



Etude expérimentale des conditions initiales de l'instabilité de Rayleigh-Taylor au front d'ablation en fusion par confinement inertiel

Barthélémy Delorme

► To cite this version:

Barthélémy Delorme. Etude expérimentale des conditions initiales de l'instabilité de Rayleigh-Taylor au front d'ablation en fusion par confinement inertiel. Physique des plasmas [physics.plasm-ph]. Université de Bordeaux, 2015. Français. NNT : 2015BORD0490 . tel-01289909

HAL Id: tel-01289909

<https://theses.hal.science/tel-01289909>

Submitted on 16 Oct 2017

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

présentée à

L'UNIVERSITÉ DE BORDEAUX

ÉCOLE DOCTORALE DES SCIENCES PHYSIQUES ET DE L'INGÉNIEUR

par **Barthélémy DELORME**

POUR OBTENIR LE GRADE DE

DOCTEUR

SPÉCIALITÉ : ASTROPHYSIQUE, PLASMAS, NUCLÉAIRE

Etude expérimentale des conditions initiales de l'instabilité de Rayleigh-Taylor au front d'ablation en Fusion par Confinement Inertiel

Soutenue le 21 janvier 2015 devant le jury composé de :

M. Dimitri BATANI, CELIA, Université de Bordeaux

M. Michel BOUSTIE, CNRS-ENSMA

M. Alexis CASNER, CEA Bruyères Le Châtel

Mme Catherine CHERFILS, CEA Bruyères Le Châtel

M. Michel KOENIG, LULI, Ecole Polytechnique

Mme Marina OLAZABAL-LOUME, CEA CESTA

M. Vladimir TIKHONCHUK, CELIA, Université de Bordeaux

Président

Examineur

Directeur de thèse

Rapporteur

Rapporteur

Directrice de thèse

Directeur de thèse

Remerciements

Il est compliqué d'écrire des remerciements, tout d'abord parce que je suis sûr que je vais oublier des personnes pourtant importantes pour moi et qui ont compté dans la réalisation de cette thèse ou dans mon orientation vers la recherche, et d'autre part car je ne me sens pas capable de retranscrire fidèlement, par des mots, toute l'ampleur de ma gratitude. Néanmoins, je tiens quand même à me livrer à l'exercice car il serait encore pire que je me "thèse" complètement à ce sujet.

Pour commencer, je voudrais remercier mes directeurs de thèse, Alexis, Marina et Vladimir, qui ont tous les trois aussi été parmi mes enseignants en master. Si j'ai choisi cette voie, c'est qu'ils ont su, ainsi que de nombreux autres professeurs, me transmettre, au cours de mes études, un fort intérêt pour les sciences et une véritable passion pour le grand jeu de réflexion qu'est la recherche. Alexis a eu le rôle principal dans la réalisation de cette thèse : merci à lui de m'avoir proposé ce sujet, d'avoir toujours fait l'effort de me suivre même à distance, d'avoir su m'encadrer tout en me laissant suffisamment de liberté et d'autonomie pour suivre mes propres idées, même quand nous n'étions pas d'accord, et surtout un grand merci d'avoir toujours fait passer mon intérêt avant le sien. Mille mercis aussi pour les activités extra-professionnelles, notamment la traversée de Manhattan à pied, le match des Golden State Warriors (dommage, ils sont meilleurs maintenant qu'à l'époque!), les nombreuses discussions politiques ou sportives, et merci d'avance pour les conseils vélocipédiques.

Pour rendre hommage à Marina, qui m'a encadré quotidiennement au CELIA pendant plus de deux ans, je voudrais évoquer sa disponibilité après qu'elle a quitté le laboratoire. En effet, lors de cette dernière année, elle a régulièrement pris sur ses soirées pour venir en réunion au CELIA, elle a toujours été la première à répondre à mes sollicitations ou à corriger mes divers écrits ; elle a même continué à résoudre mes difficultés vis-à-vis des simulations, permettant à ma "branche numérique" de ne pas s'éteindre avec son départ. Un grand merci donc pour sa disponibilité totale pendant ces trois ans, qui, allant de pair avec sa bonne humeur, son optimisme et sa rigueur, en ont fait une directrice de thèse plus qu'idéale.

Pour clore le chapitre sur mes directeurs de thèse, je voudrais exprimer toute ma gratitude à Vladimir. Je dois m'avouer tout simplement ébahi par sa capacité à trouver du temps pour tout le monde : depuis la fin de la première année de master jusqu'à aujourd'hui, il n'a cessé de me suivre, de m'encourager, de trouver le temps à n'importe quel moment pour une réunion dont on ressort toujours motivé et plein de nouvelles idées. Et ce malgré un emploi du temps tellement chargé que seul la physique quantique peut l'expliquer raisonnablement. Merci beaucoup à Vladimir pour son mélange d'exigence et d'humanité qui m'a poussé à croire en moi et à faire plus que je ne me croyais capable au cours de cette thèse.

Outre mes directeurs de thèse, je voudrais remercier l'ensemble des permanents du CELIA qui ont su se montrer remarquablement disponibles pour répondre à la moindre de mes questions, et particulièrement Xavier, qui fut mon "encyclopédie" de la physique, m'aidant dans des domaines très divers et ce jusqu'à la veille de ma soutenance ; Philippe, qui, bien que n'étant pas mon directeur de thèse, s'est montré très présent, notamment lors des répétitions de la soutenance, ainsi que Jérôme, qui a su composer de manière pédagogique avec mes lacunes numériques.

Je tiens aussi à adresser mes remerciements aux membres du jury, Michel Boustie, qui a gentiment accepté d'en être le président, Dimitri Batani, mes trois directeurs de thèse, ainsi que Catherine Cherfils et Michel Koenig, qui m'ont fait l'honneur d'être les rapporteurs de cette thèse.

Je suis bien conscient d'avoir été un doctorant très chanceux, grâce notamment à un superbe environnement de travail. Je remercie pour cela Philippe Balcou, qui m'a accueilli au CELIA, et mes divers supérieurs à la DIF, particulièrement Eric Benso, Denis Juraszek et Thierry Garié, qui m'ont permis d'assister à un grand nombre de conférences, d'effectuer plusieurs formations et surtout de me déplacer aux Etats-Unis pour chaque campagne expérimentale, dans des conditions toujours excellentes. Merci à eux aussi de m'avoir accordé une prolongation de contrat pour clôturer les projets dans lesquels j'étais engagé.

Merci aussi à tous les collaborateurs extérieurs au CELIA qui ont rendu possible la réalisation de ce travail, notamment Vladimir Smalyuk et David Martinez, nos collègues américains sans lesquels nous n'aurions pu participer à la campagne d'étude de l'instabilité de Richtmyer-Meshkov ablative. Merci aussi à tous les membres de la collaboration internationale sur l'étude des mousses sous-denses, Nataliya et le groupe russe du Lebedev qui nous ont fourni les mousses, l'équipe japonaise d'Osaka sans laquelle nous n'aurions pas eu de feuilles de plastique, Dustin et Tomline grâce auxquels nous avons obtenu deux journées de tir sur le programme interne du Laboratory for Laser Energetics de Rochester, Gilles pour ses simulations PARAX, Wolf pour ses dépouillements de données et Mickael Grech pour son apport de connaissances sur les instabilités paramétriques. Je ne peux qu'être admiratif devant les équipes des lasers OMEGA et OMEGA EP pour leur professionnalisme ainsi que pour l'éblouissante efficacité et fiabilité de leurs installations, et les remercier pour toutes ces données que nous avons pu acquérir. Merci aussi à Laurent Masse pour les différentes discussions passionnantes sur la physique des instabilités hydrodynamiques ablatives.

Je voudrais aussi remercier de nombreux autres collègues qui ont fait qu'aller au travail n'a jamais été une contrainte, et même bien souvent un plaisir. Merci donc (et j'en oublie beaucoup) à Dario, Basile, Aurore, Xavier V., Gabriel, Amélie, Mathieu B.-G., Mathieu D., Marion, Pierre, Alexandre, Sébastien, Mickael, Emmanuelle-Robertine, Céline ainsi que mes collègues de bureau, bureau dont l'excellente ambiance était enviée dans tout

le laboratoire : Emma, Mokrane, Witold, Rémi et Alexandra. Les conférences resteront de superbes souvenirs, notamment celles où j'ai eu l'occasion de passer du temps avec Xavier V. et Mathieu B.-G. qui sont deux personnes d'une grande gentillesse. J'ai enfin une pensée pour Mickael et ces images si particulières, depuis le retour en caddie à Aix-en-Provence en master jusqu'au burger de kangourou à Oxford en septembre dernier. Et une autre pensée pour Alexandra et les moments passés ensemble, avec en point d'orgue les forums ILP et le séjour à Rochester.

Je tiens aussi à adresser mes remerciements à Ralf Richter, à son équipe et au laboratoire du CIC BiomaGUNE, dans lequel j'ai effectué mon stage de première année de master. C'est ce stage qui m'a donné envie de faire de la recherche et m'a convaincu d'entreprendre un doctorat.

Si l'adage bien connu associe un esprit sain à un corps sain, j'en ai une version personnelle qui pourrait s'exprimer sous la forme "un travail sain dans une vie saine". Je voudrais donc adresser mes plus sincères remerciements à tous ceux qui ont fait des ces trois années une période inoubliable. Merci à ma famille, à mes parents particulièrement, car ils m'ont offert des conditions optimales pour étudier et, surtout, m'ont toujours laissé une liberté totale quant au choix de mon cursus, me permettant donc de me laisser guider par le plaisir des études en physique fondamentale. Un grand merci aussi à eux, à mes frères, à Michèle et Claude, à Pauline et à Géraldine pour l'organisation de mon pot de thèse, ainsi qu'à Hélène et Xavier pour la chemise et la cravate. Merci aussi à tous mes amis, pour le génial cadeau de fin de thèse, ainsi qu'à l'équipe des Vikings du Lucifer, avec ses quelques belles victoires et ses défaites mémorables (les jokers, les jokers...), merci au club du BEC et l'ensemble de mes coéquipiers, car si le rugby m'a coûté un mois et demi de thèse, il a aussi été une soupape merveilleuse pour échapper à la pression, et m'a offert de superbes rencontres. Merci à mes colocataires, et notamment à Clément, qui m'a tenu compagnie sur sa console lors de longues soirées de rédaction et qui m'a permis de gagner la ligue des champions avec les Girondins de Bordeaux. Merci aussi à Perrine et Thibault, pour avoir été complètement inutiles.

Pour finir, merci à Géraldine, qui a énormément compté pour moi durant ces trois ans et avec qui j'ai tant partagé. Des superbes souvenirs à Helsinki, Rome ou Oxford, des moments plus difficiles comme durant la rédaction, elle a toujours été là. Merci à toi.

En guise de conclusion, je pense qu'on est en grande partie le produit de son environnement. Ainsi, cette thèse est l'accomplissement de tous ceux qui m'ont éduqué, encouragé, encadré, enseigné, entouré. Mes plus sincères et profonds remerciements vous sont donc à tous adressés.

Résumé :

Les différents dimensionnements et expériences de Fusion par Confinement Inertiel (FCI) en attaque directe comme indirecte montrent qu'une des principales limites à l'atteinte de l'ignition est l'instabilité de Rayleigh-Taylor (IRT) qui cause la rupture de la coquille de la cible en vol et potentiellement le mélange du combustible chaud du coeur avec celui, froid, de la coquille. La connaissance, la compréhension et la maîtrise des conditions initiales de ce mécanisme sont donc d'un grand intérêt.

Nous présentons ainsi une étude expérimentale et théorique des conditions initiales de l'IRT ablative en attaque directe au travers de deux campagnes expérimentales réalisées sur le laser OMEGA (LLE, Rochester). La première campagne concerne l'étude de l'instabilité de Richtmyer-Meshkov (IRM) ablative imprimée par laser ; cette instabilité commence à se développer au début de l'irradiation laser et fixe l'ensemencement de l'IRT. Nous avons mis en place une configuration expérimentale qui a permis de mesurer l'évolution temporelle de l'IRM ablative imprimée par laser pour la première fois. Nous présentons ensuite une interprétation des résultats de cette expérience par des simulations hydrodynamiques réalisées avec le code CHIC, ainsi que par un modèle théorique de l'IRM ablative imprimée par laser. Nous montrons que le moyen le plus direct de contrôler cette instabilité est de réduire l'amplitude des défauts d'intensité laser.

Ceci peut être accompli en utilisant des cibles couvertes par une couche de mousse de basse densité. Ainsi, lors de la deuxième campagne, nous avons étudié pour la première fois l'effet de mousses sous-denses sur la croissance de l'IRT ablative. Au cours de ces expériences, des feuilles de plastique recouvertes d'une couche de mousse ont été irradiées par un faisceau laser portant une perturbation d'intensité destinée à imprimer des modulations sur la cible. Différentes données expérimentales sont présentées : rétrodiffusion de l'énergie laser, dynamique de la cible obtenue par mesure de côté d'auto-émission et radiographies de face faisant apparaître l'effet des mousses sur les modulations de densité surfacique des cibles. Ces données ont ensuite été interprétées à l'aide de simulations CHIC et du code d'interaction laser-plasma PARAX. Nous montrons qu'une des mousses réduit l'amplitude des modulations de l'intensité laser d'un facteur 2.

Par conséquent, cette thèse a donné lieu au développement de configurations expérimentales et d'un ensemble d'outils de dépouillement numériques pour l'étude approfondie des instabilités hydrodynamiques en FCI.

Mots clefs : fusion par confinement inertiel, instabilités hydrodynamiques, instabilité de Rayleigh-Taylor, instabilité de Richtmyer-Meshkov, attaque directe, front d'ablation

Summary :

Numerous designs and experiments in the domain of Inertial Confinement Fusion (ICF) show that, in both direct and indirect drive approaches, one of the main limitations to reach the ignition is the Rayleigh-Taylor instability (RTI). It may lead to shell disruption and performance degradation of spherically imploding targets. Thus, the understanding and the control of the initial conditions of the RTI is of crucial importance for the ICF program.

In this thesis, we present an experimental and theoretical study of the initial conditions of the ablative RTI in direct drive, by means of two experimental campaigns performed on the OMEGA laser facility (LLE, Rochester). The first campaign consisted in studying the laser-imprinted ablative Richtmyer-Meshkov instability (RMI) which starts at the beginning of the interaction and seeds the ablative RTI. We set up an experimental configuration that allowed to measure for the first time the temporal evolution of the laser-imprinted ablative RMI. The experimental results have been interpreted by a theoretical model and numerical simulations performed with the hydrodynamic code CHIC. We show that the best way to control the ablative RMI is to reduce the laser intensity inhomogeneities.

This can be achieved with targets covered by a layer of a low density foam. Thus, in the second campaign, we studied for the first time the effect of underdense foams on the growth of the ablative RTI. A layer of low density foam was placed in front of a plastic foil, and the perturbation was imprinted by an intensity modulated laser beam. Experimental data are presented : backscattered laser energy, target dynamic obtained by side-on self-emission measurement, and face-on radiographs showing the effect of the foams on the target areal density modulations. These data were interpreted using the CHIC code and the laser-plasma interaction code PARAX. We show that the foams noticeably reduce the amplitude of the laser intensity inhomogeneities and the level of the subsequent imprinted ablation front modulations.

In conclusion, this thesis allowed us to develop an experimental platform and a suite of numerical tools for future, more detailed studies of hydrodynamic instabilities for ICF applications.

Keywords : inertial confinement fusion, hydrodynamic instabilities, Rayleigh-Taylor instability, Richtmyer-Meshkov instability, direct-drive, ablation front

Glossaire :

AD : Attaque Directe
AI : Attaque Indirecte
ASBO : Active Shock Break-Out
ASP : Alignement Sensor Package
CBET : Cross Beam Energy Transfer
CBF : Caméra à Balayage de Fente
CCD : Charge-Couple Device
CEA : Commissariat à l'Energie Atomique et aux Energies Alternatives
CESTA : Centre d'Etudes Scientifiques et Techniques d'Aquitaine
CLARA : Cristal Large-Aperture Ring Amplifier
CPA : Chirped Pulse Amplification
DAM : Direction des Applications Militaires
DER : Driver Electronic Room
DIF : DAM Ile-de-France
DPP : Distributed Phase Plate
DPR : Distributed Phase Rotator
DT : Deutérium-Tritium
ETP : Equivalent Target Plane
F-ASP : F-Alignement Sensor Package
FABS : Full-Aperture Backscatter Station
FCI : Fusion par Confinement Inertiel
FCM : Fusion par Confinement Magnétique
FFT : Fast Fourier Transform
FTM : Fonction de Transfert des Modulations
GMC : Galette de Microcanaux
GP : Gaz Parfait
HiPER : High-Power Laser Energy Research
IDL : Interactive Data Language
ILE : Institute of Laser Engineering
IR : Infra-Rouge
IRM : Instabilité de Richtmyer-Meshkov
IRT : Instabilité de Rayleigh-Taylor
ITER : International Thermonuclear Experimental Reactor
KDP : Potassium Dihydrogen Phosphate
LARA : Large-Aperture Ring Amplifier
LB : Laser Bay
LBS : Laboratory Basic Science
LIL : Ligne d'Intégration Laser

LLE : Laboratory for Laser Energetics
LLNL : Lawrence Livermore National Laboratory
LMJ : Laser Méga-Joule
NIC : National Ignition Campaign
NIF : National Ignition Facility
NLUF : National Laser User's Facility
OPA : Optical Parametric Amplification
OPCPA : Optical Parametric Chirped Pulse Amplification
PGR : Pulse Generation Room
PI : Principal Investigator
ROSS : Rochester Optical Streak System
RPP : Random Phase Plate
SG4 : Supergaussienne d'ordre 4
SRF : Shot Request Form
SSCA : SIM Streak Camera
SSD : Smoothing by Spectral Dispersion
TB : Target Bay
TC : Target Chamber
TIM : Ten-Inch Manipulator
TMS : Target Mirror Structure
TPS : Target Positionner System
UR : University of Rochester
UV : Ultra-Violet
VISAR : Velocity Interferometry System for Any Reflector
XRFC : X-Ray Framing Camera
XRPHC : X-Ray Pinhole Camera

Table des matières

1	Introduction	13
1.1	Contexte actuel de la FCI	13
1.1.1	La fusion comme source d'énergie : deux approches	13
1.1.2	Limites à l'atteinte de l'ignition en FCI	15
1.1.3	Effets délétères des instabilités hydrodynamiques dans les dernières expériences de FCI	18
1.2	Compréhension de l'ensemencement de l'IRT ablative	19
1.3	Amélioration des méthodes de contrôle des instabilités hydrodynamiques .	21
1.4	Objectifs de thèse	23
2	Physique des instabilités hydrodynamiques au front d'ablation	27
2.1	Physique du front d'ablation	27
2.1.1	Absorption de l'énergie du laser	28
2.1.2	Zone de conduction	31
2.1.3	Génération et propagation d'un choc	33
2.2	L'IRT au front d'ablation	34
2.2.1	L'IRT classique	34
2.2.2	Effet de l'ablation sur l'IRT	36
2.2.3	L'IRT ablative en phase non-linéaire	38
2.3	L'IRM ablative	42
2.3.1	Physique de l'IRM classique	42
2.3.2	L'IRM ablative	43
2.4	Bilan expérimental	47
2.4.1	Expériences d'IRM ablative	48
2.4.2	Amélioration de l'uniformité de l'intensité du laser	50
2.4.3	Lissage des perturbations	51
2.4.4	Diminution du taux de croissance de l'IRT ablative	54
2.5	Points manquants traités dans cette thèse	57
3	Matériel et méthodes	59
3.1	Installation laser OMEGA	59

3.1.1	Laser OMEGA	59
3.1.2	Préparation d'une journée de tirs	63
3.1.3	Contraintes d'exploitation sur OMEGA	64
3.2	Diagnostics utilisés	66
3.2.1	X-Ray Framing Camera (XRFC)	67
3.2.2	Imageur X couplé à une caméra à balayage de fente (SSCA)	68
3.2.3	Diagnostics laser	69
3.2.4	Diagnostics fixes	70
3.3	Programmes de dépouillement développés sous IDL	70
3.3.1	Corrélation croisée	70
3.3.2	Calcul du niveau de modulation	73
3.3.3	Calcul de la barre d'erreur	78
3.3.4	Segmentation d'image	78
3.4	Code d'hydrodynamique radiative CHIC	80
4	Mesure de la phase Richtmyer-Meshkov imprimée par laser	87
4.1	Configuration expérimentale	87
4.1.1	Contexte de conception de l'expérience	88
4.1.2	Cibles	89
4.1.3	Faisceaux laser sur cible	92
4.1.4	Radiographie de face et autres diagnostics	95
4.2	Mesure de l'évolution des modulations de densité surfacique pendant la phase Richtmyer-Meshkov	96
4.2.1	Déroulement des journées de tir	96
4.2.2	Exemples de radiographies de face	96
4.2.3	Dépouillement des radiographies de face	100
4.2.4	Données issues du dépouillement des radiographies de face	106
4.3	Détermination du niveau d'empreinte induit par les lames de phase	106
5	Interprétation des mesures l'IRM imprimée par laser à partir du code CHIC et du modèle de Goncharov	113
5.1	Réalisation de simulations CHIC de la croissance des modulations imprimées par laser et calcul du modèle de Goncharov	114
5.1.1	Intensité équivalente pour les simulations CHIC 1D	114
5.1.2	Calcul analytique avec le modèle de Goncharov	115
5.1.3	Simulation de la croissance des modulations imprimée par laser	117
5.2	Comparaison des simulations CHIC et des données expérimentales	117
5.3	Interprétation avec le modèle de Goncharov	119
5.3.1	Comparaison des simulations CHIC et du modèle de Goncharov	119

5.3.2	Interprétation des observations expérimentales par le modèle de Goncharov	121
5.4	Etude des variations des simulations CHIC en fonction des paramètres numériques et expérimentaux	127
5.4.1	Variation du niveau d'empreinte	127
5.4.2	Variation de la forme de l'impulsion	127
5.4.3	Variation de l'équation d'état et du limiteur de flux	129
5.5	Conséquences de l'étude au LLE	133
5.6	Conclusion sur les interprétations	134
6	Démonstration de la réduction par des mousses sous-denses des instabilités hydrodynamiques 2D imprimées par un motif de référence	137
6.1	Configuration expérimentale	137
6.1.1	Contexte de l'expérience	137
6.1.2	Cibles	138
6.1.3	Faisceaux laser	139
6.1.4	Diagnostics	142
6.2	Analyse des données de la caméra à balayage de fente	143
6.2.1	Interprétation des images d'auto-émission	143
6.2.2	Profils extraits des données d'auto-émission mesurées par la caméra à balayage de fente	146
6.3	Détermination du temps d'ionisation	146
6.3.1	Simulations de l'ionisation des mousses	146
6.3.2	Données FABS	149
6.4	Dépouillement des données de radiographie de face	150
6.4.1	Comparaison des radiographies de face avec différents types de cibles et de motifs d'empreinte	150
6.4.2	Evolution temporelle des modulations de densité surfacique	151
6.4.3	Analyse des distributions de bulles obtenues avec les cibles avec mousses	153
6.4.4	Interprétation de l'origine des structures 3D	155
6.5	Interprétation des données expérimentales par les simulations CHIC et PARAX	157
6.5.1	Simulations d'interaction laser-plasma PARAX	157
6.5.2	Comparaison des courbes de croissance expérimentales et numériques	158
6.5.3	Effet des mousses sur l'ensemble du spectre spatial	161
7	Conclusion et perspectives	163
7.1	Etude de l'instabilité de Richtmyer-Meshkov imprimée par des inhomogénéités de l'intensité laser	163

7.2	Etude du lissage de l’empreinte laser par des mousses sous-denses	165
A	Annexe A : Détail du modèle de Goncharov	169
B	Annexe B : Méthode alternative de dépouillement des radiographies de face	171
C	Annexe C : Laser OMEGA EP	175

Chapitre 1

Introduction

1.1 Contexte actuel de la FCI

1.1.1 La fusion comme source d'énergie : deux approches

La fusion nucléaire est un axe fondamental de la recherche de nouvelles sources d'énergie. A ce titre, la France est impliquée dans deux projets majeurs : HiPER (High-Power Laser Energy Research), sous l'égide de l'Union Européenne, et ITER (International Thermonuclear Experimental Reactor), en collaboration avec 34 pays, en cours de construction à Cadarache (Bouches-du-Rhône). Ils illustrent les deux méthodes principales envisagées pour réussir à obtenir un gain d'énergie par fusion nucléaire contrôlée : la Fusion par Confinement Inertiel (FCI) [1] et la Fusion par Confinement Magnétique (FCM) [2]. Dans les deux cas, le but est de réaliser la fusion de deux atomes légers en un atome plus lourd pour libérer de l'énergie. Pour comprendre cela, il faut se référer à la courbe d'Aston qui donne l'énergie de liaison moyenne des nucléons d'un atome en fonction de son numéro atomique Z (cf [3] p. 2-3). L'énergie de liaison correspond à l'énergie qu'il faut fournir pour casser les liaisons d'un nucléon avec le noyau de l'atome. On peut voir que les éléments légers (de l'hydrogène (H) au bore (B)) possèdent une énergie de liaison faible qui croît fortement avec Z . Cela signifie que la fusion de deux de ces atomes va créer un atome de Z plus élevé qui sera nettement plus stable. Or, un objet plus stable possède un niveau d'énergie plus bas : de l'énergie va donc être dégagée par la fusion de deux atomes légers. Dans le cadre de la FCI comme de la FCM, c'est le couple deutérium-tritium (DT) qui est principalement étudié. Ces deux isotopes de l'hydrogène ($A=2$ pour D et $A=3$ pour T) présentent la meilleure section efficace de réaction pour tous les atomes légers. La difficulté majeure à surpasser pour réaliser une réaction de fusion est la répulsion électrique des deux noyaux chargés positivement. Pour cela, il est nécessaire de fournir aux noyaux une très grande énergie cinétique et donc de les porter à une température très élevée. A cette température de l'ordre d'une centaine de millions de degrés, la matière est complètement ionisée, elle est en état de plasma. La FCM propose de créer un plasma de DT,

puis de le confiner grâce à des champs magnétiques toroïdaux (en forme d'anneau). La notion de confinement est fondamentale et correspond au fait de garder le plasma dans des conditions thermodynamiques propices à la fusion en l'isolant de l'extérieur. Le plasma est composé d'ions et d'électrons, donc de particules chargées, et est ainsi sensible aux champs magnétiques : il va être confiné au coeur du tore. Ainsi, le plasma va pouvoir être chauffé sans se refroidir par contact direct avec les parois du réacteur. Ceci va permettre d'atteindre la température de fusion et d'obtenir un gain supérieur à 1, c'est-à-dire une puissance produite supérieure à celle investie pour la création du plasma, son chauffage et le fonctionnement du réacteur.

La FCI, dans le cadre de laquelle cette thèse a été réalisée, consiste à imploder une microsphère à l'aide de faisceaux laser. Les microsphères (cibles) sont faites d'une couche extérieure appelée "ablateur", qui peut être constituée par exemple de plastique. A l'intérieur de l'ablateur se trouve une couche de DT sous forme de glace, puis le coeur de la cible est constitué de DT gazeux (cf figure 1.1). Ce mélange de DT constituera le combustible pour la réaction de fusion. Deux approches sont envisagées pour la FCI : l'attaque directe (AD) et l'attaque indirecte (AI). La première méthode consiste à diriger les faisceaux laser directement sur la cible. En AI, la cible est placée dans une cavité constituée d'un matériau de Z élevé, en général de l'or ou un mélange or-uranium, appelée hohlraum. Cette cavité, illuminée par les faisceaux lasers, génère des rayons X qui vont à leur tour irradier la capsule et la cavité. Ces deux approches sont illustrées en figure 1.1. Dans les deux cas, l'irradiation de la cible va créer une onde thermique qui se propagera vers le coeur de la cible. Le pied de cette onde thermique, où la matière est ablatée en un plasma qui se détend vers l'extérieur de la cible, est appelé front d'ablation. Depuis ce front, une série de chocs va être lancée dans l'ablateur qui servira de piston pour la compression du combustible. La théorie prévoit alors une forte compression qui permet d'atteindre de très hautes pressions et températures au coeur de la cible. Les réactions de fusion vont débiter dans le coeur du gaz comprimé, formant un point chaud. Les noyaux d'hélium (aussi appelés particules α) formés par les réactions de fusion vont alors déposer une partie de leur énergie dans le point chaud, augmentant sa température et le taux de réaction. Ceci amorce une flamme nucléaire que va se se propager à travers le combustible, transférant son énergie et créant ainsi de nouvelles réactions de fusion. Un régime de combustion est atteint, ce qui permet de brûler une partie importante de la coquille et d'obtenir un gain en énergie. Deux installations majeures existent aujourd'hui pour réaliser des expériences de FCI : le Laser MégaJoule (LMJ), situé dans le centre du CEA/CESTA au Barp, en Gironde, et le National Ignition Facility (NIF) au Lawrence Livermore National Laboratory (LLNL) à Livermore, Etats-Unis.

