (en bleu) et le score moyen de la génération (en rouge), en fonction de la génération pour un poids relatif $\alpha = 0$

Cette optimisation présente un profil similaire à celui présenté précédemment pour $\alpha = 8$, avec une optimisation du score de mérite maximal et de la moyenne de celui-ci, avec un optimum atteint à la génération 6. Le score optimal atteint une puissance radiofréquence dépassant les valeurs jusqu'ici mesurées, avec un optimum à 75 dBm. Les valeurs moyennes, quant à elles, fluctuent moins que pour l'optimisation précédente, confortant l'idée que les instabilités de mesure proviennent essentiellement des largeurs

4.4. CONVERGENCE VERS UN RÉGIME DE BLOCAGE DE MODE FONDAMENTAL À SPECTRE LARGE

spectrales DFT générées, pouvant varier d'une mesure sur l'autre pour un même individu, contrairement à la puissance RF. Cette optimisation permet donc la génération de régimes lasers à blocage de mode très stables, ne prenant pas en considération les données spectrales.

Une variation des poids relatifs α est réalisée, pour des valeurs comprises entre -4 et 20, menant à la constatation suivante : une optimisation avec un poids relatif α prédéfini permet la génération d'une largeur spectrale cible. La figure 4.12 présente les différents spectres optiques de 5 régimes générés par optimisation pour 5 valeurs d' α différents dans cette fenêtre d'étude.

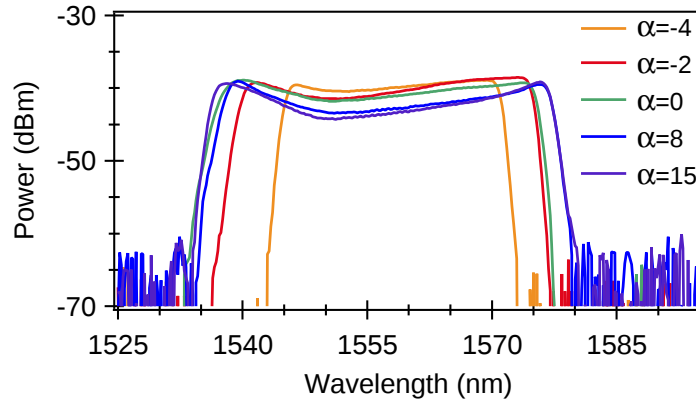


FIGURE 4.12: Graphique présentant différents spectres optimaux obtenus pour différentes valeurs d' α .

Cette figure présente des spectres ayant une largeur comprise entre 25 nm et 39 nm. Comme expliqué précédemment, et vérifié par l'expérience, une optimisation avec une valeurs $\alpha < -4$ ne converge pas vers un état de blocage de mode. Nous remarquons également que des largeurs spectrales résultantes d'optimisation avec $\alpha = 8$ et $\alpha = 15$ sont très proches (38.4 nm pour le premier et 39 nm pour le second). Cela s'explique par le fait que pour $\alpha = 10$, la largeur spectrale maximale générable par notre cavité, correspondant à la largeur du gain de l'erbium est atteinte. Pour déterminer la répétabilité du processus, 10 optimisations pour chaque poids relatif α entre -4 et 20 sont réalisées. Les résultats de cette étude sont présentés en figure 4.13, avec des barres rouges représentant les erreurs standards à 95 %. Avec ce montage, des largeurs spectrales comprises entre 25 et 39 nm sont atteignable. La largeur spectrale maximale est atteinte pour $\alpha = 10$, une valeur plus importante ne générant pas des régimes spectralement plus larges. En revanche, les fluctuations sur les scores de mérite dûs aux instabilités s'accroissent pour des valeurs du poids relatif plus importantes. Ainsi, une fenêtre d'opération entre $-4 < \alpha < 10$ semble être optimale.

Pour des valeurs $|\alpha| < 2$, les effets de la seconde contribution de la fonction de mérite sont négligeables. Nous remarquons que les erreurs sont plus importantes pour les valeurs $\alpha < 0$. Cela s'explique par le fait que dans ce cas précis, les régimes sélectionnés par l'algorithme peuvent être de deux formes différentes : soit des régimes stables avec une forte puissance RF mais un spectre peu large, soit des régimes avec un spectre plus large

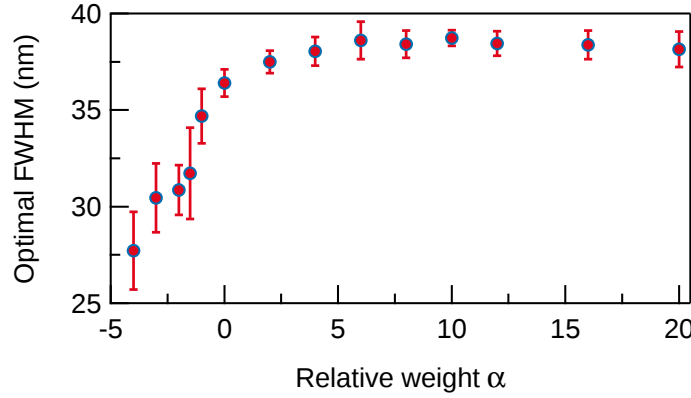


FIGURE 4.13: Graphique présentant différents spectres optimaux obtenus pour différentes valeurs d' α .

mais une puissance RF plus faible, caractérisé par un pic de quasi continu au centre du spectre. Ces deux formes de régimes auront alors des scores similaires, rendant l'erreur sur la largeur spectrale plus importante. Une fonction de mérite plus élaborée serait nécessaire pour distinguer ces différents régimes.

4.4.4 Robustesse du laser et du procédé

Pour terminer ce premier chapitre expérimental, un test sur la robustesse du laser est présenté dans cette section. Des perturbations extérieures sont appliquées sur le système pendant une longue optimisation de 70 générations. L'évolution du score au cours de l'optimisation est donné en figure 4.14 (a), tandis que les spectres optiques des optimas locaux après perturbation sont donnés en figure 4.14 (b). Cette optimisation débute avec les mêmes conditions initiales que les travaux précédents. Les actions suivantes sont réalisées à une génération donnée :

- Les fibres intracavités sont bougées de quelques centimètres, et tournées légèrement à la génération 17, en A sur la figure
- Le spectre optique “B” sur la figure 4.14 (b) est enregistré, puis la table du montage expérimental est secouée, et quelques coups sont donnés dessus à la génération 24
- Le spectre optique “C” est enregistré, puis le thermostat de l'air conditionné est diminué de 2 °C à la génération 29
- Le spectre optique “E” est enregistré à la fin du processus, en génération 70.

Nous y remarquons que le déplacement des fibres et la vibration mécanique du montage ont un impact mineur sur l'optimisation et sur la variation des scores, démontrant une bonne robustesse du montage. En revanche, le changement de température de la pièce a un impact plus important, qui fait diminuer le score de mérite. Cependant, l'algorithme retrouve un état optimal en une vingtaine de générations, jusqu'à ce que la température de la pièce retrouve une valeur quasi stationnaire. Cette étude permet de mettre en évidence la robustesse du “smart laser” (le montage et le procédé), celui-ci étant capable de s'adapter à un environnement fluctuant.

4.4. CONVERGENCE VERS UN RÉGIME DE BLOCAGE DE MODE FONDAMENTAL À SPECTRE LARGE

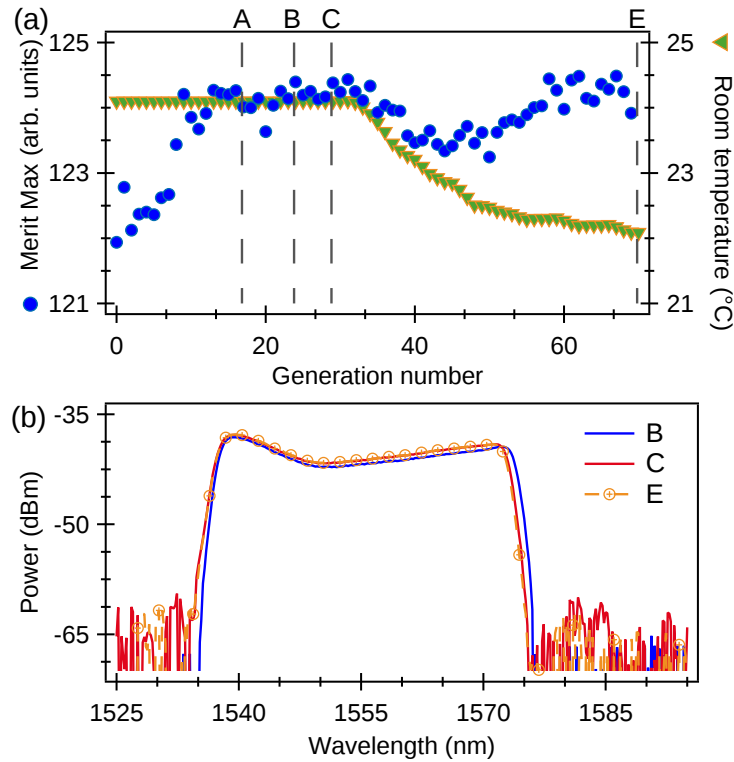


FIGURE 4.14: Graphique présentant la courbe d'évolution de notre test de robustesse

Cette dernière section conclue le premier chapitre expérimental de ce manuscrit, mettant en avant l'utilisation d'un algorithme d'évolution sur un montage laser simple, permettant la génération de dynamique impulsionnels simple. Ce procédé, couplé à l'utilisation de technique de mesures spectrales en temps réel (DFT), permet l'auto génération de régimes impulsionnel et le contrôle de leur largeur spectrales, celle ci étant prédéfinies par l'utilisateur. Cette capacité de contrôle qui n'avait jamais été réalisée, ouvre les portes à de nouvelles utilisations des algorithme d'évolution pour les travaux en photonique. Un article ([101]) ainsi qu'une présentation en conférence ([127]) ont été réalisées à partir de ces travaux. Les prochains chapitres tenteront de pousser encore plus ce concept, en auto-générant d'autres types de régimes impulsionnels plus complexes.

Chapitre 5

Augmentation de la versatilité :
génération de différentes
dynamiques impulsionsnelles

Les travaux du chapitre précédent comporte une limitation importante : le nombre de paramètres variables accessibles de la cavité (4 seulement), permettant uniquement la génération d’impulsions uniques et dont la dynamique est simple. Pour s’affranchir de cette limitation et permettre la génération de régimes lasers plus complexes, le nombre de degré de liberté doit être augmenté. Le choix s’est porté sur l’utilisation d’un façonneur d’impulsion (“pulse shaper”) intracavité. Pour la suite de ce manuscrit, ce que nous appellerons “pulse shaper” est une ligne 4f (dont le schéma sera présenté dans la section suivante), couplé à l’utilisation d’un modulateur à cristaux liquides, ou SLM pour Spatial Light Modulator, et de deux polariseurs en entrée et sortie de ligne, permettant un contrôle à la fois de la phase et de l’amplitude spectrale. La ligne 4f a été développée pendant cette thèse, ce qui permet un contrôle sur chaque élément la composant. Ce pulse shaper permet une augmentation du nombre de degré de liberté jusqu’à plusieurs centaines (324 exactement). Pour débiter ce chapitre, la première section présente le montage complet de cette ligne.

5.1 La ligne à dispersion nulle ou ligne 4f

5.1.1 Bref historique du façonnage d’impulsion

Les premiers travaux traitant de la mise en forme d’impulsions picosecondes datent du début des années 1970 [120], mais celle-ci fut fortement étudiée à partir les années 1980. En 1983, Froehly publia un chapitre décrivant toutes les techniques de façonnage d’impulsions picosecondes existantes à cette époque [128]. Puis A. M. Weiner réalisa une série d’études sur le sujet, permettant une avancée notable dans le domaine. Heritage et Weiner publiaient en 1985 leurs premier travaux sur la mise en forme d’impulsions picosecondes [129], au moyen d’un dispositif basé sur les réseaux dits de Treacy, composé de deux réseaux de diffraction comme illustré en figure 5.1.

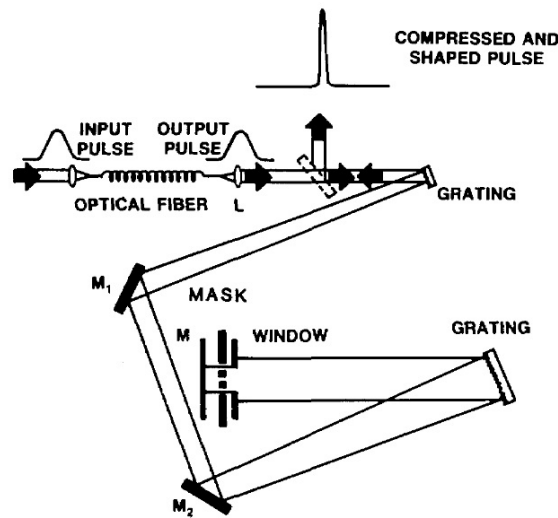


FIGURE 5.1: Illustration d’un montage à deux réseaux de diffraction utilisés pour la recompression et le façonnage d’impulsions, tirés de [129]

L'impulsion est envoyée tout d'abord sur un premier réseau, qui va la diffracter et étendre angulairement ses composantes spectrales. Celle-ci va se propager entre deux réseaux sur une distance D déterminant une dispersion induite par le dispositif. Le second réseau recombine alors toutes les composantes spectrales. Ce dispositif, utilisé à la fin des années 1960 par E. B. Treacy dans le but de comprimer des impulsions [130, 131], peut être utilisé pour le façonnage d'impulsions. En effet, une fois les composantes diffractées en champ lointain, Heritage et al utilise un masque pour éteindre certaines composantes spectrales. Le montage, étant optimisée pour la compression, ne l'est pas pour la mise en forme d'impulsion : il induit une dispersion souvent non désirée en plus du contrôle de l'amplitude spectrale. Pour la plupart des applications, il est plus adapté d'utiliser une ligne à dispersion nulle, ou ligne 4f. Weiner publia des travaux se basant sur une ligne 4f quelques années plus tard, en 1988 [132], puis les avancées technologiques ont permis d'utiliser un modulateur spatial de lumière (ou SLM pour Spatial Light Modulator) à cristaux liquides plutôt qu'un simple masque, permettant ainsi un façonnage contrôlé de la phase et de l'amplitude spectrale de l'impulsion [133, 134]. C'est ce type de dispositif qui a été implémenté pour cette thèse.

5.1.2 Principe d'une ligne à dispersion nulle

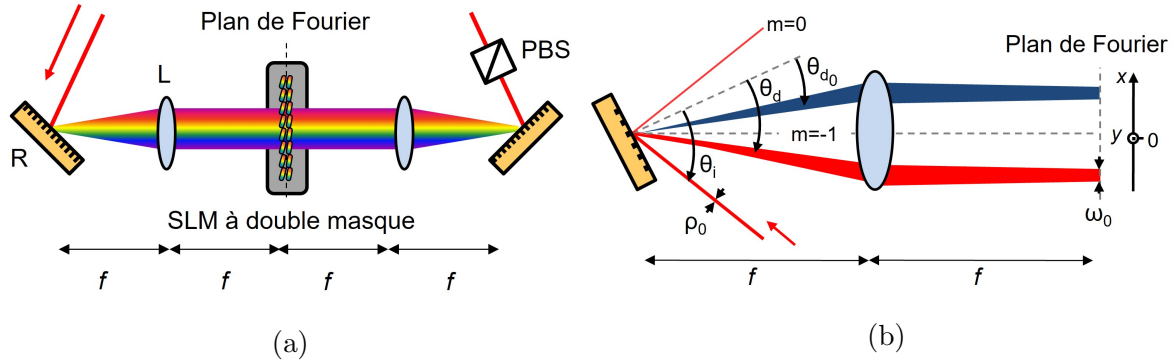


FIGURE 5.2: (a) Illustration d'un montage d'une ligne 4f composé de deux réseaux de diffraction (R), deux lentilles (L), un modulateur spatial de lumière (SLM) à double masque et un polariseur (PBS), et (b) un schéma de la ligne présentant les différentes grandeurs à considérer

Une ligne 4f à dispersion nulle (figure 5.2(a)) est un montage composé de 2 réseaux de diffractions et deux éléments de collimation (lentilles ou miroirs), séparés d'une distance f correspondant à la distance focale des éléments de collimation. Le premier réseau de diffraction va diffracter l'impulsion en champ lointain, et donc disperser angulairement les composantes spectrales de celle-ci. Ensuite, l'impulsion rencontre un premier élément de collimation. Cet élément a deux rôles. Premièrement, il collimate l'ensemble du faisceau, c'est à dire donne la même direction à toutes les composantes. Deuxièmement, il focalise chaque composante sur le plan focal image de l'élément appelé plan de Fourier 5.2(b). Pour cela, il faut que le réseau de diffraction se situe sur le plan focal objet de l'élément, soit à une distance f de celui-ci. Sur le plan de Fourier, un masque fixe

ou à cristaux liquides peut être appliqué pour le façonnage d'impulsion, puis un montage symétrique est utilisé pour recombinaison des composantes spectrales. L'impulsion en champ sera façonnée avec une forme temporelle correspondant à la transformée de Fourier du masque appliqué. Sans masque, l'impulsion en sortie reste identique à celle en entrée, permettant la vérification du bon alignement des composants.

Dans notre cas, le second réseau est monté sur une platine à translation pour pouvoir régler finement la position de celui-ci, et n'induire aucune dispersion. A noter également que l'orientation des réseaux est primordial, d'une part pour diffracter selon le bon angle sur le masque pour le premier réseau, et pour bien recombinaison des composantes de manière cohérente avec le second réseau. Le réglage des 3 angles de chaque réseau est donc primordial, et sera détaillé dans une section suivante. La distance de propagation de l'impulsion dans un tel montage est de $4f$, comme illustré sur la figure 5.2, d'où le nom "ligne $4f$ ". Un SLM est utilisé pour nos travaux plutôt qu'un simple masque, permettant un interfaçage informatique du masque de phase et d'amplitude. Ses caractéristiques sont détaillées dans la section suivante. La combinaison pas du réseau et focale des éléments détermine la dispersion du faisceau dans le plan de Fourier en fonction de la longueur d'onde, et doit donc être adapté à la longueur d'onde de travail.

5.2 Propriété du modulateur spatial de lumière à double masque

Pour la suite des travaux, dans l'optique de contrôler le plus de paramètres possibles, il est souhaitable d'avoir la possibilité de moduler le spectre optique de l'impulsion à la fois en phase mais également en amplitude. Notre choix s'est porté sur le SLM-S320d de la marque Jenoptik, à double masque disponible dans le département. Ce dernier, couplé à l'utilisation d'un polariseur en aval, permet ce double contrôle. Il comporte deux masques de 320 pixels chacun, étalés sur une largeur de 3.2 cm, soit 100 μm par pixel (97 μm par pixel et 3 μm de gap). Chaque pixel fonctionne sur le même principe que les lames à retard à cristaux liquides utilisées dans notre cavité, détaillés en section 2.3.1. Les deux lignes de pixels A et B, induisant deux phases $\varphi_A(\omega)$ et $\varphi_B(\omega)$, ont leurs axes orientés à 45° l'une de l'autre, permettant une modification de l'état de polarisation du champ en plus de l'induction d'une phase globale. Avec le second polariseur, dont l'axe est orienté horizontalement, la variation de la polarisation se traduit en variation d'amplitude. Un schéma d'illustration du dispositif pour une modulation de phase et d'amplitude est présenté en figure 5.3.

Considérons un pixel i traversé par une fréquence ω_i , et induisant deux phases φ_{Ai} et φ_{Bi} . Dans ce cas, la fréquence incidente subira une modulation $H(\omega_i)$ correspondant à

$$H(\omega_i) = \cos\left(\frac{\varphi_{Ai} - \varphi_{Bi}}{2}\right) e^{i(\varphi_{Ai} + \varphi_{Bi})/2} \quad (5.1)$$

Le premier terme en cosinus traduit la modulation de l'amplitude, tandis que le second terme en exponentielle complexe traduit la modulation de la phase. Une modulation de la phase seule est obtenue en appliquant deux phases égales sur chaque

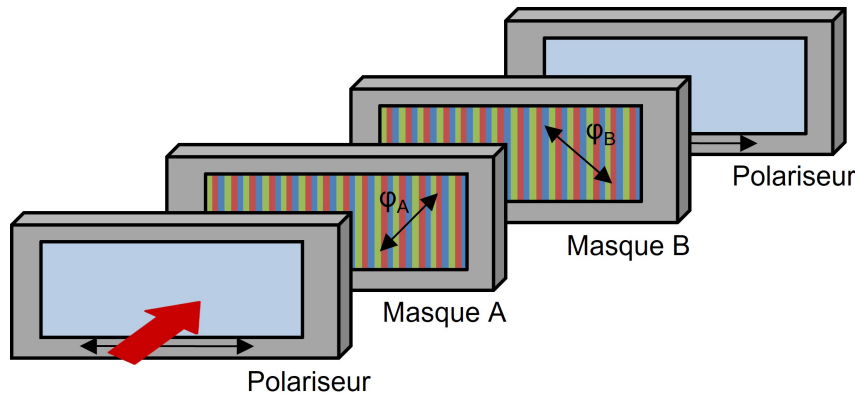


FIGURE 5.3: Schéma présentant la configuration du SLM à double masque pour l'induction de phases spectrales et pour modulations d'amplitude

masque ($\varphi_A = \varphi_B$). Dans un autre cas de figure, une modulation d'amplitude sans induction de phase peut se réaliser en appliquant deux phases opposées sur chaque masque ($\varphi_A = -\varphi_B$), conduisant à une modulation en polarisation par le SLM, qui se traduit en modulation d'amplitude avec l'utilisation d'un polariseur. On notera que chacun des 320 pixels adresse une fréquence (ou bande spectrale) différente permettant un contrôle de l'ensemble du spectre $H(\omega)$.

Dans nos premiers travaux de recompression extra cavité, présentés dans ce chapitre, une modulation de la phase spectrale seule sera appliquée. Nous précisons toutefois que le SLM utilisé n'est pas optimisé pour un travail à nos longueurs d'onde (un revêtement pour un travail autour de 1060 nm est appliqué). Les pertes à 1550 nm sont donc très importantes, autour de 3 dB uniquement pour le SLM. Ces lourdes pertes seront une vraie limitation pour nos futurs travaux. En revanche, ce revêtement n'empêche pas une modulation de phase de 0 à 2π , et permet une modulation d'amplitude sur plus de 25 dB. Il reste donc apte à nos travaux.

5.3 Développement de la ligne 4f

5.3.1 Caractéristiques des composants

Pour plus de clarté, la figure 5.2(a) présente le schéma du montage avec les grandeurs utilisées pour les calculs.

Les miroirs concaves ou lentilles

Des miroirs cylindriques sont utilisés pour le montage de la ligne. L'utilisation de miroirs cylindriques plutôt que de lentilles permet d'éviter les aberrations chromatiques dans l'une des deux dimensions transverses, dues à une focale dépendante de la longueur d'onde. Par ailleurs, les miroirs cylindriques (plutôt que sphériques) permettent d'obtenir un faisceau non focalisé selon y dans le plan de Fourier, là où est positionnée le SLM limitant ainsi le flux d'énergie le traversant. Ces miroirs ont une distance focale de 40 cm. La ligne sera alors longue de 80 cm. En théorie, pour que l'impulsion soit

recombinée parfaitement, les miroirs cylindriques en entrée de ligne et en sortie doivent être rigoureusement identiques. En pratique, un miroir cylindrique sera découpé en deux parties parallèlement à la courbure pour avoir la focale la plus proche possible entre les deux miroirs. Le reste des éléments de la ligne sera choisi en fonction de ces miroirs cylindriques pour étaler le spectre des impulsions sur une largeur adéquate.

Les réseaux

Pour un contrôle optimal du spectre, ce dernier doit être suffisamment dispersé sur le SLM pour avoir une résolution suffisante mais pas trop pour que l'ensemble du spectre soit traitable par celui-ci (maximum 3.2 cm). Pour ce faire, il faut utiliser un jeu miroir-réseau de diffraction approprié, le réseau devant également être adapté à la longueur d'onde de travail. Avec des miroirs cylindriques à focale de 400 mm, et des largeurs spectrales attendues au maximum à 50 nm, un réseau à 600 traits par millimètre est utilisé. Les différents paramètres pertinents du pulse shaper sont détaillés dans les paragraphes suivants.

La loi des réseaux dans l'air relie les angles incidents θ_i et diffractés θ_d d'ordre 1 en fonction de la longueur d'onde λ et du pas du réseau d (équation 5.2).

$$\sin(\theta_d) + \sin(\theta_i) = \lambda/d \quad (5.2)$$

Lorsque l'angle d'incidence est égal à l'angle de l'ordre premier de diffraction $\theta_i = \theta_d$, les conditions de Littrow sont réalisées. L'efficacité de diffraction est dans ce cas maximale. La relation se simplifie alors en $2 \cdot \sin(\theta_i) = \lambda/d$. En réalité, il est impossible de travailler dans ces conditions puisque l'angle d'incidence et le faisceau diffracté se confondent, mais les conditions de travail seront choisies proche des conditions de Littrow avec une déviation de 15° . En réalité, l'angle d'incidence sera autour de 35° , ce qui donnera un angle de diffraction de $\theta_d = 20^\circ$.

Les collimateurs

Après les miroirs cylindriques et les réseaux, une autre problématique est alors de choisir les collimateurs permettant le passage d'une entrée fibrée à une sortie collimatée en espace libre. Le diamètre de ceux-ci détermine celui de chaque composante spectrale dans le plan de Fourier, qu'il est souhaitable d'avoir inférieur la taille d'un pixel. En effet si chaque composante spectrale recouvre plusieurs pixels dans le plan de Fourier, le contrôle est altéré. Pour cela il convient de calculer le diamètre optimal du faisceau incident pour que chaque composante spectrale puisse être traitée par un pixel de $100 \mu m$. Pour se faire, une relation d'optique géométrique simple, basée sur la diffraction d'un faisceau par un réseau, est utilisée (équation 5.3) [134].

$$\omega_0 = \frac{1}{\pi} M^2 \frac{\lambda f \cos(\theta_i)}{\rho_0 \cos(\theta_d)} \quad (5.3)$$

Cette formule relie le waist ω_0 sur le plan de Fourier d'une composante avec f la longueur focale, ρ_0 le diamètre du faisceau incident, M^2 la qualité du faisceau (pour les lasers à fibre on considère que $M^2 = 1$) et θ_i, θ_d respectivement les angles d'incidence

et de diffraction standard de l'ordre $m = 1$ du faisceau. Ce dernier dépend de l'angle d'incidence mais également du pas du réseau.

Comme énoncé en sous section précédente, l'angle d'incidence est pris à $\theta_i = 35^\circ$, donnant un angle de diffraction de $\theta_d = 20^\circ$. Après calcul, pour un rayon dans le plan de Fourier de $\omega_0 = 50 \mu\text{m}$ afin d'avoir chaque composante traitable par un pixel, nous obtenons un rayon incident de $\rho_0 = 3.5 \text{ mm}$. Nous utiliserons donc des collimateurs de 7 mm de diamètre.