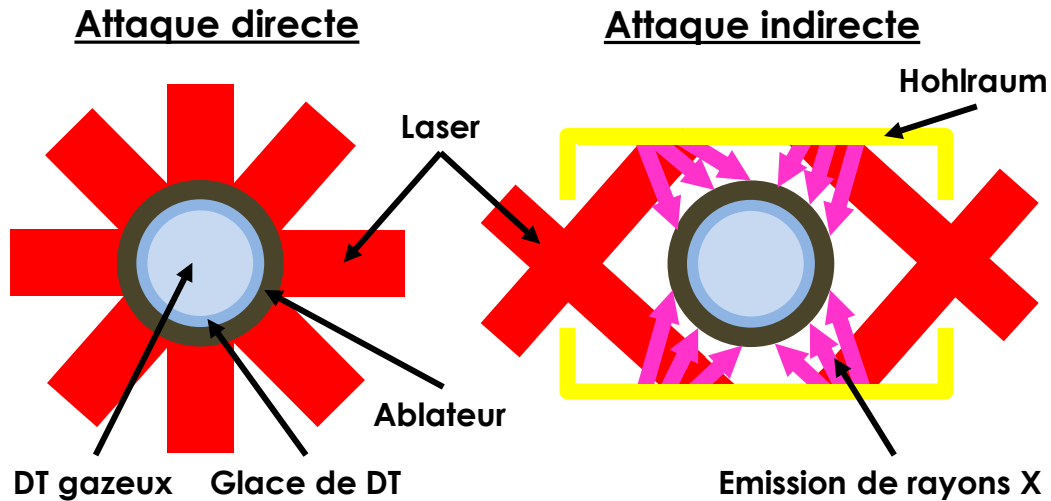


FIGURE 1.1 – Schéma des deux approches pour la FCI : l'attaque directe (à gauche) et l'attaque indirecte (à droite).

1.1.2 Limites à l'atteinte de l'ignition en FCI

Les phénomènes limitants pour l'atteinte de l'ignition sont connus de longue date. Cependant, l'échec de la campagne d'ignition conduite jusqu'à 2012 sur le NIF a remis en évidence ces différents obstacles. Pour la plupart d'entre eux, la compréhension physique et les simulations intégrées restent incapables de rendre compte des observations expérimentales faites en AI sur le NIF. Un schéma représenté en figure 1.2 aborde ces différentes limitations :

- **Couplage du laser à la cible** En AI, le schéma choisi sur le NIF et le LMJ, le laser illumine la cavité qui émet ensuite des rayons X ; or, environ 20 % de l'énergie laser n'est pas convertie en rayons X. De plus, que ce soit en AD ou en AI, l'interaction des faisceaux laser et des plasmas de cavité et/ou de cible peut provoquer l'apparition d'instabilités paramétriques du fait du couplage de l'onde laser et de modes propres du plasma. Ces instabilités entraînent la diffusion de l'énergie laser, l'éclatement du faisceau en multiples filaments ou la génération d'électrons suprathermiques. Enfin, le croisement de faisceaux peut provoquer un transfert d'énergie d'un faisceau à l'autre, appelé "Cross Beam Energy Transfer" (CBET). Ce phénomène est à même de se produire au niveau des fenêtres d'entrée des hohlraums en AI ou quand la compression de la cible a débuté en AD. Il peut induire une répartition non-homogène de l'intensité des faisceaux illuminant la cible, causant une asymétrie d'implosion.
- **Asymétrie d'implosion** Les asymétries d'implosion peuvent donc être provoquées par une illumination inhomogène de la cible, provoquée par exemple par le CBET, ou par des déformations macroscopique de cette cible. Ces asymétries sont

transmises au point chaud ; la surface de contact entre le point chaud et le combustible froid va alors dévier de sa forme sphérique idéale, accroissant les pertes thermiques et la fraction de particules α s'échappant du coeur, diminuant le gain.

- **Préchauffage de la cible** Durant la phase de compression, la coquille du combustible solide doit rester aussi froide que possible. En effet, une augmentation de sa température et donc de sa pression opposerait une résistance à la compression. Le préchauffage peut être provoqué par le dépôt d'énergie d'électrons suprathermiques ou par un saut d'entropie inconsideré dû à des chocs mal synchronisés.
- **Instabilités hydrodynamiques** Enfin, une des limitations critiques de la performance des implosions sont les instabilités hydrodynamiques, et spécialement l'instabilité de Rayleigh-Taylor (IRT). Les instabilités hydrodynamiques sont des phénomènes physiques se développant dans les fluides, qui provoquent une évolution instable de perturbations de ces fluides. L'IRT se développe à une interface quand un fluide léger accélère un fluide lourd, soit dans un champ de gravité soit du fait d'une accélération, lorsque l'accélération va dans le sens du fluide léger vers le fluide lourd. Dans le cas de la FCI, on trouve deux phases d'accélération instables : lors de la compression de la cible par le laser, l'accélération de la coquille solide par le plasma ablaté est centripète, tandis qu'à la stagnation, c'est-à-dire lorsque la pression du coeur de la cible ralentit la coquille en vol, l'accélération est alors centrifuge (décélération). Ainsi lors de la compression, le front d'ablation, qui sépare le plasma léger de l'ablateur plus lourd, va être sensible au développement d'une instabilité de type IRT alors qu'à la décélération ce type d'instabilité va se développer à la face interne de la coquille dense autour du point chaud. De plus, l'IRT peut se développer au niveau des interfaces entre les différents matériaux de la coquille. Cependant, l'instabilité nécessite une perturbation initiale pour se développer. La figure 1.3 présente les différentes origines de perturbations des interfaces des cibles, qui seront ensuite amplifiées par l'IRT. Une des sources de perturbation est la rugosité des cibles, autrement dit les irrégularités des surfaces apparaissant lors de leur usinage. Ces irrégularités peuvent se retrouver aux différentes interfaces de la cible. Dans le cas de l'AD, les inhomogénéités de l'intensité laser peuvent imprimer des modulations à la surface de l'ablateur. Les défauts peuvent aussi être transportés d'une interface à l'autre par le biais des ondes de choc et de raréfaction générées par le laser.

Aux prémices de la FCI, l'IRT était déjà suspectée de pouvoir perturber les implosions [4]. Cependant, on rappellera qu'une modélisation inadéquate de l'IRT faisait promettre l'ignition pour une énergie laser d'environ 1 kJ dès 1972 [1]. Comme on va le montrer au paragraphe suivant, les instabilités hydrodynamiques sont un problème toujours d'actualité, car elles sont un facteur limitant majeur des performances actuelles du NIF. Cette thèse porte sur l'étude des instabili-

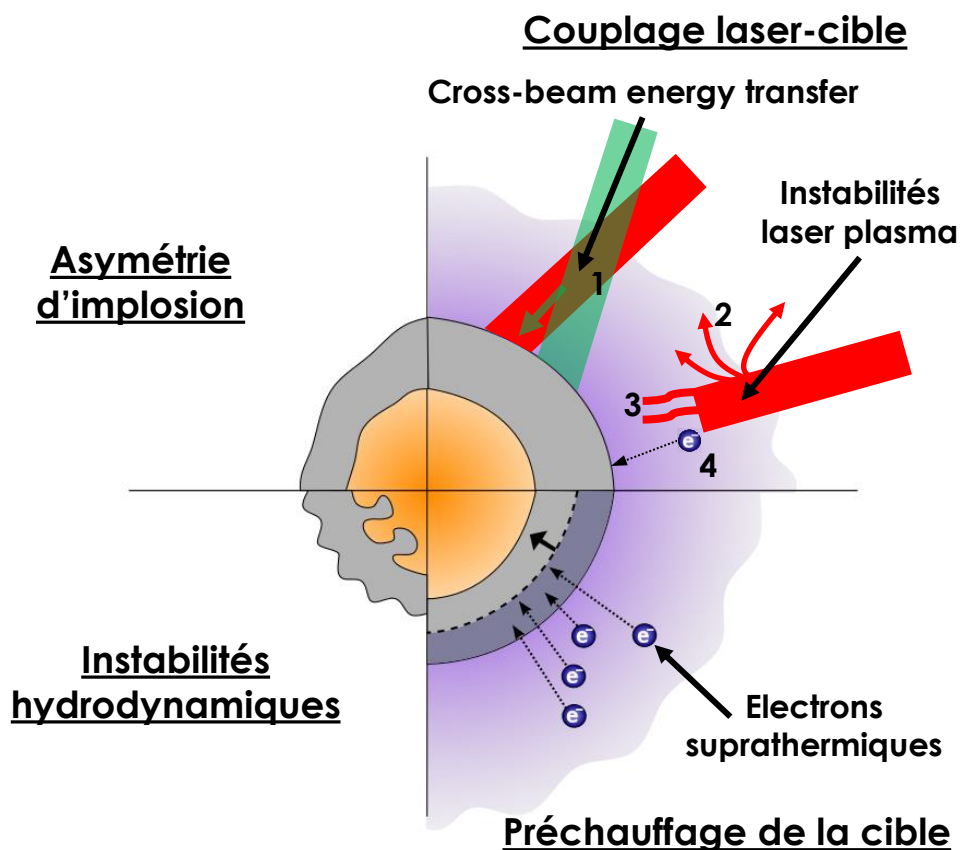


FIGURE 1.2 – Schéma des différentes limites à l'atteinte de l'ignition rencontrées en FCI.

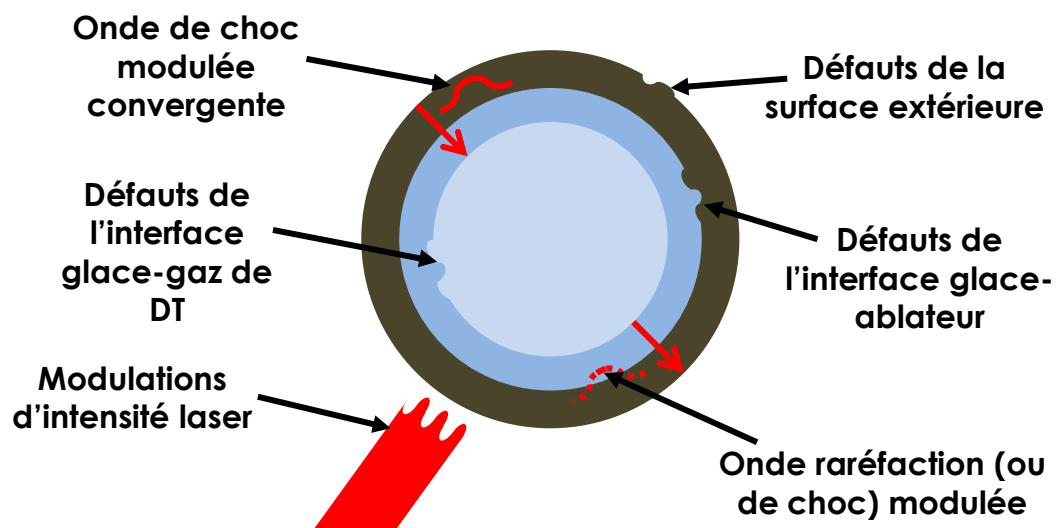


FIGURE 1.3 – Schéma représentant les origines possibles des modulations amplifiées par les instabilités hydrodynamiques en FCI.

tés hydrodynamiques au front d'ablation, particulièrement sur les perturbations initiales que causent l'ensemencement de l'IRT ablative.

1.1.3 Effets délétères des instabilités hydrodynamiques dans les dernières expériences de FCI

Pour tenter d'atteindre la fusion contrôlée, une campagne en AI appelée "National Ignition Campaign" (NIC) a été réalisée de 2009 à 2012 sur le NIF. Cette installation est composée de 192 faisceaux pour une énergie laser totale approchant les 2 MJ [5]. La réf. [6] est une revue des 3 ans de cette campagne NIC, dont l'objectif de réaliser l'ignition n'a pas été atteint. Les auteurs présentent les différentes expériences réalisées pour tenter d'obtenir du gain et pour comprendre les phénomènes qui entrent en jeu lors de l'implosion de la cible. Plusieurs pages (p. 49-54) évoquent le rôle délétère de l'IRT dans la performance des implosions. La figure 1.4 présente un graphe extrait de cet article ; la production neutronique y est représentée en fonction de la masse de matière mélangée dans le coeur comprimé de la cible pour différents tirs cryogéniques de la campagne NIC réalisés avec une énergie supérieure à 1,3 MJ. Les réactions de fusion D-T produisent des neutrons de 14,1 MeV dont le libre parcours est supérieur aux dimensions du plasma et qui peuvent être détectés : la production neutronique est donc un indicateur majeur de la performance d'une implosion. La quantité de mélange correspond à la masse d'ablateur présent dans le coeur de la cible comprimée. Ce mélange est mesuré grâce à l'émission X des matériaux de l'ablateur provenant du coeur de la cible [7]. Comme on peut le voir sur la figure 1.4, une quantité d'ablateur de l'ordre du μg présente au niveau du point chaud (lieu de démarrage des réactions de fusion dans le coeur) suffit à réduire la production de neutrons de presque un ordre de grandeur. Ce mélange est attribué à l'IRT : la croissance exponentielle des modulations des interfaces peut provoquer l'injection de matière provenant de l'ablateur dans le coeur de la cible. Cette matière, froide comparée au gaz comprimé du point chaud, provoquerait son refroidissement, réduisant ainsi les réactions de fusion. Ce refroidissement a aussi une origine radiative car les matériaux injectés ont un numéro atomique plus élevé et donc produisent une émission plus intense. Un autre point à noter dans cet article est que de gros écarts apparaissent entre les prédictions des simulations et les résultats expérimentaux. On peut d'ailleurs noter qu'à cause de prévisions numériques trop optimistes, l'ignition, finalement non réalisée, avait été prévue lors de la campagne NIC. Les auteurs insistent donc sur la nécessité de réduire l'écart entre simulations et expériences pour pouvoir prédire de manière fiable les performances des implosions et pour une conception optimale des cibles.

Les auteurs de la réf. [6] montrent que l'IRT est une cause importante de réduction des performances des implosions en AI. Dans la réf. [8], les auteurs abordent le cas de l'AD. Ils montrent que les performances des implosions sont meilleures quand on augmente l'intensité du pied de l'impulsion laser. Ceci car l'ablation - le chauffage de la surface de la cible par le laser et sa transformation en plasma - plus intense de l'ablateur provoque une stabilisation plus importante de l'IRT. Les conséquences de l'IRT sont détaillées : lors

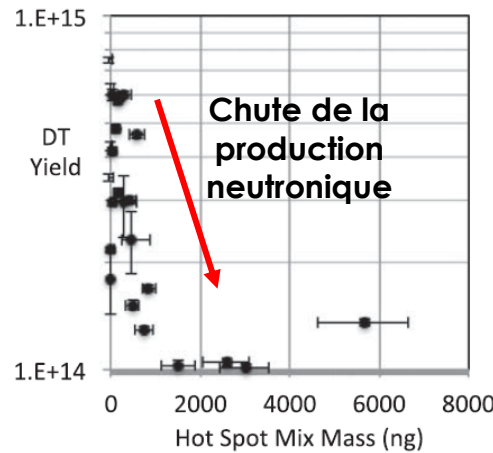


FIGURE 1.4 – Production neutronique en fonction de la quantité de mélange présent dans le coeur de la cible pour différents tirs des campagnes NIC. Graphe extrait de la réf. [6].

de la phase de compression, la croissance des modulations de la coquille solide de la cible peut provoquer sa rupture. Ainsi, du plasma d'ablation chaud et à haute pression peut se détendre dans la direction du coeur de la cible, où la pression est plus basse, augmentant ainsi sa pression et causant une difficulté de compression. Or une compression efficace est fondamentale dans les schémas d'ignition. A la stagnation, la croissance des modulations de l'interface gaz chaud (coeur)-glace froide va en augmenter la surface et donc les pertes radiatives à travers cette dernière. Le coeur va se refroidir, introduisant une diminution du taux des réactions de fusion. Une fois encore, la majeure partie des schémas présentés dans l'article [8] sont issus de simulations. On peut donc voir combien il est important d'optimiser la capacité prédictive des simulations des expériences.

1.2 Compréhension de l'ensemencement de l'IRT ablative

L'IRT qui se développe au front d'ablation est appelée IRT ablative, car ce front est continuellement ablaté par le laser ou le flux X, ce qui provoque une stabilisation partielle de cette instabilité. L'IRT ablative ne commence à se développer qu'au début de l'accélération de la coquille, ce qui nécessite un temps de l'ordre de quelques nanosecondes dans les expériences d'implosions cryogéniques. Au début de l'illumination par le laser ou par les rayons X, un choc est lancé dans la cible. Quand ce choc débouche en sortie de coquille, une onde de raréfaction est générée et remonte vers le front d'ablation. Quand l'onde de raréfaction atteint le front d'ablation, celui-ci commence à accélérer. Entre le début de l'impulsion laser et de l'accélération, l'IRT ablative ne peut pas se développer.

Cependant, la propagation d'un choc en aval du front d'ablation crée des conditions favorables pour le développement d'un phénomène appelé instabilité de Richtmyer-Meshkov (IRM) ablative. Dans le cas de l'IRM classique, les modulations d'une interface croissent linéairement après qu'un choc a traversé l'interface perturbée. Dans le cas du front d'ablation, comme pour l'IRT ablative, un effet de stabilisation apparaît. Dans la réf. [9], V. N. Goncharov et ses collaborateurs ont développé un modèle théorique de l'IRM ablative, qui considère les irrégularités issues de la rugosité de la cible. Du fait de l'ablation, les modulations du front d'ablation vont osciller en s'amortissant, effet appelé inversion de phase. Les auteurs de la réf. [10] ont réalisé des expériences d'IRM ablative avec le laser Nike KFr en géométrie plane. Ils ont mesuré l'inversion de phase (le passage des modulations par une amplitude égale à 0, caractéristique d'oscillations) de modulations monomodes usinées à la surface des cibles. Un bon accord a été trouvé avec le modèle de Goncharov ainsi qu'avec les simulations numériques sur la période des oscillations et la position des maxima d'oscillation. Certains désaccords apparaissent néanmoins avec les simulations en terme d'amplitude des pics d'oscillation notamment.

Le modèle et les expériences présentées concernent l'IRM ablative pour des modulations issues de la rugosité de la cible. Or en AD, une source importante de perturbations du front d'ablation provient de l'empreinte d'inhomogénéités de l'intensité laser. La théorie de l'IRM ablative imprimée par laser est développée dans la réf. [11] : dans un premier temps, les modulations du front d'ablation naissent et croissent du fait de l'empreinte des défauts de l'intensité laser, puis ces modulations se découplent du laser du fait de la conduction thermique dans le plasma et oscillent sous l'effet de l'IRM ablative. Il n'existe pas dans la littérature de description d'expériences d'IRM ablative imprimée à partir de défauts de l'intensité laser ; ainsi, ce modèle et les simulations n'ont jamais été comparés à des données expérimentales. Or l'IRM ablative joue un rôle fondamental : c'est elle qui fixe l'amplitude des modulations du front d'ablation au début de l'accélération ; elle ensemence l'IRT ablative. Les auteurs des refs. [12, 13] ont réalisé des expériences d'IRM ablative en AI sur des cibles comportant des bosses et comparent leurs résultats à des simulations. Ils montrent que le paramétrage des simulations est nécessaire pour obtenir un bon accord avec les données expérimentales ; dans leur cas, l'équation d'état semble jouer un rôle fondamental. Cependant ils expliquent aussi l'importance d'être capable de simuler et modéliser correctement l'évolution de l'IRM ablative : une conception d'expérience reposant sur des prédictions fiables permettrait d'utiliser les oscillations des modulations pour atteindre une amplitude minimale au moment du passage à l'IRT ablative, et donc d'en réduire l'effet.

Ce point est détaillé dans un article récent [14]. En effet, du fait des oscillations de l'IRM ablative, il existe une longueur d'onde dite "de croissance nulle", c'est-à-dire pour

laquelle les défauts seront en train de s'inverser et seront donc d'amplitude nulle au moment de début de l'accélération. Lors des implosions NIF, plusieurs chocs sont lancés. Les auteurs montrent grâce aux modèles de l'IRM ablative et à des simulations qu'en jouant sur la force et la chronométrie des chocs, il est possible de choisir quel mode aura une croissance nulle. On peut par exemple choisir de fixer une croissance nulle pour la longueur d'onde la plus fréquente dans les défauts de rugosité des capsules, ou pour les modes se situant au pic de croissance de l'IRT ablative. Un autre point mis en avant par les auteurs est que les modulations du front d'ablation jouent un rôle fondamental dans le développement des instabilités hydrodynamiques dans toute la capsule, car si le premier choc lance la phase Richtmyer-Meshkov ablative, les suivants transmettent les défauts vers le coeur de la capsule. Ces différentes propositions permettent de bien de sentir l'importance de l'IRM ablative et de sa compréhension détaillée.

En résumé, l'IRM ablative est décrite théoriquement, que ce soit pour des modulations provenant de la rugosité de la cible ou de l'empreinte d'inhomogénéités de l'intensité laser. Dans le 1er cas, de nombreuses expériences ont été réalisées et montrent la difficulté et la nécessité de prédire correctement l'évolution de cette instabilité. Il n'existe cependant pas d'expériences d'IRM ablative imprimée par laser. C'est ce qui justifie une partie de notre étude, qui consiste à étudier expérimentalement les conditions initiales de l'IRT ablative en AD. Nous avons donc réalisé des expériences pour mesurer l'IRM ablative imprimée par laser. Puis nous avons comparé les résultats obtenus à des simulations réalisées avec le code d'hydrodynamique radiative CHIC et au modèle de Goncharov.

1.3 Amélioration des méthodes de contrôle des instabilités hydrodynamiques

Même une connaissance parfaite du comportement des instabilités hydrodynamiques ablatives ne garantirait pas la disparition de leur effet délétère sur la performance des implosions cryogéniques. Un axe fondamental reste la recherche de la réduction de l'IRT, que ce soit par une diminution de l'amplitude initiale des perturbations des interfaces ou par la réduction de la croissance de l'IRT. De très nombreuses méthodes sont détaillées dans la littérature. L'utilisation, en AI, d'ablateurs laminés, c'est-à-dire alternant des couches de plastique dopé et non dopé au germanium (Ge), permet la stabilisation de la croissance de l'IRT par augmentation de la diffusion thermique transverse [15, 16]. En AD, des expériences utilisant des ablateurs dopés avec des matériaux de Z élevé ont aussi montré une réduction de la croissance de l'IRT ablative [17] par la création d'un double front d'ablation [18]. Les auteurs de la réf. [19] présentent quant à eux une méthode de lissage des défauts de l'intensité laser, particulièrement utile en AD. L'utilisation de mousses re-

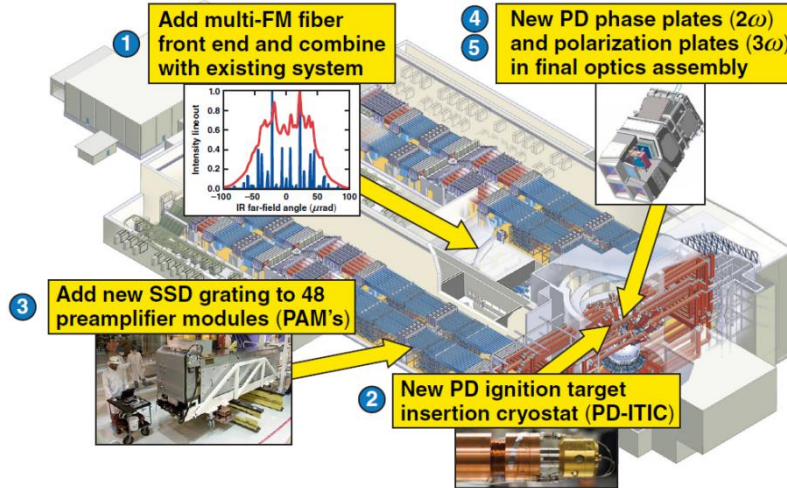


FIGURE 1.5 – Schéma des principales modifications à apporter au NIF pour réaliser des tirs en AD. Ce schéma est extrait d’une présentation de T. C. Sangster (LLE) faite au ”Management Advisory Committee Meeting” du LLNL en juin 2014.

couvrant l’ablateur a aussi été étudiée pour réduire l’empreinte du laser. Dans la réf. [20], les auteurs présentent une expérience réalisée sur le laser OMEGA [21] lors de laquelle des cibles de plastique recouvertes par une mousse de 30 mg/cm^3 sont illuminées par des faisceaux laser portant de forts défauts d’intensité. Cette mousse est considérée comme sur-dense vis-à-vis de la longueur d’onde du laser, c’est à dire que le laser est absorbé entièrement à l’entrée de la mousse et ne peut se propager dedans même quand la mousse sera complètement ionisée. Ainsi, pour obtenir un plasma qui réalisera le lissage des défauts laser par conduction thermique, la mousse est préalablement ionisée par un flash de rayons X. Les auteurs montrent que l’amplitude des modulations du front d’ablation est plus faible dans le cas de cibles recouvertes de mousse que dans le cas de cibles nues. Cependant, cette méthode comporte un gros inconvénient : le flash de rayons X provoque le préchauffage de la feuille de plastique. Or dans le cadre de la FCI, le préchauffage est problématique car il augmente l’entropie et réduit l’efficacité de compression de la cible et donc la performance des implosions.

Une alternative résiderait dans l’utilisation de mousses sous-denses, c’est-à-dire dans lesquelles le laser va se propager tout en ionisant la mousse lui-même. Une expérience a été réalisée dans cette optique sur la Ligne d’Intégration Laser (LIL) par S. Depierreux et ses collaborateurs [22]. Les auteurs montrent que les faisceaux laser sont lissés lors de la traversée de la mousse par des phénomènes d’interaction laser-plasma appelés instabilités paramétriques. L’avantage de cette méthode est qu’elle ne nécessite pas de flash de rayons X et ne risque donc pas d’entraîner de préchauffage. Cependant, pour pouvoir mesurer les défauts de l’intensité laser, les mousses n’étaient pas pourvues de feuilles de plastique

en face arrière. La réduction des défauts de l'intensité laser a donc été mesurée mais pas la conséquence de cette réduction sur l'empreinte. Ainsi, l'autre partie de notre étude expérimentale des conditions initiales de l'IRT ablative en AD consiste à étudier l'effet de mousses sous-denses sur l'empreinte et la croissance de l'IRT ablative.

Il faut noter que cette recherche de réduction des instabilités hydrodynamiques est plus d'actualité que jamais. En effet, l'échec de la campagne NIC dans l'atteinte de l'ignition a relancé le développement d'une plate-forme d'attaque directe sur le NIF dans le cadre de la Polar Ignition Campaign [23] menée par les chercheurs du LLE. Des tirs réguliers en AD sont ainsi réalisés sur le NIF à la fréquence d'une douzaine par an. Les derniers en date (novembre 2014) ont été dévolus à des mesures de la croissance de l'IRT en géométrie sphérique. Néanmoins, une campagne conséquente d'expériences en AD nécessite des modifications importantes du NIF. La figure 1.5 présente les principales modifications à apporter au NIF pour optimiser l'installation pour l'AD. On peut remarquer que les points 1, 3, 4 et 5 concernent la mise en forme des faisceaux. Plus spécifiquement, l'ajout des modules de SSD (3) et de multi-FM (1) (ces techniques seront présentées dans la section 2.4 du chapitre 2) est proposée en vue de lisser les défauts d'intensité des faisceaux laser, qui sont une source importante d'instabilités hydrodynamiques en AD, contrairement à l'AI. Cependant, de telles modifications sur une installation de la taille du NIF sont très coûteuses ; de plus, il n'est pas dit qu'elles permettent une réduction suffisante de l'empreinte des non-uniformités de l'intensité laser. Ainsi, tout schéma physique permettant la réduction de l'empreinte peut présenter un intérêt dans la mise en place de tirs en AD sur le NIF.