Un schéma de la ligne montée avec les éléments décrits dans les sous sections précédentes est présenté en figure 5.4 (a), ainsi que sa photographie en figure 5.4 (b). Ce montage à plusieurs étages, avec les réseaux et miroirs cylindriques situés sur deux étages, permet d'éviter les réflexions hors axes sur les miroirs ce qui conduit à minimiser les aberrations [135].

Fenêtre de travail avec la combinaison de paramètres choisie

La loi des réseaux $\sin(\theta_i) + \sin(\theta_d) = \lambda/d$, couplée à la relation trigonométrique simple $\theta_d \approx \tan(\theta_d) = X/f$, permet d'écrire la dépendance entre x et λ de la manière suivante :

$$d\theta_d = \frac{d\lambda}{d \cdot \cos(\theta_d)} = \frac{dX}{f} \quad (5.4)$$

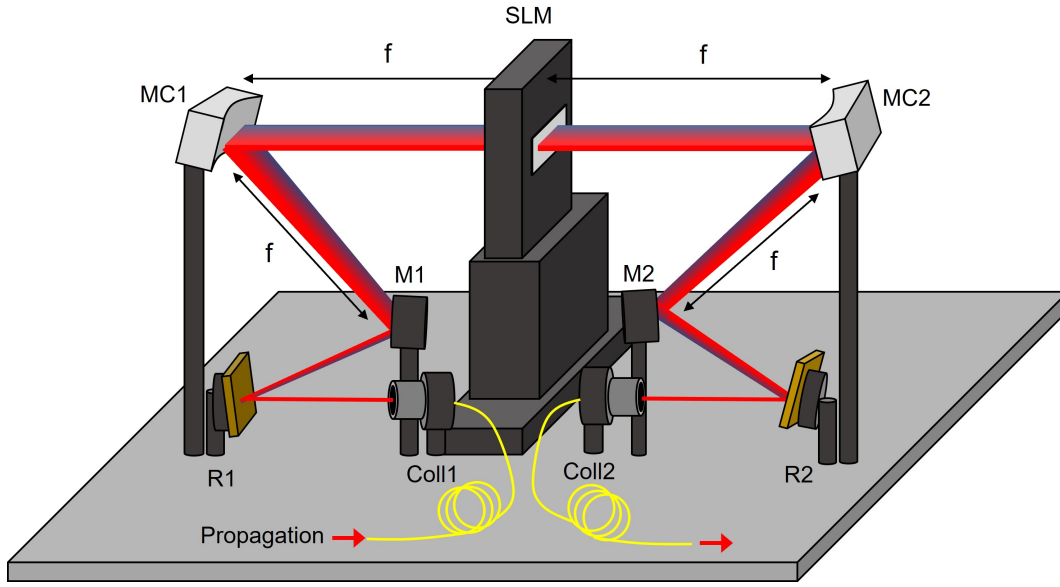
La dispersion spatiale en longueur d'onde, que nous appellerons α , s'écrit alors :

$$\alpha = \frac{dX}{d\lambda} = \frac{f}{d \cdot \cos(\theta_d)} = 2.554 \cdot 10^{-2} \text{ cm.nm}^{-1} \quad (5.5)$$

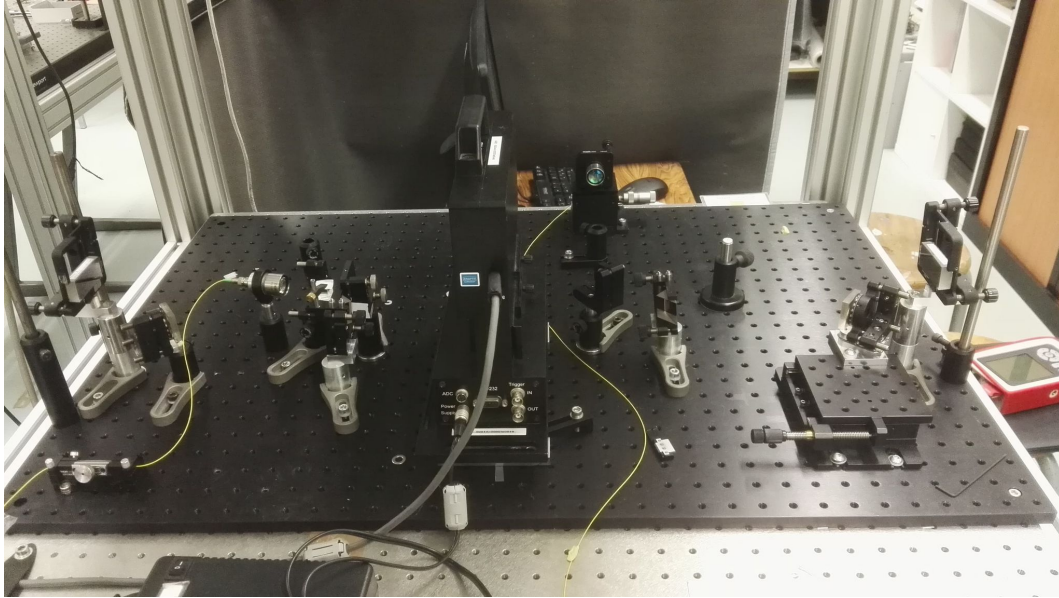
Ainsi, la fenêtre ajustable en longueur d'onde $\Delta\lambda$, calculée d'après la fenêtre spatiale $\Delta X = 3.2 \text{ cm}$, sera de $\Delta\lambda = 125 \text{ nm}$, ce qui est suffisant compte tenu des largeurs spectrales des régimes générés par notre cavité (au maximum de 50 nm). Ainsi, nous avons une marge de travail, et même si une partie du SLM est inutilisé le nombre de pixels disponibles pour le contrôle permet des modulation assez fine. De plus, la résolution en longueur d'onde $\delta\lambda$, donnée par la relation $\delta\lambda = \delta X/\alpha$ avec δX la taille d'un pixel, serait de 0.39 nm (différence de longueur d'onde entre deux pixels adjacents). Cette grandeur sera mesurée expérimentalement par la suite.

La pixelisation d'un tel dispositif induit une fenêtre temporelle de travail. En effet, pour décaler une impulsion temporellement par exemple, une phase linéaire doit être appliquée. La phase maximale pouvant être appliquée entre deux pixels successifs étant de 2π (ce qui revient à une absence de phase), le décalage temporel maximum τ s'écrit alors $\tau = 2\pi/\delta\omega$, avec $\delta\omega$ la différence de fréquence entre deux pixels consécutifs. Ceci constitue la fenêtre temporelle. Plus précisément, la forme en "porte" des pixels conduit à une fenêtre temporelle en sinus cardinal. En assumant que l'amplitude du champ ne varie pas le long d'un pixel ne varie pas le long de ce pixel, le champ à la sortie du SLM s'écrit [136] :

$$E_{out}(t) = \text{sinc}\left(\frac{\delta\omega t}{2}\right) \sum A_n E_n e^{i(\omega_n t + \varphi_n)} \quad (5.6)$$



(a)



(b)

FIGURE 5.4: (a) Schéma de la ligne montée composée de 2 collimateurs (Coll1 et Coll2), deux réseaux (R1 et R2), deux miroirs (R1 et R2), deux miroirs cylindriques (MC1 et MC2) et un SLM, et (b) sa photographie

avec A_n et φ_n respectivement la modulation d'amplitude induite par le pixel n et la phase induite par ce pixel, ω_n la pulsation centrale traversant le pixel et E_n l'amplitude du champ incident sur le pixel n . Ainsi, le terme derrière la somme induit des répliques pour des temps à $2\pi/\delta\omega$, modulées par le sinus cardinal. Compte tenu de la dispersion, le décalage en fréquence de deux pixels consécutifs $\delta\omega = 0.39$ nm induit donc une fenêtre temporelle de $\tau = \pm 21$ ps.

5.3.2 Orientation des réseaux

Une fois les composants choisis, la seconde problématique est de réaliser un bon alignement, non seulement pour réduire les pertes induites par chaque élément de la ligne, mais aussi pour recombinaison au mieux les composantes spectrales après le second réseau et éviter tout chirp non désiré. Chaque réseau est monté sur une monture cinématique rotative, permettant l'ajustement des 3 angles α , β et γ de la figure 5.5.

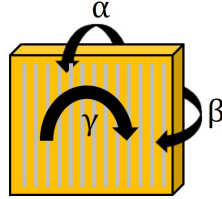


FIGURE 5.5: Schéma d'un réseau présentant les différents angles à ajuster pour l'alignement de la ligne 4f

Les angles α et γ ont une orientation particulière : ces angles sont alignés pour avoir un faisceau parallèle au plan de la table. L'angle γ représente l'orientation des traits du réseau. Il va jouer sur l'inclinaison des rayons diffractés. En théorie, il doit être nul, ce qui correspond à des traits du réseau parfaitement verticaux. Le réglage de cet angle est obtenu en s'assurant que l'ordre 1 de diffraction soit diffracté dans le plan horizontal en utilisant un iris de réglage à longue distance fixé à la hauteur du faisceau dans le plan de Fourier. Pour l'alignement de l'angle α , l'ordre 0 de diffraction est utilisé. À l'aide de l'iris fixé à la hauteur de propagation du faisceau, il est facile d'ajuster α pour que l'ordre 0 soit réfléchi dans le plan horizontal sur une longue distance. Ces deux angles sont réglés de manière identiques pour les deux réseaux. En revanche, l'ajustement de l'angle β représente un travail plus conséquent. La difficulté de cet exercice repose alors essentiellement sur l'orientation des angles β des réseaux l'un par rapport à l'autre. Pour une bonne recombinaison, le montage doit être parfaitement symétrique. En d'autres termes, l'orientation de β pour second réseau, qui permet le recouvrement des composantes spectrales doit être identique à celui du premier réseau. En pratique, pour le premier réseau, β est ajusté grossièrement pour que le faisceau ressorte dans l'axe de la ligne, et de telle sorte que la déviation soit faible. Mais pour le second réseau, un protocole d'alignement est utilisé, illustré en figure 5.6.

Pour ajuster β , le spectre est découpé en deux parties, en plaçant un masque au centre du spectre optique d'une impulsion dans le plan de Fourier. Seules les deux extrémités du spectre se propagent, laissant place à deux spots lumineux. Le faisceau en sortie de ligne devra alors avoir ces deux composantes restantes bien recouvertes en un seul spot. Si ce n'est pas le cas, deux spots seront encore visibles, et l'angle β sera encore à ajuster. Pour un bon alignement, il est préférable de visualiser le faisceau sur une longue distance (au moins 5 m) permettant des réglages plus fins. Pour cette dernière étape, l'aide d'une seconde personne est requise.

Notre alignement mène à une transmission de 40% (4.5 dB de pertes) sur l'ensemble de la ligne, soit environ 3 dB de pertes pour le SLM et 1.5 dB de pertes pour le reste

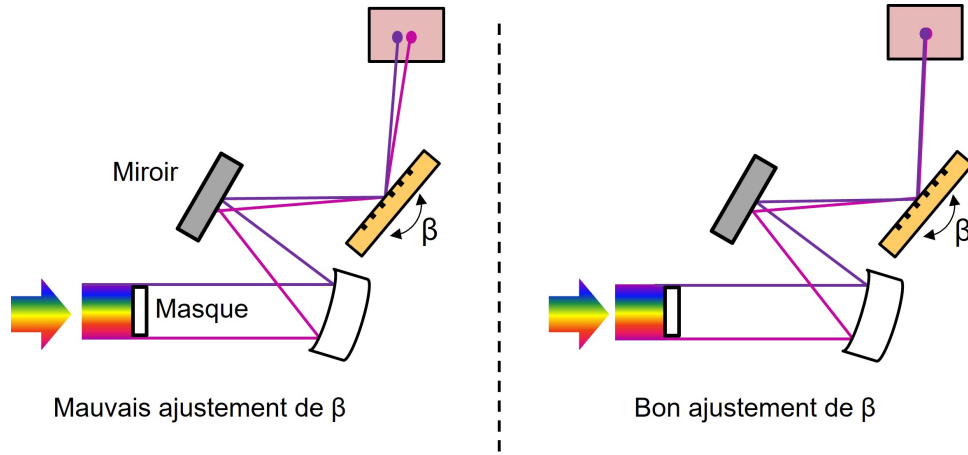


FIGURE 5.6: Schéma représentant un bon alignement de l'angle β du second réseau. La figure de gauche représente un mauvais alignement de β , visible grâce à la présence de deux spots lumineux distinct à une longue distance, tandis qu'un bon alignement est représenté à droite.

des éléments (réseaux+miroirs). La plupart des pertes proviennent donc du SLM, et l'alignement des réseaux et miroirs semble être bon. Ces pertes rendent l'utilisation de la ligne plus problématique en cavité laser. L'incorporation de cette ligne en cavité est donc un véritable challenge à relever. L'utilisation hors cavité de la ligne ne pose en revanche pas de problème majeur.

5.3.3 Mesure de la dispersion dans le plan de Fourier

Pour terminer le montage de notre ligne 4f, une mesure de la dispersion du SLM dans le plan de Fourier se doit d'être réalisée. Cette partie permet d'établir une relation entre le numéro des pixels et la longueur d'onde centrale le traversant, qui permettra d'appliquer une phase désirée. En effet, un déphasage étant une différence de phase entre deux pixels consécutifs, la valeur de la phase globale appliquée en fonction de la longueur d'onde dépendra donc de la valeur de la dispersion. Par ailleurs le SLM applique une retardance et il convient de connaître la fréquence (ou bande spectrale) adressée par un pixel afin d'appliquer le déphasage désiré. Pour mesurer la dispersion, une coupure dans le spectre à un pixel déterminé est appliquée. En déterminant la longueur d'onde correspondante à la coupure grâce à une mesure par analyseur de spectre, il devient alors possible d'associer une longueur d'onde à un pixel. Avant le pixel de coupure, la transmission est donc maximale ($T=1$), et elle est minimale après ($T=0$). L'opération est réalisée une dizaine de fois, permettant le calcul de la différence de longueur d'onde pour deux pixels consécutifs (le coefficient directeur de la droite). Les spectres coupés à différents pixels et le graphique présentant la longueur d'onde de coupure en fonction du numéro de pixel sont donnés en figure 5.7.

Cette étude permet de déterminer une dispersion de $\delta\lambda = -0.377 \text{ nm.pix}^{-1}$ (pour un champ autour de 1550 nm). Cette valeur mesurée est en bonne adéquation avec la valeur théorique calculée en section 5.3.1, au signe près. Ce changement de signe indique simplement un sens inverse du SLM dans le montage expérimental, ce qui ne pose aucun problème pour nos expériences. Il est important de constater que cette valeur

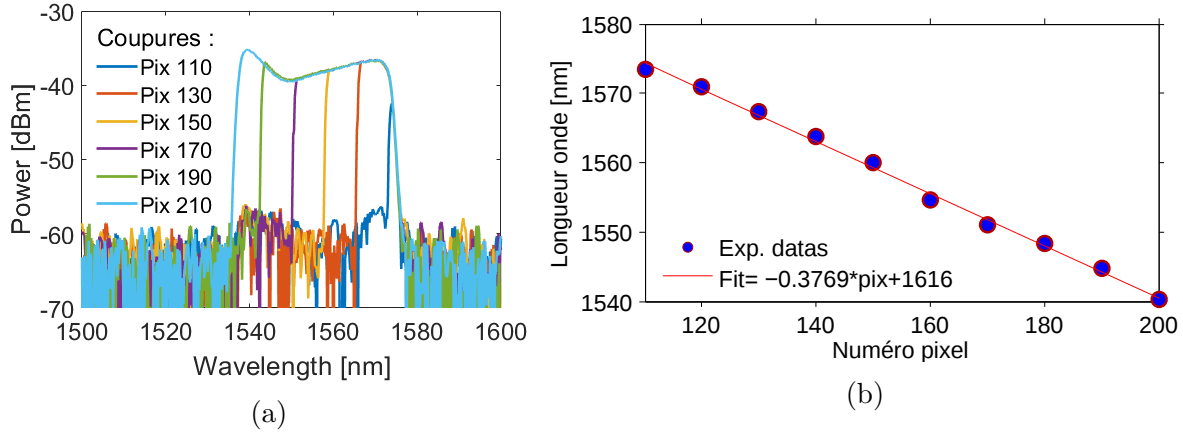


FIGURE 5.7: Graphiques présentant (a) les spectres expérimentaux coupés à chaque numéro de pixel indiqué et (b) le graphique présentant la dispersion du SLM dans le plan de Fourier, associant un numéro de pixel à une longueur d'onde

dépend de la longueur d'onde du laser, mais également de son l'angle d'incidence sur le SLM. La ligne pouvant légèrement se désaligner durant le temps, cette valeur n'est pas rigoureusement constante suivant le temps, et cette étude devra être refaite à partir de quelques semaines d'utilisation. De plus, ces valeurs sont déterminées d'après une mesure par analyseur de spectre, dont la précision peut se détériorer avec le temps (mesure à 0.5 nm près). Pour nos calculs de phase à 1550 nm, cela ne pose aucun problème majeur. Cette valeur de dispersion permet de travailler avec une fenêtre temporelle de 21 ps [137].

Le montage de la ligne et la calibration de celle ci sont réalisés. Avant la présentation des premiers résultats de recompression d'impulsion en extra cavité, certains calculs se doivent d'être réalisés. Ceux-ci seront décrits dans la prochaine section, et permettront de mieux évaluer la qualité de nos recompressions.

5.4 Calculs de caractéristiques impulsionnelles

Pour la suite de nos travaux, certaines problématiques préalables se posent en raison de la forme spectrale et donc temporelle de nos impulsions générées. Les coefficients d'autocorrélation $\Delta t = T_0/K$, avec T_0 la durée de la trace d'autocorrélation à mi hauteur, ainsi que le produit durée-largeur spectrale $\Delta\nu\Delta t = C$ (ou TBP pour time bandwidth product) ne sont renseignées dans les banques de données que pour les solitons plus conventionnels, de forme gaussienne ou sécante hyperbolique. Dans notre cas, les formes spectrales générées en fonction porte indiquent des champs électriques en sinus cardinal. Cependant, limiter la forme spectrale à une fonction porte conduit à des imprécisions notables. Nous essaierons de limiter au maximum ces imprécisions dans nos calculs par évaluation numérique, en débutant par un spectre typique d'une impulsion générée par notre cavité.

5.4.1 Calcul du TBP en limite de Fourier

La trace enregistrée par l'analyseur de spectre optique est évidemment l'intensité spectrale de celui-ci. Pour déterminer sa trace temporelle, il convient tout d'abord de travailler en champ électrique, soit avec la racine carré de l'intensité spectrale $|E(\omega)|$, pour ensuite en faire la transformée de Fourier pour obtenir le champ électrique $E(t)$ en limite de Fourier. La trace temporelle enregistrée sur des appareils de mesures correspondant à l'intensité, il convient pour finir d'élever le champ au carré.



FIGURE 5.8: Schéma présentant le chemin de calcul utilisé pour déterminer une impulsion en limite de Fourier d'après son spectre

Le chemin pris pour calculer la trace temporelle est représentée en figure 5.8. Ce cheminement n'est vrai qu'en limite de Fourier (pour des impulsions non chirpées). Les calculs suivants ne seront donc valables qu'en limite de Fourier.

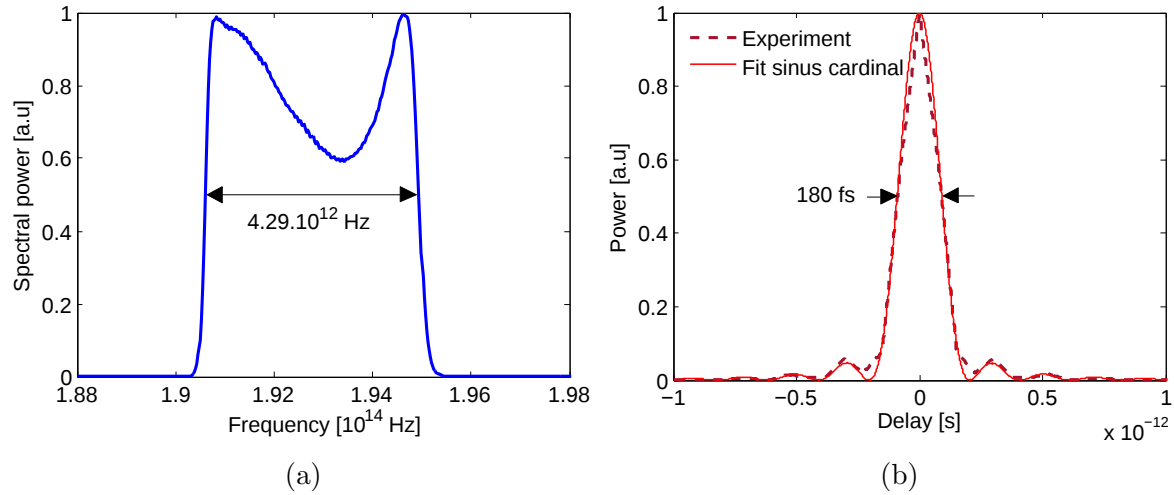


FIGURE 5.9: Graphiques présentant (a) le spectre expérimental utilisé pour cette étude, et (b) son impulsion calculée par transformée de Fourier avec son fit en sinus cardinal.

Le domaine de la trace spectrale enregistrée nous permet d'obtenir une résolution temporelle maximale de 75 fs. Les durées d'impulsions en limite de Fourier sont d'environ 200 fs. Il est alors possible de visualiser une impulsion et d'en faire quelques calculs, mais ceux-ci seront très imprécis. Une interpolation linéaire de l'impulsion puis un fit en sinus cardinal seront donc réalisés pour faire les calculs. La figure 5.9 représente le spectre utilisé pour les calculs, ainsi que l'impulsion calculée et son fit en sinus cardinal. La figure nous indique la largeur spectrale de l'impulsion : $\Delta\nu = 4.29.10^{12}$ Hz, ainsi que sa durée : $\Delta t = 180.10^{-15}$ sec. Nous pouvons facilement en déduire le coefficient en limite de Fourier $C = \Delta\nu\Delta t = 0.77$ pour ce type d'impulsions. Il est intéressant de noter que pour un spectre purement rectangulaire (fonction porte), ce coefficient serait de 0.9, d'où l'intérêt de considérer un spectre expérimental.

5.4.2 Calcul de la constante d'autocorrélation K

Pour déterminer cette constante, il suffit de faire le produit de corrélation de l'impulsion en intensité par elle-même, qui correspond à sa trace d'autocorrélation :

$$S(\tau) = \int I(t)I(t - \tau)dt \quad (5.7)$$

Le produit de corrélation du fit calculé en sous section précédente est réalisé, puis le rapport entre la durée de la trace d'autocorrélation et la durée de l'impulsion est calculé pour déterminer le coefficient K. Ces calculs sont donnés en figure 5.10. Le rapport entre la durée d'autocorrélation et la durée d'impulsion détermine la constante $K=1.32$ pour nos impulsions. Cette valeur étant uniquement dépendante de la forme de l'impulsion, elle est correcte même pour les impulsions ayant une forte dérive de fréquence.

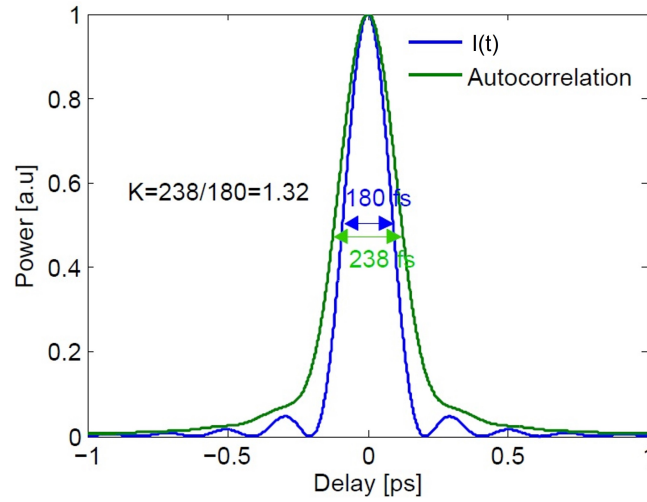


FIGURE 5.10: Graphique présentant l'impulsion calculée d'après le spectre (en bleu) et le produit de convolution de celle-ci avec elle-même (en vert), correspondant à sa trace d'autocorrélation

5.5 Recompression d'impulsions extra-cavité via un algorithme d'évolution

Traditionnellement, un faisceau optique est décrit mathématiquement par son champ électrique $E(t) = E_0(t)e^{i(\omega_0 t + \varphi(t))} = E_0(t)e^{i\phi(t)}$. La fréquence instantanée $\omega(t)$ des oscillations s'écrit alors $\omega(t) = \frac{\partial \phi(t)}{\partial t}$. Lorsque la phase $\varphi(t)$ de l'impulsion reste constante dans le temps ($\varphi(t) = a$), la fréquence instantanée est constante en fonction du temps : $\omega(t) = \omega_0$. Cependant, lorsque $\varphi(t)$ est non linéaire, on a alors $\omega = \omega(t)$. Par exemple, si la phase est quadratique ($\varphi(t) = at^2/2$), la fréquence instantanée s'écrit $\omega(t) = \omega_0 + at$. On parle alors de dérive de fréquence, ou "chirp". Notre cavité expérimentale génère principalement des impulsions avec une forte dérive de fréquence, due à son régime de dispersion normal. Ainsi, leur TBD $\Delta\nu\Delta t$ est très loin de sa valeur en limite de Fourier, indiquant une impulsion longue comparée à sa durée attendue en limite de Fourier avec un spectre large. Notre but est alors de recompresser ces impulsions grâce à la ligne

4f précédemment développée, en extra cavité, en lui appliquant une phase spectrale quadratique $\varphi(\omega)$ correspondante à l'inverse de son chirp. Pour déterminer cette phase, un algorithme d'évolution, similaire à celui utilisé aux travaux précédents est utilisé.

5.5.1 Fonction de mérite utilisée

Pour débiter cette étude, la problématique reste la même que pour les travaux du chapitre 4 : la recherche d'une fonction d'objectif décrivant les attentes de l'utilisateur. L'objectif de ce travail est de recompresser une impulsion en extracavité, il faut donc trouver une grandeur caractéristique de la durée d'une impulsion. L'intensité de doublage de fréquence (ou SHG pour second harmonic generation) répond en partie à cet objectif. En effet, cette SHG repose sur l'utilisation d'un cristal doubleur de fréquence. Lorsque son champ incident est suffisamment intense, la réponse du cristal devient non linéaire, générant alors une onde à fréquence double. La polarisation du milieu devient dans ce cas non linéaire. Cette polarisation, si la réponse du milieu est instantanée, s'écrit de la forme suivante : $P_{NL}(t) = \epsilon_0 \chi^{(2)} E^2(t)$, avec P_{NL} la polarisation non linéaire du milieu, ϵ_0 la permittivité du milieu, $\chi^{(2)}$ la susceptibilité d'ordre 2 du milieu et $E(t)$ le champ électrique vibrant à la pulsation ω . Ce champ s'écrit $E(t) = E_0 \sin(\omega t)$. La polarisation du milieu dans le domaine temporel s'écrit alors :

$$P_{NL}(t) = \epsilon_0 \chi^{(2)} E_0^2 \frac{(1 - \cos(2\omega t))}{2} = \epsilon_0 \chi^{(2)} \left(\frac{E_0^2}{2} - \frac{E_0^2 \cos(2\omega t)}{2} \right) \quad (5.8)$$

Une onde de fréquence doublée à 2ω est alors générée. Dans le cas d'un processus accordé en phase, le champ généré par SHG est proportionnelle à la polarisation non linéaire et donc à l'intensité du fondamental au carré $E_{2\omega} \propto P_{NL} \propto E_\omega^2$. En sachant que l'énergie d'une impulsion s'écrit $\varepsilon_{pulse} = P_{crete} \Delta t$ (pour une impulsion de forme carré), pour une impulsion d'énergie fixe, l'intensité du doublage de fréquence est alors proportionnelle à l'inverse de sa durée Δt au carré¹. Rechercher une optimisation de l'intensité de doublage de fréquence, dans le cas d'un travail en extra cavité, revient alors à une recompression de l'impulsion. Nous précisons que dans le cas d'une optimisation en extra cavité, le régime impulsional n'a pas besoin d'être une grandeur incluse dans la fonction de mérite, la modification de la phase induite par SLM ne modifiant pas cette caractéristique.

5.5.2 Montage développé

Un montage est développé pour permettre l'opération de recompression. Il comporte 4 sous-parties : la cavité laser développée en section 4.1.2, permettant la génération d'impulsions, la ligne 4f, un montage de doublage de fréquence servant de détection et un algorithme d'évolution, dont le fonctionnement est similaire à celui utilisé précédemment mais dont les paramètres diffèrent. Ces derniers sont donnés en tableau 5.1 . Le schéma du montage pour recompresser l'impulsion est présenté en figure 5.11.

Dans notre cas, le doublage de fréquence est réalisé par un cristal de Bêta-borate de baryum de type 2 (BBO), et détectée par un tube photomultiplicateur (PM) à gain. Le

1. $\epsilon_{2\omega} = P_{2\omega} \Delta t_{2\omega} \propto (\frac{\epsilon_\omega}{\Delta t_\omega})^2 \Delta t_{2\omega} \propto \frac{\epsilon_\omega^2}{\Delta t}$ quelle que soit la forme de l'impulsion

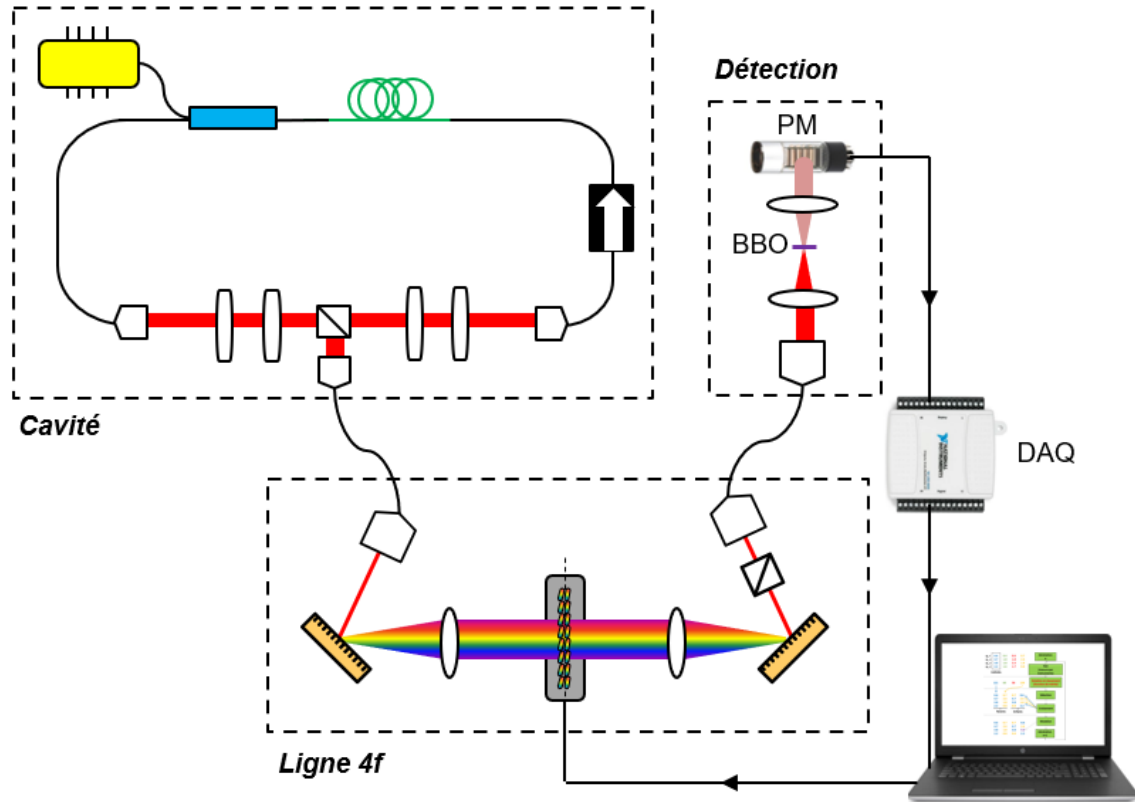


FIGURE 5.11: Schéma présentant le montage développé pour la recompression d'impulsion, avec comme élément de détection un cristal doubleur de type 2 de beta barium borate (BBO), un tube photomultiplicateur (PM) et une carte d'acquisition (DAQ)

montage de détection est similaire à celui utilisé pour l'autocorrélateur. Deux lentilles sont utilisées pour focaliser le faisceau dans le cristal doubleur puis le recollimer, ce dernier étant très fin. Le tube PM convertit une intensité optique du domaine visible en tension et, le doublage de fréquence étant un processus non linéaire, de faibles variations de la durée d'impulsion seront facilement détectées avec ce processus. Le photocourant du tube PM est ensuite mesuré par une carte DAQ, directement branchée sur l'ordinateur. La connexion entre ces deux éléments est très rapide, chaque individu est mesuré en 500 ms avec 100 mesures moyennées. Dans le cas présent, les gènes de l'algorithme seront les phases induites par chaque pixel du SLM (codées en digit), correspondant donc à 320 gènes, et ayant des valeurs comprises entre 0 et 2π . Les spectres des impulsions ne s'étalant que sur une centaine de pixels (≈ 40 nm), l'ensemble de la surface du SLM n'est pas utilisée, et il est possible de réduire le champ d'investigation à 120 gènes. Les populations sont composées de 100 individus, et les optimisations se font sur 100 générations.

5.5.3 Processus de recompression extra cavité

Les conditions initiales de l'algorithme pour le processus de recompression sont présentées dans le tableau 5.1.

Avec ces conditions initiales de travail et un démarrage sur la création d'une population aléatoire, une optimisation est réalisée donnant une évolution de score typique,

5.5. RECOMPRESSION D'IMPULSIONS EXTRA-CAVITÉ VIA UN ALGORITHME D'ÉVOLUTION

Gènes	120 pixels
Fonction de mérite	I_{SHG}
Population initiale	100
Population future	100
P_{pompe} (W)	0.6
Nombre génération	100

TABLE 5.1: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail de recompression d'une impulsion par SLM

présentée en figure 5.12. Celle ci démontre une optimisation de 600% de l'intensité du doublage de fréquence entre le meilleur individu de la génération 0 et l'individu optimal. La variation des scores est faible, notamment due à l'utilisation de valeurs moyennées sur de nombreuses mesures (100 mesures pour chaque individu). Le plateau d'optimisation est atteint autour de la génération 70, ce qui correspond à une durée d'environ une heure. La valeur FTL est déterminée en appliquant une phase constante à chaque début de génération, donnant des indications sur la stabilité du laser pendant le déroulement du procédé. Une valeur FTL n'évoluant pas indique une bonne stabilité des mesures, ce qui est le cas pour l'optimisation.

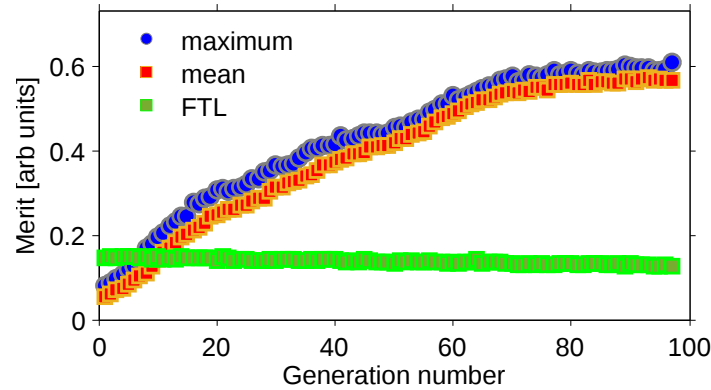


FIGURE 5.12: Graphique présentant une courbe d'optimisation typique obtenue pour la recompression d'impulsion, indiquant l'évolution du meilleur individu et le score moyen de chaque génération, ainsi que la valeur FTL d'une phase constante indiquant la stabilité des mesures

La durée d'impulsion obtenue après recompression est estimée grâce à des mesures d'autocorrélation, données en figure 5.13 (a), indiquant une nette optimisation de celle ci. Sans optimisation, l'impulsion entrante et sortante de la ligne sont similaires, avec des durées d'autocorrélation d'environ 1.2 ps. Ces traces indiquent une ligne 4f correctement alignée, n'induisant pas de dispersion sans SLM. Après recompression, la trace d'autocorrélation de l'impulsion présente une durée de 250 fs, soit environ 5 fois plus courte. Une question peut alors se poser : a-t-on atteint la limite de Fourier de l'impulsion ? Comme nous l'avons calculé précédemment, la durée d'impulsion en limite de Fourier est égale à 180 fs. D'après la trace d'autocorrélation, nous mesurons une durée d'impulsion de $250/1.32=189$ fs. La durée de l'impulsion recompressée est donc pratiquement la plus courte qu'il est possible d'atteindre, celle en limite de Fourier.

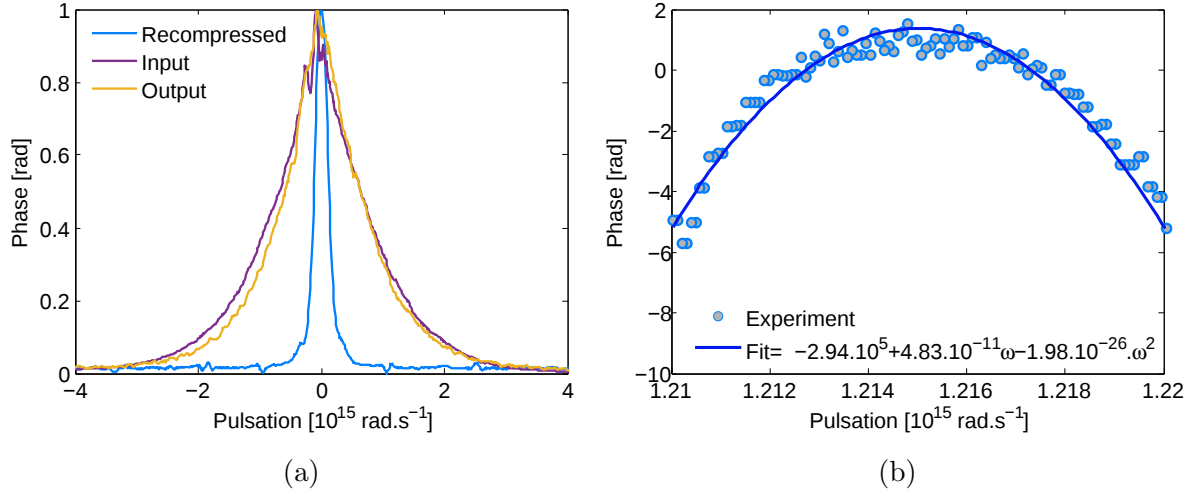


FIGURE 5.13: Graphiques présentant (a) les traces d'autocorrélation enregistrées pour plusieurs impulsions, mettant en évidence l'intérêt du processus de recompression. La trace violette 'Input' représente l'impulsion en entrée de ligne, la trace orange 'Output' représente la trace en sortie de ligne sans recompression et la trace bleue 'Recompressed' est l'impulsion recompressée. (b) l'individu optimal, c.a.d la phase induite par le SLM pour recompresser notre impulsion

La phase spectrale de l'individu optimal est présentée en figure 5.13 (b). La fenêtre de fréquence est limitée à la zone d'intérêt, entre 1530 et 1570 nm. La phase induite est principalement quadratique, comme l'attestent les données expérimentales. Le fit quadratique, de la forme polynomial au second ordre, s'écrivant

$$\phi = \phi_0 + \phi_1\omega + \frac{\phi_2}{2}\omega^2 \quad (5.9)$$

nous permet de déterminer les coefficients : $\phi_0 = 62.94.10^5$, $\phi_1 = 4.83.10^{-11}\text{s}$ et $\phi_2 = -3.96 \text{ ps}^2$, correspondant directement au chirp induit. Le centrage spectral de la phase est de $\omega_0 = 1.218.10^{15} \text{ rad} \cdot \text{s}^{-1}$. Le chirp de l'impulsion correspond donc à l'opposé du chirp induit par recompression, soit de 0.040 ps^2 . Cette optimisation sera considérée comme référence pour la sous section suivante.

5.5.4 Recompression avec ajout de fibre

Dans le but de tester notre processus à d'autres cas de figures, nous décidons d'ajouter respectivement 2m de fibre SMF 28 et 0.5m de fibre DCF entre la ligne 4f et le dispositif de détection (c'est-à-dire le dispositif de doublage de fréquence). La différence de chirp induit pour les différents tests doit correspondre à la dispersion supplémentaire ajoutée par la fibre (soit de $\beta_{2,SMF} = -0.044 \text{ ps}^2$ pour les 2m de SMF et de $\beta_{2,DCF} = 0.039 \text{ ps}^2$ pour la fibre DCF). La figure 5.14 présente les phases induites par le SLM après optimisation avec l'ajout des deux fibres citées précédemment, avec leur fit quadratique. Nous remarquons tout d'abord qu'avec l'ajout de 2m de SMF 28 (figure 5.14 a), la phase induite est presque linéaire, et n'induit donc qu'un faible chirp : $C_{SMF} = 0.0029 \text{ ps}^2$. Cette faible valeur s'explique par le fait qu'avec cette fibre,

5.5. RECOMPRESSION D'IMPULSIONS EXTRA-CAVITÉ VIA UN ALGORITHME D'ÉVOLUTION

l'impulsion est proche de son point de recompression optimal. En effet, la recompression de l'impulsion sans ajout de fibre (la référence) requiert l'induction d'un chirp de $C_{ref} = -0.040 \text{ ps}^2$, soit environ la dispersion induite par les 2m de fibre SMF ajoutés.

Pour l'ajout de fibre DCF (figure 5.14 b), le chirp à induire pour recompression est de $C_{DCF} = -0.080 \text{ ps}^2$. La différence de chirp à induire pour recompression entre la référence et ce montage avec DCF est donc de $\Delta C_{DCF} = -0.040 \text{ ps}^2$. Cette valeur est encore logique, très proche de la dispersion induite par la fibre DCF. Un tel dispositif pourrait alors permettre de calculer par exemple la dispersion d'une fibre inconnue, tant que celle-ci induit une transmission suffisante pour réaliser un doublage de fréquence.

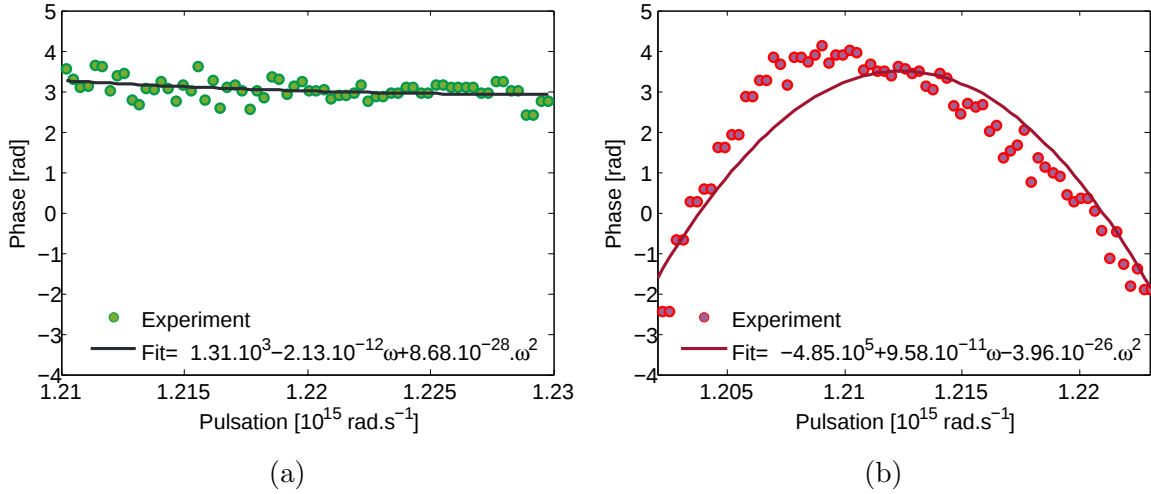


FIGURE 5.14: Graphiques présentant (a) la phase optimale induite pour recompression d'une impulsion avec ajout de 2m de fibre SMF an amont de l'autocorrélateur, et (b) la phase optimale induite pour recompression d'une impulsion avec ajout de 0.5m de fibre DCF an amont de l'autocorrélateur

Il est intéressant de remarquer que la pulsation centrale ω_0 de la phase à appliquer varie selon les optimisations, parfois pour un même dispositif expérimental. Cette variation de fréquence se traduit temporellement par un retard. En effet, pour une phase quadratique $\phi(\omega) = (\omega - \omega_0)^2$, une variation δ de la pulsation centrale donnerait une phase de $\phi(\omega) = (\omega - \omega_0 + \delta)^2 = (\omega - \omega_0)^2 + \delta^2 + \delta(\omega - \omega_0)$. La première contribution correspond au chirp, la seconde à la phase absolue ou CEP (carrier envelope phase) et la dernière contribution est linéaire, ce qui implique un retard temporel. N'utilisant pas d'instrument de détection permettant la mesure de retards temporels ni de CEP, la variation de δ ne joue aucun rôle dans nos optimisations.

En revanche, nous notons que les impulsions recompressées par ces optimisations sont identiques d'un point de vue temporel. Les traces d'autocorrélation de ces impulsions sont données en figure 5.15, où nous remarquons une durée d'impulsion proche de celle en limite de Fourier (environ 190 fs). Cela indique une optimisation complète quels que soient les montages expérimentaux.

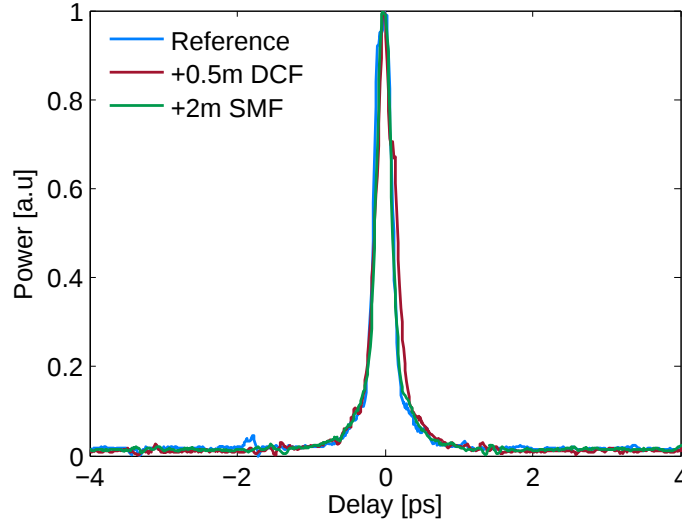


FIGURE 5.15: Graphique présentant les différentes traces d'autocorrélation enregistrées après recompression, pour les montages précédemment évoqués.

Cette étude cloture la section sur la recompression d'impulsions extra cavité. Ainsi, des recompressions d'impulsion ont pu être réalisées par notre dispositif jusqu'à leur durée en limite de Fourier. L'utilisation de ce dispositif pourrait permettre de retrouver la dispersion inconnue d'une fibre grâce à la mesure de la phase à appliquer pour recompression. Il pourrait permettre également de recompresser des impulsions dont la phase spectrale ne serait pas quadratique, et donc plus difficile à compenser par l'utilisation de fibres. Dans notre cas, les phases appliquées sont quadratiques. Il serait alors possible de simplifier des futures optimisations en appliquant directement des phases quadratiques et en les paramétrisant $\phi = \phi_0 + \phi_1\omega + \frac{\phi_2}{2}\omega^2$. La prochaine section présente le véritable objectif du développement de la ligne 4f, c'est à dire son inclusion dans une cavité fibrée comme dans les travaux de Iegorov [138] pour façonner une impulsion dès sa création, augmentant considérablement le nombre de degré de liberté de notre cavité par rapport aux travaux du premier chapitre.

Chapitre 6

Optimisations intégrant la modulation de phases spectrales par le SLM

de 0.6 m pour l'amplificateur. Le blocage de mode est toujours obtenu par rotation non linéaire de polarisation, et le montage précédent utilisant des lames à retard à cristaux liquide et un PBS est utilisé pour créer le mécanisme d'absorbant saturable. La seconde cavité développée pour ces travaux est présentée en figure 6.1. La dispersion globale de cette nouvelle cavité est similaire à celle présentée en figure 4.3, avec une dispersion globale de $\beta_{2,net} = 0.06 \text{ ps}^2$ mais avec un taux de répétition plus faible car la cavité est plus longue : $f_{rate} = 39.9 \text{ MHz}$. Cette configuration permet l'induction d'une phase spectrale par le SLM dans la cavité, ainsi qu'une modulation en intensité, permettant le façonnage des impulsions générées dès leur création (fonction de transfert spectrale du SLM $H(\omega)$ de l'équation 5.1 directement dans la cavité).

L'utilisation du SLM pour augmenter le nombre de degré de liberté de la cavité nous a permis de réaliser plusieurs travaux. Dans les prochaines sections, ceux-ci seront présentés dans l'ordre chronologique, permettant de présenter également le cheminement scientifique et pratique menant des travaux simples aux travaux plus complexes. Mais pour débiter cette étude, nous nous intéresserons uniquement à l'application de phases spectrales et des optimisation sur des caractéristiques familières, comme la largeur spectrale.

6.2 Optimisation sur la largeur spectrale

6.2.1 Paramètres de l'algorithme

Les 4 phases induites par les LCs restent des paramètres importants à contrôler, façonnant la fonction de transfert de l'absorbant saturable. Cependant, les degrés de liberté provenant du SLM demandent réflexion quant à leur nature et leur importance dans le processus d'optimisation. Pour éviter de travailler avec un nombre de gènes trop élevé, impliquant un nombre de génération très important pour obtenir un optimum, la phase induite par le SLM peut être paramétrisée. Un paramètre intéressant à contrôler pour des cavités laser pourrait être le régime dispersion. En effet, le contrôle de la GVD de la cavité permet la génération de nouvelles familles d'impulsions, dont la forme, la durée ou la largeur spectrale sont différentes. Dans cette première section, la phase induite par le SLM φ_{SLM} sera donc une phase quadratique, ce qui correspond à une dispersion, paramétrisée avec sa GVD. La phase induite par SLM est donnée en équation suivante, avec $\beta_{2,SLM}$ la GVD induite par SLM :

$$\varphi_{SLM} = \frac{\beta_{2,SLM}}{2}(\omega - \omega_0)^2 \quad (6.1)$$

Ce dernier paramètre constitue le dernier gène de l'algorithme. Ainsi, 5 gènes seront utilisés dans cette première étude intracavité. Pour se familiariser avec ce nouveau dispositif, un premier objectif similaire aux travaux de la section 4.1, à savoir l'optimisation des régimes de blocages de mode en fonction de leur largeur spectrale est réalisé. La fonction de mérite de l'algorithme reste donc identique à celle de la section 4.1.1 : $Merite = P_{RF} - P_{bck} + \alpha \Delta \lambda_{DFT}$. Les caractéristiques du fonctionnement de l'algorithme sont résumées dans le tableau 6.1.

Gènes	$\varphi_{1,2,3,4} + \beta_{2,SLM}$
Fonction de mérite	$P_{RF} - P_{bck} + \alpha \Delta \lambda_{DFT}$
Population initiale	100
Population finale	20
P_{pompe} (W)	1.0
Nombre génération	15
Taux d'enfants mutés	70 %

TABLE 6.1: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail sur l'optimisation de la largeur spectrale avec 4 LCs et la GVD induite par SLM comme gènes

6.2.2 Convergence vers un blocage de mode à largeur spectrale contrôlable

Une optimisation typique avec cette configuration donne une évolution des scores similaire aux travaux précédents, montrant une augmentation du meilleur score sur quelques générations, puis une stagnation du résultat jusqu'à la fin du processus. Le score moyen de chaque génération augmente également jusqu'à atteindre une valeur asymptotique proche du maximum. Une courbe d'optimisation typique pour $\alpha = 20$ est donnée en figure 6.2. Une optimisation de 80% du score optimal de la génération 0 et une stabilisation de ce score à partir de la génération 5 sont observés dans ce cas. Une optimisation du score moyen et stabilisation à partir de la génération 11 sont également observés, indiquant une convergence de toute la population.

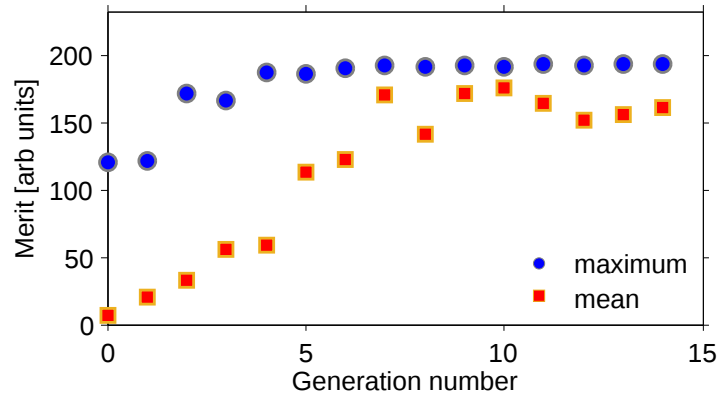


FIGURE 6.2: Graphique présentant l'évolution du meilleur score de chaque génération pour une optimisation typique avec la GVD comme gène, et avec le poids relatif $\alpha = 20$

De même que dans les travaux de la section 4.1.2 et 4.1.3, les états générés sont stables, spectralement larges et sont constitués de solitons uniques. La même méthode de contrôle de la largeur spectrale est opérée, avec un ajustement du poids relatif α , permettant la génération d'impulsions larges spectralement de 20 à 60 nm avec des poids relatifs compris entre $-2 < \alpha < 40$, comme en atteste les figures 6.3 (a) et (b).