1.4 Objectifs de thèse

Dans cette thèse, nous nous proposons d'étudier les conditions initiales de l'IRT ablative en AD. Cette étude est réalisée par une approche expérimentale, à travers des expériences sur le laser OMEGA. Les expériences sont comparées et interprétées à l'aide de simulations CHIC. Deux objectifs ont été recherchés dans cette thèse :

- Mesurer l'IRM ablative imprimée par laser pour la première fois, et interpréter ces mesures à l'aide du modèle de Goncharov et des simulations CHIC.
- Evaluer la capacité de mousses sous-denses à réduire l'IRT ablative en diminuant l'empreinte des défauts de l'intensité du laser.

Au cours du chapitre 2, nous allons détailler la physique du front d'ablation : absorption du laser, transport de la chaleur dans la zone de conduction et génération d'un choc. Nous allons ensuite présenter l'IRT et l'IRM ablatives, qui sont étudiées dans cette thèse, et aborder certaines méthodes de réduction de ces instabilités. Le chapitre 3 présentera le laser OMEGA qui fut le vecteur des expériences décrites dans ce manuscrit. Nous

expliquerons ensuite le fonctionnement des diagnostics (appareils de mesure) que nous avons utilisés. L'installation et les diagnostics induisent des contraintes d'exploitation et de conception des expériences ; ces contraintes seront abordées. Grâce aux diagnostics, nous pouvons effectuer des mesures telles que des radiographies de face des modulations du front d'ablation. Il a donc fallu développer des programmes pour dépouiller et analyser ces données ; nous les présenterons.

Le chapitre 4 présente notre démarche expérimentale de la mesure de l'IRM ablative imprimée par laser. Pour cela, nous avons défini en collaboration avec D. Martinez et V. Smalyuk du LLNL et I. Igumenshev du LLE une configuration expérimentale permettant de mesurer ce phénomène d'amplitude réduite. Nous avons notamment dû définir des conditions d'empreinte et de radiographie particulières, la phase Richtmyer-Meshkov induisant des amplitudes de modulations proche du niveau de bruit et donc du seuil de mesure. Une fois les mesures effectuées, nous avons utilisé les programmes de dépouillement présentés dans le chapitre précédent pour analyser les données et obtenir des courbes de croissance. Le dernier point a consisté à déterminer le plus précisément possible le niveau d'empreinte utilisé dans les expériences, paramètre essentiel pour pouvoir effectuer le travail d'interprétation des instabilités.

Dans le chapitre 5, nous présentons l'analyse de l'expérience décrite précédemment. Nous modélisons l'IRM ablative imprimée par laser à partir du code CHIC et du modèle de Goncharov. Une fois les simulations réalisées, nous les avons comparées aux données expérimentales. Le modèle de Goncharov a été utilisé pour interpréter les différents résultats et a été comparé aux simulations et aux données expérimentales. Enfin, nous avons étudié l'effet des variations de différents paramètres comme le niveau d'empreinte ou l'équation d'état sur les simulations.

Le chapitre 6 présente des expériences dont j'ai assuré la conception, tant numérique qu'expérimentale, en me basant sur la plate-forme développée avec nos collaborateurs du LLNL pour la mesure de l'IRM ablative imprimée par laser. Nous présentons les mesures d'instabilités hydrodynamiques ablatives sur des feuilles de CH recouvertes de mousses sous-denses, utilisées pour réduire le niveau d'empreinte. Afin d'effectuer plusieurs mesures différentes de manière concomitante, la configuration expérimentale a d'abord été optimisée. Puis nous avons dépouillé les données issues de caméra à balayage de fente pour déterminer la dynamique des cibles. Ensuite, nous avons comparé les mesures d'énergie rétrodiffusée et des simulations CHIC afin de déterminer la durée d'ionisation des mousses. Enfin, nous avons traité les radiographies de face des modulations du front d'ablation pour évaluer l'effet lissant des mousses. Le chapitre 7 permet de conclure sur l'ensemble des résultats obtenus pendant cette thèse et d'aborder de nouvelles expériences

pour poursuivre et compléter le travail réalisé.

Chapitre 2

Physique des instabilités hydrodynamiques au front d’ablation

Les expériences décrites ultérieurement portent sur les instabilités hydrodynamiques au front d’ablation en attaque directe. Ce chapitre pose les bases théoriques nécessaires à leur compréhension. La physique du front d’ablation est ainsi présentée : ablation de matière de la cible par interaction laser-matière, formation du plasma en expansion et de la zone de conduction, génération d’un choc dans la cible et accélération subséquente du front d’ablation après le retour de l’onde de raréfaction. Ces différents phénomènes interviennent directement dans le développement des instabilités hydrodynamiques. Les instabilités de Rayleigh-Taylor (IRT) et de Richtmyer-Meshkov (IRM) ablatives seront ensuite abordées par la présentation des principaux modèles et la description des phénomènes physiques, particulièrement ceux liés aux effets d’ablation. Les expériences réalisées au cours de cette thèse l’ont été en géométrie plane (interaction du laser sur des cibles planes) : les modèles décrits dans ce chapitre seront eux aussi présentés dans cette géométrie. Les phénomènes physiques sont de toute manière semblables en géométrie sphérique (propre aux expériences d’implosion), aux effets de convergence type Bell-Plesset près.

2.1 Physique du front d’ablation

Lors des expériences de FCI, un flux lumineux intense (généralement entre 10^{12} W/cm² et 10^{16} W/cm²) irradie une cible solide. Plus spécifiquement, en attaque directe, un ou plusieurs faisceaux laser sont focalisés sur la cible créant ainsi un plasma d’ablation. Plusieurs zones, représentées sur la figure 2.1, correspondant à différents phénomènes sont mises en place. Le rayonnement laser est absorbé dans le plasma d’ablation près de la densité critique, l’énergie est ensuite transportée jusqu’au front d’ablation par conduction

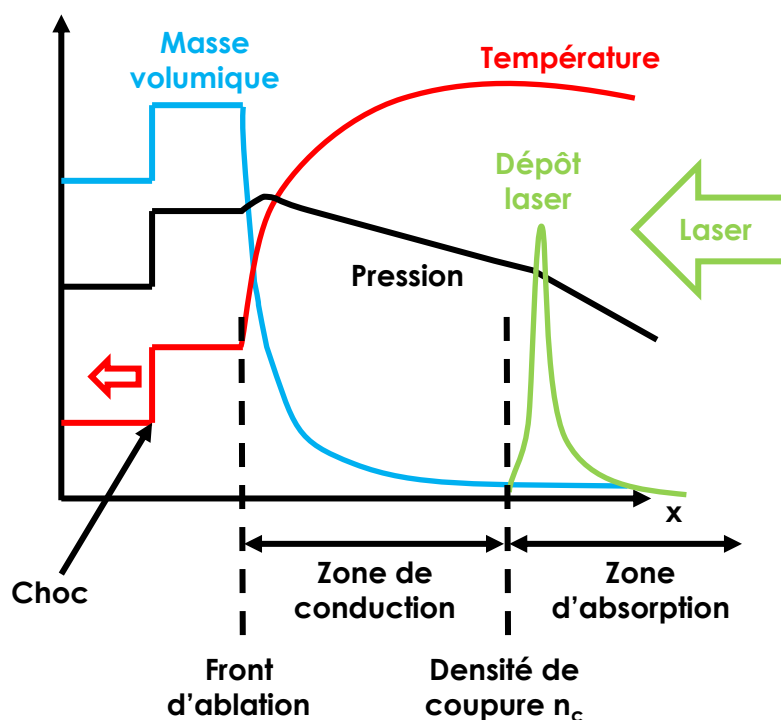


FIGURE 2.1 – Représentation schématique de la dynamique de l'interaction laser-plasma.

thermique dans la zone de conduction. La réaction à la détente du plasma crée un choc qui se propage dans la cible. Ces différents phénomènes sont détaillés dans les sous-parties suivantes.

2.1.1 Absorption de l'énergie du laser

Au moment où le laser commence à irradier la cible, il n'existe pas encore de plasma d'ablation. Pour un laser au verre de néodyme triplé en fréquence (technologie prépondérante des lasers de puissance actuels NIF, LMJ, OMEGA, etc, ...), la longueur d'onde est $\lambda_L = 351$ nm. L'énergie d'un photon du laser est d'environ 3,5 eV. Pour des cibles composées de carbone et d'hydrogène telles qu'utilisées dans nos expériences, l'énergie d'ionisation minimale est de 5-7 eV. La création des premiers électrons libres du plasma ne peut se faire que lorsqu'un électron de la bande de valence absorbe deux ou plusieurs photons laser simultanément. Ceci nécessite une intensité laser supérieure à 10^{12} W/cm².

Une fois les prémices du plasma créées, les électrons du plasma oscillent sous l'effet du champ électromagnétique du laser. Ils entrent alors en collision avec les ions et les noyaux des atomes et chauffent le plasma. Ce phénomène d'absorption d'énergie lumineuse est appelé absorption collisionnelle ou Bremsstrahlung inverse. Le coefficient d'absorption

collisionnelle est donné par

$$K = \frac{\nu_{ce}}{c} \left(\frac{n_e}{n_c} \right)^2 / \sqrt{1 - \frac{n_e}{n_c}} \quad (2.1)$$

avec c la vitesse de la lumière, ν_{ce} la fréquence de collision des électrons à la densité critique, n_e la densité électronique et n_c la densité critique. L'efficacité d'absorption croît avec la densité électronique du plasma jusqu'à la densité critique

$$n_c = \frac{\pi m_e c^2}{e^2 \lambda_L^2} = \frac{1,1 \cdot 10^{21}}{\lambda_L^2} \text{ cm}^{-3} \quad (2.2)$$

au voisinage de laquelle l'absorption est maximale. On considère généralement que l'absorption par Bremsstrahlung inverse a lieu entre $n_c/4$ et n_c . Elle est négligeable en deçà de $n_c/4$ car le coefficient d'absorption varie comme le carré de n_e . Le mouvement des électrons du plasma et leur réponse à une stimulation extérieure par le champ électromagnétique d'un laser sont caractérisés par la fréquence plasma électronique

$$\omega_{pe} = \sqrt{\frac{4\pi n_e e^2}{m_e}}, \quad (2.3)$$

laquelle croît avec la densité électronique : pour une densité électronique trop faible, les électrons ne peuvent suivre les oscillations du champ laser d'où l'absence d'absorption collisionnelle. Plus on se rapproche de la densité critique, plus les électrons du plasma arrivent à suivre le champ laser donc plus l'efficacité d'absorption croît. A la surface critique, la fréquence plasma électronique et celle du laser sont égales ; les électrons du plasma vont alors compenser exactement le champ du laser. Ainsi au delà de la densité critique, le plasma va avoir les propriétés d'un métal : le laser est réfléchi au niveau de la surface critique et il ne peut se propager dans le plasma dont la densité est supérieure à n_c (plasma surdense).

En supposant que le rayonnement laser se propage de $+\infty$ dans la direction opposée de l'axe x le long du gradient de densité, la fraction totale d'absorption collisionnelle est définie par

$$A = 1 - \exp\left(-2 \int_{x_c}^{\infty} K(x) dx\right) \quad (2.4)$$

avec x_c la position de la surface critique. Le facteur 2 vient du fait que le laser parcourt deux fois le plasma, avant et après sa réflexion. En faisant l'hypothèse d'un profil de densité du plasma exponentiel $n_e(x) = n_c \exp(-x/L)$ avec L la longueur caractéristique du gradient de densité électronique, on peut trouver que la fraction totale d'absorption collisionnelle suit l'équation suivante [3] :

$$A_L = 1 - \exp\left(-\frac{8\nu_{ec}L}{3c}\right) \quad (2.5)$$

avec ν_{ec} la fréquence de collision électron-ion à la densité critique définie par

$$\nu_{ec} = \frac{4\sqrt{2\pi}}{3} \frac{n_c Z_i e^4 \ln \Lambda_e}{\sqrt{m_e} (k_B T_e)^{3/2}} \quad (2.6)$$

avec e la charge élémentaire d'un électron, $\ln \Lambda_e$ le logarithme coulombien à la densité critique, m_e la masse d'un électron, k_B la constante de Boltzmann, T_e la température électronique et Z_i l'état d'ionisation, c'est-à-dire la charge moyenne des ions du plasma. En ne gardant que les paramètres du plasma, on peut écrire [3]

$$A_L = 1 - \exp\left(-0,007 \frac{Z_i L}{\lambda_L^2 T_e^{3/2}}\right) \quad (2.7)$$

avec L et λ_L en μm et T_e en keV. Pour un plasma en expansion, la longueur de gradient augmente au cours du temps et l'absorption collisionnelle devient très rapidement majoritaire. Par exemple, pour un plasma de CH_2 complètement ionisé ($Z_i = 3, 5$), dont la longueur $L=9 \mu\text{m}$ et la température $T_e=300 \text{ eV}$ ont été extraites d'une simulation CHIC (le code d'hydrodynamique radiative du CELIA, cf chapitre 3) après 100 ps d'irradiation laser, on trouve $A_L = 99,5 \%$.

Il faut cependant noter que l'on s'est placé dans le cas d'une incidence normale : prendre en compte un angle θ d'incidence par rapport au gradient de densité rajoute un facteur $\cos^3 \theta$ pour la fraction totale d'absorption collisionnelle. Cela s'explique par le fait que l'indice de réfraction du plasma dépend de la densité comme

$$n = \sqrt{1 - \frac{\omega_{pe}^2}{\omega_L^2}} \quad (2.8)$$

Ainsi le faisceau laser traverse un milieu dont l'indice de réfraction diminue ; cela provoque la croissance de l'angle du faisceau avec la normale au fur et à mesure de sa propagation dans le plasma car la partie du vecteur $\vec{k}$ parallèle au gradient de densité est affectée par la réfraction. La réflexion aura lieu avant d'atteindre la densité critique, à la densité électronique $n_e = n_c \cos^2 \theta$. Plus l'angle d'incidence du laser sera grand, plus sa réflexion aura lieu loin de la surface critique donc moins la fraction totale d'absorption collisionnelle sera importante. Pour des angles d'incidence de 23° et 48° (les deux géométries utilisées lors de nos expériences), le coefficient d'absorption A_L est diminué de 20% et 70% respectivement.

Un autre type d'absorption est couramment rencontré : l'absorption résonante. Elle correspond à l'excitation d'oscillations des électrons du plasma (appelées onde de Langmuir) à la densité critique par le champ laser en incidence oblique et de polarisation dans le plan d'incidence. Bien que le faisceau laser n'atteint pas la surface critique, le couplage entre la surface de réflexion et la surface critique se fait par effet tunnel. Pour que la distance de l'amortissement du champ laser évanescant ne soit pas trop importante, le

gradient de densité doit être assez fort. D'après [24], la fraction d'absorption résonante est non négligeable pour

$$0 < (kL)^{2/3} \sin^2 \theta < 2,5 \quad (2.9)$$

avec $k = 2\pi/\lambda_L$. En se plaçant dans le même cas que dans le paragraphe précédent, pour des angles d'incidence de 23° et 48° , on trouve respectivement des valeurs du paramètre $(kL)^{2/3} \sin^2 \theta$ d'environ 11 et 40. L'absorption résonante est donc considérée comme négligeable dans le cadre de nos expériences; elle prend de l'importance pour des expériences avec impulsions courtes (de l'ordre de quelques ps) où le gradient de densité au voisinage de la densité critique reste très fort tout au long de l'interaction laser-matière.

D'autres phénomènes de couplage entre le laser et le plasma sont dus à des effets non-linéaires; c'est le cas des instabilités paramétriques. Ces instabilités se développent dans le plasma sous-dense pour des intensités supérieures à quelques 10^{14} W/cm² voire 10^{15} W/cm² pour une longueur d'onde laser de 351 nm. Les instabilités de type Brillouin et Raman stimulées provoquent la diffusion d'une partie du faisceau laser par couplage résonant avec respectivement une onde sonore et une onde électronique du plasma. L'instabilité de filamentation provoque, elle, des phénomènes d'autofocalisation du faisceau laser, le fragmentant en plusieurs filaments plus intenses, du fait de la variation de l'indice de réfraction du plasma avec l'intensité du faisceau laser incident. Certains effets de ces instabilités paramétriques seront envisagés lors de l'analyse des expériences présentée dans le chapitre 6, lors desquelles une intensité supérieure à quelques 10^{14} W/cm² a été utilisée. Dans ce cas, les instabilités paramétriques provoquent un élargissement angulaire du faisceau laser du fait de la filamentation et une diffusion stimulée de Brillouin vers l'avant.

2.1.2 Zone de conduction

L'énergie déposée par le laser chauffe la matière. Il en résulte une onde thermique qui se déplace du plasma vers la cible. L'extrémité, ou pied, de cette onde thermique est appelée front d'ablation : c'est la surface où la matière de la cible est ablatée, c'est-à-dire là où elle commence à être chauffée par conduction électronique et transformée en plasma. La zone située entre le front d'ablation et la surface critique est appelée zone de conduction. Dans cette zone, le transfert de l'énergie déposée par le laser dans le plasma avant la densité critique se fait par conduction électronique, c'est-à-dire par diffusion des électrons chauffés du plasma vers la cible dense et froide.

L'ablation est caractérisée par un taux de masse ablatée $\dot{m}_a$ qui s'exprime en g.cm⁻².s⁻¹ et qui correspond à la masse de matière ablatée par unité de temps et de surface ($\dot{m}_a$ peut ainsi être assimilée à une "intensité" d'ablation, par analogie avec l'intensité d'irradiation qui est la quantité d'énergie lumineuse par unité de temps et de surface). On en tire la

vitesse d'ablation, c'est-à-dire la vitesse de pénétration du front d'ablation dans la cible,

$$V_a = \frac{\dot{m}_a}{\rho_s} \quad (2.10)$$

avec ρ_s la densité de la cible avant le front d'ablation. On voit que pour un taux de masse ablatée constant, le laser "creusera" moins vite la cible si celle-ci est plus dense. La matière qui s'expand du front d'ablation sous forme de plasma pousse par "effet fusée" le front d'ablation et donc la cible. Cet effet est caractérisé par une pression d'ablation P_a . Dans les cas d'une ablation stationnaire en AD, des relations basées sur la conservation de l'énergie et de l'impulsion permettent de retrouver ces quantités sous la forme de lois d'échelle [3] :

$$\dot{m}_a = \left(\frac{\rho_c}{2}\right)^{2/3} I^{1/3} \quad (2.11)$$

$$P_a = \left(\frac{\rho_c}{2}\right)^{1/3} I^{2/3} \quad (2.12)$$

avec $\rho_c = A m_n n_c / Z_i$ la densité à la surface critique où A est le nombre de masse et $m_n \approx 1,7 \cdot 10^{-21}$ mg est la masse d'un nucléon. D'après l'équation (2.2) :

$$\rho_c = \frac{1,8A}{Z\lambda_L^2} \text{ mg/cm}^3 \quad (2.13)$$

En injectant (2.13) dans (2.11) et (2.12), on obtient

$$\dot{m}_a = 0,93 \left(\frac{A}{Z}\right)^{2/3} \left(\frac{I}{\lambda_L^4}\right)^{1/3} \text{ g.s}^{-1}.\text{m}^{-2} \quad (2.14)$$

$$P_a = 0,97 \left(\frac{A}{Z}\right)^{1/3} \left(\frac{I}{\lambda_L}\right)^{2/3} \text{ Pa} \quad (2.15)$$

avec λ_L en μm . On peut alors évaluer ces grandeurs importantes dans la physique du front d'ablation - et par conséquent dans celle des instabilités hydrodynamiques ablatives - à partir du type de matériau de la cible et de l'intensité et de la longueur d'onde du laser. Un autre point à noter est qu'au début de l'interaction entre le laser et la cible, la surface critique s'éloigne progressivement du front d'ablation jusqu'à obtenir un régime d'ablation stationnaire. Au début de l'interaction, la taille de la zone de conduction D_c croît au cours du temps, de manière linéaire d'après [11] $D_c = V_c t$. Au bout d'un temps allant de quelques centaines de ps à quelques ns, un régime stationnaire s'installe et D_c reste constante, d'une taille de l'ordre du diamètre de la tache focale du laser en géométrie plane et de l'ordre du rayon de la cible en géométrie sphérique. Dans la zone de conduction, le transfert de chaleur se fait par diffusion thermique et selon la loi de Fourier $Q = -\kappa \nabla T$. Le coefficient de conductivité thermique κ dans un plasma est donné par les calculs de Spitzer-Härm [25]. Cependant, pour les intensités laser élevées, ce flux thermique possède

une limite : c'est le flux d'énergie transporté par tous les électrons se propageant dans la direction du gradient de température à leur vitesse thermique. Ce type de cas limite se présente quand la longueur caractéristique de variation de la température atteint l'ordre de grandeur du libre parcours moyen des électrons. Cependant, la comparaison entre les expériences et les simulations a montré que la valeur du flux limite se situait bien en dessous de ce flux limite théorique. Les codes hydrodynamiques de FCI utilisent un limiteur de flux pour tenir compte de cet effet, généralement égal à 3 à 10 % du flux limite maximal. Les effets cinétiques conduisant à la modification de la fonction de distribution des électrons sont responsables de cette limitation.

Outre la conduction électronique, l'énergie est transportée dans la cible par le rayonnement X. Cependant, dans le cas de matériaux de Z peu élevé (comme C et H lors de nos expériences), la conduction radiative est peu importante ; elle peut conduire à un léger préchauffage de la matière (cf réf. [26] partie 2 p. 109). On notera qu'en attaque indirecte où un plasma d'or est créé dans le hohlraum, le transport radiatif est d'une importance fondamentale.

2.1.3 Génération et propagation d'un choc

Du fait de la pression d'ablation, la matière non chauffée située juste après le front d'ablation est poussée contre les couches voisines non perturbées, et ainsi de suite : un choc est mis en place, si la vitesse d'ablation est et reste subsonique. La matière choquée est mise en mouvement à une vitesse constante. Les relations de Rankine-Hugoniot développées à partir des équations de conservation au niveau du front de choc définissent les grandeurs caractéristiques de la matière post-choc : vitesse, densité, pression. Pour une cible plane, lorsque le choc atteint la face arrière du matériau, il débouche dans le vide ; la face arrière ne rencontre plus de résistance et se détend dans le vide en accélérant. Sa densité chute, et les couches précédentes peuvent elles aussi se détendre en accélérant : une onde de raréfaction remonte de la face arrière vers le front d'ablation. Lorsqu'elle atteint le front d'ablation, celui-ci se propage plus vite ; la cible sera accélérée tant qu'elle sera illuminée par le laser. En considérant les relations de Rankine-Hugoniot dans l'hypothèse d'un gaz parfait mono-atomique et d'un choc fort (pression avant le choc négligeable devant la pression post-choc), on obtient une vitesse de choc u définie par (cf réf. [26] partie 2 p. 137)

$$u = 2\sqrt{\frac{P_a}{3\rho_0}} \quad (2.16)$$

avec ρ_0 la densité du matériau avant le choc. La vitesse de choc augmente avec la pression d'ablation qui augmente elle même avec l'intensité laser à la puissance $2/3$. On peut donc

voir qu'une augmentation de l'intensité d'un facteur 8 n'augmentera ainsi la vitesse du choc que d'un facteur 2.

2.2 L'IRT au front d'ablation

Il existe une diversité notable d'instabilités hydrodynamiques : l'IRT, l'IRM, mais aussi l'instabilité de Plateau-Rayleigh, qui provoque la séparation d'un jet de liquide en gouttelettes, l'instabilité de Kelvin-Helmoltz, provoquée par le cisaillement de deux fluides, etc, ... La section qui suit présente l'IRT et l'IRM classiques ainsi que l'IRT et l'IRM ablatives.

2.2.1 L'IRT classique

Pour comprendre l'origine de l'IRT [27, 28], envisageons deux fluides non miscibles de densités différentes, séparés par une interface plane, et placés dans un champ de gravitation comme celui de la pesanteur terrestre. La gravité équivaut à une accélération dans le sens opposé (dans un ascenseur qui monte et donc accélère vers le haut, on a l'impression de peser plus lourd, donc qu'une composante de gravité supplémentaire dirigée vers le bas apparaît). On prend souvent l'exemple de l'eau et de l'huile : si l'huile -moins dense- se trouve au dessus de l'eau, le système est stable, et toute perturbation de l'interface -par exemple des oscillations parce qu'on secoue le récipient- sera atténuée et finalement annulée. En revanche, si l'eau est placée sur l'huile, le système est instable car son énergie potentielle est plus importante que dans le cas précédent. Cependant, si l'interface (cas idéal impossible à reproduire) est parfaitement plane, le système est en équilibre : l'eau ne peut pas pénétrer l'huile à un endroit ou un autre de l'interface. Cette situation équivaut à celle d'une sphère immobile placée en haut d'un pic : la moindre perturbation fera dégringoler la sphère en bas. Ainsi toute perturbation de l'interface sera amplifiée par l'IRT, le système minimisant ainsi son énergie potentielle en échangeant le fluide lourd avec le fluide léger ; l'huile pénétrera dans l'eau et vice-versa, jusqu'à ce que l'eau soit au fond et l'huile au dessus. On retrouvera alors la configuration stable précédemment évoquée. Ce cas peut être généralisé à n'importe quels fluides accélérés lorsque l'accélération et le gradient de densité vont dans la même direction dans le référentiel du laboratoire. On peut aussi interpréter ce phénomène par la plus grande inertie du fluide le plus lourd qui va donc opposer plus de résistance à l'accélération et donc se retrouver "au fond" tandis que le fluide léger se retrouvera "devant", comme dans une centrifugeuse.

L'IRT passe par différentes phases. Prenons une perturbation sinusoïdale de l'interface d'amplitude notée η très inférieure à la longueur d'onde λ . Tant que $\eta < 0,1\lambda$, on se trouve en phase linéaire [29] comme représenté en figure 2.2 (a). En supposant des fluides

incompressibles, on obtient alors une croissance exponentielle de la perturbation de type

$$\eta(x, y, t) = \eta(x, y, 0) \exp \sigma t \quad (2.17)$$

avec le taux de croissance σ défini par

$$\sigma = \sqrt{g A_t k} \quad (2.18)$$

où g est l'accélération, $k = 2\pi/\lambda$ le nombre d'onde et A_t le nombre d'Atwood tel que

$$A_t = \frac{\rho_H - \rho_L}{\rho_H + \rho_L} \quad (2.19)$$

où ρ_H et ρ_L sont les densités respectivement du matériau le plus dense et le moins dense. On voit ainsi que le taux de croissance augmente quand l'accélération croît ou que la longueur d'onde diminue. Les petites longueurs d'onde croissent donc plus vite que les grandes et atteignent donc plus vite des stades non-linéaires. Le nombre d'Atwood quantifie le fait que plus la différence de densité entre les deux fluides est importante, plus le taux de croissance est important. On peut aussi remarquer qu'une perturbation arbitraire de petite amplitude pourra être traitée par décomposition de Fourier comme une somme de perturbations sinusoïdales croissant à des vitesses différentes.

Quand l'amplitude atteint quelques dixièmes de λ , l'IRT entre en phase non-linéaire (cf figure 2.2 (b)). La perturbation perd alors la symétrie de son profil sinusoïdal : la pénétration du fluide dense dans le fluide léger se produit plus vite et prend la forme d'aiguilles tandis que la pénétration du fluide léger dans le fluide lourd se ralentit et prend la forme de bulles. En phase fortement non-linéaire représentée en figure 2.2 (c), quand l'amplitude des perturbations dépasse la longueur d'onde, deux phénomènes supplémentaires apparaissent. Du fait de l'interpénétration des bulles et des aiguilles, une composante transverse de vitesse apparaît à l'interface au niveau des bords des aiguilles. Une instabilité hydrodynamique de type Kelvin-Helmholtz [30, 31] se développe alors. Ce type d'instabilité crée des vagues à l'interface entre deux fluides lors qu'il y a cisaillement, c'est-à-dire que les composantes tangentielles des vitesses de deux fluides à l'interface sont différentes. C'est typiquement ce qui arrive lorsque le vent souffle à la surface de l'eau : des vagues se forment car l'eau est immobile tandis que l'air se déplace à une certaine vitesse parallèlement à l'interface. Cette instabilité provoque la formation de structures sur les bords des aiguilles qui prennent ainsi une allure de champignon. L'autre phénomène qui apparaît en phase fortement non-linéaire est appelé compétition et mélange de bulles. Les bulles vont fusionner entre elles en produisant des bulles plus grosses. Ce phénomène se répète, éliminant les plus petites bulles et créant des bulles de plus en plus grosses. Si l'IRT a suffisamment de temps pour se développer, on obtiendra à la fin une seule grosse bulle de la largeur du système. Le fluide léger peut alors finir par remonter au-dessus du fluide lourd et le système atteint la configuration stable.