Il serait alors intéressant alors d'étudier les résultats de cette étude sous le point de vue de la GVD induite par SLM : existe-t-il des régimes de dispersion générant des largeurs spectrales toujours faible, et d'autres pour lesquels elles sont toujours larges ?

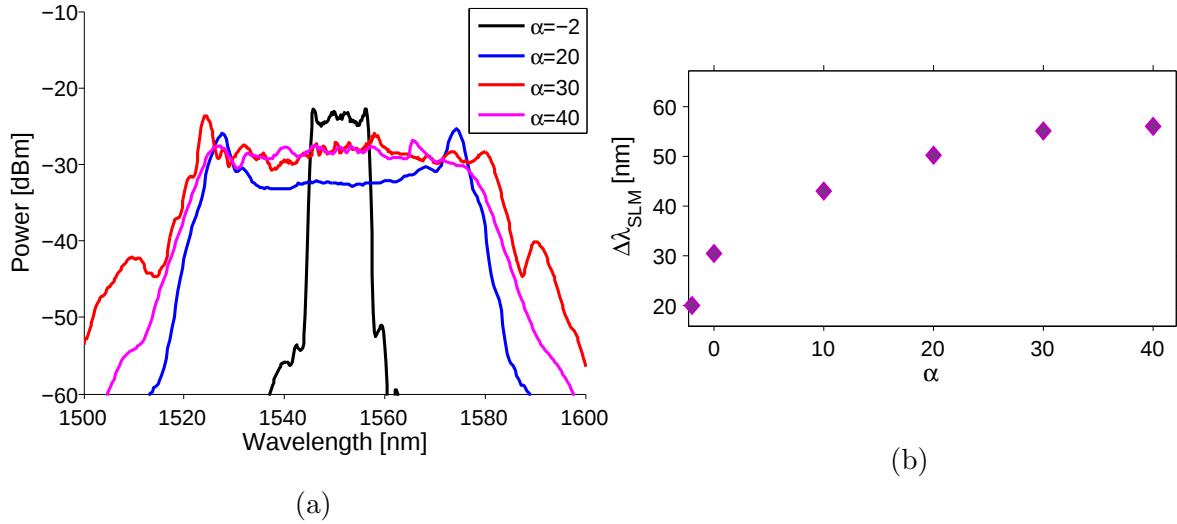


FIGURE 6.3: Graphiques présentant (a) les spectres optimaux obtenus après optimisation en fonction du poids relatif α dans la fonction de mérite, (b) la largeur spectrale optimale obtenue après optimisation en fonction du poids relatif α dans la fonction de mérite

Pour répondre à cette question, nous nous intéressons aux GVDs optimales globales, obtenues après optimisation, en ajustant le poids relatif α pour jouer sur les largeurs spectrales. Les résultats optimaux se situent autour de la dispersion nulle, sauf pour

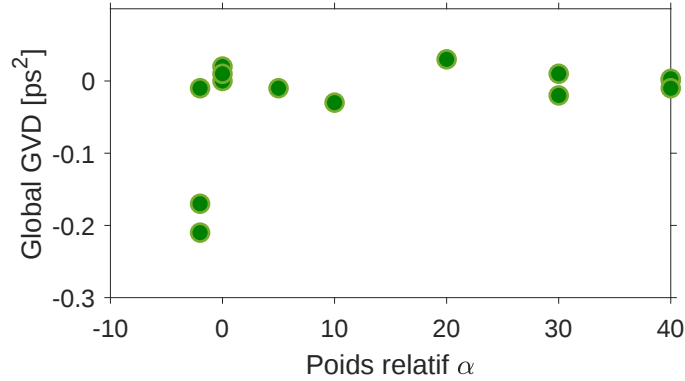


FIGURE 6.4: Graphique présentant l'évolution de la GVD globale optimale en fonction du poids relatif α

quelques optimisations à poids relatif négatif. Cela s'explique par une très grande diversité des régimes potentiellement générés par ces régimes de dispersion caractéristiques, tant par leurs formes spectrales que temporelles. Quelles que soient les largeurs spectrales désirées, un ajustement de la dispersion globale vers 0 est bénéfique, et le simple ajustement des lames à retard LCs est suffisant pour générer des régimes de toutes les largeurs spectrales. Pour des dispersion très anormales, les largeurs spectrales sont généralement très faibles, expliquant les points à GVD fortement anormales pour des poids relatif liés à des largeurs spectrales faibles, c'est à dire $\alpha < 0$. La conclusion de ce travail demande alors l'intérêt d'inclure la GVD comme gène de l'algorithme. Une simple gestion de dispersion par l'utilisation d'une fibre DCF et un ajustement des LCs

est alors suffisant, et plus rapide. L'utilisation du SLM devient alors superflue pour des optimisations sur la largeur spectrale avec seule la GVD induite comme gène. Une phase spectrale plus complexe aurait pu être implémentée comme gène.

Hormis une fenêtre de largeurs spectrales atteignables plus importante que pour les travaux précédents, ce travail ne présente qu'un intérêt limité. Dans le but d'exploiter le contrôle des phases spectrales par le SLM au mieux, nous sélectionnons une nouvelle fonction d'objectif pour générer de nouveaux états où la phase spectrale pourrait jouer un grand rôle. Nous décidons pour cela d'optimiser les régimes sur leurs intensités de second harmonique SHG, mais cette fois-ci avec le pulse shaper directement en intra-cavité.

6.3 Optimisation sur la SHG

6.3.1 Paramètres de l'algorithme

L'objectif de cette section est de générer une impulsion stable, à forte énergie et la plus brève possible. En optimisant sur la SHG, on souhaite optimiser son énergie et sa durée. La première problématique à se poser pour motiver ce travail est l'effet de la phase spectrale, et plus précisément de la GVD sur la SHG d'un régime. En d'autres termes, quel est l'impact de la GVD sur la durée d'une impulsion, ou sur son énergie. Une étude simple a été réalisée au préalable pour visualiser cet effet. Hors algorithme, la GVD est ajustée manuellement sur un état de blocage de mode stable préalablement optimisé, et la réponse en SHG est enregistrée. Le dispositif de détection utilisé en extra cavité (figure 5.11) est réutilisé pour la SHG. Le résultat de cette étude est présenté en figure 6.5 pour deux régimes de blocage de mode différents. Cette figure démontre deux comportements

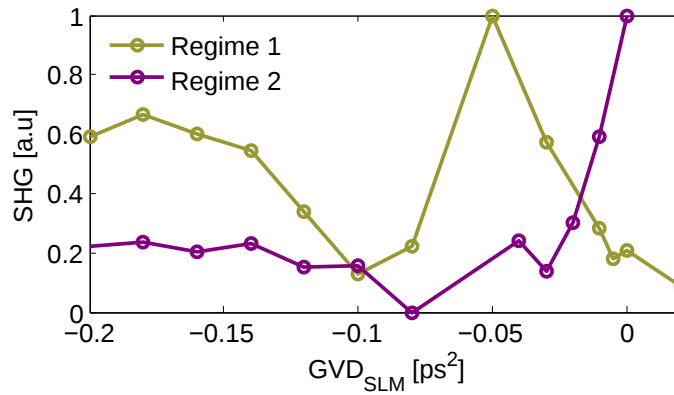


FIGURE 6.5: Graphique présentant l'intensité de SHG d'un régime de blocage de mode en fonction de la GVD induite par SLM

distincts selon les régimes déterminés par deux combinaisons différentes $\varphi_{1,2,3,4}$ des LCs. Ainsi, une combinaison non prévisible de $\varphi_{1,2,3,4}$ et de $\beta_{2,SLM}$ est nécessaire pour atteindre une génération de SHG maximale. Autrement dit le comportement de la SHG en fonction de la GVD n'est pas monotone. Un algorithme d'évolution semble donc appropriée pour trouver une combinaison optimale de ces paramètres.

Des optimisations via l'algorithme vont donc être réalisées avec des conditions initiales de fonctionnement identiques à celles des travaux de la section précédente à l'exception de la fonction de mérite qui va différer (tableau 6.2). La fonction de mérite

Gènes	$\varphi_{1,2,3,4} + \beta_{2,SLM}$
Fonction de mérite	$P_{RF} - P_{bck} + \gamma I_{SHG}$
Population initiale	100
Population finale	20
P_{pompe} (W)	1.0
Nombre génération	15
Taux d'enfants mutés	70 %

TABLE 6.2: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail sur l'optimisation de la SHG avec 4 LCs et la dispersion induite par SLM comme gènes

utilisée est de la même famille que celle présentée en section 3.5, comportant deux contributions. La première contribution reste la puissance RF, la différence se situe dans la nature de la seconde contribution : l'intensité de seconde harmonique SHG. Elle permet donc d'optimiser un régime impulsionnel selon sa stabilité et sa SHG, le poids relatif noté γ dans cette configuration (pour le différencier de α se référant à l'optimisation de la largeur spectrale), permet de favoriser une contribution par rapport à l'autre. Ce poids relatif n'étant pas associé à la même grandeur, sa valeur va également être différente que pour l'optimisation de la largeur spectrale.

6.3.2 Convergence vers un blocage de mode à impulsion brève

Une courbe d'optimisation typique pour avec de telles conditions initiales, avec un poids relatif $\gamma = 15$, présente une évolution des scores donnée en figure 6.6.

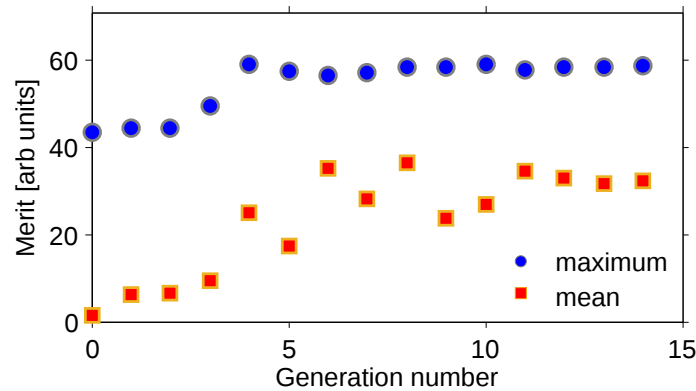


FIGURE 6.6: Graphique présentant l'évolution du meilleur score de chaque génération pour une optimisation typique avec la GVD comme gène pour $\gamma = 15$

Le comportement est similaire aux optimisation précédentes, avec une évolution du meilleur score sur 5 générations puis une stagnation de celui-ci jusqu'à la fin de l'optimisation, montrant une réelle convergence de l'algorithme vers un état optimal.

L'optimisation par rapport au meilleur individu de la génération initiale est de 30%. La moyenne des scores va également converger vers son maximum, moins important que le score optimal, et avec plus de fluctuations. Les caractéristiques de l'individu optimal sont présentées dans la figure 6.7.

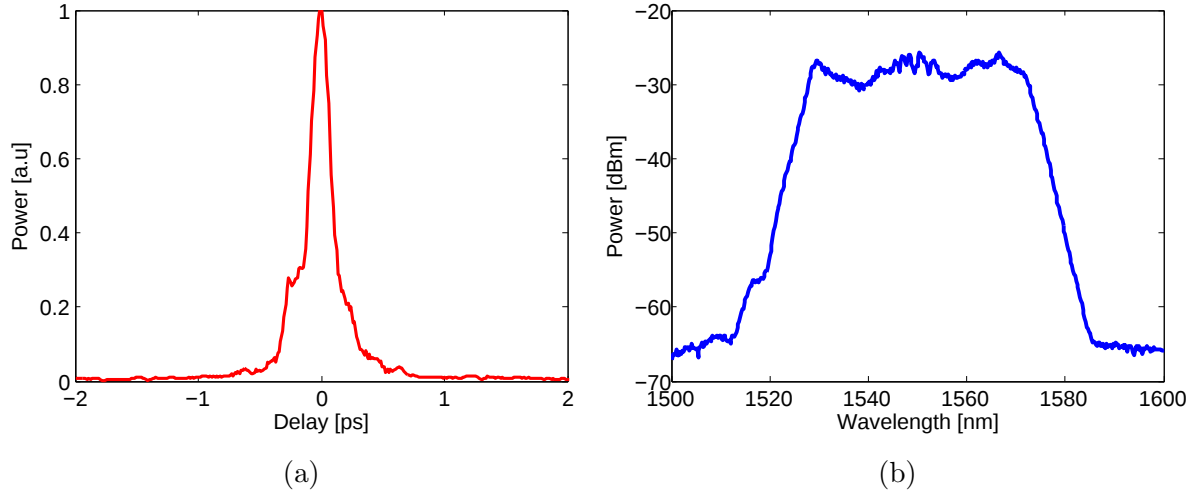


FIGURE 6.7: Graphiques présentant (a) la trace d'autocorrélation de l'individu obtenu avec une optimisation avec un poids relatif $\gamma = 15$, et (b) son spectre.

Il s'agit d'un blocage de mode stable et conventionnel, mais dont l'impulsion a une durée de 150 fs, calculée d'après sa trace d'autocorrélation (durée d'autocorrélation de 190 fs). Cette impulsion est donc plus brève que celles générée par recompression extra cavité (section 5.3.3, dont l'impulsion recompressée avait une durée de 189 fs). De plus, compte tenu de la largeur spectrale de l'impulsion (45 nm), elle aurait en limite de Fourier une durée de 135 fs. Le procédé permet donc de générer des impulsions très courtes, et très proche de la limite de Fourier.

6.3.3 Ajustement de la durée d'impulsion

Une variation du poids relatif γ , promouvant le deuxième objectif de la fonction par rapport au premier, permet la génération automatique d'impulsions à durée variables. Trois traces d'autocorrélation, provenant de trois régimes générés par le processus avec des valeurs de $\gamma = -5$, 0, et 15 sont présentées en figure 6.8(a). Pour une optimisation avec $\gamma = -5$, les fluctuations autour du délai nul laisse penser que l'individu optimal est un régime multisolitonique, constitué d'un long paquet d'ondes. La variation des poids relatifs sur cette plage démontre un ajustement de la durée des impulsions, avec un comportement monotone. Plus ce poids est élevé, plus les impulsions générées seront brèves (figure 6.8(b)). Ainsi, la plage de durée d'impulsion atteignable par ce procédé est de 150 fs à 2.7 ps, pour des poids relatifs $-5 < \gamma < 15$. Les largeurs spectrales, dont la variation selon γ est présentée en figure 6.8(b), ne varie pas de manière monotone selon celui-ci même si une augmentation globale est démontrée. Le produit temps-fréquence TBP $\Delta\nu\Delta t$, calculé d'après ces deux grandeurs (figure 6.8(c)), diminue avec une augmentation du poids relatif de manière monotone. Celui-ci tend vers une asymptote horizontale d'impulsions en durée de Fourier pour $\gamma > 5$ (on rappelle qu'en

limite de Fourier, $\Delta\nu\Delta t = 0.77$). Nous pouvons en conclure que pour des poids relatifs $\gamma < 5$, l'algorithme génère des régimes avec un fort chirp.

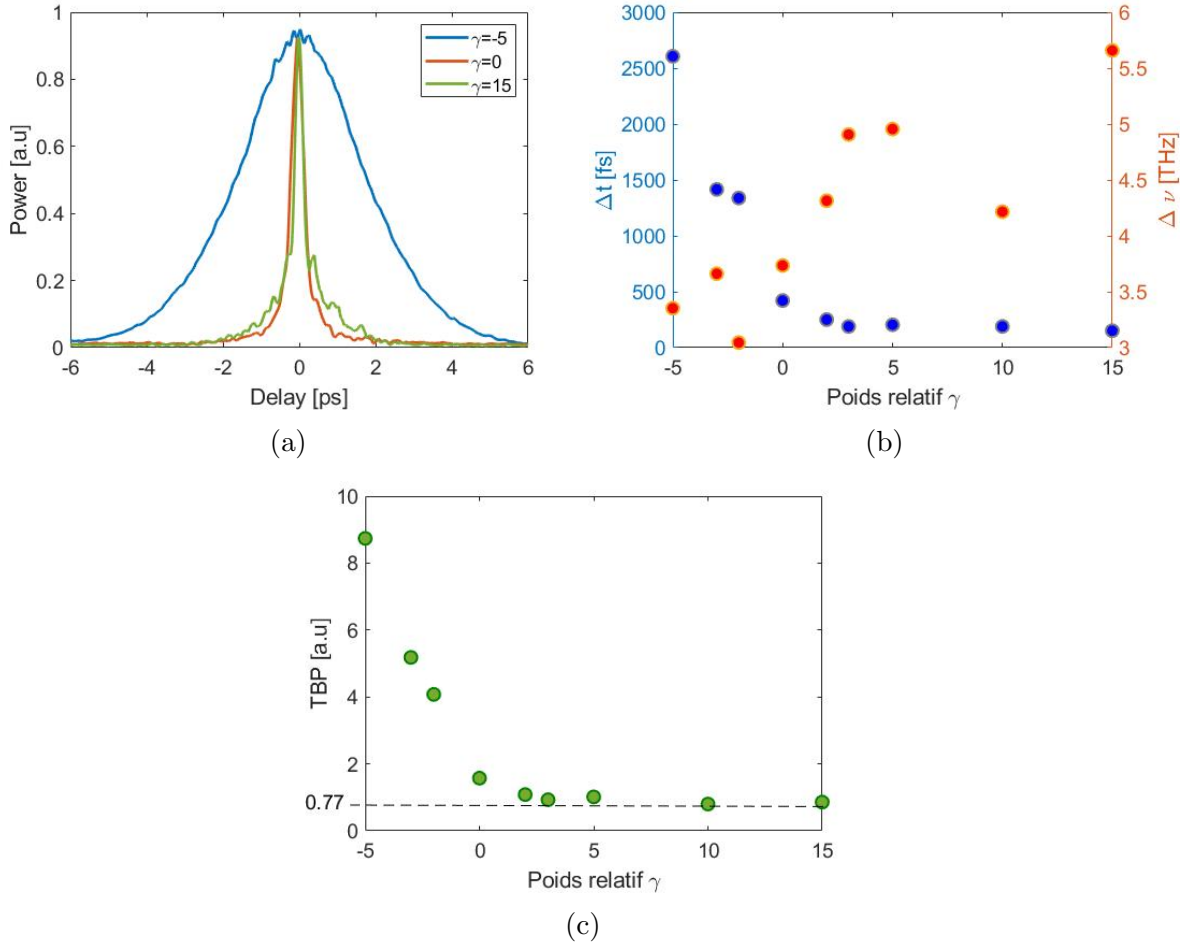


FIGURE 6.8: Graphiques présentant (a) les traces d'autocorrélation d'individus obtenus par optimisation avec des poids relatif γ différents, et (b) La durée impulsionnelle optimale Δt et la largeur spectrale à mi hauteur $\Delta\nu$ en fonction du poids relatif γ de la fonction de mérite, et (c) le TBP $\Delta\nu\Delta t$ calculé d'après les deux valeurs précédentes

En terme de dispersion, il est impossible de conclure quant à une éventuelle convergence de l'algorithme vers un régime de dispersion bien défini. La figure 6.9 présente les différentes GVD globales de individus optimaux $\beta_{2,opt}$, en fonction du poids relatif γ de la fonction de mérite. Le nuage de point est clairsemé, avec des points disséminés sur toute la surface de travail. Ces résultats rejoignent l'étude préliminaire de cette sous section, démontrant une combinaison de $\varphi_{1,2,3,4}$ et de β_2 non prédictible. L'algorithme est donc bien nécessaire à la détermination de ces combinaisons, et son utilisation fait sens dans ce cas. En revanche, ce nuage clairsemé ne permet aucune conclusion sur l'étude. A noter que cette étude se concentre sur la durée des impulsions or, comme énoncé en section précédente, une optimisation sur la SHG maximise à la fois la durée d'impulsion mais également son énergie. Ceci peut donner un début d'explication sur l'aspect clairsemé du nuage, et une étude incluant les énergies d'impulsion pourrait être

réalisée pour plus de compréhension.

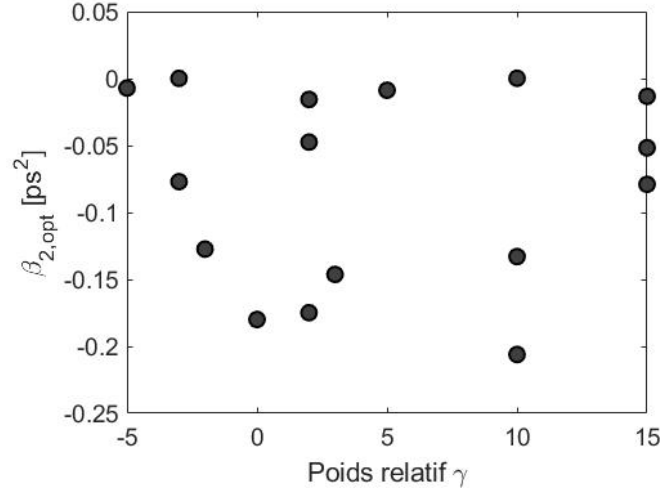


FIGURE 6.9: Graphique présentant l'évolution de la GVD globale des individus optimaux $\beta_{2,opt}$, en fonction du poids relatif γ de la fonction de mérite

Ces premiers résultats sont prometteurs, et montrent l'intérêt d'utiliser le SLM en cavité. Cependant, son potentiel d'utilisation ne semble pas complètement exploité. Il serait intéressant d'étudier l'effet des dispersions d'ordres plus élevées (d'ordre 3 ou d'ordre 4 comme pour les travaux de Runge [139]) sur la génération d'impulsions ultracourtes. Nous choisirons une continuité différente, avec le but d'exploiter les capacités du SLM pour générer des régimes impulsions plus complexes, ou pour des applications plus concrètes. Ainsi, ces résultats concluent la section présentant uniquement le contrôle de phase spectrale comme gène pour l'optimisation de deux fonctions de mérites différentes. Pour mieux utiliser le potentiel du pulse shaper, des travaux sur la modulation spectrale en intensité, et non plus en phase, sont réalisés et présentés dans la section suivante.

Chapitre 7

Optimisations intégrant la modulation de l'intensité spectrale par le SLM

Pour une modulation de l'intensité spectrale, l'ajout d'un second polariseur en fin de ligne 4f est requis, comme expliqué en section 5.1.2. Dans notre montage expérimental, un cube séparateur de polarisation, identique à celui utilisé pour créer le mécanisme de rotation non linéaire de polarisation, sera utilisé juste avant les lames de phase à cristaux liquide 3 et 4. La partie "pulse shaping" de la cavité, comprenant la ligne 4f et le SLM, est donc bordé par deux cubes séparateurs de polarisation : le premier est utilisé pour injecter une polarisation horizontale dans la ligne 4f, permettant une modulation optimale par les SLM, et le second polariseur placé en fin de ligne est utilisé pour moduler l'amplitude spectrale.

Pour tenter d'exploiter le potentiel du façonnage d'impulsion intracavité, une première étude portant sur la génération d'impulsion ajustable en énergie et en largeur spectrale est réalisée, et présentée dans la section suivante.

7.1 Optimisation sur l'énergie de l'impulsion et ajustement de la largeur spectrale

7.1.1 Paramètres de l'algorithme

Dans cette étude, le régime laser sera contrôlé sur 3 critères : la stabilité du régime impulsionnel, la largeur spectrale de celui-ci et son énergie. Ce travail présente donc un pas en avant par rapport aux études précédentes, ne permettant le contrôle de régimes impulsionnels que sur deux critères. Les optimisations par algorithme seront réalisées selon la même famille de fonction de mérite à double objectifs : la puissance RF et la puissance moyenne des régimes (proportionnelle à l'énergie d'impulsion), et un filtre optique à bande passant variable et prédéfinie sera appliqué par le pulse shaper.

Premièrement, le souhait est de contrôler l'énergie d'une impulsion E_{pulse} . Celle-ci est liée directement à la puissance moyenne du régime : en régime impulsionnel, elle peut s'écrire sous la forme : $P_{moyenne} = E_{pulse} \cdot f_{rate}$, avec $f_{rate} = 39.9$ MHz. La mesure de la puissance moyenne d'un régime impulsionnel, avec un simple puissance-mètre permet donc une détermination facile de l'énergie d'une impulsion. Cette mesure sera utilisée dans la fonction de mérite de notre algorithme, en tant que second objectif. Le premier objectif reste la puissance RF au taux de répétition fondamental de la cavité, et un poids relatif que nous nommerons σ dans ce cas précis pour le différencier des autres est utilisé pour favoriser une contribution par rapport à l'autre. Ainsi, comme pour les études précédentes, un contrôle de l'énergie des impulsions pourra être réalisé grâce à cette configuration. Cependant, un contrôle de la largeur spectrale est également souhaité dans ces travaux.

L'idée serait de générer automatiquement des impulsions d'une énergie prédéfinie, avec une largeur spectrale définie également. Le contrôle de l'énergie d'une impulsion se fait par l'ajustement du poids relatif σ , et le contrôle de la largeur spectrale sera réalisée par la modulation spectrale appliquée par le SLM. Ainsi, la transmission du SLM sera une fonction porte, centrée autour d'une longueur d'onde centrale $\lambda_{0, filtre}$ et

avec une bande passante BP_{filtre} . Ces deux paramètres sont fixés au préalable, avant le commencement de l'optimisation par algorithme, et seules les lames à retard de phase contrôlant la rotation non linéaire de polarisation sont implémentés comme gènes. Les nombres d'individus pour chaque génération de l'algorithme, ainsi que le taux d'enfants mutés restent les mêmes que pour les études précédentes. La configuration initiale des optimisations est résumée dans le tableau 7.1

Gènes	$\varphi_{1,2,3,4}$
Fonction de mérite	$P_{RF} - P_{bck} + \sigma P_{moy}$
Population initiale	100
Population finale	20
P_{pompe} (W)	1.0
Nombre génération	15
Taux d'enfants mutés	70 %

TABLE 7.1: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail sur l'optimisation de la puissance moyenne, avec 4 LCs comme gènes et une transmission spectrale prédéfinie

7.1.2 Convergence vers un régime de blocage de mode

La première campagne de test réalisée pour cette étude visait à déterminer la relation entre la puissance moyenne de l'individu optimal et la valeur du poids relatif σ . Une bande passante test à $BP_{filtre} = 3$ nm est tout d'abord imposée par le SLM, avec une longueur d'onde centrale à $\lambda_{0,filtre} = 1550$ nm. Une optimisation est réalisée, avec $\sigma = 300$, menant à la courbe d'optimisation présentée en figure 7.1. Cette courbe présente une optimisation du meilleur individu de 20%, (de 40 à 50) sur 3 générations. Le processus d'optimisation atteint une population optimale à la génération 6, visible sur l'évolution du score moyen. Ces résultats restent similaires à ceux des précédentes études, traduisant des optimisation rapides (moins de 10 minutes). Après ce seuil, le score optimal varie très peu, indiquant un régime stable.

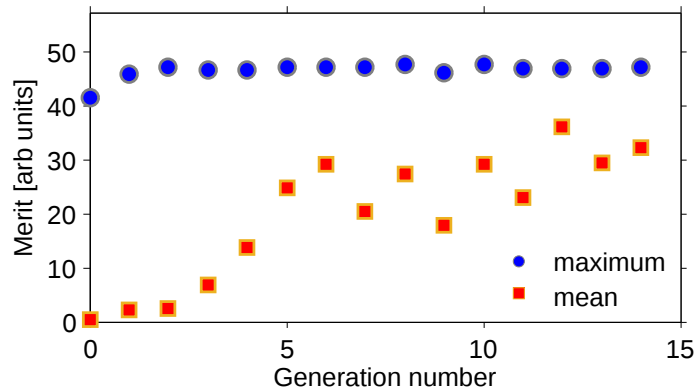


FIGURE 7.1: Graphique présentant l'évolution du meilleur score de chaque génération pour une optimisation typique avec une bande passante du filtre $BP_{filtre} = 10$ nm

L'optimisation mène au régime dont le spectre et la trace d'autocorrélation sont présentées en figure 7.2. Le spectre, mesuré au niveau des pertes du PBS, est fin (6 nm de large). Cette largeur, étant supérieure à la bande passante du filtre, peut être surprenante au premier abord. Cette différence s'explique par un élargissement spectral induit par les milieux à gain, situés entre le filtre et le PBS. Ainsi, pour les bandes passantes fines, on observe un élargissement du spectre dans le milieu à gain, ce qui n'est pas le cas pour les régimes spectralement larges. La relation entre la largeur spectrale mesurée et la bande passante du filtre n'est pas directe, et un travail se doit d'être réalisé pour la mesurer. Au niveau temporel, la trace d'autocorrélation est très longue (2.7 ps) par rapport aux précédents travaux sans modulation d'amplitude, dont les durées d'impulsions avoisinaient la centaine de femtoseconde. Cette relation est logique, une impulsion spectralement large est souvent très brève. En revanche, la trace d'autocorrélation présente un régime multi-impulsionnel, avec plusieurs fluctuations au centre. Il semblerait donc que l'utilisation d'un filtre à bande passante fine a tendance à générer des paquets d'ondes, malgré la dispersion normale de la cavité.

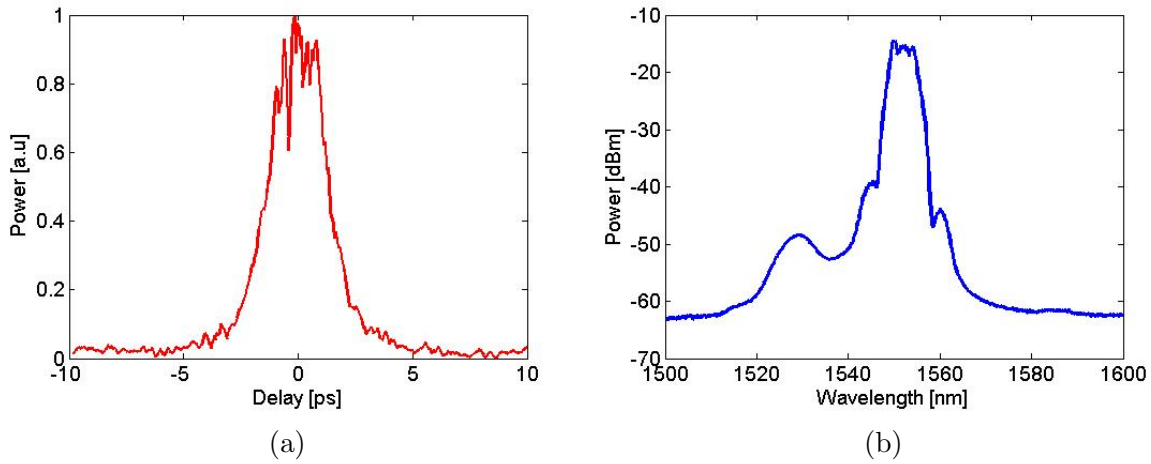


FIGURE 7.2: Graphiques présentant (a) la trace d'autocorrélation du meilleur individu obtenu par optimisation sur la puissance moyenne, et (b) son spectre

Focalisons nous sur la puissance moyenne à présent. L'évolution de la puissance moyenne au cours du processus, ainsi que celle de la puissance radiofréquence (figure 7.3) indique une optimisation de ces deux grandeurs. La puissance RF est optimisée de 25% et la puissance moyenne de 23%. La faible augmentation de la puissance RF s'explique par le poids relatif σ favorisant fortement l'optimisation de la puissance moyenne. Après optimisation, elle est de 26 mW, ce qui correspond à une énergie de $E_{pulse} = 0.7$ nJ.

Pour déterminer l'énergie des impulsions dans la cavité, un rapide calcul est réalisé.

- La puissance du champ extra cavité est mesurée après tout un dispositif de diagnostic. Le faisceau, juste après le PBS, est recouplé dans un collimateur (environ 20% de pertes), et 90% de ce signal est récupéré pour la mesure de la puissance moyenne (les 10% restant sont envoyés sur l'oscilloscope pour calculer la puissance RF). Compte tenu de la puissance moyenne mesurée ($E_{mesure} = 26$ mW),

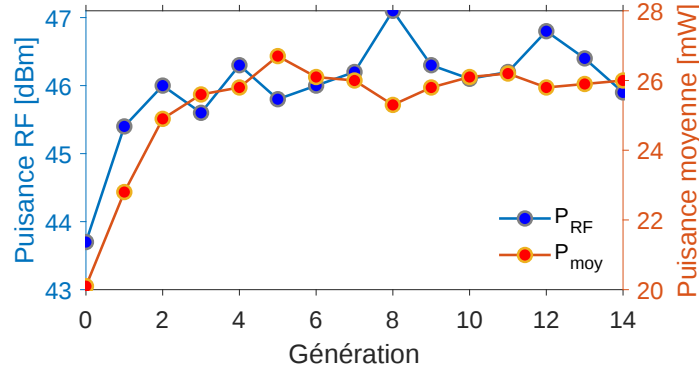


FIGURE 7.3: Graphique présentant l'évolution de la puissance radiofréquence et de la puissance moyenne au cours de l'optimisation avec un poids relatif $\sigma = 300$, en fonction du numéro de génération

la puissance des pertes induites par le PBS juste avant recouplage extracavité est de 36 mW.

- La pente d'efficacité du laser, présentant le rapport $\frac{P_{Laser}}{P_{Pompe}}$ est mesuré à 11 %. La puissance de pompe étant fixée à $P_{Pompe} = 1$ W, la puissance totale générée par le laser est de $P_{Laser} = 110$ mW.
- Ainsi, la puissance du laser restante dans la cavité est de $P_{Laser,intra} = 110 - 36 = 74$ mW, ce qui correspond à une énergie d'impulsion intracavité de $\epsilon_{intra} = 1.8$ nJ.

7.1.3 Contrôle de la largeur spectrale et de l'énergie de l'impulsion

La variation de la bande passante du filtre, lorsqu'elle est réalisée directement avec un régime de blocage de mode, est la plupart du temps synonyme de perte de l'impulsion. Pour obtenir un blocage de mode stable, avec une puissance moyenne optimale, une optimisation par l'algorithme est donc requise. Une pré détermination de BP_{filtre} avant une optimisation par l'algorithme permet alors la génération de régimes solitoniques avec une largeur spectrale proche de BP_{filtre} , mais également une puissance moyenne optimale. Pour chaque optimisation, cette puissance moyenne optimale reste constante quelle que soit la bande passante du filtre, comme illustré en figure 7.4 (a). L'intérêt de l'utilisation de l'algorithme est donc encore démontré. En revanche, sans surprise, la largeur spectrale évolue, et peut facilement être contrôlée. Pour générer automatiquement une impulsion de puissance moyenne de 29 mW (une énergie de 0.73 nJ), avec une certaine largeur spectrale désirée, il suffit alors de pré déterminer la bande passante du filtre optique, et de laisser l'algorithme optimiser le régime impulsionnel selon sa puissance moyenne. Ce procédé permet donc un contrôle de la largeur spectrale comme pour les travaux de la section 4.4, mais cette fois ci avec une optimisation sur la puissance moyenne. Comme énoncé dans le paragraphe précédent, la largeur spectrale n'est pas exactement la bande passante du filtre. Elle est plus importante pour $BP_{filtre} < 20$ nm, mais est légèrement plus importante au delà de ce seuil. Avec ce dispositif, des largeurs spectrales entre 5 et 31 nm sont atteignables (comparaison avec 4.4 : de 25 à 38 nm).

De plus, une variation du poids relatif σ dans la fonction de mérite de l'algorithme est également synonyme de variation de la puissance moyenne optimale (figure 7.4(b)).

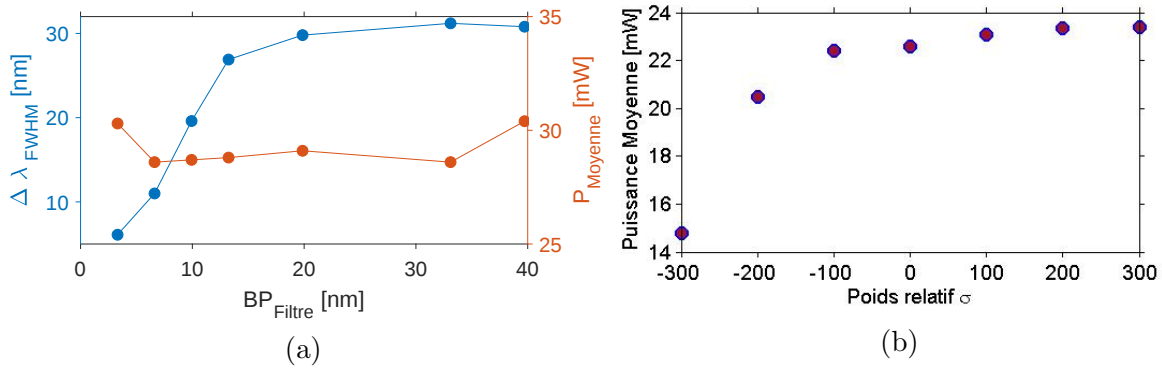


FIGURE 7.4: Graphiques présentant (a) l'évolution de la largeur spectrale à mi hauteur et de la puissance moyenne optimale selon la bande passante du filtre optique, et pour des poids relatifs $\sigma = 300$, et (b) l'évolution de la puissance moyenne optimale des régimes générés selon le poids relatif et pour une bande passante du filtre fixée à $BP_{filtre} = 30$ nm et centrée à 1550 nm.

Ainsi, il devient possible de générer automatiquement des régimes impulsionnels dont la puissance moyenne varie entre 15 et 24 mW. Pour un état à largeur spectrale prédéfinie, son énergie peut être contrôlée en modifiant le poids relatif σ . En résumé, la génération automatique d'un régime avec un couple énergie-largeur spectrale peut être réalisée avec ce montage, en une durée raisonnable car le SLM ne possède aucun gène implémenté dans l'algorithme. Nous notons que le désalignement fréquent du montage extracavité (pour la récupération du faisceau) fait varier ces valeurs, d'où la différence de valeur maximale entre les deux sous figures 7.4 (a) et (b). Chaque étude est réalisée sur quelques heures, le désalignement n'a donc qu'un impact mineur sur une seule étude. En revanche, une comparaison entre celles ci, réalisées à plusieurs jours de différence, n'a pas de sens. Seule la tendance est démontrée ici.

Ce travail représente une étape supplémentaire en complexité, proposant un contrôle de 2 paramètres impulsionnels, contre un seul précédemment. Cependant, ces fonctionnalités ne requièrent pas fondamentalement l'utilisation d'un filtre programmable par ordinateur, un simple filtre optique serait suffisant. Une manipulation spectrale plus poussée, conduisant à des formes d'impulsions plus complexes, requiert par contre l'utilisation d'un SLM. Dans la prochaine section, un travail répondant à ces attentes sera présenté, pour la génération de régimes lasers plus complexes que des solitons uniques : les molécules de solitons.

7.2 Génération de molécule de soliton à retard temporel pré-déterminé

7.2.1 Paramètres de l'algorithme

Pour finaliser ce chapitre et cette thèse, nous démontrons la génération de régimes impulsionnels stables plus complexes, rendues possible par l'incorporation du SLM en cavité. Parmi les régimes complexes présentés en section 2.2.3, un régime nous paraissait le plus intéressant à auto-générer : la molécule de soliton (figure 2.18). En effet, aucun travail d'autogénération de ce type de régime complexe, très localisé dans l'espace de recherche et très stable, n'avait été reporté jusqu'alors. Ce travail visait ainsi à une plus grande versatilité du "smart laser", permettant une génération de nouvelles dynamiques, qui constitue un véritable défi à relever. Plusieurs approches sont alors à considérer pour réaliser cet objectif : doit-on optimiser selon un critère spécifique aux molécules de solitons par rapport aux solitons uniques, et donc ajouter ce critère à la fonction de mérite, ou peut-on utiliser les fonctionnalisés du SLM pour s'en affranchir ? La seconde option sera choisie pour sa plus grande facilité, et présentée dans cette dernière section.

Une molécule de soliton est une structure de plusieurs solitons en interaction, stabilisée par la présence d'un attracteur [29]. Ce dernier résulte d'un équilibre non linéaire dissipatif. Une molécule de soliton est donc une structure stable et robuste. Dans le cas de notre étude, les molécules recherchées seront composées de 2 solitons, ayant entre eux une certaine relation de phase et de retard temporel. L'attracteur à rechercher est illustré en figure 7.5. L'évolution du système amène le régime vers l'attracteur (au centre de la figure) à relation de phase ($\Delta\varphi$) et retard temporel (τ) fixée. Le spectre d'un tel régime résulte des interférences à deux ondes avec une modulation périodique en $\cos^2(\omega\tau/2)$ (en intensité), dont la période est directement liée au retard temporel $\Delta\omega = 2\pi/\tau$. En effet, pour deux impulsions simplifiées en Diracs, séparés par un retard τ , le champ s'écrit :

$$E(t) = E_0 \otimes \left[\delta(t - \frac{\tau}{2}) + \delta(t + \frac{\tau}{2}) \right] \quad (7.1)$$

La transformée de Fourier de ce champ permet d'écrire :

$$\tilde{E}(\omega) = \tilde{E}_0 \left[e^{i\omega\frac{\tau}{2}} + e^{-i\omega\frac{\tau}{2}} \right] = 2\tilde{E}_0 \cos(\omega\frac{\tau}{2}) \quad (7.2)$$

ce qui permet d'écrire l'intensité spectrale :

$$\tilde{I}(\omega) = 4\tilde{E}_0^2 \cos^2(\omega\frac{\tau}{2}) \quad (7.3)$$

En terme de longueur d'onde, la relation devient $\Delta\lambda[\text{nm}] = \lambda_0^2/c\tau \approx 8/\tau[\text{ps}]$. En induisant cette modulation au spectre, il est possible de scinder une impulsion en deux et d'en générer une molécule. De plus, en jouant sur la période spectrale, il est possible de modifier le retard temporel entre les deux solitons, permettant un contrôle de celui-ci. En revanche, toutes les valeurs de τ ne sont à priori pas atteignables, la génération d'attracteur dépendant de la fonction de transfert de l'absorbant saturable mais également des propriétés de la cavité laser. L'espace de travail sur le retard temporel est donc discret, ne permettant pas la génération de n'importe quelle molécule

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

de soliton. Le contrôle de la phase relative entre les deux solitons n'est pas un objectif dans ces travaux, et ne sera pas mesuré.

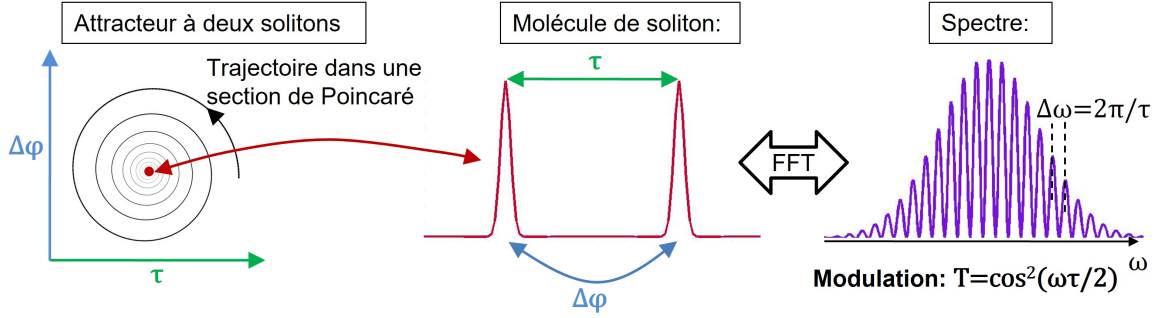


FIGURE 7.5: Illustration de la structure d'une molécule de soliton basé sur un attracteur à deux solitons. La relation phase/retard temporel entre les solitons donne une modulation spectrale périodique

Pour générer un molécule à deux solitons à partir d'un soliton unique en dehors d'une cavité laser, une modulation spectrale en intensité par le SLM en $\cos^2(\omega\tau/2)$ doit être appliquée. Cependant, pour un contrôle intracavité, où l'impulsion retransverse le pulse shaper à chaque tour, la question de la modulation à appliquer est plus complexe. Quelle transmission par le SLM doit-on appliquer pour générer un régime stable ? Comment évolue le champ laser tour à tour ? Autant de questions auxquelles il est impossible de répondre a priori. De plus, la propagation de l'impulsion à travers les fibres apportant également une modulation en phase et en amplitude, le façonnage est moins précis et moins direct que pour des contrôles en extracavité. Notre première idée est alors d'appliquer une modulation spectrale en intensité en $\cos^2(\omega\tau/2)$ sans phase. Les intérêts et les limitations de cette transmission seront présentés et discutés dans une section suivante.

En ce qui concerne les optimisations par algorithme, une question doit alors se poser : quels gènes seront implémentés ? En d'autres termes, quels paramètres de cette modulation est-il intéressant de contrôler pour atteindre l'objectif fixé dans les meilleures conditions ? Quelques expériences préliminaires nous ont montré que l'application d'une modulation spectrale fixe et donc non optimisée, ne menait que rarement à des molécules de soliton stables. Il devient donc indispensable d'implémenter quelques paramètres du SLM comme gène. Plusieurs paramètres de la modulation, présentés en figure 7.6, peuvent être envisagés. Nous distinguons 3 principaux paramètres importants sur lesquels jouer : la période de modulation $\Delta\omega_{SLM}$, le centrage spectral $\omega_{0,SLM}$, et la transmission minimale T_{min} . Ces trois paramètres définissent totalement la transmission du SLM en intensité (la transmission maximale est fixée à 1). Parmi ces paramètres, deux seront utilisés comme gènes. Tout d'abord, $\omega_{0,SLM}$, lié à la relation de phase entre les deux impulsions, n'est pas une caractéristique de la molécule que nous souhaitons contrôler, son implémentation comme gène et donc son investigation n'est pas problématique. Cependant, elle permet d'étendre le champ d'investigation et donc de permettre à l'algorithme de trouver plus de solutions (des régimes de molécule de solitons). En revanche, $\Delta\omega_{SLM}$ est une caractéristique de la molécule intéressante

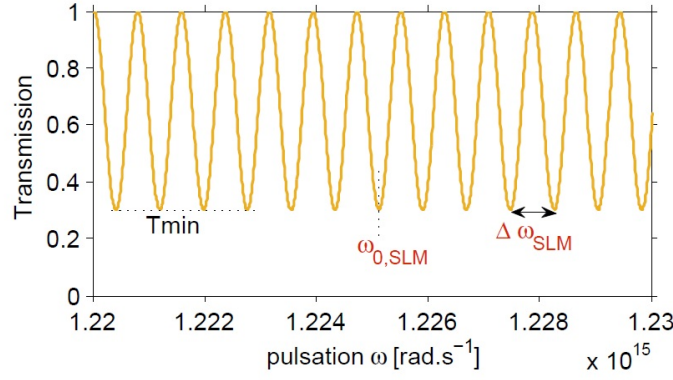


FIGURE 7.6: Graphique présentant la transmission du dispositif ligne 4f, SLM et polariseur en fonction de la pulsation. La période $\omega_{0,SLM}$ et la phase relative $\Delta\omega_{SLM}$ sont implémentés comme gènes pour l'algorithme, mais pas la transmission minimale T_{min} .

à contrôler. En effet contrôler cette grandeur revient à contrôler le retard temporel τ entre les deux impulsions. Cependant, compte tenu de la discrétisation de l'espace de travail due à la complexité des interactions intra-solitons, une légère flexibilité sur cette grandeur est judicieuse pour permettre à l'algorithme de trouver suffisamment de solutions. Un léger compromis entre valeur désirée et valeur à appliquer pour créer l'attracteur est donc opéré pour favoriser les optimisations par l'algorithme. Dans cette étude, la période des oscillations sera donc implémentée comme gène mais dans une fenêtre d'investigation autour d'une valeur cible : $\Delta\omega_{Cible} - 5\% < \Delta\omega_{SLM} < \Delta\omega_{Cible} + 5\%$. Pour terminer, l'expérience nous a montré que lorsque $T_{min} = 0$, tous les régimes de blocages de mode générés par la cavité sont des molécules à 2 solitons. Dans ce cas, il est inutile d'ajouter un critère de sélection des molécules de solitons au détriment des régimes d'impulsions uniques dans la fonction de mérite, la sélection étant indirectement faite par le SLM. En revanche, pour tout $T_{min} > 0$, il n'est pas garanti de générer que des molécules. Dans le présent travail, nous faisons le choix d'utiliser cette propriété plutôt qu'un critère de mérite, et T_{min} ne sera donc pas utilisé comme gène dans l'algorithme. Des études purement pratiques montrent qu'une valeur de T_{min} fixée à 0 dès le début de l'optimisation mène à un taux de succès (c.a.d le taux de génération de molécules de soliton stable par le procédé) plus faible que lorsque celui-ci décroît progressivement pendant l'optimisation. Pour cette étude, une transmission de $T_{min}(\text{numéro de génération}) = 0.3(1 - \text{numéro de génération}/\text{génération maximale})$ est testée. Ainsi, $T_{min}(0) = 0.3$ et $T_{min}(\text{génération maximale}) = 0$. Le contraste augmente de manière linéaire avec le nombre de génération. Autrement dit, en début d'optimisation, les pertes induites par le montage SLM+PBS sont faibles pour trouver facilement des régimes de blocage de mode (à ce moment mono-impulsionnels), puis ces pertes indispensables pour sélectionner les molécules de solitons augmentent continuellement pour trouver plus "doucement" la solution finale optimale. Cette configuration de fonctionnement sera donc utilisée tout au long de ce travail. Les études préliminaires menant à cette configuration d'utilisation optimale sont résumées en tableau 7.2.

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

Configuration	(i)	(ii)	(iii)	(iv)
Taux de succès	10%	60%	60 %	90%

TABLE 7.2: Tableau présentant les différents taux de succès (génération de molécules stables) par l'algorithme selon sa configuration de fonctionnement. Configuration **(i)** transmission fixée au préalable, optimisation sur 4 LCs, $T_{\min}$ constant, **(ii)** optimisation sur 4 LCs+ $\Delta\omega_{SLM}$, $T_{\min}$ constant, **(iii)** optimisation sur 4 LCs+ $\Delta\omega_{SLM} + \omega_{0,SLM}$, $T_{\min}$ constant, **(iv)** optimisation sur 4 LCs+ $\Delta\omega_{SLM} + \omega_{0,SLM}$, $T_{\min}$ décroissant

Ainsi, dans ces travaux, les optimisations seront réalisées sur 6 gènes : $\varphi_{1,2,3,4}$, $\Delta\omega_{SLM}$ et $\omega_{0,SLM}$. La fonction de mérite se doit de sélectionner les régimes de blocage de mode les plus stables, compatibles avec la transmission spectrale induite par le SLM. Comme énoncé précédemment, il est inutile d'ajouter un critère promouvant les molécules de solitons par rapport aux impulsions uniques, celle ci étant réalisée par le SLM. Ainsi, la fonction de mérite utilisée sera simplement la puissance RF au taux de répétition fondamental de la cavité (39.9 MHz). Les optimisations seront réalisées sur 15 générations, avec un contraste grandissant, avec une nombre d'individu pour chaque génération future similaires aux précédents travaux. En revanche, le nombre d'individu de la génération initiale est augmenté à 140 pour un plus grand taux de succès. Cette configuration initiale est présentée en tableau 7.3.

Gènes	$\varphi_{1,2,3,4}$, $\Delta\omega_{SLM}$, $\omega_{0,SLM}$
Fonction de mérite	$P_{RF} - P_{bck}$
Population initiale	140
Population finale	20
P_{pompe} (W)	1.0
Nombre génération	15
Taux d'enfants mutés	70 %

TABLE 7.3: Tableau présentant les caractéristiques de fonctionnement de l'algorithme pour le travail sur la génération à la demande de molécules de solitons, avec retard temporel ajustable

7.2.2 Convergence vers un régime de molécule de soliton

Pour débiter cette série d'optimisations, une génération de molécule de solitons avec un retard temporel de $\tau_{Cible} = 4$ ps entre les impulsions est testé. La période spectrale à appliquer est donc $\Delta\lambda_{Cible} = 2$ nm. Cette valeur est utilisée comme valeur cible pour notre algorithme, qui investiguera donc des périodes spectrales comprises entre 1.9 et 2.1 nm. Les autres gènes sont investigués sur des plages complètes. L'évolution des meilleurs et scores moyens est présenté en figure 7.7.