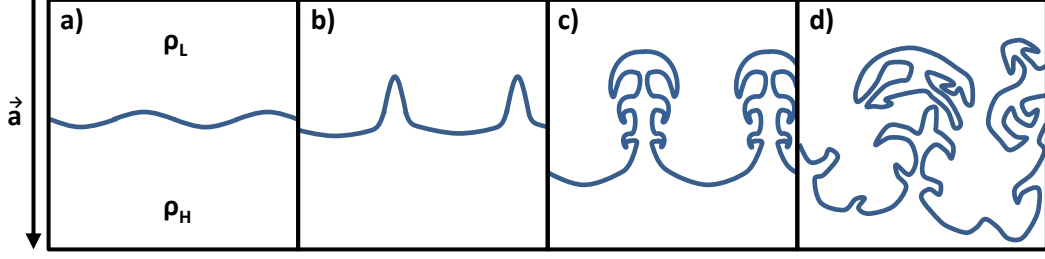


FIGURE 2.2 – Evolution d’une perturbation sous l’effet de l’IRT à l’interface entre deux fluides de densité (léger de densité ρ_L et lourd de densité ρ_H) a) en phase linéaire, b) en phase faiblement non-linéaire, c) en phase fortement non-linéaire et d) lors de la turbulence.

Cependant, à ce stade, si l’épaisseur du fluide lourd est suffisamment grande, un autre phénomène apparaît : l’IRT entre dans une phase chaotique appelée phase de turbulence (cf figure 2.2 (d)), où les aiguilles se cassent sous forme de gouttes. L’instabilité de Kelvin-Helmholtz sur les bords des aiguilles entre en phase non-linéaire et provoque un mélange entre les deux fluides. L’écoulement est alors extrêmement complexe et ne peut plus être traité de manière complètement prédictive.

2.2.2 Effet de l’ablation sur l’IRT

Dans le cadre de la FCI, le front d’ablation sépare le matériau dense de la cible qui a subi le choc et le plasma en expansion, peu dense. On a pu voir dans les sections précédentes que le front d’ablation est accéléré par l’effet fusée : on se retrouve donc dans le cas d’une accélération et d’un gradient de densité de mêmes directions, propice au développement de l’IRT. L’IRT qui se développe au front d’ablation est appelée IRT ablative. L’IRT ablative a été largement étudiée théoriquement, en phase linéaire [32, 33, 34, 35] comme non linéaire [36, 37, 38]. L’approche linéaire commence de manière similaire dans ces différents articles, par la linéarisation des équations d’Euler,

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \vec{v}) = 0, \quad (2.20)$$

$$\rho \frac{\partial \vec{v}}{\partial t} + \rho (\vec{v} \cdot \nabla) \vec{v} = -\nabla P + \rho \vec{g}, \quad (2.21)$$

$$\frac{\rho(\epsilon + \vec{v}^2/2)}{\partial t} + \nabla \cdot [(\rho(\epsilon + \vec{v}^2/2) + P)\vec{v} + \vec{Q}_e] = \rho \vec{g} \cdot \vec{v}, \quad (2.22)$$

avec $\vec{v}$ la vitesse, ρ la densité, P la pression, ϵ l’énergie interne, $\vec{g}$ l’accélération dans le repère du front d’ablation et $\vec{Q}_e$ le flux de chaleur. Ce système est fermé par l’équation

d'état des gaz parfaits $P = (\gamma - 1)\rho\epsilon$ avec γ le rapport des chaleurs spécifiques.

Les hypothèses d'un écoulement fortement subsonique et quasi-isobare permettent de négliger les variations de pression par rapport aux gradients de température dans l'équation (2.22). En prenant la pression égale à la pression d'ablation P_a et en supposant un gaz monoatomique ($\gamma = 5/3$), l'équation (2.22) devient l'équation de l'énergie stationnaire

$$\nabla \cdot \left(\frac{5}{2} P_a \vec{v} + \vec{Q}_e \right) \approx 0. \quad (2.23)$$

On exprime le flux de chaleur par $\vec{Q}_e = -\kappa \nabla T_e$ avec κ la conductivité thermique.

On linéarise ensuite le système en utilisant un développement à l'ordre 1 de chaque quantité de l'écoulement, qui est considéré subir de petites perturbations par rapport à sa valeur moyenne. En général, le système perturbé est écrit dans le référentiel du front d'ablation sous la forme d'un système aux valeurs propres pour une perturbation sinusoïdale de nombre d'onde k . Cependant, la résolution analytique de ces équations reste difficile. La célèbre "formule de Takabe" [39] propose une formulation empirique du taux de croissance sous la forme

$$\gamma(k) = \alpha \sqrt{kg} - \beta k V_a, \quad (2.24)$$

obtenue par résolution numérique des équations (2.20), (2.21) et (2.22) linéarisées. Les auteurs de la réf. [39] donnent une évaluation des constantes $\alpha \approx 0,9$ et $\beta \approx 3 - 4$, qui dépendent des caractéristiques et du matériau de l'écoulement. On voit apparaître dans l'équation (2.25) une stabilisation du taux de croissance due à l'ablation, appelée stabilisation convective. Cependant, cette formule ne permet pas de décrire tous les mécanismes de stabilisation de l'IRT au front d'ablation. Différentes approches analytiques ou semi-analytiques, comme celles présentées dans les réfs. [32, 34, 40, 41, 42, 43], ont permis de mettre en évidence ces mécanismes. Dans ces différents modèles, le problème aux valeurs propres est résolu de part et d'autre du front d'ablation puis les solutions sont raccordées au niveau de la région de transition, qui correspond au front d'ablation. Ces modèles, qui permettent d'obtenir des formulations analytiques du taux de croissance, sont valables dans certaines conditions du régime d'accélération. Le régime d'accélération est caractérisé par son nombre de Froude $Fr = V_a^2/gL_0$, où L_0 est l'épaisseur du front d'ablation ; ce nombre quantifie l'influence relative du flux d'ablation par rapport à l'accélération et à la conductivité thermique (qui augmente avec L_0). Ainsi, les différents modèles sont développés pour des petites ou grandes valeurs du nombre de Froude, et dans la limite $kL_0 \gg 1$ ou $kL_0 \ll 1$.

Pour des grands nombres de Froude, les réfs. [41, 34] font apparaître des expressions du taux de croissance qui mettent en avant des phénomènes de stabilisation supplémentaires sous la racine carrée. Ces expressions sont résumées par la formule présentée dans la réf. [44] :

$$\gamma(k) = \sqrt{kg - k^2 V_a V_{bl} + 4k^2 V_a^2} - 2k V_a, \quad (2.25)$$

avec V_{bl} la vitesse dite de "blow-off", c'est-à-dire la vitesse d'expansion du plasma. Le terme $k^2 V_a V_{bl}$ est associé au phénomène de surpression dynamique et le terme $4k^2 V_a^2$ au phénomène dit de "fire polishing". Ces phénomènes seront détaillés ultérieurement (cf section 2.3.2). Des expressions plus complexes encore du taux de croissance existent, par exemple celle présentée dans la réf. [34], qui met en évidence le rôle de la conduction thermique latérale [44].

De plus, les travaux de R. Betti et V. N. Goncharov [35] exhibent une solution générale qui combine les solutions des réfs. [34, 41] pour les grands nombres de Froude et de la réf. [42] pour les petits nombres de Froude. Cette solution théorique est néanmoins complexe ; un ajustement de celle-ci permet d'obtenir la "formule de Takabe modifiée"

$$\sigma_a = \alpha \sqrt{\frac{gk}{1 + kL_m}} - \beta k V_a \quad (2.26)$$

avec L_m la longueur minimum de gradient de densité au front d'ablation. Pour un ablateur de CH, on a $\alpha \approx 0,98$ et $\beta \approx 1,6$. On peut noter deux points : tout d'abord, le nombre d'Atwood A_t classique est remplacé par un facteur qui dépend de la longueur caractéristique du front d'ablation. Plus L_m est grand, plus le taux de croissance est petit : le fait que le changement de densité à l'interface ne soit plus brusque mais progressif amortit la croissance de l'IRT ablative. Cet effet dit de stabilisation de l'IRT est plus important pour les longueurs d'ondes plus petites que L_m . D'autre part, on peut aussi remarquer, dans les différentes formulations du taux de croissance de l'IRT ablative présentées, qu'à partir d'une longueur d'onde λ_c appelée longueur d'onde de coupure, l'IRT ablative est stabilisée.

La comparaison entre le taux de croissance de l'IRT classique et de l'IRT ablative issue de l'équation (2.26) est présentée sur la figure 2.3. Les différents paramètres ont été extraits d'une simulation CHIC à l'intensité laser 5.10^{13} W/cm^2 comme évoqué dans la partie précédente. La cible est en CH_2 , d'épaisseur $30 \mu\text{m}$ et g , V_a et L_m sont calculés à 1 ns : on trouve $g = 2,8.10^{16} \text{ cm.s}^{-2}$, $V_a = 5.10^5 \text{ cm.s}^{-1}$ et $L_m = 0,88 \mu\text{m}$. On peut voir que les perturbations de longueur d'onde inférieure à $3,5 \mu\text{m}$ sont stabilisées par l'ablation et ne croissent donc pas. En revanche, les perturbations de longueur d'onde aux alentours de $13 \mu\text{m}$ sont celles dont le taux de croissance est maximal et sont donc les plus instables.

2.2.3 L'IRT ablative en phase non-linéaire

La phase non-linéaire de l'IRT ablative joue un rôle important dans les expériences d'instabilités hydrodynamiques en FCI : la croissance exponentielle de l'IRT ablative en phase linéaire est souvent suivi par son passage rapide en phase non-linéaire. Ainsi, une partie des données mesurées l'est souvent en phase non-linéaire. L'approche la plus simple pour envisager la non-linéarité de l'IRT est la suivante : un mode sinusoïdal de longueur

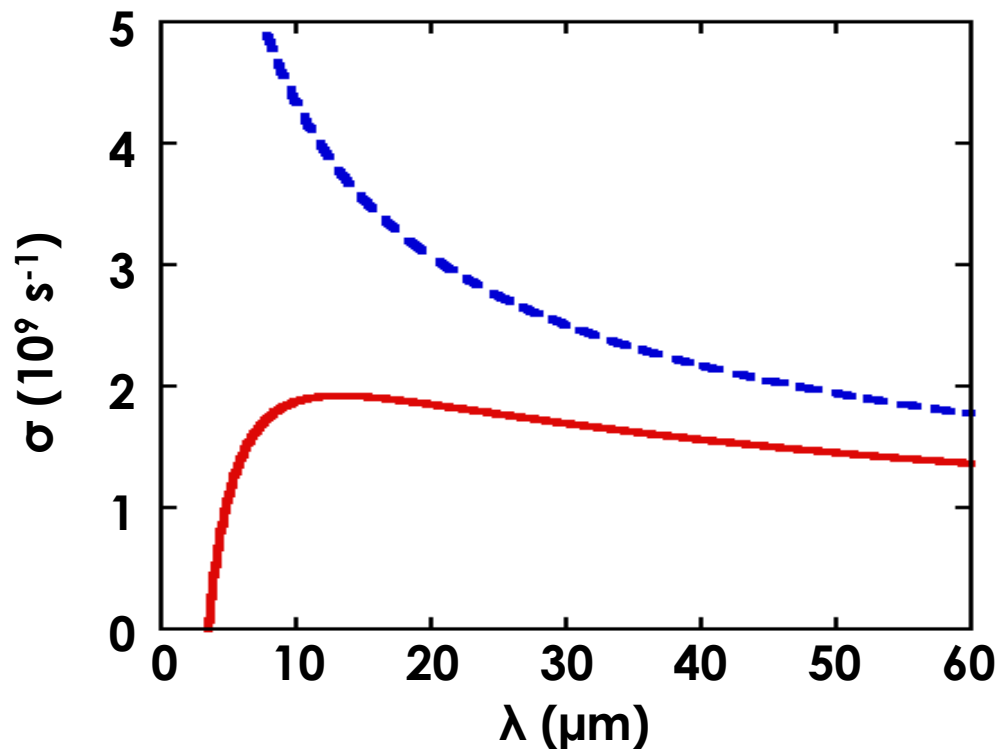


FIGURE 2.3 – Relations de dispersion de l'IRT classique (courbe bleue pointillée) et ablative (courbe rouge continue) calculées à 1 ns à partir de paramètres extraits d'une simulation CHIC à l'intensité laser $5 \cdot 10^{13} \text{ W/cm}^2$ et à la longueur d'onde laser 351 nm pour une cible de CH_2 de $30 \mu\text{m}$.

d'onde λ va croître exponentiellement et indépendamment des autres modes jusqu'à atteindre une amplitude de saturation de l'ordre de $\eta = S(\lambda) = 0,1\lambda$ [45]. Une transition pour la vitesse de bulle va alors avoir lieu : celle-ci va passer d'exponentielle à constante. De plus, comme expliqué précédemment, le mode passera de sa forme sinusoïdale initiale à des bulles et des aiguilles sous l'effet de la croissance des harmoniques de la longueur d'onde initiale. Enfin, le mode ne sera alors plus indépendant des autres modes et des phénomènes de couplage pourront intervenir. Cependant, le modèle développé par Haan [46] dépeint un tableau plus complexe pour la saturation des modes, du fait que plusieurs modes croissent en même temps. Le critère précédemment évoqué donne une saturation pour une amplitude d'un mode de l'ordre de $S = 0,1\lambda$; ce que sous-tend ce critère est que le passage à la non-linéarité s'effectue quand la pente de l'interface est de l'ordre de 10-20 %, autrement dit quand la déformation locale de l'interface n'est pas négligeable. Outre un mode unique dont l'amplitude atteint l'ordre de grandeur de la longueur d'onde, une déformation locale non négligeable de l'interface peut intervenir du fait d'une superposition de plusieurs modes d'un spectre.

Ce cas est illustré par un exemple dans [46] que l'on reprend ici : on considère une interface perturbée par une modulation de longueur d'onde λ et d'amplitude 0.1λ . On peut construire cette interface par une somme de modes divers du spectre proches du mode k (tout mode k' tel que $k - k' < \epsilon k$ est considéré comme participant à cette somme ; expérimentalement, ϵ pourrait être défini comme la résolution minimale des mesures). On aura donc construit localement un mode qui se comporte comme le mode unique k d'amplitude 0.1λ ; en s'éloignant de cette zone, on pourra discriminer les différents modes participant à la somme du fait de leur déphasage. On voit que l'amplitude de chacun des modes du spectre est bien inférieure à 0.1λ et pourtant ces modes sont saturés. D'après [46], le niveau de saturation individuel des modes du spectre est alors

$$S(k) = \frac{0.2\pi\sqrt{\pi}}{L\epsilon k^2} \quad (2.27)$$

avec L la taille du système, c'est-à-dire la taille de la zone d'analyse dans le cadre du dépouillement d'une expérience. Le fait que la taille de la zone d'analyse influe sur le niveau de saturation des modes ne semble pas intuitif, mais pourtant, plus cette zone est grande, plus nombreux sont les modes qui peuvent participer à la somme donc plus l'importance relative d'un mode donné dans la somme est faible. De la même manière, on voit que plus ϵ est grand, plus le niveau de saturation individuel d'un mode est petit car plus nombreux sont les modes qui participent à la somme. De plus, les modes de plus petite longueur d'onde saturent plus vite. Une fois qu'un mode a atteint son niveau de saturation individuel $S(k)$, il croît linéairement à la vitesse $V_s(k) = S(k)\sigma_a(k)$ [47]. Les expériences de la réf. [48] en AD ont montré un bon accord entre le modèle de Haan et l'expérience. De plus, des modèles plus récents enrichissent la théorie de l'IRT ablative

en phase non-linéaire en exhibant des comportement très spécifiques du front d'ablation pour des nombres d'onde élevés [36, 37, 38, 49].

En phase non-linéaire, la fusion de bulles apparaît. Les travaux théoriques des réfs. [50, 51, 52] ont permis de développer un modèle (valable en régime classique et ablatif, pour A_t proche de 1, ce qui est le cas au front d'ablation) pour décrire ce phénomène. Pour une bulle de taille λ soumise à une accélération g , la vitesse asymptotique de bulle est définie par

$$v_b = \sqrt{\frac{g\lambda}{6\pi}} \quad (2.28)$$

Les bulles de plus grande taille atteignent donc des vitesses supérieures. Ainsi, pour un front de bulles de diverses dimensions, les plus grosses bulles vont finir par dépasser la vitesse de leurs voisines plus petites ; le front finira par aller plus vite que ces petites bulles qui vont être emportées vers l'intérieur du fluide et ainsi disparaître du front. Les grosses bulles vont alors s'étendre dans la place laissée vacante par les petites bulles, créant des bulles encore plus grosses de diamètre égal à la somme des diamètres de la petite bulle et de la grosse bulle qui ont ainsi fusionné. Ce phénomène a deux conséquences : une augmentation de la taille moyenne des bulles et une accélération du front de bulles. Concernant l'augmentation de la taille moyenne des bulles, on va se retrouver finalement dans une configuration où une bulle beaucoup plus grosse que le reste de la distribution, qualifiée d'incontrôlée, va dominer le processus de fusion de bulles et s'étendre jusqu'à atteindre la largeur du système. Quant à l'amplitude du front de bulles, sa croissance accélérée est définie par

$$h_b = \alpha g t^2 \quad (2.29)$$

avec α une constante allant de 0,04 à 0,08 [53] (estimée à 0.04 dans la réf. [54]) dépendante de A_t . Il faut cependant noter ici que des simulations récentes semblent indiquer qu'en régime ablatif et en 3D, la vitesse des bulles ne sature pas (contrairement au cas classique), du fait d'une accumulation de vorticit  dans la t te de la bulle, caus e par l'ablation.

Un point important de ce mod le est que le taux de fusion des bulles (exprim  en s^{-1}) n'est fonction que d'un param tre : le rapport de la taille des deux bulles concern es. Ce taux de fusion est une fonction croissante de ce rapport. En effet, plus la diff rence de taille entre deux bulles est importante, plus la vitesse de la grosse bulle est sup rieure   celle de la petite donc plus la grosse bulle va consommer la petite rapidement. En revanche, deux bulles de m me taille vont   la m me vitesse, aucune ne prend donc le pas sur l'autre et elles ne peuvent ainsi pas fusionner. Le fait que le taux de fusion ne d pend que du rapport des tailles des deux bulles implique une  volution auto-semblable de la distribution de bulles

$$f(\lambda/\langle\lambda\rangle) = \exp[-(\lambda/\langle\lambda\rangle - 1)^2/2C_1^2]/(\sqrt{2\pi}C_1) \quad (2.30)$$

avec C_1 une constante et $\lambda = 2\sqrt{\pi/A}$ la taille de la bulle avec A l'aire de la bulle. Ainsi, au long du phénomène de coalescence de bulles, la taille moyenne des bulles croît mais la distribution et l'amplitude des bulles normalisées à la taille moyenne restent constantes. D'après réf. [55], la taille moyenne des bulles, l'amplitude moyenne des bulles et leur nombre évoluent respectivement selon

$$\langle \lambda \rangle(t) = \frac{\varpi^2 a t^2}{16} + \frac{\varpi \sqrt{a} C t}{4} + \frac{C^2}{4} \quad (2.31)$$

$$\langle h \rangle(t) = \langle h \rangle(0) + \alpha_h a t^2 \quad (2.32)$$

$$N(t) = \frac{D}{(\varpi \sqrt{a} t + 2C)^4} \quad (2.33)$$

avec $C = 2\sqrt{\langle \lambda \rangle(0)}$, $D = N(0)(2C)^4$ et ϖ une constante. Ainsi, en connaissant les caractéristiques de la distribution initiale des bulles, on peut déterminer l'évolution à tout moment futur de cette distribution. On voit aussi que l'amplitude comme la taille des bulles évoluent en t^2 tandis que le nombre de bulles décroît en t^4 . Les expériences décrites en réf. [55] montrent de bons accords avec ce modèle de fusion de bulles. D'après des fits des résultats expérimentaux, les auteurs trouvent $C_1 = 0.23$ et $\varpi = 0.83$.

2.3 L'IRM ablative

2.3.1 Physique de l'IRM classique

Comme expliqué dans les chapitres précédents, lorsqu'un laser illumine une cible de FCI, qu'elle soit plane ou sphérique, un choc est lancé dans la matière à partir du front d'ablation. Le front d'ablation ne commencera ensuite à être accéléré qu'après que le choc a traversé le matériau et que l'onde de raréfaction est revenue au front d'ablation. A partir de cet instant, l'IRT ablative pourra se développer. Mais durant ce temps de transit, le front d'ablation est le siège du développement d'une autre instabilité : l'instabilité de Richtmyer-Meshkov (IRM) ablative. On dit ainsi généralement qu'au front d'ablation, l'IRM ablative ensemence l'IRT ablative.

L'IRM classique, décrite théoriquement par Richtmyer [56] et mesurée par Meshkov [57], apparaît lorsqu'une interface initialement perturbée subit le passage d'un choc. L'IRM peut être traitée théoriquement comme un cas particulier de l'IRT. En effet, un choc correspond entre autres à une discontinuité de vitesse d'écoulement. Par conséquent, lors du passage d'un choc, l'interface subira une accélération impulsionnelle $g(t) = \Delta u \delta(t)$ avec Δu la discontinuité de vitesse induite par le choc. Le modèle impulsionnel, utilisé par Richtmyer et repris dans [3], consiste à utiliser le cas particulier d'une accélération

impulsionnelle dans les équations hydrodynamiques. On trouve alors qu'une perturbation sinusoïdale de nombre d'onde k et d'amplitude η va croître linéairement en temps selon

$$\eta(x, y, t) = \eta(x, y, 0) + \eta(x, y, 0)\sigma_{RM}t \quad (2.34)$$

avec le taux de croissance

$$\sigma_{RM} = A_t k \Delta u \quad (2.35)$$

L'IRM classique fait donc croître les perturbations de l'interface à vitesse constante. Dans le cadre de l'IRT, l'accélération fournit continuellement de l'énergie pour la croissance des modulations tandis que dans le cas de l'IRM, l'accélération impulsionnelle lance la croissance des modulations (comme le ferait l'IRT) à une vitesse $\eta(x, y, 0)\sigma_{RM}$ mais le système évolue ensuite sans être perturbé par une accélération donc la vitesse reste constante. Cette analogie possède cependant ses limites ; les mécanismes de croissance de l'IRT et de l'IRM sont différents. Ainsi, contrairement à l'IRT, l'IRM va se développer quelle que soit la disposition des fluides.

2.3.2 L'IRM ablative

Depuis la fin des années 90, des modèles ont été développés pour décrire l'IRM ablative pour des perturbations issues de la rugosité de la cible [9, 58, 59] comme pour des modulations imprimées par les inhomogénéités du laser [11, 59, 60]. Comme nous le justifierons ultérieurement dans cette section, nous allons nous appuyer sur les modèles développés par Goncharov et ses collaborateurs [9, 11] pour modéliser l'IRM ablative. Des perturbations issues de la rugosité de la cible et issues de défauts d'intensité laser diffèrent car dans le premier cas, les perturbations de l'interface sont déjà présentes et l'IRM ablative provoque leur évolution tandis que dans le deuxième cas, une forme de compétition apparaît entre le phénomène d'"empreinte" des défauts du laser sur le front d'ablation et l'effet de l'IRM ablative. Mais dans tous les cas, l'évolution induite par l'IRM est similaire : les modulations du front d'ablation oscillent et finissent par s'amortir. Bien qu'au pic de l'oscillation, les modulations atteignent une amplitude supérieure à leur niveau initial, la dénomination instabilité est abusive pour l'IRM ablative car toute perturbation du front d'ablation est à terme amortie.

L'équation obtenue grâce au modèle de Goncharov pour l'amplitude η de la perturbation permet de bien appréhender la physique de l'IRM ablative [61] :

$$d_t^2 \eta + 4kV_a d_t \eta + k^2 V_{bl} V_a \eta = \frac{2}{5} k \frac{\delta I}{I} c_s^2 e^{-kD_c} + \frac{1}{\sqrt{5}} \frac{\delta I}{I} d_t (c_s e^{-kD_c}) \quad (2.36)$$

avec k le nombre d'onde de la modulation du front d'ablation et $\delta I/I$ la modulation relative de l'intensité laser. La partie à gauche de l'équation contient la physique de l'IRM ablative.

Pour l'interpréter, envisageons les modulations du front d'ablation, qui ont été mises en mouvement initialement par le passage du choc : les pics des modulations pénètrent dans le plasma chaud de la couronne tandis que les creux s'inscrivent plus profondément dans le matériau froid de la cible. Le front d'ablation étant isotherme (même température tout au long du front d'ablation), il en résulte que le gradient de température au front d'ablation est plus prononcé au niveau des pics qu'au niveau des creux, ce qui implique un flux de chaleur plus important. La vitesse d'expansion du plasma dépendant du flux de chaleur, elle est plus importante au niveau des pics, ce qui induit une force de réaction par effet fusée plus importante. Le membre $k^2 V_{bl} V_a \eta$ correspond à ce phénomène dit de "dynamic overpressure" (ou surpression dynamique) qui est comparable à la force de rappel d'un ressort. Ceci permet de définir une pulsation d'oscillation du front d'ablation

$$\omega_{RM} = k \sqrt{V_a V_{bl}} \quad (2.37)$$

Ainsi, plus la longueur d'onde d'une modulation du front d'ablation est importante, plus sa période d'oscillation sous l'effet de l'IRM ablative est petite. De plus, cette période diminue du fait d'une ablation plus forte et donc d'un éclaircissement laser plus important. On peut noter ici que dans la réf. [58], les auteurs utilisent une discontinuité au front d'ablation de type déflagration de Chapman-Jouguet [62] : la vitesse d'expansion du plasma derrière le front d'ablation est considérée comme étant la vitesse acoustique, et ce tout le long du front d'ablation. L'effet stabilisant de surpression dynamique n'est pas présent dans ce modèle.

Le membre $4k V_a d_t \eta$ est identifié comme correspondant à un phénomène dit de "fire polishing" par les auteurs des réfs. [11, 61] ; c'est un terme stabilisant qui amortit les oscillations. Dans la réf. [9], ce phénomène est associé au fait que le gradient de température (et donc la vitesse d'ablation) soit plus important(e) au niveau des pics qu'au niveau des creux de modulation du front d'ablation. Dans les réfs. [60, 59], l'effet de "fire polishing" n'est pas pris en compte dans le modèle du fait de conditions aux limites différentes.

Cependant, un terme n'a pas été pris en compte dans l'équation (2.36) : l'effet dit de la "vorticity convection" (ou convection de vortacité) [9]. Du fait que le choc soit lui aussi modulé, et donc que le front de choc ne soit pas perpendiculaire à sa direction de propagation, le front d'ablation progresse dans une matière dont la vitesse n'est pas uniformément orientée dans la direction de propagation du choc. Cet effet de vortacité va apporter une contribution à l'oscillation des modulations du front d'ablation. Comme la "vorticity convection" s'amortit plus lentement que la "dynamic overpressure", cet phénomène prendra au bout d'un certain temps le relais pour assurer les oscillations du front d'ablation [63].