Nous y remarquons un score maximum constant pendant l'optimisation, et un score moyen à comportement non monotone. Cette évolution est différente de celles présentées dans les sections précédentes, et est plus compliquée à interpréter. Tout d'abord, la constance du score maximal de chaque génération peut s'expliquer par une raison évidente : le fort nombre d'individu dans la génération initiale, menant à une simple

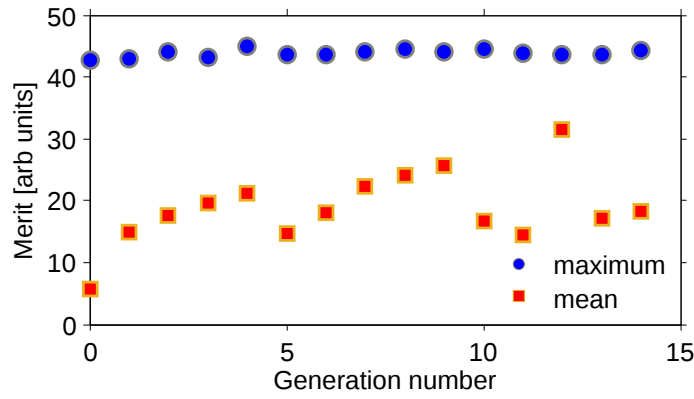


FIGURE 7.7: Graphique présentant l'évolution du meilleur score de chaque génération pour une optimisation typique avec une transmission sinusoïdale du SLM

sélection d'un individu optimal, sans optimisation. Cependant, dans cette étude, l'individu optimal final n'est pas celui de la génération initiale. L'hypothèse d'un individu optimal dès la génération initiale est donc fausse. Dans ces travaux, les conditions de travail changent pendant le procédé : le SLM induit une transmission de plus en plus contrastée. Un individu testé à la génération 1 et à la génération 10 n'aura donc pas nécessairement les mêmes caractéristiques. L'algorithme semble alors capable de retrouver un individu optimal pour chaque génération, malgré ces contraintes variables. De plus, le score moyen de chaque génération varie beaucoup durant l'optimisation, sans présenter de plateau de convergence. Cela démontre un comportement plus instable de chaque individu durant le processus. En analysant les valeurs des individus pour chaque génération, nous remarquons que les parents retestés ont souvent un score légèrement plus faible que celui obtenu à la génération précédente, mais celui-ci reste élevé. Cela signifie que l'augmentation du contraste pendant l'optimisation est suffisamment "douce", pour permettre à des régimes de blocage de mode de le rester d'une génération sur la suivante. En revanche, le meilleur individu de chaque génération n'est généralement pas un parent, le croisement et la légère mutation permettant de retrouver un individu avec un meilleur score. Cela démontre l'intérêt de l'utilisation d'un tel algorithme dans ces conditions d'utilisation.

Les caractéristiques de l'individu optimal avec sa trace d'autocorrélation et son spectre sont présentées en figure 7.8. Pour une molécule de soliton composée de deux impulsions identiques, la trace d'autocorrélation intensimétrique présente 3 pics, dont deux pics latéraux avec une amplitude à 50% du pic central. Ce n'est pas exactement le cas ici, en figure 7.8(a). Premièrement, les pics latéraux ont une amplitude à 60% du pic central, et de légers rebonds sont présents à $\pm 2\tau$ (ici à -8 et 8 ps). Cette trace d'autocorrélation est plutôt caractéristique d'un régime à 2 solitons identiques séparés de τ et deux germes d'amplitudes réduites à 2τ .

La détermination de l'origine des germes temporels demande de nouvelles investigations. Il est possible que ceux-ci proviennent de la transmission appliquée par le SLM, comme illustré en figure 7.9. En effet, on devrait appliquer un cosinus, avec une modulation d'amplitude égale à $t(\omega) = \sqrt{T(\omega)}$ en champ et une modulation de phase lorsque

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

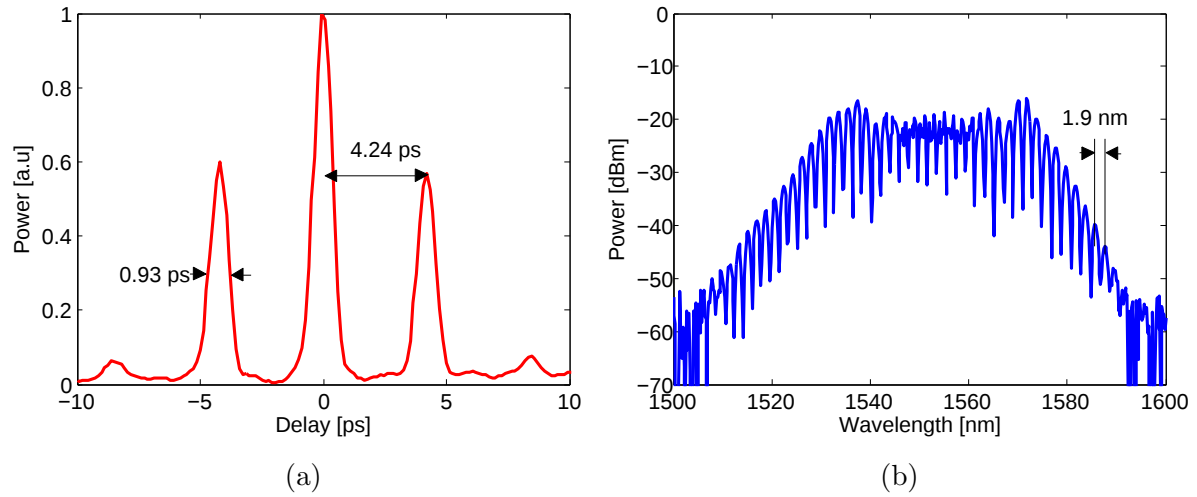


FIGURE 7.8: Graphiques présentant (a) la trace d'autocorrélation du meilleur individu obtenu par optimisation générant des molécules de solitons, et (b) son spectre

celui-ci devient inférieur à 0. En réalité, on applique une modulation d'amplitude en valeur absolue de cosinus sans modulation de phase. Sans modulation de phase (figure 7.9 (a)), le signal théorique généré est un pic central à forte amplitude, mais des rebonds situés à chaque multiple du retard fondamental, et dont l'intensité décroît (figure 7.9 (b)). En revanche, une modulation en cosinus, c'est à dire une modulation d'amplitude en valeur absolue de cosinus avec saut de phase de π (figure 7.9 (c)) correspond exactement à une molécule à deux solitons (figure 7.9 (d)). Il aurait été donc plus logique à première vue d'induire une modulation en amplitude et en phase. Cependant, pour un travail intracavité, le champ laser traverse le SLM à chaque tour de cavité. Si un saut de phase de π est appliqué par le SLM celui-ci sera appliqué à chaque tour de cavité. Dans ce cas, pour un tour de cavité N , la phase présenterait des sauts d'amplitude $N\pi$, conduisant à une phase constante pour chaque valeur de N paire. On aurait alors une alternance entre le cas avec une phase constante et celui avec saut de phase sur deux tours de cavités successifs. Ce cas de figure semble alors trop variable et chaotique pour générer des molécules de solitons stables. Dans notre cas expérimental (figure 7.8), la situation semble se situer entre les deux cas simulés avec et sans saut de phase. Le plus probable est que la relation de phase est apportée par la propagation de l'impulsion dans la partie fibrée du montage expérimental car la situation semble plus favorable, et seule une modulation d'amplitude est appliquée par le SLM. La relation de phase n'étant probablement pas exactement des saut de π comme souhaité, ces légers germes sont générés. D'un point de vue temporel, la molécule se formera à partir de la dynamique dissipative non linéaire de la cavité au cours des passages successifs à travers la cavité et le SLM d'après les germes.

D'un point de vue spectral, le retard temporel τ entre les impulsions correspond bien à la période de modulation du spectre ($1.9 \approx 8/4.24$), et le spectre est large d'environ 40 nm (FWHM), démontrant une certaine performance du montage laser. Cependant, là encore, le spectre enregistré ne correspond pas complètement à celui d'une molécule à deux solitons (figure 7.8(b)) dont le spectre serait modulé en $\cos^2(\omega\tau)$. En effet, au centre du spectre, autour de 1550 nm, des surmodulations apparaissent,

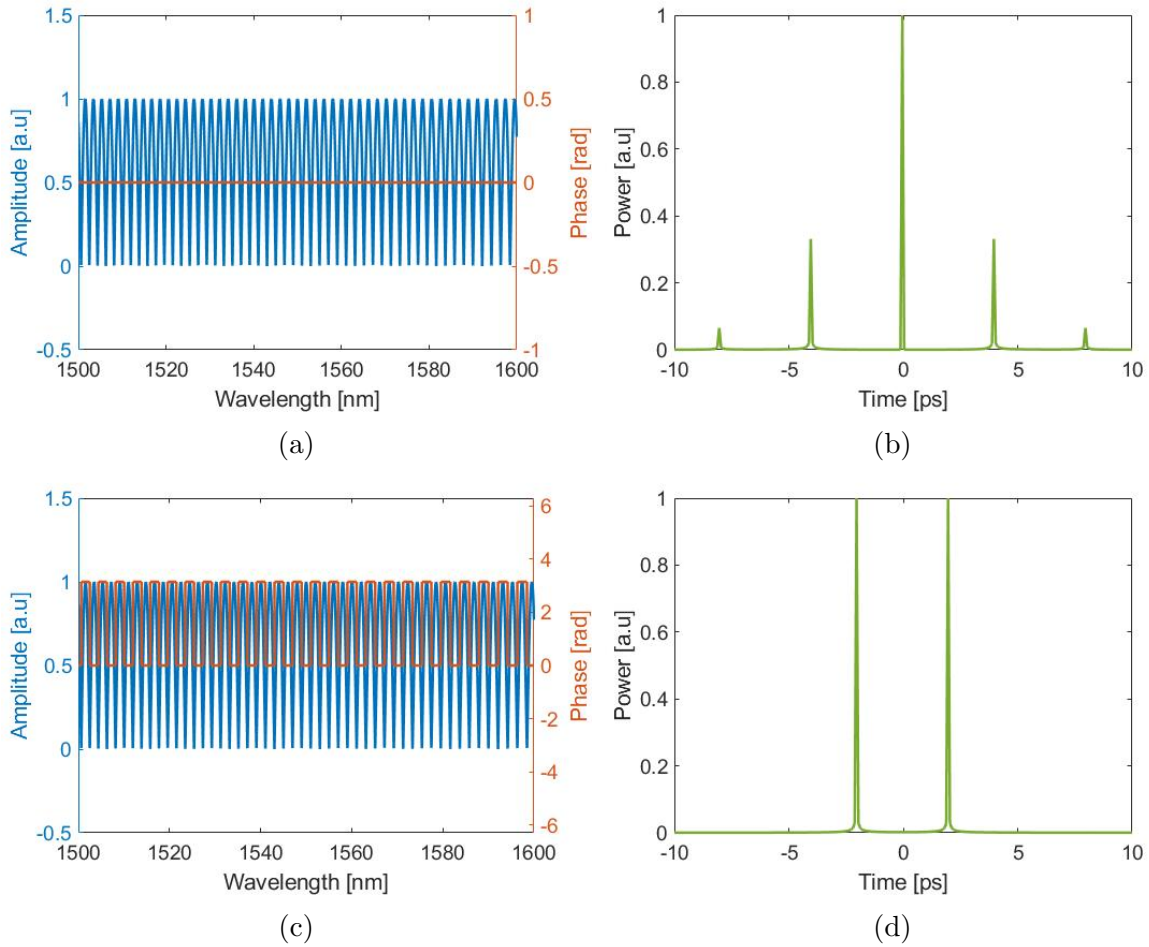


FIGURE 7.9: Graphiques présentant (a) une transmission simple passage en $|\cos(\omega\tau/2)|$ avec une phase plate et $\tau = 4$ ps, avec (b) sa trace temporelle calculée par fft, et (c) une transmission en $|\cos(\omega\tau/2)|$ avec des sauts de phases de π avec (d) sa trace temporelle calculée par fft.

rendant le spectre moins régulier. L'origine de ce phénomène se trouve probablement dans la définition des modulations spectrales, jusqu'ici définies en longueurs d'onde, non conforme à décrire un faisceau réel dans le plan de Fourier. Une correction sera apportée avec un travail directement en fréquence pour les prochaines études. Ces surmodulations peuvent également apporter un élément de réponse quant aux germes visibles sur la trace d'autocorrélation. Un travail de réalignement de la ligne 4f pourrait également être réalisé pour rendre le spectre plus lisse.

Enfin, pour terminer, la durée de chaque impulsion ($930/1.32 = 704$ fs sur la figure 7.8(a)) est bien plus longue que les valeurs attendues en limites de Fourier pour cette largeur spectrale (avec les coefficients calculés en section 4.2.2, $\Delta t = 0.77/\Delta\nu = 154$ fs). Une question peut alors se poser : avons nous réellement deux solitons simples ou bien deux solitons composites comportant des sous-structures ? Pour y répondre, une recompression des impulsions sera réalisée et présentée dans une section suivante. En effet, s'il existe des sous-structures masquées par l'élargissement dispersif, elles apparaîtront

après recompression.

7.2.3 Correction des surmodulations

La précédente sous section montre des instabilités de modulation au centre du spectre, autour de 1550 nm. Pour tenter de corriger ces défauts, un travail de réaligement est effectué. Premièrement, la cavité laser est coupée, et une émission spontanée ASE (Amplified Spontaneous Emission) est envoyée dans la ligne 4f pour déterminer l'origine des surmodulations. La figure 7.10 présente le résultat de ces travaux pour de l'ASE modulée ou non.

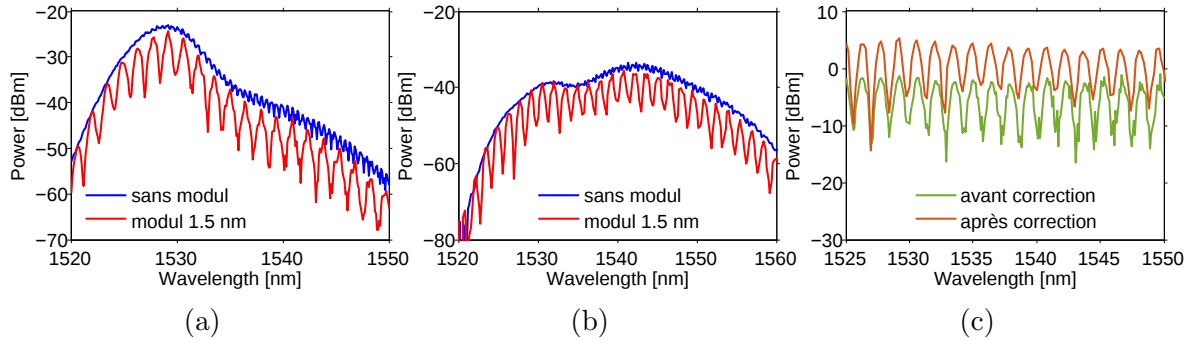


FIGURE 7.10: Graphiques présentant (a) les spectres des signaux avec et sans modulation en cosinus carré à période de 1.5 nm par le SLM avant réaligement, (b) après le réaligement, et (c) la modulation appliquée par le SLM (soustraction du spectre sans modulations au spectre modulé en dB), avant et après réaligement

En l'état initial de ce travail, celui de la section précédente, nous remarquons que même sans modulation induite par SLM, des surmodulations spectrales entre 1535 et 1550 nm (figure 7.10 (a)) sont présentes. Ces surmodulations, relativement intenses (3 dB d'amplitude), peuvent être à l'origine des surmodulations détectées en figure 7.8. Nous pouvons en déduire qu'elles ne proviennent pas entièrement d'une mauvaise calibration du SLM mais aussi d'un mauvais alignement de la ligne. Cependant, les modulations induites par SLM, calculées par soustraction du spectre sans modulation et du spectre modulé, démontrent malgré tout des désaccords vis à vis du cosinus carré dans la zone à problème. Les problèmes de surmodulation ont donc deux origines : un mauvais alignement de la ligne et un problème de calibration avec notamment des modulations appliquées en longueur d'onde. Un réaligement et une recalibration du SLM ont donc été réalisés, en essayant de minimiser les surmodulations. Les spectres avec et sans modulation par SLM, après les corrections, sont donnés en figure 7.10 (b). Nous y remarquons que les surmodulations ne disparaissent pas complètement, mais sont plus lisses (amplitude < 1 dB). L'orientation du SLM, ainsi que l'angle d'incidence du faisceau sur le second PBS, sont les composants qui ont été corrigés. Les modulations induites par le SLM (figure 7.10 (c)), après recalibration, définition des modulations en fréquence et réaligement, présentent une forme en cosinus carré attendue et propre. Ces corrections vont permettre de générer des molécules de soliton plus propres, qui se traduiront par des traces d'autocorrélation plus propres. Un problème reste encore à corriger : la longue durée des impulsions chirpées. Un travail de recompression de ces impulsions en extra cavité est de fait réalisé et présenté en prochaine sous section.

7.2.4 Recompression des molécules générées

L'objectif de ce travail est de déterminer si la longue durée de la trace d'autocorrélation des régimes générés par le dispositif est due à un fort chirp, et un élargissement extra cavité dans les fibres entre la cavité et l'autocorrélateur, ou à la présence de plusieurs impulsions. Pour cela, nous tentons de recompresser les impulsions proches de leur durée en limite de Fourier. Une fibre DCF à dispersion de $-60 \text{ ps.nm}^{-1}.\text{km}^{-1}$ sera utilisée, à longueur judicieusement choisie par méthode de coupe progressive. Pour ce travail, nous considérons que le chirp des impulsions générées est similaire quel que soit le régime impulsif. Une impulsion unique est alors générée par la cavité et la fibre DCF est utilisée dans la partie extra cavité juste avant l'autocorrélateur. Une trace d'autocorrélation est enregistrée pour chaque longueur de fibre DCF utilisée, sa longueur diminuant au cours des tests. Le graphique présentant ce travail est présenté en figure 7.11.

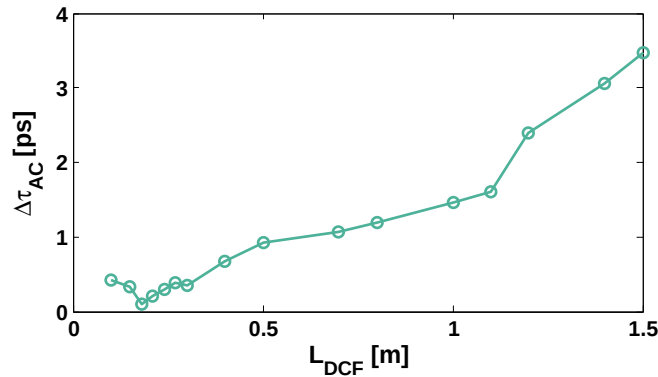


FIGURE 7.11: Graphique présentant la durée de la trace d'autocorrélation d'une impulsion générée en fonction de la longueur de fibre DCF utilisée pour recompression

Cette figure présente une durée d'autocorrélation et donc de l'impulsion minimale pour une longueur de fibre DCF de 18 cm. Nous notons que la durée d'impulsion pour cette recompression (90 fs) correspond à sa durée en limite de Fourier pour un spectre large de 30 nm. Cependant, la dispersion de la fibre étant importante, et la durée des impulsions recompressées varie rapidement avec la longueur de fibre DCF. En d'autres termes, pour une recompression optimale, la longueur de fibre se doit d'être précise au cm, ce qui n'est pas le cas dans notre étude, et ce qui est difficilement le cas dans la pratique (clivage de la fibre, soudure...). La longueur de fibre utilisée pour la recompression ne sera probablement pas exactement de 18 cm, ce qui ne permettra pas d'avoir des régimes recompressés jusqu'à la limite de Fourier. Cette limitation ne pose toutefois pas de problème majeur pour la conclusion de cette étude.

7.2.5 Convergence vers une molécule de soliton avec paramètres corrigés

Une nouvelle optimisation avec les mêmes paramètres que ceux de la section 7.2.2, mais avec les corrections apportées en sous section 7.2.3 et 7.2.4, est donc réalisée. La courbe d'évolution, similaire à la précédente, est donnée en figure 7.12, tandis que la

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

trace d'autocorrélation du champ recompressé et le spectre de l'individu optimal sont donnés en figure 7.13.

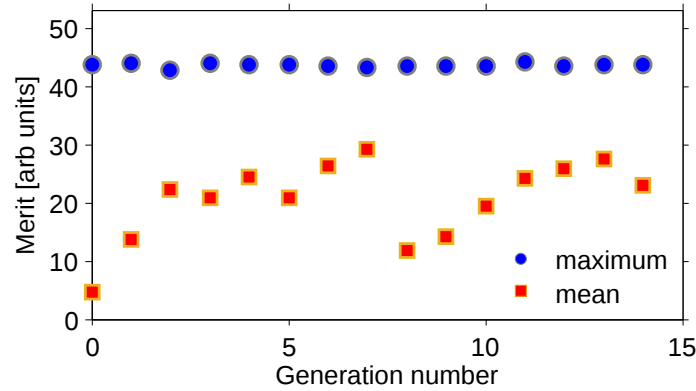


FIGURE 7.12: Graphique présentant l'évolution du meilleur score de chaque génération pour une optimisation typique avec une transmission sinusoidale du SLM

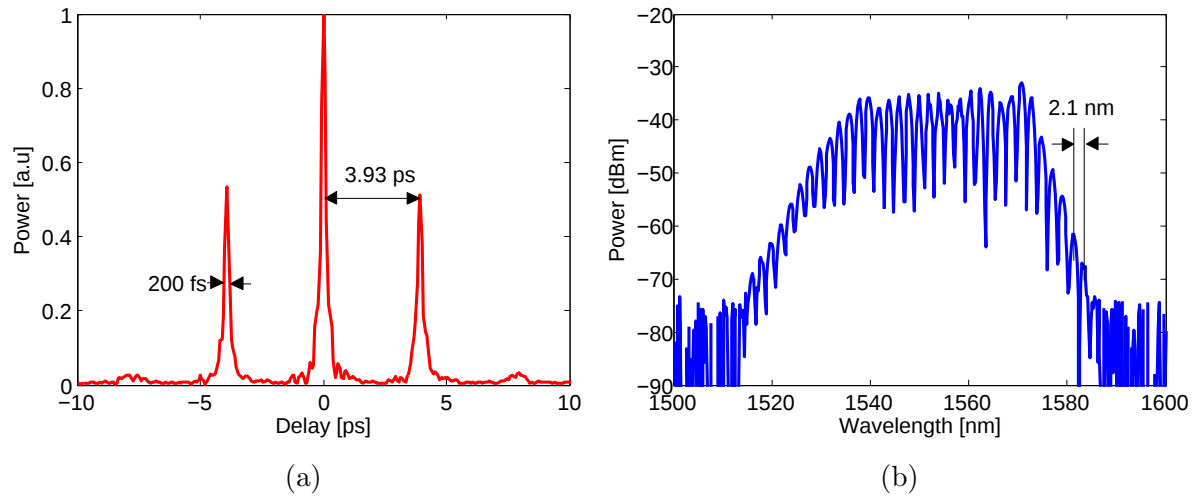


FIGURE 7.13: Graphiques présentant (a) la trace d'autocorrélation du meilleur individu obtenu par optimisation générant des molécules de solitons, et (b) son spectre

La trace d'autocorrélation (figure 7.13(a), présente 3 pics principaux brefs, de largeur à mi hauteur de 200 fs, correspondantes à des impulsions d'une durée de 150 fs. Cette durée est proche de la limite de Fourier (autour de 90 fs) mais est légèrement plus importante. Cette différence s'explique par les limitations dans la recompression évoquées plus haut. Cette durée d'impulsion étant suffisamment proche de la limite de Fourier, nous concluons que ce régime est bien composé d'une structure principale de molécule à deux solitons, sans sous structures internes. De plus, dans cette étude, le ratio entre les amplitudes des pics latéraux et le central approche les 50% (contre 60% pour l'étude précédente), et les germes à $\pm 2\tau$ sont moins intenses. Nous notons qu'aucun rebond n'est visible à $\pm 3\tau$ sur la trace d'autocorrélation.

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

D'un point de vue temporel, le régime généré semble alors correspondre à une structure principale à 2 solitons séparés de $\tau = 4$ ps, mais avec deux germes très légers, comme illustré en figure 7.14(a), correspondant à une trace d'autocorrélation illustrée en figure 7.14(b) en bonne adéquation avec la trace expérimentale figure 7.13(a).

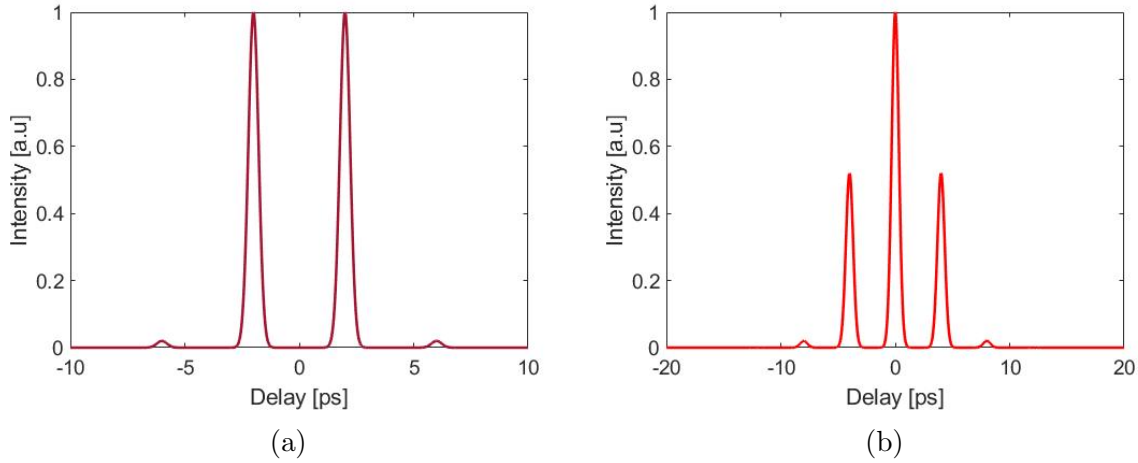


FIGURE 7.14: Graphiques présentant (a) une illustration de la trace temporelle en intensité générée par le processus, avec (b) sa trace d'autocorrélation calculée par produit de corrélation

Cette trace est bien compatible avec le spectre optique donné en figure 7.13 (b), qui présente une période de modulation plus régulière de 2.1 nm. La régularité de cette période spectrale, apportée par la définition des modulations en fréquence plutôt qu'en longueur d'onde, a donc permis une atténuation des germes temporels, mais ceux-ci n'ont pas disparus. Une partie de leur origine s'explique donc aussi par la relation de phase nécessaire à la production de molécules à 2 solitons comme expliqué en section précédente, et qui est induite par la dynamique de propagation. La période de modulation spectrale reste dans un intervalle proche de la valeur cible, avec un écart de moins de 5%. Ce spectre est moins large que pour le travail précédent (≈ 35 nm), mais reste malgré tout relativement large pour laisser penser à une bonne stabilité du régime.

La stabilité du régime optimal est testée en visualisant l'évolution du spectre pour chaque tour de cavité, sur 20000 tours correspondant à une durée d'environ 1 ms. Pour des vérifications plus longues, plusieurs enregistrements de la trace DFT sont réalisés. Une de ces traces est présentée en figure 7.15.

Ces traces présentent des impulsions stables sur 20000 tours de cavités, avec un spectre cannelé non fluctuant (figure 7.15 (b)). La différence entre l'évolution de ces spectres au tour n et $n-1$ est présentée en figure 7.15 (c), la différence pour le tour 1 n'étant pas représentée. Cette trace montre un phénomène cyclique de battement, avec des cycles stables d'un tour de cavité sur l'autre (la différence est très faible), puis des cycles plus instables où la différence spectrale sur deux tours de cavité peut atteindre 15 %. Ce phénomène se répète sur toutes les traces DFT enregistrées pour nos individus

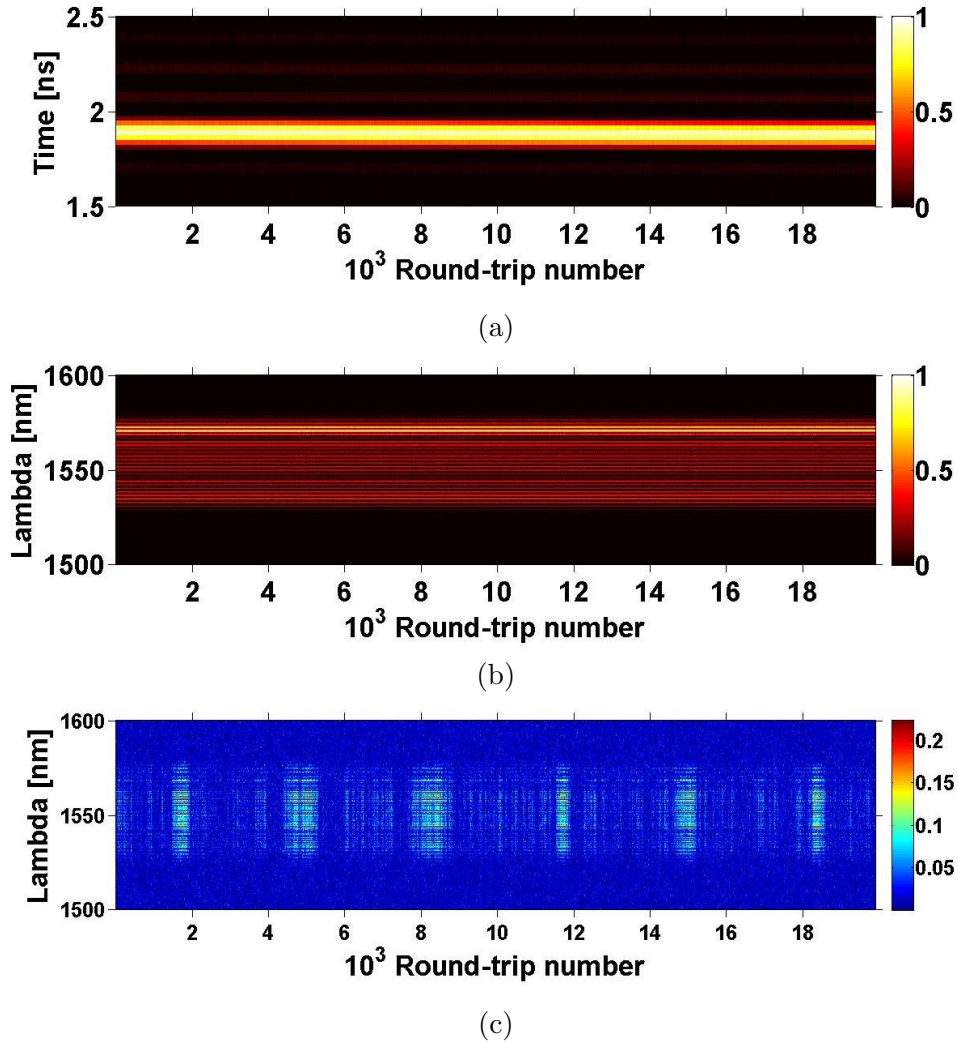


FIGURE 7.15: Graphiques présentant (a) l'évolution de la trace d'oscilloscope pour chaque tour de cavité, (b) l'évolution de son spectre enregistré par DFT et (c) l'évolution de la différence entre le spectre du tour de cavité n et $n-1$.

optimaux. Si la valeur de 15 % semble relativement importante, la répétition des cycles démontre bien une stabilité du régime malgré tout.

7.2.6 Contrôle du retard temporel entre les solitons

Le procédé est réutilisé avec différents retards temporels cibles. Ainsi, des optimisations avec des cibles comprises entre 1 et 8 ps ont été conduites. L'algorithme est capable générer des solutions correspondantes à des molécules de soliton pour des retard temporels compris entre 3.5 et 8.0 ps, mais pas pour des retard inférieurs (entre 1 et 3.5 ps) . Des résultats d'optimisation avec des retard temporels cibles respectivement de 5.3 et 7.3 ps sont donnés en figure 7.16.

Nous notons que pour chaque optimisation l'individu optimal est une molécule de soliton avec la même structure principale et des légers rebonds. Le retard temporel mesuré est proche de la valeur cible (moins de 5% de différence). Ce retard correspond

7.2. GÉNÉRATION DE MOLÉCULE DE SOLITON À RETARD TEMPOREL PRÉ-DÉTERMINÉ

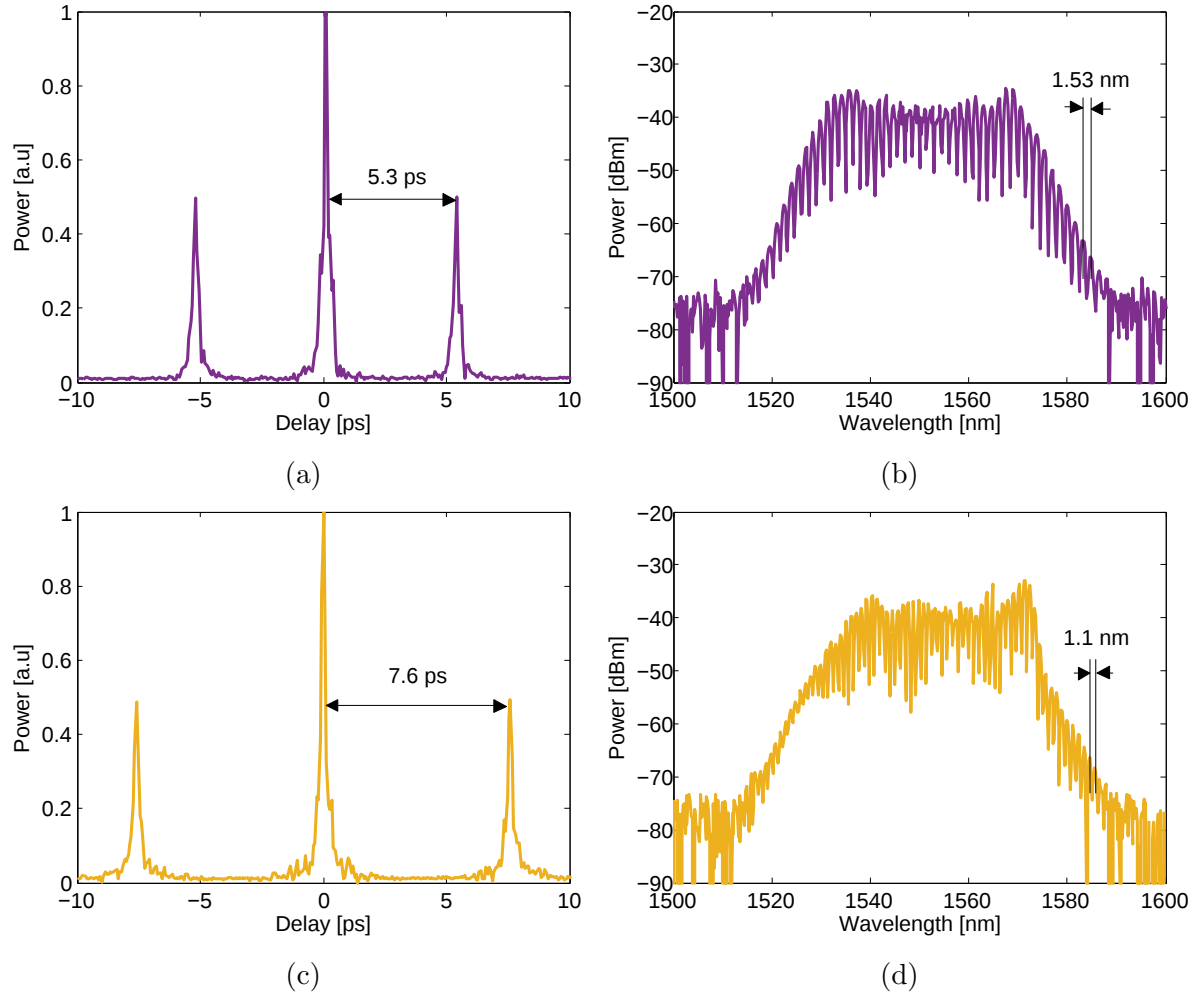


FIGURE 7.16: Graphiques présentant (a) et (c) les traces d'autocorrélation des meilleurs individus obtenus par optimisation avec un retard cible respectivement de 5.3 ps et 7.3 ps, (b) et (d) leur spectre respectif

bien à la période d'oscillation des modulations spectrales, et les impulsions sont courtes après recompression (≈ 150 fs), et ce quelle que soit la valeur du retard temporel. Les spectres sont de largeurs variables (entre 30 et 40 nm), généralement plus fins pour des périodes d'oscillations plus courtes. La figure 7.17 présente la moyenne des périodes mesurées en fonction de la période cible (la moyenne étant réalisée sur 4 optimisations). La barre rouge représente l'erreur standard à 95 % pour chaque test. Nous remarquons que cette dernière est très faible pour des valeurs cibles entre 1 et 2 nm, mais augmente pour des périodes cibles de 2.2 nm. Cela s'explique par les limites physiques du système. Nous notons que certaines optimisations ont pu être réalisées pour des périodes cibles jusqu'à 2.5 nm, mais présentaient un taux de réussite plus faible que pour la fenêtre 1-2.2 nm. Nous avons donc décidé de ne pas représenter ces optimisations sur notre graphique.

Notre système et procédé permet donc de générer à la demande des molécules de solitons à deux impulsions, à retard prédéterminé entre 3.5 et 8 ps. La fenêtre opérationnelle s'explique par deux raisons. Premièrement, le montage impose une limite technique :

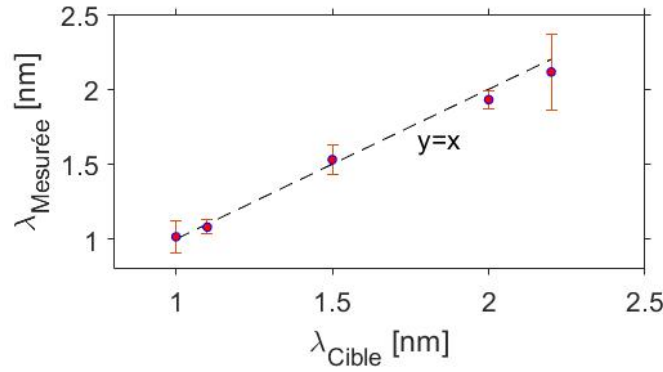


FIGURE 7.17: Graphique présentant les périodes des modulations spectrales mesurées en fonction des périodes cibles

la résolution spectrale du SLM étant de 0.3 nm dans nos conditions d'utilisation, il est impossible d'induire des modulations périodiques de moins de 1 nm (correspondant à 3 pixels). Ainsi, il est impossible de générer des retards de plus de 8 ps. Deuxièmement, une limite physique explique également la difficulté à générer des retards inférieurs à 3 ps : les impulsions étant fortement chirpées, avec des durées de 2 ps pour certaines, ils apparaît compliqué de créer des attracteurs avec deux impulsions plus proches temporellement que 3 ps. Pour plus de versatilité dans la génération de molécules de soliton, il conviendrait par exemple de gérer la dispersion de la cavité avec le SLM par exemple, pouvant permettre une réduction du chirp, ou encore travailler avec un réseau à pas plus important (autour de 800 traits.mm⁻¹), augmentant la dispersion dans le plan de Fourier et donc la résolution spectrale du SLM tout en gardant une fenêtre d'étude suffisamment grande en longueur d'onde. Les travaux de cette dernière section ont été présentés dans deux conférences : CLEO Europe [140] et EOSAM [141], et ont fait l'objet d'une publication [142]. Pour mieux comprendre les effets physiques permettant la génération de ces molécules, et l'origine des rebonds sur les traces d'autocorrélation, des travaux de simulation seront réalisés. L'enregistrement de l'établissement du régime optimal à chaque tour de cavité permettrait d'expliquer la différence entre la modulation appliquée et les oscillations mesurées. Enfin, un enregistrement des meilleurs individus de chaque génération permettrait de comprendre comment l'algorithme permet la création pas à pas d'une molécule à 2 solitons stable.

Cette section finalise la présentation des travaux expérimentaux de ce manuscrit, démontrant différentes stratégies afin d'auto-générer des dynamiques laser complexes. Ce travail constitue une illustration des possibilités apportées par l'implémentation d'un algorithme d'évolution. Ces travaux ne sont donc pas une finalité, mais posent les bases de travaux plus ambitieux.

Chapitre 8

Conclusion

Cette thèse présente l'implémentation d'algorithmes sur des systèmes laser fibrés. L'objectif à long terme est une automatisation complète de lasers fibrés, avec une grande versatilité des régimes auto-générables et une adaptation du système à tous les environnements. Si les régimes lasers continus sont faciles à obtenir, et présentent un intérêt limité, la génération de régimes impulsions ultracourts est plus intéressante et plus complexe, nécessitant l'implémentation d'algorithmes pour tenter de générer ces régimes. Ce projet a débuté par la thèse d'Ugo Andral en 2013, donnant les premiers résultats d'auto-génération de régimes impulsions à soliton unique. Les travaux présentés dans ce manuscrit s'inscrivent dans la continuité du projet en augmentant la stabilité des régimes, la durée des optimisations ainsi que la versatilité des régimes générés. Cet objectif permet une approche pluridisciplinaire des travaux, avec une partie purement photonique (développement de cavités laser versatiles et robustes) et une partie programmation (adaptation des algorithmes pré-développés aux nouvelles cavités). Chaque nouvelle étape présentée requiert un ajustement et un développement de ces deux volets.

Ce manuscrit est composé de deux parties expérimentales. La première partie, le chapitre 3, présente la suite directe des travaux d'Ugo Andral. L'objectif de ce premier chapitre expérimental est d'auto-générer des régimes lasers impulsions stables, robustes, à soliton unique et à largeur spectrale pré-déterminée. Par rapport à ces travaux, un autre objectif est de diminuer sensiblement les durées d'optimisations. Une nouvelle cavité laser a donc dû être développée. Tout d'abord, une nouvelle fibre fortement dopée, nouvellement utilisée dans la laboratoire, est caractérisée pour déterminer la longueur à utiliser afin d'avoir à la fois un gain et un régime de dispersion satisfaisants. Une réduction de la longueur de fibre passive est également réalisée, pour réduire les instabilités. Cette étape m'a permis d'obtenir des compétences pratiques en photoniques, avec l'utilisation de soudeuses, la manipulation et le nettoyage de fibres optiques. Pour créer l'effet d'absorbant saturable, et avec l'idée de réduire à la fois le temps d'optimisation et les longueurs de fibres, des lames à retard de phase à cristaux liquides sont utilisées dans une partie en espace libre. Le développement de mes capacités de manipulation des éléments optiques en espace libre a donc été réalisé lors de cette étape, puis encore plus lors de la seconde partie expérimentale. Les lames à retard à cristaux liquides permettent également l'interfaçage de la cavité et l'implémentation de l'algorithme. La seconde discipline commence à ce moment, demandant une adaptation des algorithmes développés à cette nouvelle cavité, ainsi qu'à l'instrument de mesure utilisé. En effet, dans ce chapitre, avec l'idée de rendre plus compact le montage expérimental, un seul oscilloscope ultrarapide, travaillant sur deux voies, sera utilisé pour le diagnostic de chaque régime généré. La première voie permettra de calculer la puissance radiofréquence, bon indicateur de la stabilité d'un régime de blocage de mode, et la seconde voie permettra, grâce à un processus de mesure spectrale en temps réel (DFT), d'obtenir des informations spectrales avec un oscilloscope. Combinant cette mesure avec l'implémentation d'un algorithme d'évolution, une nouvelle méthode pour générer des régimes impulsions à largeur spectrale prédéterminée a été développée. Cette stratégie présente des durées d'optimisation convenables (7 minutes), et permet la génération de régimes stables et robustes. De plus, la répétabilité des résultats a été démontrée, prouvant la robustesse du montage et du procédé. Cependant, la principale limitation est son manque de diversité dans les régimes générés. En effet, elle ne per-

met l’auto génération que de régimes impulsionnels uniques, avec un seul contrôle de la largeur spectrale. La suite du travail a donc été de développer une nouvelle cavité laser, incluant plus de degrés de liberté pour générer d’autres régimes impulsionnels.

Cette seconde partie expérimentale s’articule donc sur une nouvelle cavité incluant un “pulse shaper” intracavité, permettant un contrôle plus important des régimes de propagation. La première étape, tout comme pour la première partie expérimentale, a été le développement de cette nouvelle cavité, avec l’inclusion de deux amplificateurs pour augmenter le gain, ainsi que le développement du “pulse shaper”. Celui-ci est en réalité une ligne 4f, incluant deux réseaux et un modulateur de lumière spatial. Le développement de cette ligne a été un travail compliqué et minutieux, tant l’alignement des optiques se doit d’être fin. Des compétences sur la manipulation des éléments optiques en espace libre sont donc fortement requises pour monter ce dispositif. Un ajout du SLM comme gène dans l’algorithme demande également un réajustement du programme informatique, et l’utilisation d’autres diagnostics. Pour commencer, un travail de recompression d’une impulsion extra cavité par contrôle de la phase spectrale est réalisé pour poser les bases de travaux plus complexes, ainsi que pour tester l’efficacité du dispositif. Ensuite, le “shaper” est introduit dans la cavité et le même travail d’optimisation sur la largeur spectrale d’un régime, avec induction de phases spectrales (la GVD), est réalisé. Ce travail a permis de vérifier les capacités de la source laser à générer des états impulsionnels simples, mais apporte peu de nouveauté par rapport aux travaux précédents. Ensuite, d’autres diagnostics ont pu être testés (intensité de seconde harmonique, puissance moyenne...), permettant, en jouant sur la modulation d’intensité spectrale, la génération d’état impulsionnels avec largeur spectrale et énergie prédéterminés. Pour la première fois, deux caractéristiques impulsionnelles sont contrôlées grâce à l’algorithme, contre une seule dans les travaux préalables. L’inclusion du pulse shaper permet cela. Cependant, l’utilisation d’un tel dispositif, nous a poussé à aller encore plus loin dans l’autogénération de régimes laser complexes avec l’auto génération de molécules à 2 solitons avec retard temporel prédéterminé. Ce type de régime, à dynamique et structure plus complexes que les régimes mono-impulsionnels, semble nécessiter l’utilisation de composants laser plus élaborés par rapport aux premiers travaux présentés dans ce manuscrit. De plus, le contrôle facile et interfaçable du retard temporel entre les deux solitons ne permet pas l’utilisation de composés simples. Le SLM semble donc être un moyen efficace d’atteindre l’objectif. Après adaptation des paramètres de fonctionnement de l’algorithme, il a été possible d’auto-générer des molécules de 2 solitons, à retard temporel pré-déterminé entre 3 et 8 ps. C’est la première fois qu’une telle performance est réalisée dans ce domaine. La répétabilité des résultats, également présentée, démontre le potentiel des algorithmes d’évolution à générer des régimes lasers complexes à la demande. Ce chapitre est donc une illustration des capacités de tels algorithmes dans le domaine de la photonique, appelant de nouvelles utilisations académiques et industrielles. Ces deux parties expérimentales ont toutes deux débouché sur l’écriture d’un article suivie de publication (IEEE [101] et Optic Letters [142]), ainsi que sur la présentation des travaux dans 4 conférences : Euro-photon 2020 [127] (présentation orale), CLEO Europe 2021 [140] (présentation orale), Optique Dijon 2021 (poster) et EOSAM 2021 [141] (présentation orale).

Les voies d'amélioration de ces travaux sont multiples. Tout d'abord, la continuité directe de ces travaux serait l'utilisation du montage avec pulse shaper pour tenter d'auto-générer des états impulsionnels multiples et versatiles. Une identification des ces régimes, avec leurs caractéristiques propres, permettant de construire une fonction de mérite capable des les favoriser, et menant à l'écriture d'une bibliothèque de ces fonctions de mérite serait la dernière étape du projet. Dans ce cas, l'utilisation d'algorithmes d'évolution ne serait pas remis en doute pour l'auto-génération de ces régimes, ces travaux démontrant son efficacité pour cette tâche. La compacité du laser pourra être sensiblement augmentée par l'utilisation de puces informatiques comme des FPGA, permettant de réduire drastiquement la taille de la boucle de rétroaction et de gagner en compacité. Pour terminer, l'utilisation de fibre à maintien de polarisation (PM), dans un laser en "figure 9", permettrait une stabilité de propagation optimale, tout en augmentant le contrôle des effets de polarisation. L'utilisation de deux lames à cristaux liquides pourrait alors être suffisante pour contrôler toute la fonction de transfert de l'absorbant saturable, rendant le montage encore plus compact et moins coûteux.

Bibliographie

- [1] A. L. Schawlow and C. H. Townes, "Infrared and optical maser," *Phys. Rev.*, vol. 112, No. 6, pp. 1940-1949, 15 Décembre 1958.
- [2] T. H. Maiman, "Stimulated Optical Radiation in Ruby," *Nature*, vol. 187, pp. 493-494, 1960
- [3] C. K. N. Patel, "Optical Maser Experiments," *Physical Review Letters*, vol. 12, no. 21, pp. 588–590, 1964.
- [4] A. J. De Maria, D. A. Stetser, and H. Heynau, "Self mode-locking of lasers with saturable absorbers," *APL*, vol. 8, no. 7, pp. 174, 1966.
- [5] C. J. Koester and E. Snitzer, "Amplification in a Fiber Laser," *Applied Optics*, vol. 3, no. 10, 1964
- [6] S. B. Poole, D. N. Payne, and M. E. Fermann, "Fabrication of low-loss optical fibres containing rare-earth ions," *Elec. Lett.*, vol. 21, no. 17, pp. 737, 1985.
- [7] R. J. Mears, L. Reekie, S. B. Poole, and D. Payne, "Low-Threshold Tunable CW and Q-Switched Fibre Laser operating at 1.55 μm ," *Elec. Lett.*, vol. 22, No. 3, pp. 159-160, 1986.
- [8] L. Reekie, R. J. Mears, S. B. Poole, and D. N. Payne, "Tunable single-mode fiber lasers," *J. Lightwave. Technol.*, vol. 4, no. 7, pp. 956-960, 1986.
- [9] B. S. Kawasaki and K. O. Hill, "Low-loss access coupler for multimode optical fiber distribution networks," *Appl. Opt.*, vol. 16, pp. 1794-1795, 1977.
- [10] M. Digonnet and H. J. Shaw, "Wavelength multiplexing in single-mode fiber couplers," *Appl. Opt.*, vol. 22, pp. 484-491, 1983.
- [11] T. Bricheno, and V. Baker, "All-fibre polarisation splitter/combiner, " *Elec. Lett.*, vol. 21, no. 6, pp. 251-252, 1985.
- [12] M. Schneider, "Perçage profond par laser : Analyse des processus physiques. Physique Atomique," Physique. Université Pierre et Marie Curie - Paris VI, 2006.
- [13] C. Mas, "Modélisation physique du procédé de découpe de métaux par laser," Physique. Université Pierre et Marie Curie - Paris VI, 2003.
- [14] M. Mostafa, "Etude du perçage et du soudage laser : dynamique du capillaire," Matériaux. Université de Bourgogne, 2011.
- [15] W. Shi, Q. Fang, X. Zhu, R. A. Norwood, and N. Peyghambarian, "Fiber lasers and their applications," *Appl. Opt.*, vol. 53, No. 8, pp. 6554-6568, 2014.
- [16] P. Steyer, S. Valette, R. Diemiaszonek, B. Forest, J.-P. Millet, C. Donnet and E. Audouard, "Influence du marquage laser d'instruments chirurgicaux en acier inoxydable sur leur résistance à la corrosion," *Matériaux et Techniques*, vol. 94, pp.111–119, 2006.
- [17] D. Gubencu, and A. Han, "Marquage Laser des Emballages de Matière Plastique dans l'Industrie Alimentaire," *Int. Sys. Agri. Food. Prod.*, vol. 11, pp. 214-218, 2003.

- [18] W. Denk, K.R. Delaney, A. Gelperin, D. Kleinfeld, B.W. Strowbridge, D.W. Tank, and R. Yuste, "Anatomical and functional imaging of neurons using 2-photon laser scanning microscopy," vol. 54, no. 2, pp. 151–162, 1994.
- [19] R. T. H. Collis, "Lidar," *Appl. Opt.*, vol. 9, pp. 1782-1788, 1970.
- [20] S. E. Reutebuch, R. J. McGaughey, H-E. Andersen, and W. Carson, "Accuracy of a high-resolution lidar terrain model under a conifer forest canopy," *Can. J. Remote. Sens.*, vol. 29, no. 5, pp. 527–535, 2003.
- [21] T. A. Birks, W. J. Wadsworth, and P. St. J. Russell, "Supercontinuum generation in tapered fibers," *Opt. Lett.*, vol. 25, pp. 1415-1417, 2000.
- [22] J. Dudley, and J. R. Taylor, "Supercontinuum Generation in Optical Fibers," Cambridge University Press, 2010.
- [23] F. Forghieri, R. W. Trach, and A. R. Chraplyvy, "Fiber nonlinearities and their impact on transmission systems," in T. L. Koch, "Optical Fiber Telecommunications III," Academic Press, 1997.
- [24] Irl N. Duling, "All-fiber ring soliton laser mode locked with a nonlinear mirror," *Opt. Lett.*, vol. 16, pp. 539-541, 1991.
- [25] M. E. Fermann, F. Haberl, M. Hofer, and H. Hochreiter, "Nonlinear amplifying loop mirror," *Opt. Lett.*, vol. 15, pp. 752-754, 1990.
- [26] M. Zirngibl, L.W. Stulz, J. Stone, J. Hugi, D. DiGiovanni, and P. B. Hansan, "1.2 ps pulses from passively mode-locked laser diode pumped Er-doped fibre ring laser," *Elec. Lett.*, vol. 27, no. 19, pp. 1734, 1991.
- [27] H. Haus, "Theory of mode locking with a fast saturable absorber," *J. Appl. Phys.*, vol. 46, no. 7, pp. 3049–, 1975.
- [28] H. Haus, "Theory of mode locking with a slow saturable absorber," *IEEE J. Quantum. Electron.*, vol. 11, no. 9, pp. 736-746, September 1975.
- [29] Ph. Grelu, and J. Soto-Crespo, "Temporal Soliton "Molecules" in Mode-Locked Lasers : Collisions, Pulsations, and Vibrations. In : Dissipative Solitons : From Optics to Biology and Medicine," *Lecture Notes in Physics*, vol. 751. Springer, Berlin, Heidelberg, 2008.
- [30] A. Hideur, T. Chartier, M. Brunel, M. Salhi, C. Özkul, and F. Sanchez, "Mode-lock, Q-switch and CW operation of an Yb-doped double-clad fiber ring laser," *Opt. Comm.*, vol. 198, pp. 141–146, 2001.
- [31] X. Feng, H. Tam, and P. K. A. Wai, "Stable and uniform multiwavelength erbium-doped fiber laser using nonlinear polarization rotation," *Opt. Express*, vol. 14, pp. 8205-8210, 2006.
- [32] S. Hamdi, A. Coillet, and Ph. Grelu, "Real-time characterization of optical soliton molecule dynamics in an ultrafast thulium fiber laser," *Opt. Lett.*, vol. 43, pp. 4965-4968, 2018.
- [33] U. Andral, 2016, "Algorithme d'évolution pour laser à fibre optique en régime d'impulsions courtes," [Thèse de doctorat, Université de Bourgogne].
- [34] K. Tamura, H. A. Haus, and E. P. Ippen, "Self-starting additive pulse mode-locked erbium fibre ring laser", *Elec. Lett.*, vol. 28, no. 24, pp. 2226, 1991.
- [35] R. A. Betts, T. Tjugiarto, Y. L. Xue, and P. L. Chu, "Nonlinear refractive index in erbium doped optical fiber : Theory and experiment," *IEEE J Quantum Electron*, vol. 27, no. 4, pp. 908–913, 1991.