Dans le cas de perturbations issues d'une rugosité de la cible, $\delta I/I = 0$ et la partie à droite de l'équation (2.36) s'annule. Les perturbations évoluent donc librement sous l'effet

de l'IRM ablative, ne dépendant que de l'amplitude initiale et de la longueur d'onde des perturbations. Dans le cadre d'inhomogénéités de l'intensité laser imprimées sur la cible, la partie à droite de l'équation (2.36) n'est pas nulle et dépend du temps. Le phénomène d'empreinte laser est donc dynamique et se trouve dans une forme de concurrence avec l'IRM ablative. On peut voir que le terme d'empreinte laser décroît selon $\exp(-kD_c)$. Ce terme tient compte de l'amortissement des perturbations dans la zone de conduction thermique. Comme l'énergie laser est déposée au niveau de la surface critique et que l'IRM ablative se développe au niveau du front d'ablation, plus la zone de conduction sera grande, plus la diffusion thermique latérale dans cette zone atténuera les modulations de l'intensité laser, moins l'empreinte sera efficace. Or, comme on l'a vu précédemment, au début de l'impulsion laser $D_c = V_c t$. On peut donc définir un temps de découplage des perturbations laser du front d'ablation [11]

$$t_D = \frac{1}{kV_c} \quad (2.38)$$

Avant ce temps, le phénomène d'empreinte domine et les modulations du front d'ablation de longueur d'onde supérieure à $2\pi/k$ croissent sous l'influence des perturbations de l'intensité laser. Après ce temps, les perturbations laser et le front d'ablation sont découplés et les modulations du front d'ablation évoluent du fait de l'IRM ablative, comme si l'éclairement laser était uniforme à cette longueur d'onde. Cet amortissement exponentiel des perturbations laser dans la zone de conduction est décrit par le modèle dit du "cloudy day" [64]. Dans la réf. [58], les auteurs tentent d'appliquer leur modèle développé pour des modulations issues de la rugosité de la cible à l'empreinte de défauts laser. Or, comme la zone de conduction n'est pas prise en compte dans leur modèle, les défauts laser ne sont jamais découplés du front d'ablation et les modulations du front d'ablation croissent sans saturer ni osciller.

Pour comparer à nos simulations et à nos mesures, nous utiliserons ultérieurement la solution pour l'amplitude d'une modulation de nombre d'onde k donnée par Goncharov et al [11] dans la limite $kc_s t \gg 1$ et pour $D_c = V_c t$

$$\eta(t) = \frac{2}{3\gamma} \frac{\delta I}{I} \left[\frac{c_s^2}{V_c^2 + V_a V_{bl}} e^{-kV_c t} + (A \cos \omega_{RM} t + B \sin \omega_{RM} t) e^{-2kV_a t} \right] + \eta_v(t) \quad (2.39)$$

où les constantes A et B et la fonction η_v sont définis en annexe A. Le terme η_v correspond au phénomène de "vorticity convection", le terme $e^{-2kV_a t}$ à l'amortissement dû au "fire polishing", le terme $A \cos \omega_{RM} t + B \sin \omega_{RM} t$ à la "dynamic overpressure" et le terme restant à l'empreinte laser. Pour un temps $t < t_D$ donné par (2.38), on obtient [11]

$$\eta(t) = \frac{2}{3\gamma} \frac{\delta I}{I} \left(c_s t + \frac{k c_s^2 t^2}{2} \right) \quad (2.40)$$

Cette croissance en t^2 correspond au fait que les perturbations de l'intensité laser créent et amplifient les modulations du front d'ablation avant leur découplage.

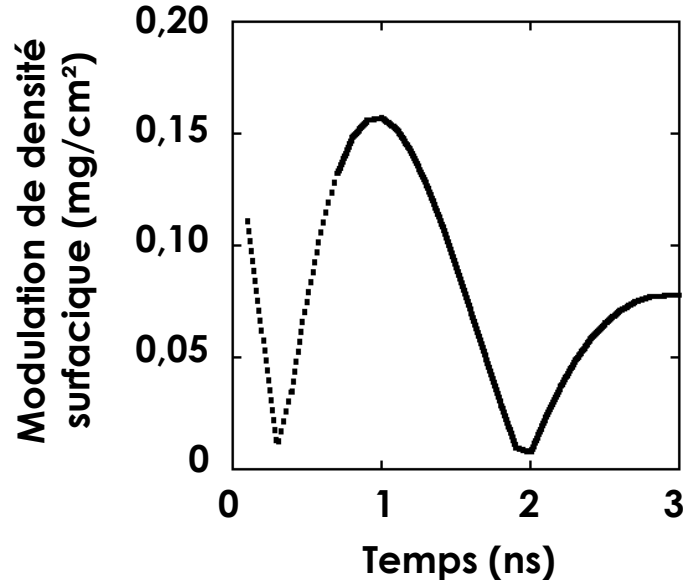


FIGURE 2.4 – Evolution temporelle de la valeur absolue de l'amplitude de modulations de densité surfacique de longueur d'onde $30 \mu\text{m}$ sous l'effet de l'IRM ablative. La zone représentée en pointillés correspond à la période où le modèle n'est pas encore valable car le temps ne satisfait pas $kc_s t \gg 1$.

Ce modèle permet de calculer l'évolution de l'IRM ablative à partir de paramètres accessibles dans les simulations 1D : V_a , V_{bl} , V_c , c_s et u . Un exemple de l'évolution de l'amplitude $\rho\eta(t)$ calculé à partir de ce modèle est représenté en figure 2.4. Des modulations de $30 \mu\text{m}$ sont imprimées à une intensité de $5 \cdot 10^{13} \text{ W/cm}^2$. On peut remarquer des inversions de phases ainsi qu'un amortissement de l'amplitude des modulations.

La figure 2.5 résume les points abordés dans ce chapitre concernant l'hydrodynamique d'une cible plane illuminée par un laser portant des modulations d'intensité de longueur d'onde λ . A $t=0$, le laser est allumé. L'ablation du matériau de la cible commence ; la pression d'ablation est modulée, proportionnellement aux modulations de l'intensité laser. Le front d'ablation se creuse donc plus au niveau des zones d'intensité supérieures. Une zone de conduction, dont la taille augmente au cours du temps, se forme dans le plasma d'ablation. Quand cette zone de conduction atteint une taille de l'ordre de la longueur d'onde, le front d'ablation est découplé des modulations de l'intensité laser à $t=t_D$. A partir de cet instant, le front d'ablation "voit" une illumination laser uniforme. Les modulations oscillent alors librement sous l'effet de l'IRM ablative. En parallèle, au début de l'illumination, un choc a été lancé dans la cible. Ce choc traverse la cible jusqu'à déboucher en face arrière. La face arrière commence donc à accélérer dans le vide et une onde de raréfaction parcourt la matière compressée de la face arrière vers le front d'ablation. Une fois l'onde de raréfaction passée, la matière de la cible commence elle aussi à

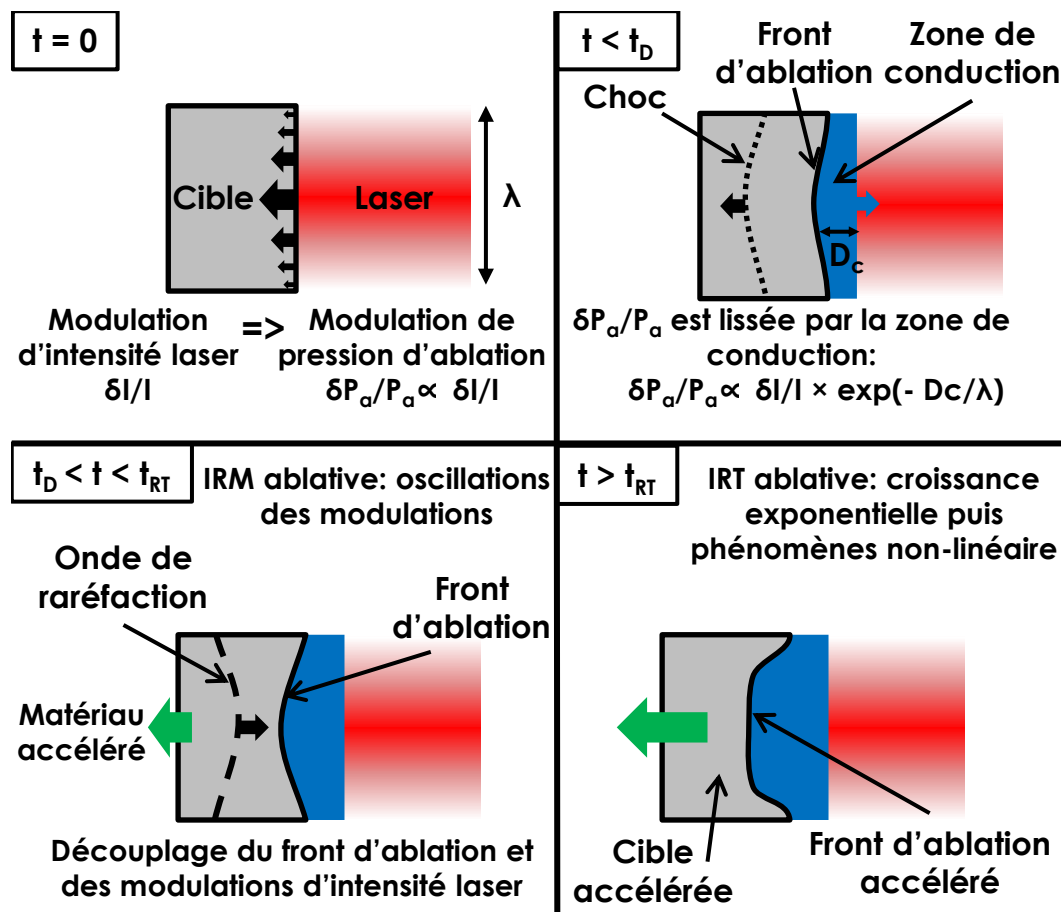


FIGURE 2.5 – Schéma résumant l'hydrodynamique d'une cible plane illuminée par un laser portant des modulations d'intensité de longueur d'onde λ . Les temps t_D et t_{RT} représentent respectivement le temps de découplage des modulations de l'intensité laser et du front d'ablation, et le temps de transition de l'IRM ablative à l'IRT ablative.

accélérer. Ainsi, quand l'onde de raréfaction atteint le front d'ablation, celui-ci subit une accélération, qui signe la transition de l'IRM ablative à l'IRT ablative. Les modulations du front d'ablation croissent alors exponentiellement, jusqu'à la phase non-linéaire où les différents phénomènes présentés dans la sous-section 2.2.3 se développent.

2.4 Bilan expérimental

Après ces rappels théoriques, nous allons présenter un bilan expérimental non exhaustif. Nous allons tout d'abord évoquer des expériences de mesure de l'IRM ablative, puis dans une deuxième partie, nous nous intéresserons aux méthodes de réduction des instabilités hydrodynamiques ablatives en AD, que ce soit par l'amélioration de l'uniformité des faisceaux laser, le lissage des perturbations d'intensité ou la diminution du taux de croissance de l'IRT ablative. Cette liste de méthodes de réductions des instabilités

hydrodynamiques n'est absolument pas exhaustive ; la réduction de l'IRT ablative est un domaine de recherche actuel et les publications à ce sujet sont régulières.

2.4.1 Expériences d'IRM ablative

Agglitskiy et al [10] ont effectué une série de mesures de l'IRM ablative sur le laser Nike KrF [65]. Les tirs ont été réalisés à une intensité de 5.10^{13} W/cm² par une impulsion laser carrée de 4 ns sur des cibles planes de CH où des modulations 2D (rayures) avaient été usinées. Les auteurs ont fait varier différents paramètres : l'amplitude crête-à-crête des modulations de 1 à 3 μm , l'épaisseur de cible de 40 à 100 μm et la longueur d'onde des modulations qui était de 30 ou 45 μm . La mesure de l'évolution des modulations a été réalisée au moyen d'une caméra à balayage de fente (ce diagnostic sera abordé dans le chapitre 3) couplée à une source de radiographie de Si. Un résultat majeur de cet article est que l'inversion de phase, c'est-à-dire le fait que les modulations du front d'ablation s'annulent pour croître ensuite avec un déphasage de π , a été mesuré, confirmant ainsi la théorie des oscillations dues à l'IRM ablative. Un autre point intéressant est montré en figure 2.6, qui représente l'évolution temporelle de l'amplitude des modulations pour une amplitude pic-vallée de 3 μm , une longueur d'onde de 45 μm et une épaisseur de cible de 65 μm . Les calculs théoriques, comme les simulations, prédisent un retour de l'onde de raréfaction au front d'ablation et donc un début d'accélération de celui-ci à 2,4 ns. Cependant, la croissance des modulations ne commence qu'à 3,1 ns. Ce décalage de 0,7 ns est interprété comme étant le temps nécessaire pour que l'IRT ablative prennent le pas sur le phénomène de "vorticity convection". Deux variations de paramètres montrent la validité des prédictions théoriques. Tout d'abord, pour une même longueur d'onde (45 μm) et une même amplitude initiale (3 μm), les oscillations causées par l'IRM ablative sont observées plus longtemps pour une cible plus épaisse (93 μm comparée à une de 65 μm) car pour une telle cible, le transit de choc ainsi que la remontée de l'onde de raréfaction dureront plus longtemps. L'évolution temporelle comparée de l'amplitude de modulations de 30 μm et 45 μm de même amplitude initiale (3 μm) montre que les modulations de plus grande longueur d'onde oscillent plus lentement, conformément à l'équation (2.37). Cependant, les auteurs montrent aussi un désaccord quantitatif inexpliqué entre leurs simulations et leurs résultats expérimentaux, avec des différences dans les périodes d'oscillation comme dans le temps de début de croissance de l'IRT ablative.

Gotchev et al [66] ont réalisé une expérience assez proche sur le laser OMEGA (cf chapitre 3). Des cibles de CH de 40 μm d'épaisseur portant des modulations de longueur d'onde 20 et 30 μm d'amplitude crête-à-crête 1,65 μm sont éclairées par 10 faisceaux laser créant une impulsion carrée de 1,5 ns et d'intensité $4,2.10^{14}$ W/cm². La mesure de l'évolution temporelle des modulations est aussi effectuée par une caméra à balayage de fente couplée à une cible de radiographie en U. Un ASBO [67], diagnostic similaire au VISAR [68] qui

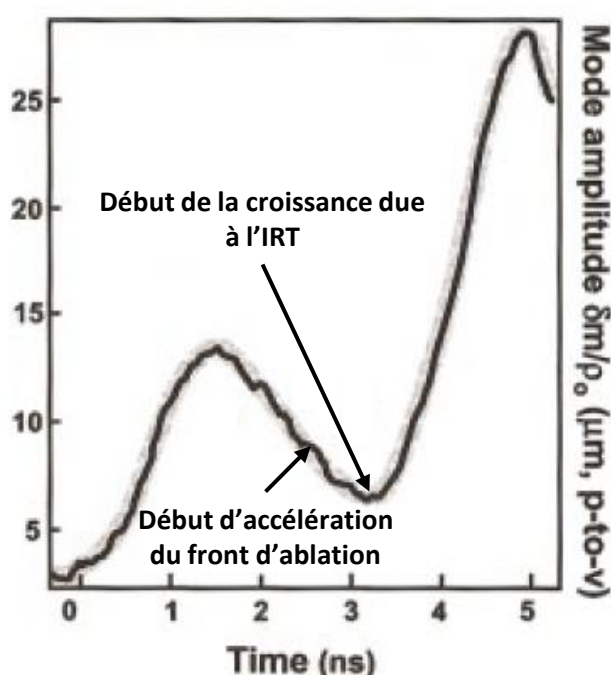


FIGURE 2.6 – Evolution temporelle de l’amplitude d’une modulation de longueur d’onde $45 \mu\text{m}$, d’amplitude initiale $3 \mu\text{m}$ pour une cible de $65 \mu\text{m}$ d’épaisseur d’après [10].

permet d’effectuer des mesures de dynamique des cibles par interférométrie, a été ajouté qui permet de mesurer le temps de débouché de choc en face arrière de la cible. Une inversion de phase est observée pour la longueur d’onde de $20 \mu\text{m}$. Cette inversion a lieu à $1,2 \text{ ns}$, bien plus tôt que celle montrée dans la référence précédente ($\approx 3,5 \text{ ns}$) car la longueur d’onde des modulations est plus petite et l’intensité plus grande. Les auteurs comparent leurs simulations et les résultats expérimentaux. Ils montrent que le limiteur de flux utilisé dans les simulations est un paramètre critique : un limiteur de flux plus petit implique une taille de la zone de conduction plus petite et donc un effet moindre de la stabilisation par la "dynamic overpressure". Un limiteur de flux trop petit peut même rendre négligeable cette stabilisation et la simulation donnera alors une croissance à vitesse constante des modulations comme dans le cas de l’IRM classique. Dans le cadre de la comparaison des simulations et des résultats expérimentaux, un limiteur de flux de $0,1$ donne un bon accord pour l’évolution de l’amplitude des modulations, ce qui n’est pas le cas des simulations avec limiteurs de flux de $0,04$ et $0,06$. Cependant, pour les temps de débouché de choc, les simulations avec limiteur de flux de $0,1$ sous-estiment les résultats expérimentaux tandis qu’un bon accord est trouvé avec $0,06$. Ceci explique les désaccords expérimentaux observés en réf. [10] : il n’y a pas de limiteur de flux constant qui permette de simuler parfaitement l’expérience. Les auteurs utilisent donc un modèle du transport de chaleur appelé non-local [63] qui leur permet grâce à un limiteur de flux variable en

temps de trouver une bonne cohérence avec l'ensemble de leurs résultats expérimentaux. Cependant, malgré ces différents exemples, nous n'avons pas trouvé dans la littérature d'étude spécifique de l'IRM ablative imprimée par laser.

2.4.2 Amélioration de l'uniformité de l'intensité du laser

Un des axes de recherche dans le domaine de la FCI en attaque directe est d'essayer d'obtenir des profils d'intensité des faisceaux laser les plus uniformes possible pour limiter l'empreinte des défauts. Une méthode pour ce faire est d'utiliser des "Random Phase Plates" (RPP) [69]. Ces lames de phase sont composées de nombreuses zones différentes qui divisent le faisceau laser en sous-faisceaux de phase variable. Ces sous-faisceaux se recouvrent sur la cible formant une enveloppe d'intensité lisse. Cependant, la superposition des sous-faisceaux crée des structures d'interférence qui s'impriment à la surface de la cible. Pour réduire l'effet de ces interférences, une méthode appelée "Smoothing by Spectral Dispersion" (SSD) a été développée et implémentée sur OMEGA par les auteurs de la réf. [19]. Le principe est d'irradier chaque zone de la RPP par une lumière de fréquence différente. Les auteurs montrent que la structure d'interférence formée par deux sous-faisceaux de fréquences f_1 et f_2 varie selon $(f_2 - f_1)t$. Le faisceau initial devra donc avoir une certaine largeur spectrale car plus la largeur spectrale est importante plus la variation des interférences peut être rapide. Cet élargissement spectral du faisceau est obtenu par un modulateur de fréquence tel un cristal électro-optique. Cependant, la longueur d'onde laser de 351 nm utilisée sur OMEGA est obtenue à partir d'une méthode de triplement en fréquence du spectre grâce à des cristaux de conversion. L'efficacité de ces cristaux de conversion est très sensible : chaque fréquence possède un angle optimal pour la conversion. Ainsi, chaque fréquence devra avoir un angle d'incidence sur le cristal de conversion différent. Cet accord est réalisé par l'utilisation d'un réseau de diffraction. Pour garder une bonne efficacité de conversion, la largeur de spectre est limitée à 0,2 nm. Cette méthode de SSD permet donc de faire varier temporellement les figures d'interférence des RPP, sur des temps inférieurs aux temps caractéristiques de l'hydrodynamique de la cible, ce qui empêche l'impression de ces figures sur la cible. Cependant, la fréquence des sous-faisceaux est aléatoire (dans la limite de la largeur du spectre). Ainsi, deux sous-faisceaux peuvent avoir la même fréquence : quelques structures d'interférence seront donc toujours présentes malgré le SSD.

Plus récemment ([70]), une méthode appelée "Multiple Frequency Modulators" (Multi-FM) a été développée. Elle permet de supprimer les structures d'interférence persistantes par l'utilisation de plusieurs modulateurs de fréquence. A largeur de spectre égale, le Multi-FM est plus efficace que le SSD pour lisser les défauts laser de petite longueur d'onde (cf réf. [71] p. 73-80). Il faut noter que cette technique de lissage a été mise en place sur

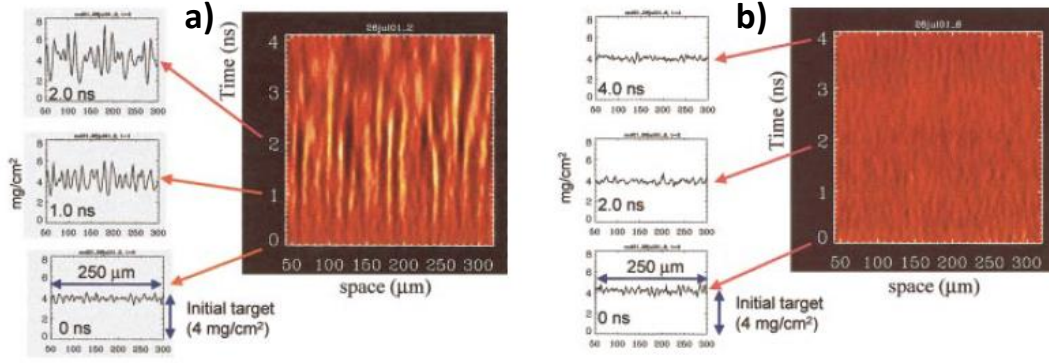


FIGURE 2.7 – Image mesurée par une caméra à balayage de fente pour a) une cible de CH seul et b) une cible de CH recouverte de 120 nm de Pd. Le temps $t=0$ correspond au début de l'impulsion principale. Cette figure est extraite de la réf. [72].

OMEGA EP. Cependant, son coût élevé est un frein à sa mise en place systématique.

2.4.3 Lissage des perturbations

Malgré l'amélioration de l'uniformité de l'intensité des faisceaux laser, des inhomogénéités subsistent toujours et impriment des défauts sur les cibles. Dans la réf. [72], les auteurs utilisent des cibles de CH recouvertes d'une fine couche de matériau de Z élevé. L'expérience est réalisée sur le laser Nike KrF. Les cibles de CH de 40 μm d'épaisseur sont recouvertes (sauf les cibles pour les mesures de référence) d'une couche d'épaisseur variant de la dizaine à la centaine de nm de Pd ou d'Au. 40 faisceaux réalisent l'irradiation des cibles; l'impulsion est constituée d'un pied de 4 ns d'intensité $2 - 3 \cdot 10^{12} \text{ W/cm}^2$ réalisé par un faisceau pour compresser la cible tandis que l'impulsion principale de 4 ns réalisée par les 39 autres faisceaux à $7 \cdot 10^{13} \text{ W/cm}^2$ accélère la cible. La rugosité des cibles est mesurée à moins de 5 nm tandis que les faisceaux sont lissés optiquement au maximum. L'évolution des modulations est mesurée par radiographie de face à l'aide d'une caméra à balayage de fente couplée à une source de radiographie en Si (émission à 1,86 keV). Le principe des radiographies de face et de côté est présenté en figure 2.8. Le résultat majeur de cet article est montré en figure 2.7 : les modulations imprimées par le laser sur la cible de CH seul croissent rapidement, à un niveau non négligeable dès la 1ère ns, tandis qu'aucune modulation ni croissance n'est observée pour la cible recouverte de 120 nm de Pd durant les 4 ns de l'impulsion principale. Un résultat similaire est observé pour une cible recouverte de 60 nm d'Au. Un autre point est que plus la couche de matériau est fine, plus l'effet de réduction des modulations diminue jusqu'à une épaisseur limite (10 nm pour l'or) où l'amplitude des modulations est même renforcée par la présence de la couche de matériau.

Pour interpréter ces résultats, les auteurs ont quantifié les différents effets de la couche de matériau de Z élevé. Tout d'abord, la présence de cette couche supplémentaire de matériau de la cible retarde le début de l'accélération jusqu'à 0,5 ns. Des mesures ont aussi été effectuées pour des cibles portant des modulations préimposées de longueur d'onde 30 μm et d'amplitude 0,125 μm , recouvertes par 42 nm d'Au. Un retard de croissance allant jusqu'à 1,2 ns est mesuré : ce retard est dû en partie à l'accélération décalée mais aussi au fait que l'émission X de la couche d'or augmente la longueur minimale de gradient de densité au front d'ablation, réduisant le taux de croissance de l'IRT ablative. Cependant, ces effets ne suffisent pas à expliquer qu'aucune croissance ne soit observée en figure 2.7 pour la cible recouverte de Pd. Les auteurs estiment que cet effet est dû à la suppression de l'empreinte laser : lors du pied d'impulsion, le laser chauffe la couche métallique qui émet des rayons X qui chauffent la surface du plastique, créant donc une zone de conduction. Lorsque l'impulsion principale atteint la cible, la zone de conduction est suffisamment large pour lisser les inhomogénéités du laser : il n'y a donc pas de perturbations du front d'ablation pour croître sous l'effet de l'IRT ablative lors de l'accélération de la cible. Une couche plus fine en surface crée moins de rayons X et donc une zone de conduction moins large, permettant une empreinte des défauts laser. L'acroissement du niveau des modulations pour une couche d'or de moins de 10 nm est expliqué par le fait que cette fine couche est instable vis-à-vis de l'IRT quand elle est poussée par le plasma de CH en expansion qui se trouve derrière. On pourra noter qu'un inconvénient de cette méthode dans le cadre d'une application à la FCI est le préchauffage du CH causé par les rayons X émis par la couche de Z élevé, ce qui limite la compression et peut être un frein dans le cadre de la FCI où des compressions maximales sont recherchées.

Watt et al [20] utilisent une méthode différente pour lisser les défauts du laser. Dans leur expérience réalisée sur le laser OMEGA, une cible de CH de densité 1,05 g/cm³ et d'environ 20 μm d'épaisseur est recouverte par une mousse de CH de 30 mg/cm³ de densité et de 100 μm d'épaisseur. Cette mousse est elle-même recouverte de 15 nm d'or. Les cibles couvertes comme les cibles de référence de CH seul sont irradiées par 5 faisceaux formant une impulsion carrée de 3 ns à une intensité de 2.10^{14} W/cm². Les mesures d'accélération montrent des comportements similaires pour les deux types cibles, avec un décalage de 500 ps pour le début d'accélération de la cible couverte de mousse.

Les radiographies de face représentées en figure 2.9 (a-b) montrent (en tenant compte du décalage de 500 ps) une réduction du niveau des modulations pour une cible couverte de mousse, ce qui est confirmé par les spectres de Fourier représentés en figure 2.9 c). Les modulations de longueurs d'onde inférieures à 50 μm sont supprimées tandis que celles de longueurs d'onde comprises entre 50 et 100 μm sont réduites d'un facteur ≈ 2 . Le mécanisme de lissage induit par les mousses est le suivant : lors de son irradiation, la couche d'Au émet des rayons X qui ionisent la mousse de manière supersonique (sans

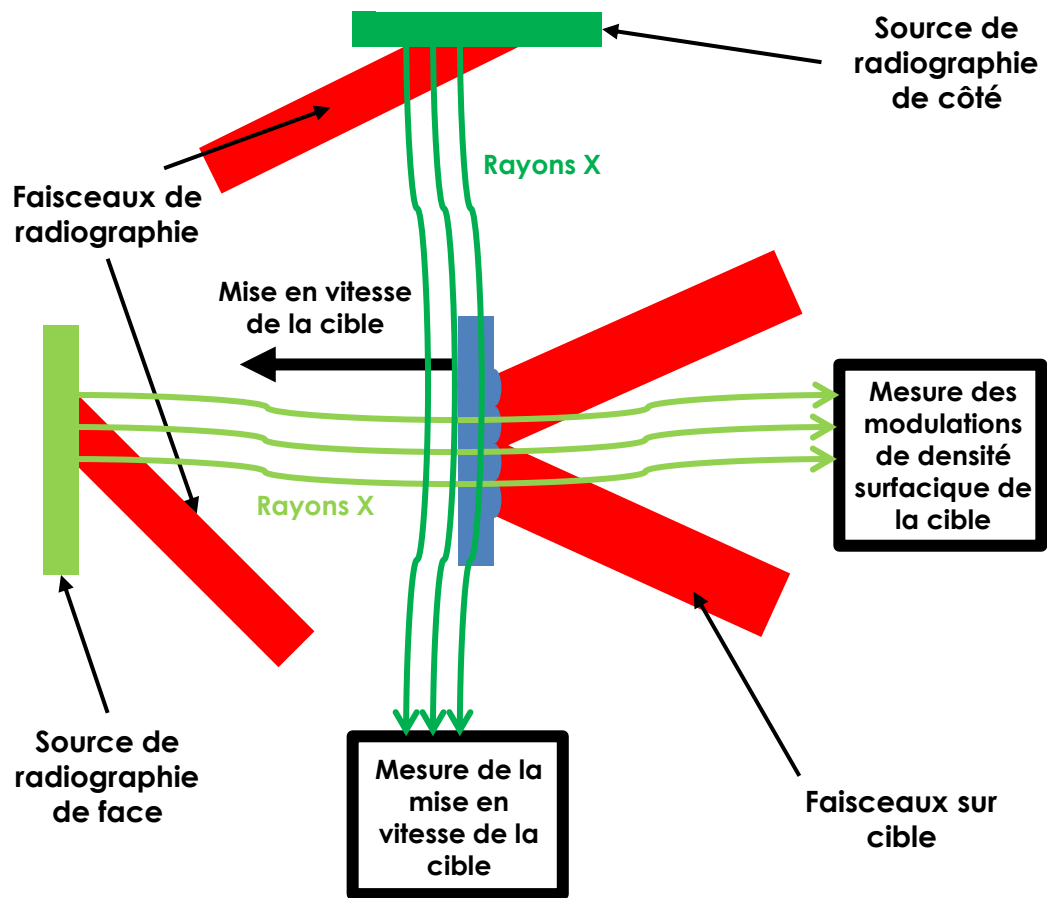


FIGURE 2.8 – Schéma d'une configuration expérimentale type pour mesurer la mise en vitesse et l'évolution des modulations de densité surfacique d'une cible plane.

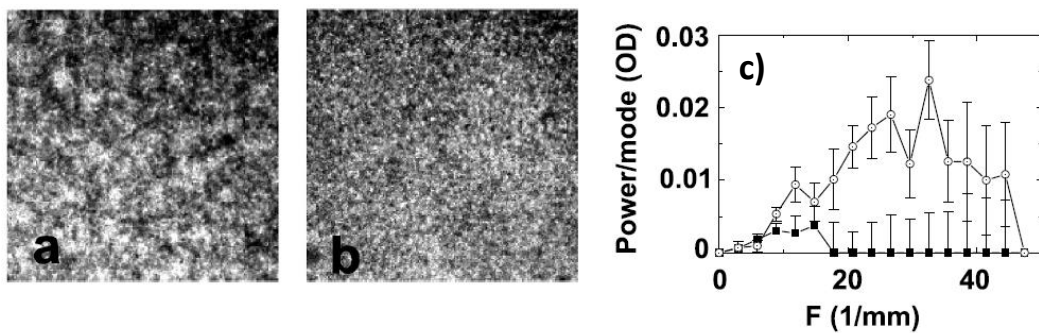


FIGURE 2.9 – Radiographie de face d'une cible a) de CH seul à 2,2 ns et b) de CH recouverte de mousse à 2,9 ns. c) Spectres de Fourier des radiographies de face pour une cible de CH seul à 2,45 ns (ronds blancs) et une cible recouverte de mousse à 2,95 ns (carrés noirs). Cette figure est extraite de la réf. [20].

lancer un choc), la transformant en plasma. Cela crée une large zone de conduction : quand le laser irradiera le CH, cette zone va lisser les défauts du laser. Pour une longueur de zone de conduction donnée, plus la longueur d'onde des modulations du laser est petite, plus le lissage est efficace (cf section 2.3.2). Ceci explique pourquoi les petites longueurs d'onde sont complètement supprimées dans le spectre de Fourier. L'efficacité de ce design de cible recouverte d'une mousse de basse densité est donc démontrée. Il se pose néanmoins aussi le problème du préchauffage de la cible par les rayons X émis par la couche d'or, comme dans l'expérience décrite précédemment.

2.4.4 Diminution du taux de croissance de l'IRT ablative

Malgré les différentes méthodes pour réduire les défauts laser et leur empreinte, il subsiste toujours des perturbations du front d'ablation. Un moyen de limiter leur croissance est de diminuer le taux de croissance de l'IRT ablative. Fujioka et al [17] présentent une expérience dans laquelle des cibles de plastique (CH) dopées (ou non, le dopage d'un matériau consistant à ajouter une petite quantité d'un autre matériau pour modifier ses propriétés) au Br et portant des modulations 2D sont accélérées sur le laser GEKKO XII [73]. Les simulations et la théorie prédisent que pour ce type de cible, une structure de double front d'ablation va se former, comme représentée en figure 2.10. Le premier front d'ablation (1) correspond à celui présenté dans la première partie de ce chapitre, créé par la conduction électronique de l'énergie déposée par le laser à la surface critique. Il est appelé front d'ablation de conduction électronique. L'émission X intense du Br de la couronne de plasma pénètre dans la cible et crée un second front d'ablation, appelé front d'ablation radiatif, au niveau où ces radiations sont absorbées (2). Le plateau entre les deux fronts d'ablation est de faible densité, le front d'ablation de conduction électronique est donc stable vis-à-vis de l'IRT. L'IRT ne pourra donc se développer qu'au niveau du deuxième front d'ablation. Les auteurs montrent par des simulations que la vitesse d'ablation est augmentée d'un facteur 3 au front d'ablation radiatif, comparé au cas de la cible de CH non dopée. La longueur minimale de gradient de densité est aussi augmentée. Le taux de croissance de l'IRT d'ablative est donc réduit d'après (2.26). Cependant, le pic de densité du matériau est plus petit dans le cas du CH dopé : la forte émission de la couronne provoque le préchauffage de la cible et rend donc sa compression plus difficile.

Les cibles d'épaisseur $25\text{ }\mu\text{m}$ portent des modulations de longueur d'onde $80\text{ }\mu\text{m}$ et d'amplitude $0,8\text{ }\mu\text{m}$ crête-à-crête. L'impulsion sur cible, formée par 3 puis 9 faisceaux, est faite d'un pied de 2 ns à 10^{12} W/cm^2 suivi de la partie principale de $2,5\text{ ns}$ à environ $1,5 \cdot 10^{14}\text{ W/cm}^2$. Tout d'abord, des mesures de profils de densité ont été effectuées par ombroscopie 1D [74] à l'aide d'une caméra à balayage de fente. Un double front d'ablation est alors observé. De plus, l'évolution des modulations du front d'ablation a été mesurée en radiographie de face par l'utilisation couplée d'une source de radiographie de Zn

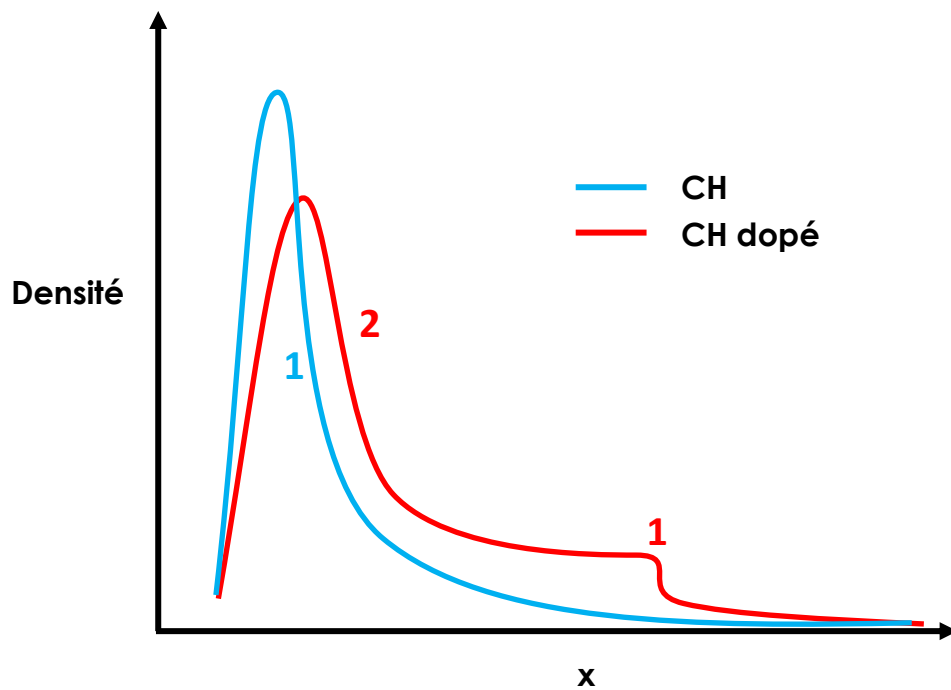


FIGURE 2.10 – Schéma des profils de densité obtenus dans le cas d’un front d’ablation classique avec une cible de CH seul (en bleu) et dans le cas d’un double front d’ablation avec une cible de CH dopé (en rouge), où l’on peut voir le front d’ablation créé par conduction électronique (1) et le front d’ablation radiatif (2).

(émission à 1,5 keV) et d'une caméra à balayage de fente. Parmi les résultats, les auteurs montrent que le taux de croissance mesuré à $1,7 \text{ ns}^{-1}$ pour le CH non dopé est abaissé à $1,2 \text{ ns}^{-1}$ pour le CHBr. Cela démontre expérimentalement que l'utilisation d'ablateurs dopés par des matériaux de Z élevé permet, par la création d'un double front d'ablation, la réduction de l'IRT ablative. Un dernier point abordé par les auteurs est que moins de 1 % du rayonnement de la couronne de plasma traverse l'ablateur de CHBr ; le rayonnement mesuré en sortie se compose essentiellement de rayons X d'énergie supérieure à quelques dixièmes de keV. Ce fait est important dans le cadre d'une application à la FCI, car il écarte le risque d'un préchauffage des isotopes d'hydrogène d'une cible cryogénique, ces atomes étant quasi-transparents aux rayons X à ces énergies.

Otani et al [75] étudient eux l'effet sur l'IRT ablative de l'irradiation d'une cible par plusieurs longueurs d'onde simultanément. De nombreux lasers de puissance utilisent une longueur d'onde laser de $0,35 \mu\text{m}$ car l'efficacité de conversion de l'énergie laser en énergie cinétique, et donc en compression, est meilleure qu'avec de plus grandes longueurs d'onde. Cependant, en irradiant une même cible avec un laser de plus grande longueur d'onde, plus d'électrons suprathérmiqes, aussi appelés électrons chauds seront produits. Ces électrons correspondent à la queue de la distribution énergétique maxwellienne des électrons, c'est-à-dire à la relativement petite fraction des électrons les plus énergétiques produits lors de l'interaction laser-plasma. Ces électrons pénètrent plus profondément dans la cible, élargissant la zone d'ablation ; la longueur minimale de gradient de densité devrait donc être augmentée et donc le taux de croissance de l'IRT ablative réduit.

L'expérience réalisée sur GEKKO XII consiste donc à irradier des cibles de CH de $25 \mu\text{m}$ d'épaisseur portant des modulations de longueur d'onde $20 \mu\text{m}$ et d'amplitude $0,2 \mu\text{m}$ par un mélange de longueur d'onde laser de $0,35 \mu\text{m}$ et $0,53 \mu\text{m}$ (ou de longueur d'onde $0,35 \mu\text{m}$ seule pour avoir un étalon). Ces lumières correspondent respectivement à un triplement et un doublement de la fréquence de la lumière initiale ($\lambda = 1,05 \mu\text{m}$), nous les qualifierons donc de lumière à 3ω et 2ω . Trois faisceaux laser servent à former un pied d'impulsion pour pré-comprimer la cible pendant 2,3 ns à $8 \cdot 10^{11} \text{ W/cm}^2$, puis la cible est irradiée par 9 faisceaux, dont 3 à 2ω , formant une impulsion carrée de 2,5 ns. Les mesures de l'IRT ablative sont effectuées par l'utilisation couplée d'une caméra à balayage de fente et d'une source de radiographie de Zn. L'intensité sur cible est de $1,3 \cdot 10^{14} \text{ W/cm}^2$ pour l'irradiation à 3ω seule, et de $7,7 \cdot 10^{13} \text{ W/cm}^2$ à 3ω et $5,1 \cdot 10^{13} \text{ W/cm}^2$ à 2ω pour le cocktail de couleur. L'accélération mesurée dans les deux cas par radiographie de côté (caméra à balayage de fente et source de radiographie en Al) est similaire : $8 \cdot 10^{15} \text{ cm/s}^2$. Les taux de croissance mesurés sont de $3,3 \text{ ns}^{-1}$ pour l'irradiation à 3ω et de $2,2 \text{ ns}^{-1}$ pour le cocktail de couleur, ce qui démontre donc son efficacité pour la réduction de l'IRT ablative. Un défaut est cependant partagé entre cette méthode et les ablateurs dopés par

des matériaux de Z élevé : une compression plus faible. Ici, ce défaut de compression est dû au fait qu'une partie de l'énergie laser est utilisée pour l'irradiation à 2ω qui est moins efficace pour la compression que la lumière à 3ω .

2.5 Points manquants traités dans cette thèse

Au cours de ce chapitre, nous avons présenté la physique du front d'ablation et des instabilités hydrodynamiques qui s'y développent. La modélisation de l'IRM ablative est relativement récente, que ce soit pour des modulations de la surface d'une cible [9, 58] ou dans le cas de l'empreinte de défauts laser [11, 60]. L'IRM ablative a été largement étudiée expérimentalement pour des cibles sur lesquelles des modulations avaient été usinées. En revanche, aucune expérience n'a été réalisée pour l'IRM ablative imprimée par laser. Nous allons donc nous attacher à mesurer la phase de Richtmyer-Meshkov ablative imprimée par laser et à interpréter ces mesures à l'aide du modèle le plus complet présenté dans la réf. [11].

D'autre part, on a pu voir que de nombreuses méthodes ont été développées pour réduire l'effet des instabilités hydrodynamiques ablatives, que ce soit en réduisant le niveau initial de modulations ou en diminuant la croissance. Les méthodes induisant une modification de la cible (couverture de mousse [20] ou métallique [72], ablateur dopé [17], ...) ont montré des résultats probants mais possèdent certains inconvénients, comme le préchauffage de la cible. C'est la raison pour laquelle ces méthodes avaient été un peu mises de côté ces dernières années. Cependant, les expériences récentes de FCI [6, 8] montrent que l'effet délétère des instabilités hydrodynamiques avait été sous-estimé ; les différentes méthodes de réduction de ces instabilités connaissent donc un regain d'intérêt. Une voie prometteuse semble être l'utilisation de mousses sous-denses présentée dans l'introduction [22]. Des expériences utilisant ces mousses ont donc été préparées et interprétées pendant cette thèse. La phase Rayleigh-Taylor ablative d'un ablateur placé après une mousse sous-dense a été mesurée pour la première fois.

Chapitre 3

Matériel et méthodes

Dans ce chapitre, nous allons présenter les détails techniques de ce qui a été mis en oeuvre pour réaliser et analyser nos expériences. En premier lieu, l'installation laser utilisée, située à Rochester (E-U), est présentée, tout comme dans une deuxième partie les diagnostics employés sur cette installation. Une fois que les mesures ont été effectuées, il faut dépouiller et analyser les résultats : les programmes développés à cet effet sous IDL [76] seront donc expliqués. Enfin, pour la préparation des expériences comme pour leur interprétation, nous nous sommes appuyés sur des simulations réalisées à l'aide des codes CHIC et PARAX. Le fonctionnement de ces codes sera brièvement présenté en dernière partie.

3.1 Installation laser OMEGA

3.1.1 Laser OMEGA

Le laser OMEGA se situe au Laboratory for Laser Energetics (LLE) de l'Université de Rochester (UR) à Rochester (Etats-Unis). Il délivre, par le biais de 60 faisceaux, des impulsions nanosecondes pour une énergie totale pouvant atteindre 30 kJ. Il est constitué de deux parties : la "Laser Bay" (LB) et la "Target Bay" (TB), séparées par un mur de protection. Nous allons tout d'abord présenter la chaîne laser qui crée les conditions nécessaires aux expériences de FCI à partir d'une impulsion nJ originelle.

Une chaîne laser d'OMEGA est schématisée en figure 3.1. Une source laser produit des impulsions de longueur d'onde 1054 nm (infra-rouge, IR), de durée 80 ps et d'énergie 1 nJ à 76 MHz dans la "Driver Electronic Room" (DER). Ces impulsions sont mises en forme puis transférées à la "Pulse Generation Room" (PGR) par une fibre optique. Une impulsion est sélectionnée, puis amplifiée dans une cavité laser appelée amplificateur régénératif. Au bout d'environ 100 aller-retours dans la cavité, l'impulsion sort avec une énergie de 0,1 mJ. Si requis, l'impulsion passe ensuite dans le système de SSD (cf chapitre

2). Il est à noter qu'à ce moment, le faisceau laser fait 40 mm de diamètre. Puis l'impulsion quitte la PGR, pour être amplifiée par un "Large-Aperture Ring Amplifier" (LARA) ; en 4 tours de l'anneau, l'impulsion est amplifiée 10 000 fois et atteint donc environ 1 J. Enfin, un pré-amplificateur cylindrique monte l'énergie de l'impulsion à 4,5 J pour un diamètre de faisceau de 64 mm. L'impulsion quitte ensuite la partie "driver" pour la partie "amplification" de la chaîne. Elle va rencontrer 6 niveaux d'amplification, du niveau A au niveau F. Les amplificateurs A à D sont des cylindres, tandis que les niveaux E et F sont un enchaînement de 4 disques.

Tous les amplificateurs sont constitués de verre dopé au Néodyme pompé par lampe flash. Le fonctionnement est le suivant : le flash des lampes éclaire les amplificateurs peu avant le passage du faisceau laser. Les électrons sont excités et passent donc sur un niveau d'énergie plus élevé, causant une inversion de population. Lorsque les photons du faisceau laser traversent le matériau amplificateur, les électrons se désexcitent par émission stimulée : les électrons retrouvent leur niveau d'énergie initial en émettant un photon qui a les mêmes caractéristiques que le photon incident, en énergie comme en direction.

Le diamètre du faisceau est progressivement augmenté au cours de l'amplification : de 64 mm à 90 mm après le niveau B, puis 150 mm après le niveau D, 200 mm après le niveau E et enfin 280 mm après le niveau F. Cet accroissement est nécessaire pour diminuer la fluence (rapport énergie sur surface) du faisceau, car dans le cas contraire les optiques de la chaîne seraient détériorées. Les lampes flash sont refroidies par de l'eau froide tandis que les amplificateurs le sont par une solution de diéthylène glycol (le diéthylène glycol permet d'abaisser la température de solidification de l'eau). L'écart de 45 min entre 2 tirs est dicté par le temps nécessaire au refroidissement des amplicateurs.

C'est aussi dans la partie amplification que l'on passe à 60 faisceaux : un seul faisceau quitte la partie "driver". Il est séparé en 3 avant l'amplificateur A. Puis chaque faisceau est séparé en 5 avant l'amplificateur B ; enfin, chacun des 15 faisceaux est divisé en 4 après l'amplificateur C. Après cette division, on obtient 60 faisceaux. La partie qui permet d'ajuster le timing relatif des faisceaux est placée juste après la dernière division. En allongeant le chemin des faisceaux voulus, on peut ainsi obtenir un écart entre les impulsions allant jusqu'à 9 ns. De plus, à chaque niveau d'amplificateur se trouve un filtre spatial. Ainsi, après chaque amplification, le faisceau est nettoyé de la lumière mal collimatée.

Après l'amplification, la conversion de fréquence est réalisée. Le faisceau IR ($\lambda = 1054$ nm) qu'on appelle faisceau à 1ω est converti en UV ($\lambda = 351$ nm) à 3ω . Pour ce faire, on utilise des cristaux de phosphate de dihydrogène de potassium appelés cristaux de KDP. Sous certaines conditions (polarisation, intensité, ...), le KDP permet d'additionner les fréquences de deux ondes incidentes. Ainsi, le faisceau incident à 1ω traverse un polariseur, puis une partie de ce faisceau est prélevée pour traverser un premier cristal de KDP, ce qui permet d'obtenir une onde à 2ω . Puis cette onde à 2ω et le reste du faisceau à 1ω traversent un

deuxième cristal de KDP, ce qui permet d'obtenir en sortie un faisceau à 3ω . Enfin, un 3ème cristal est utilisé pour optimiser la conversion dans le cas de l'utilisation du lissage SSD, car on a alors un spectre de fréquence. Le phénomène de conversion de fréquence est complexe et dépend entre autres de l'intensité initiale; ainsi, la conversion est maximale dans le cas d'une impulsion carrée de 1 ns. On obtient alors au maximum une impulsion de 500 J en sortie de conversion.

Le faisceau laser traverse alors le mur de protection et entre dans la TB. Après un passage par le F-ASP, le "Alignement Sensor Package" du niveau F (nous reviendrons sur ce point ultérieurement), le faisceau est dirigé vers la cible par deux miroirs. Avant de pénétrer dans la chambre de tir ("Target chamber", TC), le faisceau passe à travers diverses lames de phase puis une lentille de focale 1,8 m qui le focalise sur la cible. La lentille est placée sur une monture de précision qui permet de la déplacer et donc de faire varier la position de focalisation du faisceau. Le faisceau entre alors dans la TC par un hublot protégé par un bouclier à débris. La TC est une sphère d'aluminium de 3,3 m de diamètre et de 9 cm d'épaisseur. Lors des tirs, l'intérieur de la sphère se trouve à une pression inférieure à 10^{-3} Pa. La TC est soutenue par une structure à trois niveaux, la "Target Mirror Structure" (TMS), qui permet la manutention des diagnostics et des lames de phase. Plusieurs diagnostics fixes sont insérés dans la TC. De plus, 6 inserteurs de diagnostics, appelés "Ten-Inch Manipulators" (TIM) sont disponibles. Ils correspondent à des ouvertures dans la TC qui permettent d'insérer et de retirer facilement des diagnostics. Les TIMs possèdent différentes entrées et sorties qui permettent d'alimenter électriquement le diagnostic, de le contrôler et d'y faire le vide.

Pour effectuer l'alignement des faisceaux laser, un faisceau d'alignement IR est inséré au niveau du séparateur A et un faisceau UV au niveau du F-ASP. Plusieurs ASP, non représentés en figure 3.1, sont disposés le long de la ligne laser : deux dans la partie "driver", un dans le niveau d'amplification A, un dans le niveau C et un, donc, après le niveau F. Ces ASP contiennent des capteurs qui permettent de détecter les faisceaux d'alignement et donc de corriger l'alignement de la chaîne. Un autre point à noter est que 3 "drivers" différents peuvent alimenter les faisceaux laser : le "main", le "SSD" et le "backlighter". "Main" et "SSD" sont des drivers similaires qui peuvent alimenter les 60 faisceaux, à la différence près que le driver "SSD" possède le système de lissage SSD sur sa ligne. Un miroir mobile permet de sélectionner un driver ou l'autre pour alimenter les faisceaux d'OMEGA. Quant au driver "backlighter", il ne possède pas de préamplificateur après son LARA, il ne peut donc alimenter que 20 faisceaux. On peut donc utiliser le "SSD" ou le "main" pour alimenter seul les 60 faisceaux, ou l'utiliser pour 40 faisceaux et le "backlighter" pour les 20 faisceaux restants. Ceci permet d'utiliser deux types d'impulsion différentes lors d'un même tir, un par driver (on ne peut utiliser les 3 drivers en même temps). Les 60 faisceaux d'OMEGA sont numérotés de 10 à 69, et séparés en trois groupes appelés "legs" (jambes) : leg 1 pour les faisceaux 10-19 et 40-49, leg 2 pour les faisceaux

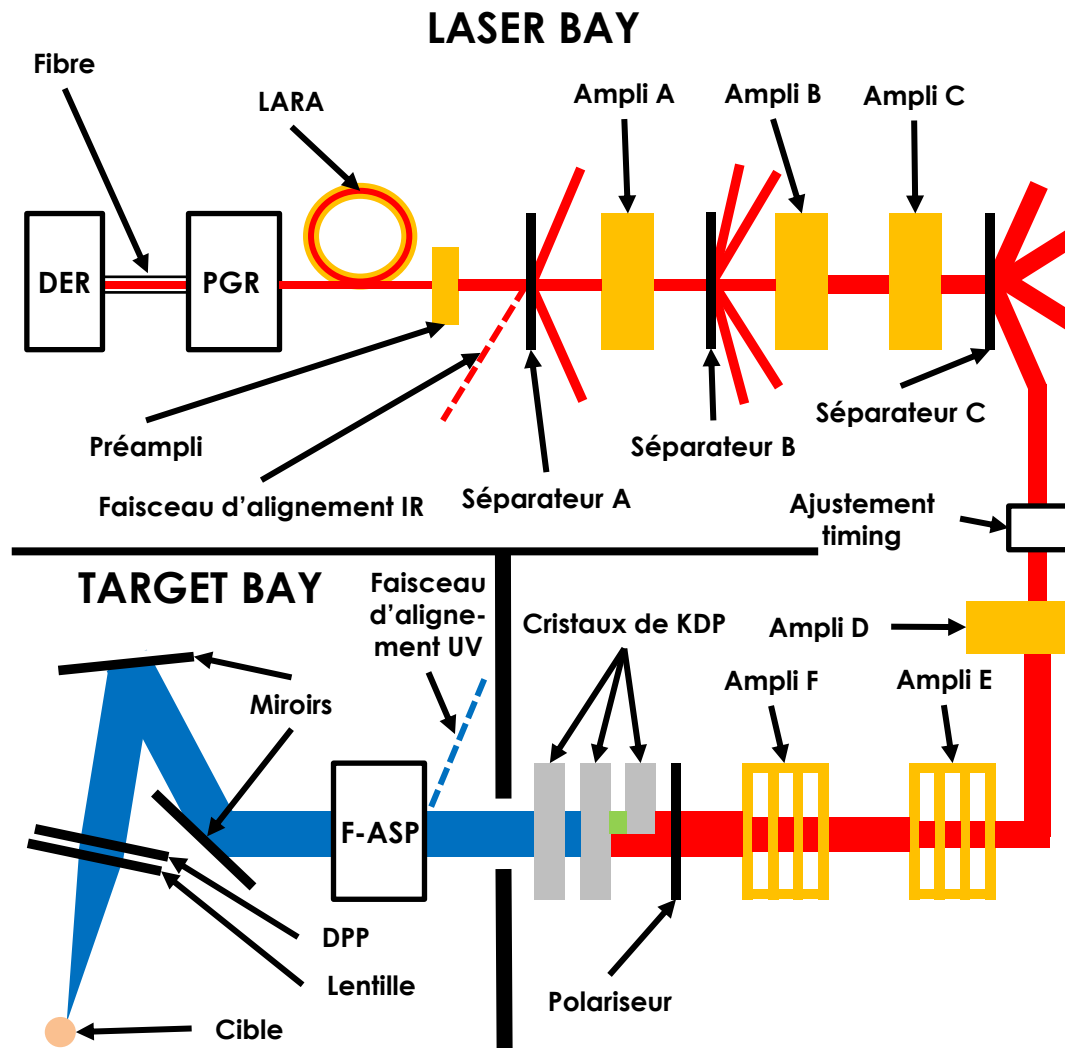


FIGURE 3.1 – Schéma d'une chaîne laser OMEGA. Tous les composants ne sont pas représentés.

20-29 et 50-59 et leg 3 pour les faisceaux 30-39 et 60-69. Tous les faisceaux d'un "leg" sont alimentés par le même driver ; ainsi, deux faisceaux pourront avoir des impulsions de formes différentes s'ils n'appartiennent pas au même "leg".

Deux types de lames de phase sont utilisés le long des chaînes laser. Tout d'abord, une "Distributed Phase Rotator" (DPR) insérée dans le F-ASP permet lisser la polarisation du faisceau par un fonctionnement similaire aux RPP présentées dans le chapitre 2. D'autre part, une "Distributed Phase Plate" (DPP) placée juste avant la lentille de focalisation permet de donner la forme désirée au spot du faisceau laser sur la cible.

On peut enfin noter que la source d'impulsions initiale dans la DER permet d'alimenter les faisceaux dits de "fiducial" qui permettent d'obtenir une référence temporelle pour certains diagnostics tels que les caméras à balayage de fente.

L'installation laser OMEGA EP est présentée en annexe C. Nous n'avons effectué qu'une seule journée de tir sur cette installation ; ces expériences ne sont pas présentées dans ce manuscrit mais sont abordées dans la conclusion.