- [36] L.E. Nelson, D.J. Jones, K. Tamura, H.A. Haus, and E.P. Ippen, "Ultrashort-pulse fiber ring lasers," *Appl. Phys. B*, vol. 65, pp. 277–294, 1997.
- [37] C. Schwob, and L. Julien, "Le laser : principe de fonctionnement," *Reflète phys.*, no. 21, 2010.
- [38] S. M. J. Kelly, "Characteristic sideband instability of periodically amplified average soliton," *Electron. Lett.*, vol. 28, no. 8, pp. 806–807, 1992.
- [39] C. Fabre, "Les lasers - Principes fondamentaux," *Collection SFO*, vol. 4, pp. 1-39, 1996.
- [40] P. Le Boudec, F. Sanchez, P-L. François, J-F Bayon, and G. M. Stephan, "Dynamique des lasers à fibre dopée à l'erbium," *Ann. Télécommun.*, vol. 49, no. 3, pp. 178-192, 1994.
- [41] G. A. Ball, and W. H. Glenn, "Design of a single-mode linear-cavity erbium fiber laser utilizing Bragg reflectors," *J. Light. Tech.*, vol. 10, no.10, pp. 1338–1343, 1992.
- [42] Z. Jianluo, Y. Chao-Yu, G. W. Schinn, W. R. L. Clements, and J.W.Y. Lit, "Stable single-mode compound-ring erbium-doped fiber laser," *J. Light. Tech.*, vol. 14, no. 1, pp. 104–109, 1996.
- [43] H. C. Lefevre, "Single-Mode Fibre Fractional Wave Devices and Polarization Controllers," *Elec. Lett.*, 1980.
- [44] L. J. Aplet, and J. W. Carson, "A Faraday Effect Optical Isolator," *Appl. Opt.*, vol. 3, pp. 544-545, 1964.
- [45] J. Fujita, M. Levy, R. M. Osgood, L. Wilkens, and H. Dotsch, "Polarization-independent waveguide optical isolator based on nonreciprocal phase shift," *IEEE Photonics. Technol. Lett.*, vol. 12, no. 11, pp. 1510–1512, 2000.
- [46] Y. Fujii, "High-isolation polarization-independent optical circulator," *J. Light. Tech.*, vol. 9, no. 10, pp. 1238-1243, Oct. 1991.
- [47] N. Sugimoto, T. Shintaku, A. Tate, H. Terui, M. Shimokozono, E. Kubota, M. Ishii, and Y. Inoue, "Waveguide polarization-independent optical circulator," *IEEE Photonics. Technol. Lett.*, vol. 11, no. 3, pp. 355-357, March 1999.
- [48] M. Dantus, M. J. Rosker, and A. H. Zewail, "Real-time femtosecond probing of "transition states" in chemical reactions," *J. Chem. Phys.*, vol. 87, no. 4, pp. 2395-2397, 1987.
- [49] A. H. Zewail, "Laser Femtochemistry," *Science*, vol. 242, no. 4886, pp. 1645-1653, 1988.
- [50] H. A. Haus, "Mode-locking of lasers," *IEEE J Sel Top Quantum Electron*, vol. 6, no. 6, pp. 1173-1185, Nov.-Dec. 2000.
- [51] J. Kafka, D. Hall, and T. Baer, "Mode-locked erbium-doped fiber laser with soliton pulse shaping," *Opt. Lett.*, vol. 14, No. 22, pp. 1269-1271, 1989.
- [52] D. C. Hanna, A. Kazer, M. Philips, D. Shepherd, and P. Suni, "Active Mode-Locking of an Yb :Er Fibre Laser," *Elec. Lett.*, vol. 25, No. 2, pp. 95-96, 1989.
- [53] F. X. Kurtner, J. A. der Au, and U. Keller, "Mode-locking with slow and fast saturable absorbers-what's the difference?," *IEEE J Sel Top Quantum Electron*, vol. 4, no. 2, pp. 159-168, March-April 1998.
- [54] U. Keller, K.J. Weingarten, F.X. Kartner, D. Kopf, B. Braun, I.D. Jung, R. Fluck, C. Honninger, N. Matuschek, and J. Aus der Au, "Semiconductor saturable absorber mirrors (SESAM's) for femtosecond to nanosecond pulse generation in solid-state lasers," *IEEE J Sel Top Quantum Electron*, vol. 2, no. 3, pp. 435-453, Sept. 1996.
- [55] O. Okhotnikov, A. Grudinin, and M. Pessa, "Ultra-fast fibre laser systems based on SESAM technology : new horizons and applications," *New J. Phys.*, vol. 6, no. 177, 2004.

- [56] J.C. Mialocq, and P. Goujon, "Pinacyanol as a mode-locking saturable absorber for Rhodamine dye lasers," *Optics Communications*, vol. 20, no. 3, pp. 342-343, 1977.
- [57] Jin-Long Xu, Xian-Lei Li, Yong-Zhong Wu, Xiao-Peng Hao, Jing-Liang He, and Ke-Jian Yang, "Graphene saturable absorber mirror for ultra-fast-pulse solid-state laser," *Opt. Lett.*, vol. 36, pp. 1948-1950, 2011.
- [58] Xinhuan Feng, Hwa-yaw Tam, Heliang Liu, and P.K.A. Wai, "Multiwavelength erbium-doped fiber laser employing a nonlinear optical loop mirror," *Opt. Communication*, vol. 268, no. 2, pp. 278-281, 2006.
- [59] H. J. S. Dorren, D. Lenstra, Yong Liu, M. T. Hill, and G. -. Khoe, "Nonlinear polarization rotation in semiconductor optical amplifiers : theory and application to all-optical flip-flop memories," *IEEE J Sel Top Quantum Electron*, vol. 39, no. 1, pp. 141-148, Jan. 2003.
- [60] M. E. Fermann, M. J. Andrejco, Y. Silberberg, and M. L. Stock, "Passive mode locking by using nonlinear polarization evolution in a polarization-maintaining erbium-doped fiber," *Opt. Lett.*, vol. 18, pp. 894-896, 1993.
- [61] R. K. Bullough, and P. J. Caudrey, "The Soliton and Its History". In : R. K. Bullough, P. J. Caudrey, "Solitons," Topics in Current Physics, vol 17, Springer, Berlin, Heidelberg, 1980.
- [62] R. H. Goodman, Ph. J. Holmes, and M. I. Weinstein, "Strong NLS soliton-defect interactions," *Physica D : Nonlinear Phenomena*, vol. 192, pp. 215-248, 2004.
- [63] T. R. Taha, and M. I. Ablowitz, "Analytical and numerical aspects of certain nonlinear evolution equations. II. Numerical, nonlinear Schrödinger equation," *J. Comput. Ph.*, vol. 55, no. 2, pp. 203-230, 1984.
- [64] H. A. Haus, J. G. Fujimoto, and E. P. Ippen, "Structures for additive pulse mode locking," *J. Opt. Soc. Am. B*, vol. 8, No. 10, pp. 2068-2076, 1991.
- [65] O. E. Martinez, R. L. Fork, and J. P. Gordon, "Theory of passively mode-locked lasers for the case of a nonlinear complex-propagation coefficient," *J. Opt. Soc. Am. B*, vol. 2, pp. 753-760, 1985.
- [66] Ph. Grelu, and N. Akhmediev, "Dissipative solitons for mode-locked lasers," *Nat. Photonics*, vol. 6, pp. 84-92, 2012.
- [67] J. M. Soto Crespo, N. Akhmediev, and V. V. Afanasjev, "Pulse solutions of the cubic-quintic complex Ginzburg-Landau equation in the case of normal dispersion," *Phys. Rev. E*, vol. 55, No. 4, pp. 4783-4796, Avril 1997.
- [68] J. P. Gordon, "Dispersive perturbations of solitons of the nonlinear Schrödinger equation," *J. Opt. Soc. Am. B*, vol. 9, pp. 91-97, 1992.
- [69] D. U. Noske, N. Pandit, and J. R. Taylor, "Source of spectral and temporal instability in soliton fiber lasers," *Opt. Lett.*, vol. 17, pp. 1515-1517, 1992.
- [70] C. Honninger, R. Paschotta, F. Morier-Genoud, M. Moser, and U. Keller, "Q-switching stability limits of continuous-wave passive mode locking," *J. Opt. Soc. Am. B*, vol. 16, No. 1, pp. 46-56, Janvier 1999.
- [71] M. Horowitz, Y. Barad, and Y. Silberberg, "Noiselike pulses with a broadband spectrum generated from an erbium-doped fiber laser," *Opt. Lett.*, vol. 22, No. 11, pp. 799-801, 1 Juin 1997.
- [72] A. Runge, C. Agüergaray, N. G. R. Broderick, and M. Erkintalo, "Coherence and shot-to-shot spectral fluctuations in noise-like ultrafast fiber lasers," *Opt. Lett.*, vol. 38, pp. 4327, 2013.

- [73] M. Horowitz, and Y. Silberberg, "Nonlinear filtering by use of intensity-dependent polarization rotation in birefringent fibers," *Opt. Lett.*, no. 22, pp. 1760-1762, 1997.
- [74] M. Horowitz, and Y. Silberberg, "Control of noiselike pulse generation in erbium-doped fiber lasers," *IEEE Photonics Technol. Lett.*, vol. 10, no. 10, pp. 1389-1391, Oct. 1998.
- [75] D. Y. Tang, L. M. Zhao, and B. Zhao, "Soliton collapse and bunched noise-like pulse generation in a passively mode-locked fiber ring laser," *Opt. Express*, vol. 13, pp. 2289-2294, 2005.
- [76] Z.Q. Wang, K. Nithyanandan, A. Coillet, P. Tchofo-Dinda, and Ph. Grelu, "Buildup of incoherent dissipative solitons in ultrafast fiber lasers," *Phys. Rev. Res.*, vol. 2, 2020.
- [77] A. Boucon, B. Barviau, J. Fatome, C. Finot, T. Sylvesre, M.W. Lee, Ph. Grelu and G. Millot, "Noise-like pulses generated at high harmonics in partially mode-locked km-long Raman fiber laser," *Appl Phys B*, vol. 106, pp. 283-287, 2012.
- [78] J. M. Soto-Crespo, Ph. Grelu, and N.I Akhmediev, "Dissipative rogue waves : Extreme pulses generated by passively mode-locked lasers," *Phys. Rev. E*, vol. 84, no. 1, Juillet 2011.
- [79] K. Krupa, K. Nithyanandan, and Ph. Grelu, "Vector dynamics of incoherent dissipative optical solitons," *Optica*, vol. 4, pp. 1239-1244, 2017.
- [80] K. Tamura, E. P. Ippen, H. A. Haus, and L. E. Nelson, "77-fs pulse generation from a stretched-pulse mode-locked all-fiber ring laser," *Opt. Lett.*, vol. 18, pp. 1080-1082, 1993.
- [81] A. B. Grudinin, and S. Gray, "Passive harmonic mode locking in soliton fiber lasers," *J. Opt. Soc. Am. B*, vol. 14, No. 1, pp. 144-153, Janvier 1997.
- [82] A. B. Grudinin, D. J. Richardson, and D. Payne, "Passive harmonic modelocking of a fibre soliton ring lasers," *Elec. Lett.*, vol. 29, No. 21, pp. 1860-1861, 1993.
- [83] C. Lecaplain and Ph. Grelu, "Multi-gigahertz repetition-rate-selectable passive harmonic mode locking of a fiber laser," *Opt. Express*, vol. 21, pp.10897-10902, 2013.
- [84] A. N. Pilipetskii, E. A. Golovchenko, and C. R. Menyuk, "Acoustic effect in passively mode-locked fiber ring lasers," *Opt. Lett.*, vol. 20, pp. 907-909, 15 Avril 1995.
- [85] J. N. Kutz, B. C. Collings, K. Bergman, and W. H. Knox, "Stabilized Pulse Spacing in Soliton Lasers Due to Gain Depletion and Recovery," *J. Quant. Elec.*, vol. 34, No. 9, pp. 1749-1756, Septembre 1998.
- [86] G. Herink, F. Kurtz, D. R. Solli, and C. Ropers, "Real-time spectral interferometry probes the internal dynamics of femtosecond soliton molecules," *Science*, vol. 356, pp. 50-54, 7 Avril 2017.
- [87] B. A. Malomed, "Bound solitons in the nonlinear Schrodinger —Ginzburg-Landau equation," *Phys. Rev. A*, vol. 44, No. 10, pp. 288-294, 15 Novembre 1991.
- [88] J. P. Gordon, "Interaction forces among solitons in optical fibers," *Opt. Lett.*, vol. 8, No. 11, pp. 596-598, Novembre 1983.
- [89] F. M. Mitschke, and L. F. Mollenauer, "Experimental observation of interaction forces between solitons in optical fibers," *Opt. Lett.*, vol. 12, No. 5, pp. 355-357, Mai 1987.
- [90] Ph. Grelu, F. Belhache, F. Guty, and J. M. Soto-Crespo, "Phase-locked soliton pairs in a stretched-pulse fiber laser," *Opt. Lett.*, vol. 27, No. 11, pp. 966-968, 2002.
- [91] S. Chouli, and Ph. Grelu, "Rains of solitons in a fiber laser," *Opt. Express*, vol. 17, pp. 11776-11781, 2009.
- [92] S. Chouli, and Ph. Grelu, "Soliton rains in a fiber laser : An experimental study," *Phys. Rev. A*, vol. 81, no. 6, Juin 2010.

- [93] F. Sanchez, Ph. Grelu, H. Leblond, A. Komarov, K. Komarov, M. Salhi, A. Niang, F. Amrani, C. Lecaplain, and S. Chouli, "Manipulating dissipative soliton ensembles in passively mode-locked fiber lasers," *Opt. Fib. Technol.*, vol. 20, pp. 562-574, 2014.
- [94] U. Andral, R. Si Fodil, F. Amrani, F. Billard, E. Hertz, and Ph. Grelu, "Fiber laser mode locked through an evolutionary algorithm," *Optica*, vol. 2, pp. 275-278, 2015.
- [95] U. Andral, J. Buguet, R. Si Fodil, F. Amrani, F. Billard, E. Hertz, and Ph. Grelu, "Toward an autotuning mode-locked fiber laser cavity," *J. Opt. Soc. Am. B*, vol. 33, No. 5, pp. 825-833, 2016.
- [96] R. I. Woodward and E. J. R. Kelleher, "Toward 'smart lasers' : self-optimization of an ultrafast pulse source using a genetic algorithm," *Sci. Rep.*, vol. 6, 37616, 2016.
- [97] R. I. Woodward and E. J. R. Kelleher, "Genetic algorithm-based control of birefringent filtering for self-tuning, self-pulsing fiber laser," *Opt. Lett.*, vol. 42, pp. 2952-2955, 2017.
- [98] D. G. Winters, M. S. Kirchner, S. J. Backus, and H. C. Kapteyn. "Electronic initiation and optimization of nonlinear polarization evolution mode-locking in a fiber laser," *Opt. Express*, vol. 25, No. 26, pp. 102-112, 2017.
- [99] Z. Zhuang, S. W. Suh, and J. S. Patel, "Polarization controller using nematic liquid crystals," *Opt. Lett.*, vol. 24, No. 10, pp. 694-696, 1999.
- [100] M. Nikodem, K. Krzempek, K. Zygodlo, G. Dudzik, A. Waz, K. Abramski, and K. Komorowska, "Intracavity polarization control in mode-locked Er-doped fibre lasers using liquid crystals," *Opto-electron. Rev.*, vol. 22, No. 2, pp. 113-117, 2014.
- [101] J. Girardot, F. Billard, A. Coillet, E. Hertz and Ph. Grelu, "Autotuning Mode-Locked Laser Using an Evolutionary Algorithm and Time-Stretch Spectral Characterization," *IEEE Journal of Selected Topics in Quantum Electronics* , vol. 26, no. 5, pp. 1-8, Sept.-Oct. 2020.
- [102] A. Komarov, H. Leblond, and F. Sanchez, "Multistability and hysteresis phenomena in passively mode-locked fiber lasers," *Phys. Rev. A*, vol. 71, no. 5, pp. 053809, 2005.
- [103] X. Liu, "Hysteresis phenomena and multipulse formation of a dissipative system in a passively mode-locked fiber laser," *Phys. Rev. A*, vol. 81, no. 2, pp. 023811, 2010.
- [104] H. B. Winful, "Self-induced polarization changes in birefringent optical fibers," *Appl. Phys. Lett.* , vol. 47, No. 3, pp. 213, 1985.
- [105] J. H. Holland, "Outline for Logical Theory of Adaptive Systems", *Journal of the ACM*, vol. 9, no. 3, pp. 297-314, 1962.
- [106] W. Q. Zhang, S. Afshar V., and T. M. Monro, "A genetic algorithm based approach to fiber design for high coherence and large bandwidth supercontinuum generation," *Opt. Express*, vol. 17, pp. 19311-19327, 2009.
- [107] B. Wetzal, M. Kues, P. Roztock, C. Reimer, P-L. Godin, M. Rowley, B. E. Little, S. T. Chu, E. A. Viktorov, D. J. Moss, A. Pasquazi, M. Peccianti and R. Morandotti, "Customizing supercontinuum generation via on-chip adaptive temporal pulse-splitting," *Nat. Commun.*, vol. 9, pp. 4884, 2018.
- [108] I. Malkiel, M. Mrejen, A. Nagler, U. Arieli, L. Wolf, and H. Suchowski, "Plasmonic nanostructure design and characterization via deep learning," *Light Sci. Appl.*, vol. 7, no. 60, 2018.
- [109] R. S. Hegde, "Deep learning : a new tool for photonic nanostructure design," *Nanoscale Adv.*, vol. 2, pp. 1007-1023, 2020.
- [110] T. Zahavy, A. Dikopoltsev, D. Moss, G. I. Haham, O. Cohen, S. Mannor, and M. Segev, "Deep learning reconstruction of ultrashort pulses," *Optica*, vol. 5, pp. 666-673, 2018.

- [111] S. Kleinert, A. Tajalli, T. Nagy, and U. Morgner, “Rapid phase retrieval of ultrashort pulses from dispersion scan traces using deep neural networks,” *Opt. Lett.*, vol. 44, pp. 979-982, 2019.
- [112] Li Gao, Yang Chai, Darko Zibar, and Zongfu Yu, “Deep learning in photonics : introduction,” *Photon. Res.*, vol. 9, no. 8, 2021.
- [113] G. Genty, L. Salmela, J.M. Dudley, D. Brunner, A. Kokhanovskiy, S. Kobtsev and S. K. Turitsyn , “Machine learning and applications in ultrafast photonics,” *Nat. Photonics*, vol. 15, pp. 91–101, 2021.
- [114] J. N. Kutz, and S. L. Brunton, “Intelligent systems for stabilizing mode-locked lasers and frequency combs : machine learning and equation-free control paradigms for self-tuning optics,” *Nanophotonics*, vol. 4, pp. 459–471, 2015.
- [115] S. L. Brunton, X. Fu, and J. N. Kutz, “Self-Tuning Fiber Lasers,” *IEEE J. Selec. Top. Quantum Electron.*, vol. 20, No. 5, 1101408, 2014.
- [116] T. Baumeister, S. L. Brunton, and J. N. Kutz, “Deep learning and model predictive control for self-tuning mode-locked lasers,” *J. Opt. Soc. Am. B*, vol. 35, No. 3, pp. 617-626, 2018.
- [117] U. Mahlab, J. Shamir, and H. J. Caulfield, “Genetic algorithm for optical pattern recognition,” *Opt. Lett.*, vol. 16, pp. 648–650, 1991.
- [118] K. D. Kihm, and D. P. Lyons, “Optical tomography using a genetic algorithm,” *Opt. Lett.*, vol. 21, pp. 1327–1329, 1996.
- [119] F. Meng, and J. M. Dudley, “Toward a self-driving ultrafast fiber laser,” *Light : Science and Applications*, vol. 9, No. 26, 2020.
- [120] J. Desbois, F. Gires and P. Turnois, “A new approach to picosecond laser pulse analysis shaping and coding,” *IEEE Journal of Quantum Electronics*, vol. 9, no. 2, pp. 213-218, February 1973.
- [121] K. Goda and B. Jalali, “Dispersive Fourier transformation for fast continuous single-shot measurements,” *Nat. Photonics*, vol. 7, pp. 102-112, 2013.
- [122] P. D. Maker, R. W. Terhune, and C. M. Savage, “Intensity-Dependent Changes in the Refractive Index of Liquids,” *Phys. Rev. Lett.*, vol. 12, no. 18, pp. 507–509, 1964.
- [123] D. Radnatarov, S. Khripunov, S. Kobtsev, A. Ivanenko, and S. Kukarin, “Automatic electronic-controlled mode locking self start in fibre lasers with nonlinear polarization evolution,” *Opt. Express*, vol. 21, No. 18, pp. 20626-20631, 2013.
- [124] B. Oktem, C. Ülgüdür, and F. Ö. Ilday, “Soliton-similariton fibre laser,” *Nat. Photonics*, vol. 4, pp. 307-311, 2010.
- [125] J. A. C. Weideman, and B. M. Herbst, “Split-Step Methods for the Solution of the Nonlinear Schrödinger Equation,” *SIAM J. Numer. Anal.*, vol. 23, no. 3, pp. 485–507, 1984.
- [126] O. V. Sinkin, R. Holzlohner, J. Zweck, and C. R. Menyuk, “Optimization of the split-step Fourier method in modeling optical-fiber communications systems,” *JJ. Light. Tech.*, vol. 21, no. 1, pp. 61-68, Jan. 2003.
- [127] J. Girardot, F. Billard, A. Coillet, E. Hertz, and P. Grelu, “Smart auto-setting mode-locked laser using an evolutionary algorithm,” in 2020 EPS-QEOD Europhoton Conference, (European Physical Society, 2020).
- [128] C. Froehly, B. Colombeau, and M. Vampouille, “Shaping and Analysis of Picosecond Light Pulses,” E. Wolf, *Progress in Optics*, Elsevier, vol. 20, pp. 63-153, 1983.

- [129] J. P. Heritage, A. M. Weiner, and R. M. Thurston, "Picosecond pulse shaping by spectral phase and amplitude manipulation," *Opt. Lett.*, vol. 10, No. 12, pp. 609-611, Décembre 1985.
- [130] E. B. Treacy, "Compression of picosecond light pulses," *Phys. Lett.*, vol. 28A, No. 1, pp. 34-35, 21 Octobre 1968.
- [131] E. Treacy, "Optical pulse compression with diffraction gratings," *J. Quant. Elec.*, vol. 5, no. 9, pp. 454-458, September 1969.
- [132] A. M. Weiner, J. P. Heritage, and E. M. Kirshner, "High resolution femtosecond pulse shaping," *J. Opt. Soc. Am. B*, vol. 5, No. 8, pp. 1563-1572, Aout 1988.
- [133] A. M. Weiner, D. E. Leaird, J. S. Patel, and J. R. Wullert, "Programmable femtosecond pulse shaping by use of a multielement liquid-crystal phase modulator," *Opt. Lett.*, vol. 15, No. 6, pp. 326-328, 15 Mars 1990.
- [134] A. M. Weiner, "Femtosecond pulse shaping using spatial light modulators," *Rev. Sc. Instr.*, vol. 71, No. 5, pp. 1929-1960, Mai 2000.
- [135] A. Monmayrant, S. J. Weber, and B. Chatel, "A newcomer's guide to ultrashort pulse shaping and characterization," *J Phys B : At Mol Opt Phys*, vol. 43, no. 10, 2010.
- [136] C. Vaughan, T. Feurer, Katherine W. Stone, and Keith A. Nelson, "Analysis of replica pulses in femtosecond pulse shaping with pixelated devices," *Opt. Express*, vol. 14, pp. 1314-1328, 2006.
- [137] E. Hertz, F. Billard, G. Karras, P. Béjot, B. Lavorel, and O. Faucher, "Shaping of ultra-violet femtosecond laser pulses by Fourier domain harmonic generation," *Opt. Express*, vol. 24, pp. 27702-27714, 2016.
- [138] R. Iegorov, T. Teamir, G. Makey, and F. Ö. Ilday, "Direct control of mode-locking states of a fiber laser," *Optica*, vol. 3, No. 12, pp. 1312-1315, 2016.
- [139] A.F.J. Runge, D. D. Hudson, K. K. K. Tam, C. Martijn de Sterke and A. Blanco-Redondo, "The pure-quartic soliton laser," *Nat. Photonics*, vol. 14, pp. 492-497, 2020.
- [140] J. Girardot, F. Billard, A. Coillet, M. Nafa, E. Hertz, and P. Grelu, "Autosetting Mode-locked Laser with Genetic Algorithm Optimization and Advanced Intracavity Controls," in 2021 Conference on Lasers and Electro-Optics Europe and European Quantum Electronics Conference, OSA Technical Digest (Optical Society of America, 2021).
- [141] J. Girardot, F. Billard, A. Coillet, M. Nafa, E. Hertz, and Ph. Grelu, "On-demand ultrafast soliton molecules through genetic algorithm optimization," in 2021 European Optical Society Annual Meeting.
- [142] J. Girardot, A. Coillet, M. Nafa, F. Billard, E. Hertz, and Ph. Grelu, "On-demand generation of soliton molecules through evolutionary algorithm optimization," *Opt. Lett.*, vol. 47, pp. 134-137, 2022.