3.1.2 Préparation d'une journée de tirs

Il existe plusieurs moyens d'obtenir une journée de tirs sur OMEGA (ou OMEGA EP) : en achetant une journée au LLE, sur le quota de tirs interne au LLE en collaborant avec des membres du laboratoire ou par le biais de l'appel à projet annuel "Laboratory Basic Science" (LBS). Une proposition LBS ne peut néanmoins être soumise que par un membre d'un laboratoire américain du "National Laser User's Facility" (NLUF) ; il faut donc aussi être en collaboration pour obtenir des tirs par ce biais. La proposition LBS doit établir clairement plusieurs points : socle théorique/numérique, forte probabilité de résultats, approvisionnement des cibles, diagnostics nécessaires... Si une proposition est acceptée, deux réunions auront lieu, trois mois et un mois avant l'expérience. Lors de la réunion à 3 mois, une configuration expérimentale détaillée devra être présentée, ainsi qu'une fiche de tir ("Shot Request Form", SRF) typique de la journée de tir. Lors de la réunion à 1 mois, un plan de tir prévisionnel devra être présenté, ainsi qu'une SRF pour chacun de ces tirs. Les derniers détails devront avoir été réglés et les cibles reçues par le LLE pour être assemblées.

Une fiche de tir (SRF) contient tous les détails nécessaires à la réalisation d'un tir. Sur OMEGA, elle contient 7 parties : "General", "Drivers", "Target", "Beams", "TIM", "Fixed" et "Neutronics". La partie "General" sert à définir le contexte de l'expérience : date, "Principal Investigators" (PI), numéro de tir, ... La partie "Driver" permet de choisir

quel "leg" sera alimenté par quel driver. Typiquement, on peut choisir soit d'alimenter les 3 "legs" et donc les 60 faisceaux avec le driver "SSD" (le "main" est devenu obsolète depuis le développement du driver "SSD"), soit d'alimenter 2 "legs" avec le driver "SSD" et un avec le "backlighter". Les expérimentateurs peuvent aussi décider d'activer ou non le système de SSD sur le driver "SSD", et doivent donner le temps d'activation du driver par rapport au t_0 d'OMEGA. La partie "Target" permet de définir le nom et le type des cibles utilisées, leur taille ainsi que leur forme. Toutes les cibles comme les sources de radiographie doivent porter un nom différent pour que l'équipe expérimentale d'OMEGA sache laquelle insérer. De plus, les expérimentateurs doivent aussi choisir par où la cible sera insérée dans la chambre. Des systèmes appelés "Target Positionner Systems" (TPS) sont prévus à cet effet ; un TPS est situé en permanence dans l'un des ports de la TC d'OMEGA (port H2) tandis que 2 autres TPS peuvent être insérés dans les TIM. Dans la partie suivante, nommée "Beams", les PI doivent définir les caractéristiques de chacun des faisceaux qui seront utilisés lors de la journée : le numéro du faisceau, le type d'impulsion, le pointage, la focalisation (si l'on veut agrandir le spot à la meilleure focalisation par exemple), le délai des faisceaux par rapport au temps d'activation du driver, s'il faut une DPP et laquelle et s'il faut une DPR. La partie "TIM" permet aux PI de choisir ce qu'ils veulent insérer dans les TIMs. De très nombreux diagnostics sont disponibles (XRFC, caméra à balayage de fente, spectromètres, diagnostics protoniques, ...) ainsi que les deux TPS. Si les PI choisissent d'utiliser des diagnostics, il faudra ensuite en définir les réglages. La partie "Fixed" permet de choisir d'activer ou non une liste de diagnostics installés de manière permanente dans la TC d'OMEGA. La partie "Neutronics" permet d'en faire de même pour les diagnostics neutroniques lors de tirs avec dégagement neutronique.

En annexe C, les fiches de tirs utilisées sur OMEGA EP sont décrites.

3.1.3 Contraintes d'exploitation sur OMEGA

Au cours de cette thèse, nous avons obtenu au total cinq journées de tirs sur OMEGA, quatre dans le cadre de l'étude de l'IRM ablative imprimé par laser et une pour l'étude des mousses sous-denses. Ceci représente 52 tirs dans le premier cas (soit une moyenne de 13 tirs par journée) et 14 tirs dans le second. Sur ces 66 tirs, 35 ont permis d'extraire les données les plus pertinentes présentées dans les chapitres suivants. Chaque tir revêt donc une grande importance. Ainsi, lors d'une campagne expérimentale sur le laser OMEGA, un des objectifs à garder à l'esprit est d'essayer de maximiser le nombre de tirs exploitables, et donc d'équilibrer la balance entre la nouveauté de l'expérience réalisée et les risques pris dans la conception de cette expérience.

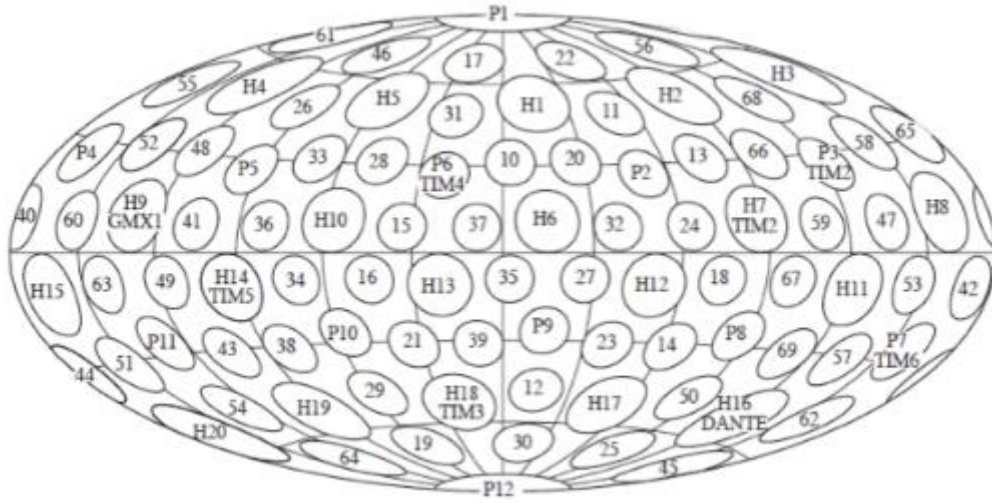


FIGURE 3.2 – Vue des différents hublots et inserteurs de la TC d'OMEGA.

Le point fondamental à prendre en considération lors de la conception d'une expérience est que tout changement de configuration au cours d'une journée de tirs peut s'avérer rédhibitoire en terme de nombre de tirs obtenus. Le changement du diagnostic principal, par exemple, coûte dans le meilleur des cas deux tirs, car il faut casser le vide de l'inserteur, changer le diagnostic, refaire le vide et repointer le diagnostic ; de même pour le changement d'une lame de phase suivi du repointage d'un faisceau. Certains changements sont en outre impossibles : par exemple, on ne peut ajouter en cours de journée des faisceaux qui n'avaient pas été prévu préalablement. Il est donc nécessaire de soigneusement préparer chaque journée de tir, en prévoyant les difficultés qui peuvent être rencontrées et en effectuant des simulations qui vont permettre, par exemple, de déterminer les instants auxquels les mesures doivent être faites.

Lors de chaque journée de tirs que nous avons effectuée, notre objectif principal était de réaliser une radiographie de face de la cible. Le diagnostic utilisé pour cela devait être inséré dans un TIM. Comme on peut le voir sur le schéma de la chambre d'OMEGA en figure 3.2, les six TIMs d'OMEGA forment trois axes de deux TIMs face-à-face : l'axe P6-P7 (TIMs 4 et 6), l'axe H3-H18 (TIMs 1 et 3) et l'axe H7-H14 (TIMs 2 et 5). Lors de toutes nos campagnes, l'axe H3-H18 a été choisi. En effet, les cônes de faisceaux à 23° et 48° possèdent 6 faisceaux sur cet axe et 5 seulement sur les autres axes. Ceci permet donc d'augmenter le nombre de faisceaux disponibles ainsi que l'intensité maximale sur cible que l'on peut atteindre. De plus, les faisceaux 25 et 30, qui sont les seuls de l'installation à porter des diagnostics de rétrodiffusion, font partie des cônes d'illumination de l'axe H3-H18. Or ces diagnostics étaient nécessaires dans la campagne d'étude des mousses

sous-denses.

Lors de nos campagnes, il aurait été intéressant de pouvoir mesurer la dynamique de la feuille, le débouché de choc et sa mise en vitesse. Pour cela, il aurait fallu réaliser, dans l'idéal, une radiographie de côté, ou utiliser un VISAR. Pour effectuer une radiographie de côté, il faut insérer une caméra à balayage de fente à la perpendiculaire de la direction de mise en vitesse de la feuille, donc de l'axe de radiographie de face. Or les axes formés par les TIMs ne sont pas perpendiculaires entre eux. L'utilisation d'un VISAR doit quant à elle se faire dans le même axe que la radiographie de face. Or, d'un côté de la cible se trouve le diagnostic de radiographie, et de l'autre côté la source de radiographie, qui masquerait la cible pour l'utilisation d'un VISAR. Ainsi, on ne peut effectuer simultanément une radiographie de face et une mesure de la dynamique de la feuille. Or un changement de configuration en cours de journée pour passer d'un type de mesure à l'autre aurait coûté plusieurs tirs. La possibilité restante aurait été de consacrer une journée entière aux mesures de la dynamique de la feuille. En gardant à l'esprit que la partie novatrice des expériences résidait dans les données mesurées par radiographie de face, il a été décidé, en accord avec nos collègues américains, de ne pas effectuer de mesure de dynamique de la feuille (à l'exception des mesures d'auto-émission de la campagne sur les mousses sous-denses qui sera abordée ultérieurement).

Dans la partie suivante, nous allons détailler les diagnostics disponibles sur OMEGA qui auront été utilisés dans nos expériences. Il faut noter que le nombre et la variété de diagnostics sont très importants sur ces installations, supérieur à une centaine.

3.2 Diagnostics utilisés

Dans le domaine des lasers de puissance, les appareils de mesure sont appelés diagnostics. Ils sont le fruit d'un développement poussé car ils doivent fonctionner dans des conditions difficiles : projection de débris, présence de rayonnement X, émission de neutrons... Le diagnostic fondamental de toutes nos expériences est un imageur X multi-sténopés, appelé "X-ray framing camera" en anglais. C'est à l'aide de cet appareil couplé à une source de radiographie que nous avons mesuré les variations de densité surfacique des cibles induites par les instabilités hydrodynamiques. La partie qui suit présente son fonctionnement ainsi que celui des autres principaux diagnostics utilisés lors de nos expériences.

3.2.1 X-Ray Framing Camera (XRFC)

Le diagnostic principal de nos expériences était un imageur X multi-sténopé ou "X-Ray Framing Camera" (XRFC). Le détail de la composition des XRFC est représenté en figure 3.3. Cette caméra est utilisée couplée à une source de radiographie : les rayons X émis par la source de radiographie traversent l'objet d'intérêt puis sont mesurés par la XRFC. Un ensemble de sténopés permet de sélectionner plusieurs images de l'objet avec une résolution qui dépendra du diamètre des sténopés. En revanche, un sténopé plus petit réduit la quantité de rayons X et donc le niveau du signal. Les rayons X sont ensuite absorbés par une photocathode qui les convertit en électrons. Ces photoélectrons atteignent alors une galette de microcanaux (GMC) [77]. La GMC est composée d'un ensemble de petits tubes, les microcanaux, qui la traversent ; ces microcanaux sont recouverts en entrée et en sortie d'une fine couche d'or. On applique une différence de potentiel entre ces deux couches ; ainsi, lorsqu'un photoélectron frappe la couche d'entrée, il produit plusieurs électrons qui sont accélérés le long du microcanal. En sortie, ces électrons frappent à nouveau une couche d'or et un nombre plus grand encore d'électrons est donc produit. On désire mesurer l'évolution temporelle des phénomènes d'instabilités hydrodynamiques : il faut donc que l'activation électrique de l'amplification de la GMC soit limitée dans le temps, à des périodes pas plus grandes que les temps caractéristiques d'évolution des phénomènes observés (de l'ordre d'une centaine de ps pour les instabilités hydrodynamiques). Dans le cas contraire, on obtiendrait des images de l'évolution des phénomènes intégrées en temps et donc quasiment impossibles à analyser. C'est donc la GMC qui fixe la résolution temporelle. Pour ce faire, des impulsions de quelques dizaines ou quelques centaines de ps (créées par avalanche dans une jonction p-n polarisée en inverse [78]) sont envoyées dans la GMC. Pour les XRFC utilisées sur OMEGA, les GMC sont séparées en 4 bandes, reliées à une source d'impulsion différente chacune, ce qui permet de réaliser des mesures à 4 instants différents indépendants. Dans le cadre de nos expériences, 2 ou 4 images peuvent être formées sur chaque bande, selon l'utilisation d'un réseau de 8 ou 16 sténopés. Les électrons en sortie de la GMC atteignent ensuite une plaque de phosphore, sur laquelle l'image électronique formée est convertie en lumière, qui forme enfin l'image définitive sur un film ou une CCD. Quatre XRFC sont disponibles sur OMEGA et OMEGA EP : une XRFC rapide avec une résolution temporelle courte de 50 ps (XRFC 1), 3 autres possédant des résolutions temporelles plus longue de 200 ps à 1 ns (XRFC 3, 4 et 5). Les résolutions temporelles correspondent à la durée de l'impulsion qui parcourt la MCP. Divers filtres peuvent être utilisés pour optimiser la bande spectrale que l'on veut détecter. Différents nez d'imagerie sont aussi disponibles, permettant d'obtenir jusqu'à plus de 30 images et d'atteindre des grandissements de 25. Lors de nos expériences, l'énergie des rayons X mesurés par les XRFC était d'environ 1,3 keV à 2,5 keV.

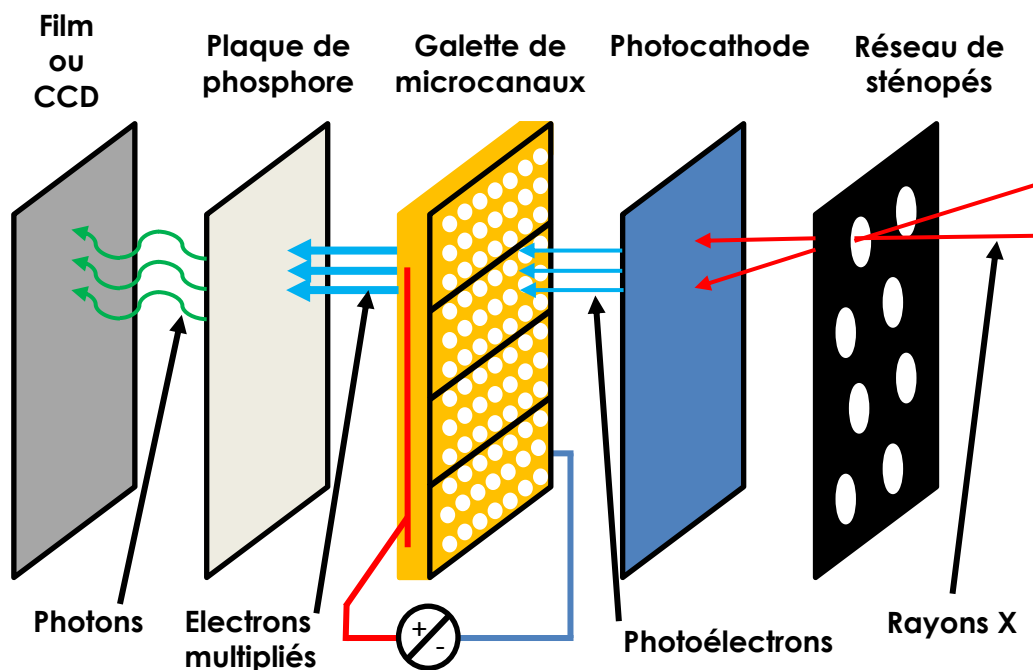


FIGURE 3.3 – Détail en vue éclatée de la composition d'une XRFC.

3.2.2 Imageur X couplé à une caméra à balayage de fente (SSCA)

Lors de certaines de nos expériences, nous avons utilisé des caméras à balayage de fente (CBF). Ce type de caméra permet d'obtenir des images 1D résolues continuellement en temps, contrairement aux caméras à images intégrales comme la XRFC qui mesurent une image 2D sur un intervalle de temps donné. Sur une image obtenue par une CBF, le temps est en abscisse et la quantité mesurée (espace, puissance laser, énergie de photons, ...) en ordonnée. Le détail du fonctionnement d'une CBF utilisée comme imageur est représenté en figure 3.4. Les rayons X, provenant d'une source de radiographie ou de l'émission propre d'une cible, sont filtrés par deux fentes, une de résolution spatiale, ce qui donne une image mono-dimensionnelle, et une de résolution temporelle. Comme pour la XRFC, les rayons X rencontrent ensuite une photocathode qui produit des photoélectrons. Ces électrons sont accélérés puis passent ensuite entre deux plaques de déviation ; une différence de potentiel variable temporellement $V(t)$ est imposée entre les deux plaques. La tension va décroître linéairement, modifiant continuellement la déviation des photoélectrons. Si ces derniers ont été émis à des temps différents, ils iront frapper une plaque de phosphore à des endroits différents. La plaque de phosphore émettra des photons qui formeront une image sur le système de détection final couplé à la CBF. La CBF que nous avons utilisé sur OMEGA est la "X-Ray Streak Camera A", nommée SSCA. Plusieurs paramètres sont à choisir sur la SSCA : le grandissement (fixé par l'écart entre la fente et la photocathode), la résolution spatiale (fixée par la taille de la fente), la vitesse de balayage (qui dépend de la vitesse de décroissance du voltage des plaques de déviation), le matériau de la photocathode,

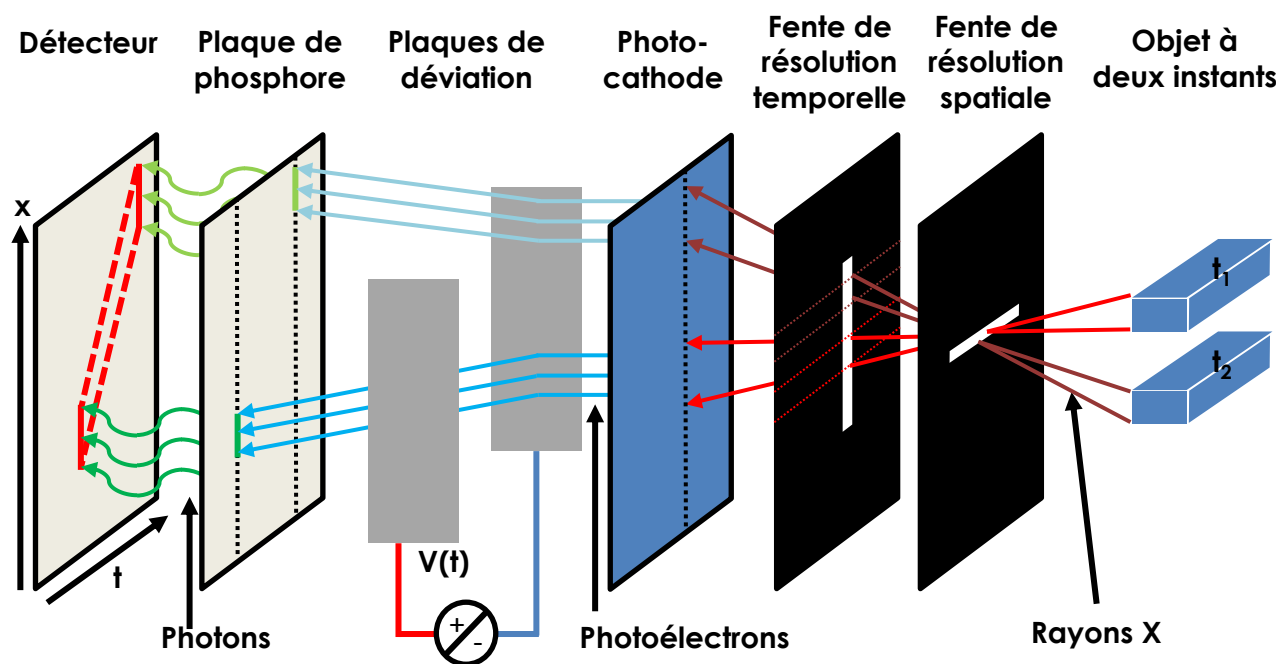


FIGURE 3.4 – Détail en vue éclatée de la composition d'une caméra à balayage de fente.

différents filtrages si besoin, ... Par exemple, le réglage de la vitesse de balayage permet d'avoir une fenêtre temporelle de 2, 4 ou 9 ns. La SSCA peut être utilisée avec un faisceau de "fiducial" à 4ω (263 nm) qui permet une calibration temporelle absolue. Lors de nos expériences, nous avons utilisé la SSCA pour détecter des rayons X d'énergie comprise entre 1 et 2 keV.

3.2.3 Diagnostics laser

Sur OMEGA, au niveau du F-ASP, 4 % de l'énergie laser est prélevée pour être analysée par un ensemble de diagnostics. 4 % de l'énergie prélevée est ensuite dirigée vers une sphère d'intégration, couplée à un spectromètre et une caméra CCD. Cet ensemble permet de mesurer précisément la puissance des 0,16 % d'énergie prélevée et donc la puissance totale du faisceau. Une autre partie de l'énergie prélevée est dirigée vers une CBF appelée P510 qui permet de mesurer la forme temporelle de l'impulsion. Pour finir, un système appelé "OMEGA Transport Instrumentation System" permet de mesurer précisément la transmission des optiques de fin de chaîne (miroirs, DPP, lentille de focalisation, hublot à vide et bouclier à débris). A l'aide de tous ces diagnostics, l'énergie qui atteint la cible et la forme de l'impulsion peuvent être mesurées précisément.

3.2.4 Diagnostics fixes

Deux types de diagnostics fixes ont été particulièrement utiles lors de nos expériences, le premier étant le "Full-Aperture Backscatter Station" (FABS). Deux FABS sont installés sur OMEGA, au niveau des faisceaux 25 et 30. Ils récupèrent l'énergie émise depuis la TC à travers les hublots des deux faisceaux. Ces diagnostics montrent leur utilité dans les expériences où se développent des instabilités paramétriques. Ils comprennent une partie dédiée à la diffusion Raman stimulée ($400 \text{ nm} < \lambda < 700 \text{ nm}$) et une à la diffusion Brillouin stimulée ($350 \text{ nm} < \lambda < 352 \text{ nm}$). Chaque partie est composée d'un calorimètre qui permet de mesurer la quantité d'énergie rétrodiffusée dans la gamme spectrale correspondant et d'une CBF ROSS utilisée en spectromètre. Les CBF permettent donc de mesurer l'évolution temporelle du spectre d'énergie rétrodiffusée. Les deux spectres sont ensuite accolés l'un à l'autre. Les résolutions spectrales et temporelles sont respectivement 0,04 nm et 80 ps pour le Brillouin et de 9 nm et 100 ps pour le Raman. La barre d'erreur sur l'énergie mesurée par les calorimètres est de 5 à 10 %.

L'autre type de diagnostic fixe utilisé lors de tous nos tirs est la "X-Ray Pinhole Camera" (XRPHC). Ces caméras sont au nombre de 6 sur OMEGA et 3 sur OMEGA EP. Elles permettent d'obtenir une image 2D intégrée en temps du centre de la TC. Elles montrent une grande utilité pour imager le spot des faisceaux qui émet sur les différentes cibles, et ainsi d'en vérifier la bonne illumination. Les XRPHC sont composées d'un sténopé de $10 \mu\text{m}$ situé à 17 cm de l'objet pointé. La distance entre le sténopé et le détecteur est de 68 cm induisant un grandissement d'un facteur 4.

3.3 Programmes de dépouillement développés sous IDL

Au cours des diverses journées de tir au LLE, les XRFC nous ont permis de mesurer des centaines d'images de modulations de densité surfacique par radiographie de face. Pour analyser ces images et gagner du temps, plusieurs programmes de dépouillement ont été mis au point sur le logiciel IDL, en utilisant le programme additionnel TSIGAN développé par le CEA. Cette partie décrira donc les principaux programmes développés au cours de cette thèse.

3.3.1 Corrélacion croisée

La figure 3.5 présente une radiographie effectuée par une XRFC. La radiographie a été réalisée de face, avec pour objet une cible de CH portant des modulations 3D au front d'ablation. De haut en bas, on voit 4 pistes qui portent chacune deux images de la cible,

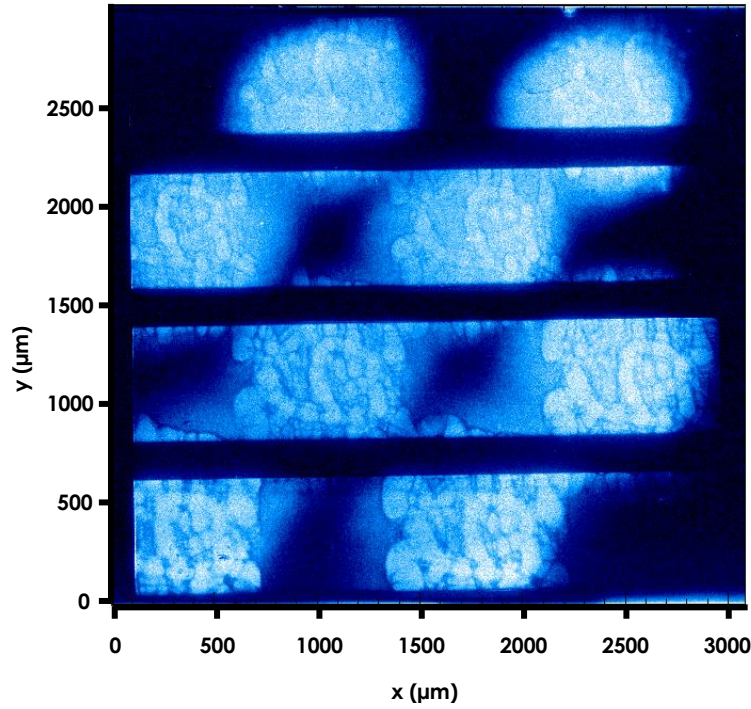


FIGURE 3.5 – Image type d’une radiographie XRF de face d’une cible présentant des modulations 3D du front d’ablation.

soit en tout 8 images. Chaque piste, comme expliqué dans la section 3.2, correspond à un instant de déclenchement de la mesure différent. On distingue d’ailleurs des variations des structures observées entre deux bandes. On peut voir aussi que d’une part les spots lumineux ne sont pas centrés au même endroit d’une piste à l’autre, ce qui est dû à un effet de parallaxe du spot de la source de radiographie. D’autre part, les modulations sont toujours situées au même endroit d’une piste à l’autre. Il y a donc un décalage de la zone de la cible qui est radiographiée entre les différentes pistes. Or pour effectuer une analyse rigoureuse d’une radiographie, il faut comparer systématiquement les mêmes modulations d’une image de la cible à l’autre. C’est dans ce but que l’algorithme de corrélation croisée a été développé. La corrélation croisée temporelle notée CC de deux signaux f et g peut s’écrire

$$CC_{fg}(\tau) = \int_{-\infty}^{+\infty} f^*(t)g(t + \tau)dt \quad (3.1)$$

avec τ le décalage entre les deux signaux. Le maximum de cette fonction de corrélation croisée donne la valeur de τ , soit le décalage temporel pour lequel les deux signaux se ressemblent le plus, qu’on peut interpréter par exemple comme le retard du signal g sur le signal f . Un décalage négatif donnerait bien entendu un retard de f sur g .

Dans notre cas, nous avons deux images de la cible Im_1 et Im_2 à des instants différents, qui sont décalées spatialement l’une par rapport à l’autre. Nous choisirons, pour la cohérence de l’analyse, des images carrées de même taille. Une différence par rapport au cas défini

dans l'équation (3.1) réside dans le fait que les images sont des signaux 2D discrets et non continus car constitués d'un certain nombre de pixels. On peut donc définir la corrélation croisée discrète entre Im_1 et Im_2 ainsi :

$$CC_{12}(m, n) = \sum_{i=0}^N \sum_{j=0}^N \sum_{k=0}^N \sum_{l=0}^N Im_1(i, j) Im_2(k + m, l + n) \quad (3.2)$$

avec N le nombre de pixel des images dans une direction. Cependant, ce calcul sera lourd car pour chaque couple (m, n) , le programme devra effectuer N^4 opérations. On pourra donc s'appuyer sur le théorème de la corrélation croisée dans l'espace de Fourier qui donne

$$CC_{fg} = F^{-1} [F^*(f)F(g)] \quad (3.3)$$

$$CC_{12} = F^{-1} [F^*(Im_1)F(Im_2)] \quad (3.4)$$

avec F la transformée de Fourier directe et F^{-1} la transformée de Fourier inverse. Nous utiliserons donc l'équation (3.4) dans notre programme de corrélation croisée. Le détail du programme est représenté en figure 3.6 à l'aide d'un exemple. Trois signaux sont utilisés en entrée : la radiographie qui comporte les 8 images de la cible, et deux sous-images positionnées sur des images de la cible. Une des sous-images correspond à l'image de référence (1) : l'autre sous-image (2) va être redimensionnée pour atteindre la taille de l'image de référence. On notera qu'il faut choisir une image de référence de forme carrée : ainsi, une analyse par transformée de Fourier rapide ("Fast Fourier Transform", FFT) se fera sur le même nombre de pixels quelle que soit la direction, et donc la décomposition en fréquence du signal se fera sur le même ensemble. Une fois la 2ème image redimensionnée (3), la corrélation croisée est réalisée avec l'image de référence. On obtient une matrice de corrélation de la même taille que les deux images (ici 250*250 pixels). La position du maximum de cette matrice donne le décalage à effectuer sur l'image 3. Si le maximum se situe en $(0,0)$, les deux images sont parfaitement corrélées. Dans notre exemple, le maximum est en $x=73$, $y=244$. Cependant, cela peut aussi signifier qu'un décalage de $x-250=-177$ et/ou $y-250=-6$ doit être effectué. Le programme décale donc l'image 3 de 73 vers la droite et 244 vers le haut, puis l'image obtenue (4) est corrélée avec l'image de référence. On décale ensuite la position des éléments de la matrice de corrélation de vérification ainsi obtenue de 125 pixels dans les deux directions. Si le maximum obtenu est au milieu de la matrice de corrélation de vérification, cela signifie que la corrélation a bien fonctionné. Ce n'est pas le cas dans notre exemple. On peut d'ailleurs voir que l'image 4 ne ressemble en rien à l'image 1. Le programme va donc essayer avec de décaler l'image 3 de $(x-250, y)$, $(x, y-250)$ et $(x-250, y-250)$, et va calculer la matrice de corrélation de vérification dans chaque cas. Si aucune de ces matrices ne possède un maximum en son milieu, le programme envoie en sortie un message d'erreur. Dans notre cas, un décalage de l'image 3 de $(x, y-250)$ (image 5) donne une matrice de corrélation de vérification

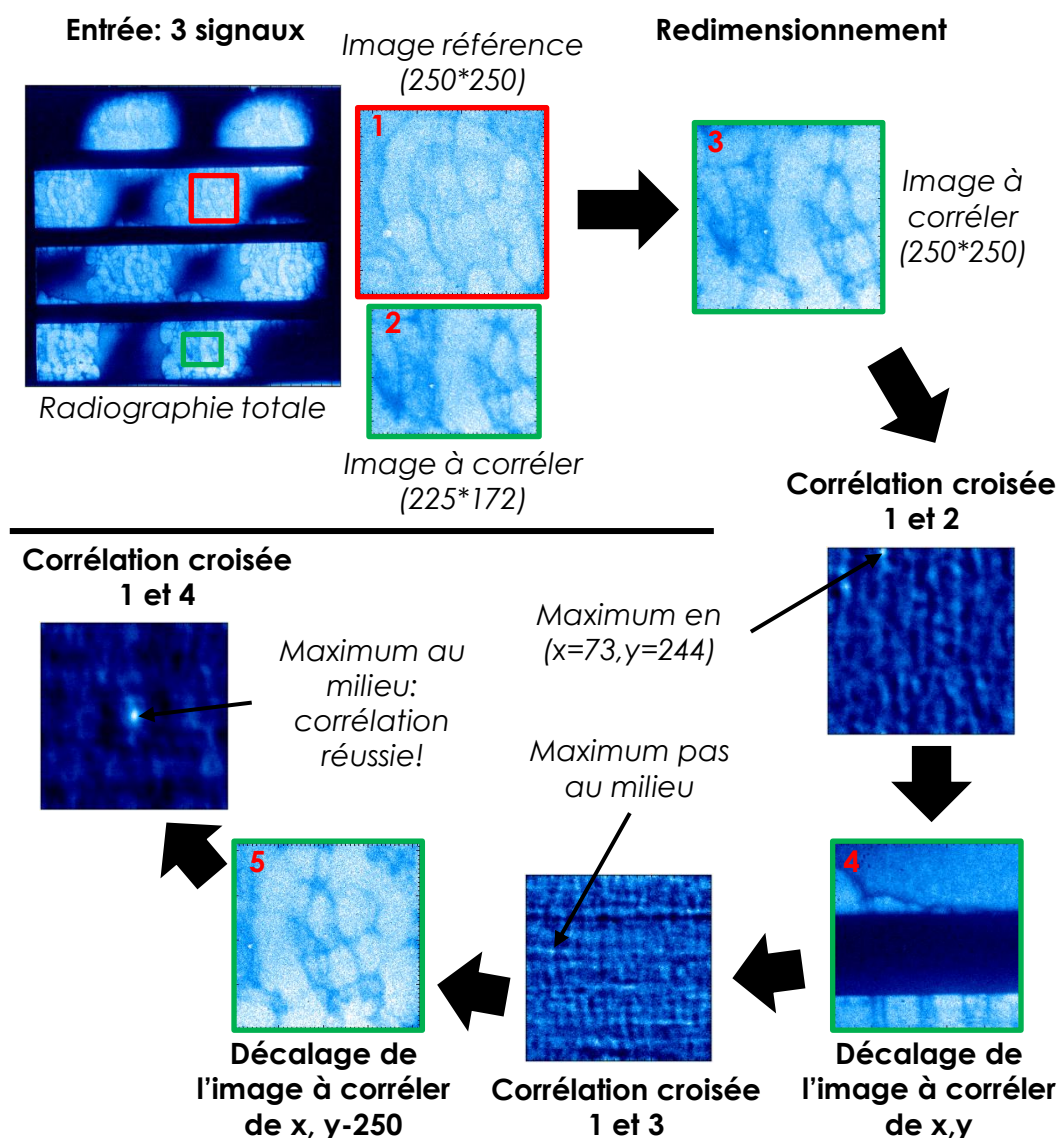


FIGURE 3.6 – Exemple du fonctionnement du programme de corrélation croisée appliqué à une radiographie XRF de face provenant d'expériences non décrites dans cette thèse réalisées sur OMEGA EP en 2011.

dont le maximum apparaît très nettement en milieu d'image (cf figure 3.6). La corrélation croisée a donc bien fonctionné ; l'image 5 et l'image 1 sont de même taille et les structures observées dessus présentent une grande ressemblance, à l'évolution due aux instabilités hydrodynamiques près.

3.3.2 Calcul du niveau de modulation

Pour analyser les images sélectionnées par corrélation croisée, on peut utiliser plusieurs méthodes. L'une d'elle est l'analyse de Fourier, en utilisant des algorithmes de FFT. Elle s'appuie sur le fait que tout signal peut être décomposé en une somme de sinusoïdes

de fréquences différentes. La transformation de Fourier permet donc de passer de l'espace réel (signal) à l'espace de Fourier (fréquences et amplitudes des sinusoïdales). Les algorithmes de FFT permettent de réaliser les calculs de transformation de Fourier rapidement en s'appuyant sur diverses méthodes (Cooley-Turkey, Good-Thomas, ...) et sont donc particulièrement adaptés au monde numérique.

Dans le chapitre 2, nous avons vu par exemple qu'en phase linéaire de l'IRT, les différentes longueurs d'onde évoluent indépendamment les unes des autres. Cela illustre bien l'intérêt de décomposer les différentes fréquences des images de la cible obtenues par les XRFC. Nous utiliserons deux types de FFT : FFT 1D et FFT 2D. Pour réaliser un traitement par FFT 1D d'une image XRFC, plusieurs étapes sont nécessaires. Le fonctionnement du programme est détaillé en figure 3.7 à l'aide d'un exemple de radiographie portant des modulations 2D. Le programme utilise en entrée l'image à analyser, la longueur d'onde supposée des modulations et l'angle de rotation pour que les modulations soient horizontales. L'image à analyser (1) est donc tournée (2), ici d'un angle de 30° . Puis un profil orthogonal aux modulations est calculé (3), en moyennant le profil sur la moitié de la largeur de l'image, ce qui permet de supprimer une partie du bruit. La largeur du profil n'est que de la moitié de l'image pour ne pas prendre en compte les coins de l'image tournée. Un fit polynomial de coefficient 4 du profil est réalisé (4) : le faible coefficient permet de ne prendre en compte que les très grandes structures d'intensité [79]. Ces très grandes structures correspondent à la tâche d'éclairement de la source de radiographie ; les plus petites structures dues aux modulations du front d'ablation ne sont pas prises en compte dans le fit. Le profil (3) est ensuite divisé par le fit (4) : on obtient ainsi une courbe (5) dont les variations autour de 1 sont directement liées aux modulations de densité surfacique de la cible (détaillé dans le chapitre suivant). On prend alors une partie de cette courbe (6) de taille égale à un nombre entier de fois la longueur d'onde supposée des modulations nommée λ_0 ($\lambda_0 = 70 \mu\text{m}$ dans cet exemple). En effet, le pas en fréquence spatiale (appelé df) du spectre de Fourier qui sera obtenu après FFT est défini par

$$df = \frac{1}{n_{pas} pas} \quad (3.5)$$

le pas étant la taille d'un pixel et n_{pas} le nombre de pixels dans le signal 6 analysé par la FFT. L'abscisse de ce signal est définie sur plusieurs longueurs d'onde ; on a donc

$$n_{pas} = \frac{N\lambda_0}{pas} \quad (3.6)$$

avec N un entier. En introduisant (3.6) dans (3.5), on trouve

$$df = \frac{1}{N\lambda_0} = \frac{f_0}{N} \quad (3.7)$$

avec f_0 la fréquence spatiale des modulations de longueur d'onde λ_0 . La fréquence spatiale supposée des modulations apparaît donc dans le spectre de Fourier obtenu par la FFT

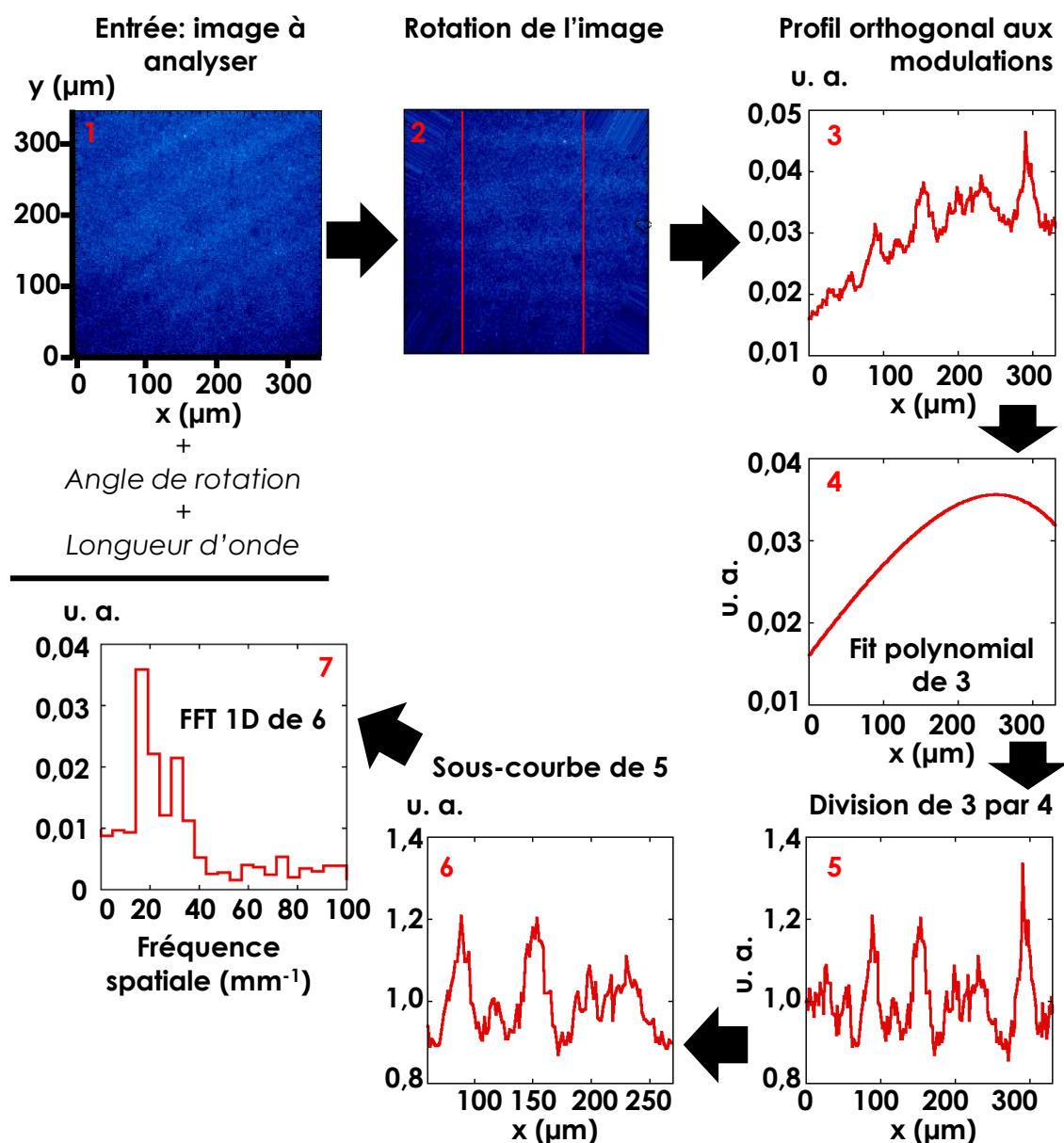


FIGURE 3.7 – Exemple du fonctionnement du programme de FFT 1D appliqué à une radiographie XRF de face provenant d'expériences décrites au chapitre suivant.

car elle est égale à un nombre entier fois le pas en fréquence.

Un algorithme de FFT existe déjà dans IDL : il est donc utilisé pour calculer le spectre de Fourier (7) du signal (6). On peut voir notamment deux pics sur ce spectre : le plus grand correspond à la longueur d'onde des modulations observées sur l'image de la cible ($70 \mu\text{m}$). Le pic secondaire, de longueur d'onde $35 \mu\text{m}$, correspond au second harmonique de ces modulations. Le programme va donner deux éléments en sortie : le spectre de Fourier et l'amplitude du pic de Fourier correspondant à la longueur d'onde supposée des modulations.

Le programme de FFT 1D est très adapté pour l'analyse de modulations 2D telles que présentées en figure 3.6. Pour une image comportant des modulations 3D de type bulles et aiguilles, une analyse par FFT 2D peut s'avérer utile. Un programme simple de FFT 2D a été développé ; il est utilisé sur deux exemples en figure 3.8. Comme pour le programme de FFT 1D, ce programme divise l'image par un fit d'elle-même. Puis la FFT 2D de l'image est effectuée grâce à l'algorithme de FFT existant dans IDL. La FFT 2D revient à faire la FFT de toutes les lignes et toutes les colonnes de l'image. Les images 1 et 4 présentent respectivement des modulations 3D et 2D. Les spectres de Fourier 2D correspondants sont représentés en 2 et 5. Un spectre de Fourier 2D se lit de la manière suivante : le centre de l'image correspond à la fréquence 0, puis la fréquence augmente linéairement vers l'extérieur de l'image dans toutes les directions. Pour un pixel donné, on obtient 3 informations : sa distance par rapport au centre de l'image donne la fréquence de la composante de Fourier qu'il représente, sa valeur donne l'amplitude de cette composante et la direction donnée par la droite qui passe par ce pixel et le centre de l'image correspond à la direction de cette composante, donc de la sinusoïde de cette fréquence et de cette amplitude. Ainsi, le spectre de Fourier 2 fait apparaître une sorte de "halo" autour du centre de l'image, avec une intensité qui paraît légèrement supérieure dans la direction gauche-droite. La halo traduit des composantes dans toutes les directions, donc un signal typiquement 3D, ce qui est le cas avec les bulles de l'image 1. La sur-intensité dans la direction gauche-droite correspond au fait que les bulles paraissent incluses dans des structures verticales. Le spectre de Fourier 5 ne fait apparaître de signal non négligeable que dans une direction, du coin haut gauche vers le coin bas droit de l'image : cela correspond à la direction des modulations 2D observées sur l'image 4. A partir des spectres de Fourier 2D, le programme réalise ensuite une moyenne azimuthale du signal. Les spectres équivalents 1D sont représentés en images 3 et 6. On peut voir que les deux spectres tendent vers une valeur constante de 0,001 u. a. vers les grandes fréquences spatiales : cela correspond au niveau de bruit. Pour le spectre moyenné 3, la plupart du signal se trouve à des fréquences spatiales inférieures à 30 mm^{-1} donc à des longueurs d'onde supérieures à $30 \mu\text{m}$. De plus, le signal est assez constant en deçà de cette fréquence : les structures de bulles sont donc de tailles variables mais supérieures à $30 \mu\text{m}$. Le pic observé à très basse fréquence n'a pas vraiment de sens physique : il correspond à une composante de la taille de l'image. Le spectre de Fourier 6 fait apparaître un pic correspondant à une longueur d'onde d'environ $70 \mu\text{m}$, comme avec le programme de FFT 1D. Cependant, ce pic est beaucoup moins net par rapport aux autres composantes du spectre qu'avec la méthode de FFT 1D ; sa valeur est divisée d'un facteur 10. Ce programme n'est pas adapté à l'analyse de modulations 2D ; il permet uniquement de comparer des radiographies portant des modulations 3D entre elles. De même, le programme de FFT 1D permet de comparer des radiographies de modulations 2D entre elles.

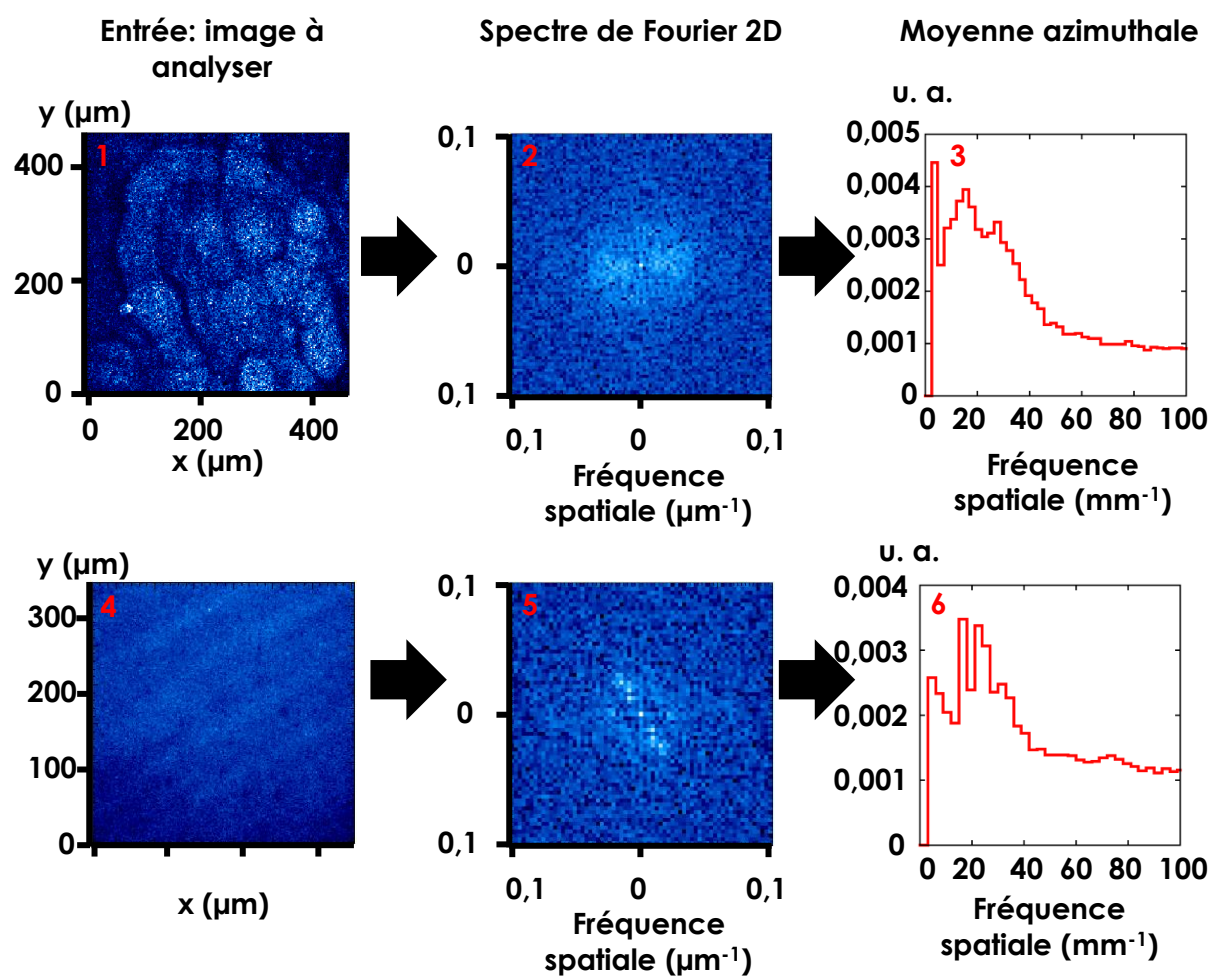


FIGURE 3.8 – Exemple de FFT 2D sur une image de modulations 2D et une image de modulations 3D.

3.3.3 Calcul de la barre d'erreur

Le programme de FFT 1D permet d'obtenir des mesures d'amplitude de modulations à un temps donné. Un programme a donc été développé pour calculer la barre d'incertitude sur cette amplitude. Ce programme fonctionne grâce à la combinaison de la corrélation croisée et de la FFT 1D. Il effectue tout d'abord la corrélation croisée de l'image d'intérêt et de l'image de référence pour récupérer la position de l'image décalée après corrélation. Le programme utilise en entrée, outre les signaux à corrélérer, la dimension de la zone de la cible autour de l'image décalée dans laquelle le programme va effectuer des prélèvements d'images, ainsi que l'angle de rotation et la longueur d'onde à laquelle on s'intéresse pour la partie du programme qui réalise les FFT 1D. Après avoir effectué la corrélation, le programme prélève un certain nombre d'images dans chaque direction, à intervalles réguliers, autour de l'image corrélée, jusqu'à atteindre les limites fixées par l'utilisateur. La limitation à une certaine zone par l'utilisateur permet que le programme ne prélève pas d'image dans une zone qui ne correspond pas à l'image de la cible, par exemple la bande sombre entre deux pistes. Chaque image prélevée va être analysée par FFT 1D ; il y aura alors une valeur d'amplitude de Fourier à la longueur d'onde d'intérêt pour chaque image prélevée. Le programme va ensuite calculer l'amplitude moyenne et l'écart-type à cette moyenne : c'est l'écart-type qui sera utilisé comme barre d'incertitude. Cette méthode est par exemple utilisée dans la réf. [79].

3.3.4 Segmentation d'image

Une méthode d'analyse particulièrement adaptée dans le cas de modulations 3D telles que représentées en figure 3.5 est la segmentation d'image. Cette méthode consiste à séparer chaque structure, à les compter et à en mesurer la taille. Un programme développé à cet effet est décrit en figure 3.9. Il s'appuie sur la même image que celle utilisée en figure 3.6. L'image à segmenter (1) est d'abord traitée : elle est divisée par un fit d'elle-même pour que le niveau moyen de l'image soit constant (et ne varie pas avec l'intensité de la tâche d'éclairement de la source de radiographie). L'image est ensuite lissée pour supprimer le bruit et le niveau d'intensité est multiplié par un coefficient choisi par l'utilisateur (2), avant que l'image soit "inversée" (3), les zones claires (sombres) devenant sombres (claires). Toutes ces opérations ont pour but de préparer l'image à la segmentation, qui est réalisée par un algorithme dit de "watershed" (ligne de partage de eaux, en anglais). Cet algorithme est déjà existant dans IDL, ce qui a permis de gagner du temps. Cependant, il n'est pas possible de modifier cet algorithme, qui possède donc une sensibilité au contraste de l'image qui lui est propre ; c'est pour cela que le niveau d'intensité doit être multiplié par un facteur qui varie en fonction de l'image à analyser, dont le contraste peut varier. Le fonctionnement de l'algorithme de "watershed" [80] est détaillé en figure 3.10, à l'aide d'une coupe de l'image traitée précédemment. La méthode décrite sur la coupe est appliquée

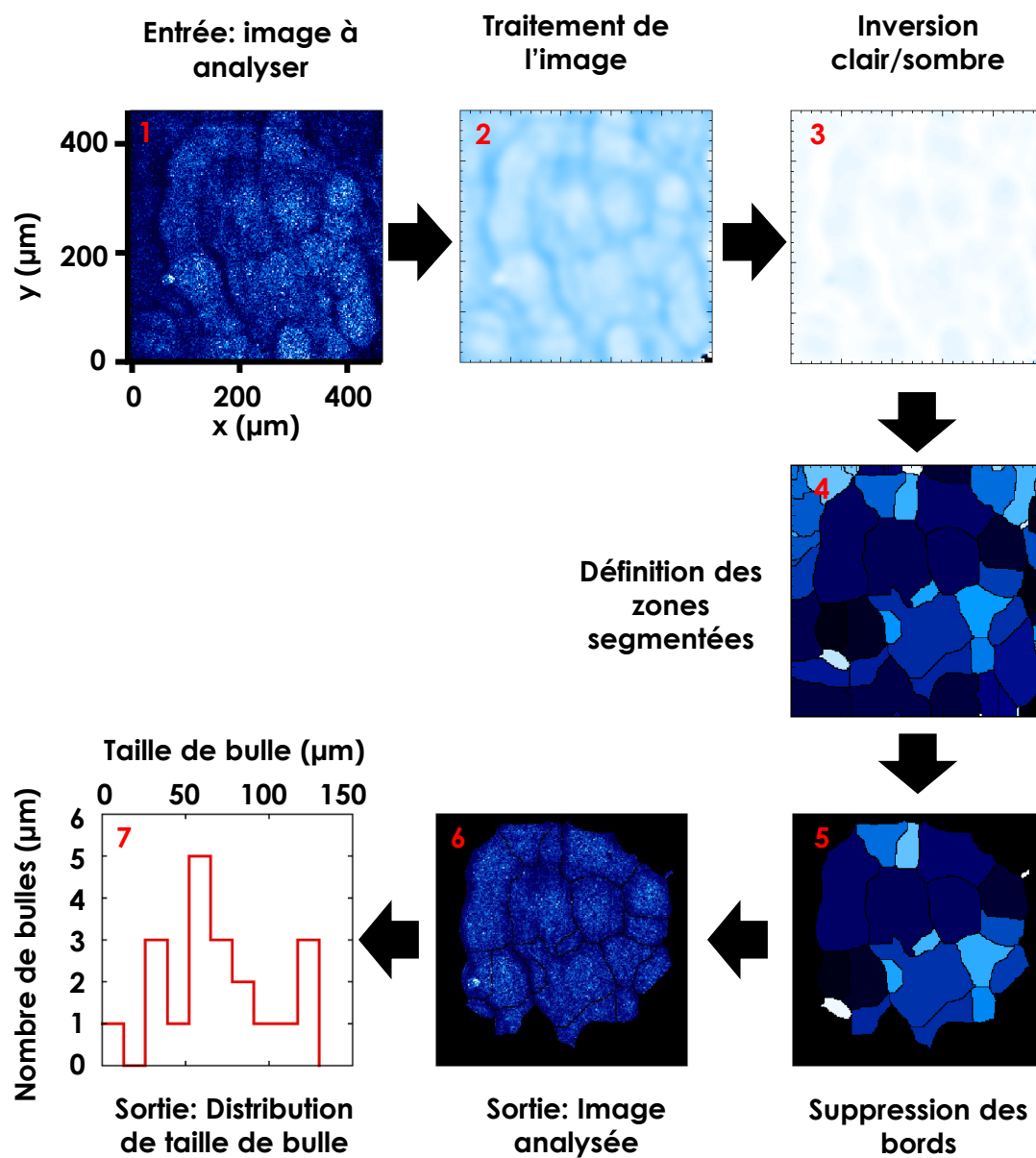


FIGURE 3.9 – Exemple du fonctionnement du programme de segmentation d'image en utilisant l'image de référence de la figure 3.6.

à toute l'image. La segmentation par "watershed" peut être envisagée par une analogie géographique : l'image peut être vue comme une surface terrestre vue de dessus. Plus une zone est sombre, plus elle est profonde. Les minima locaux définissent des bassins (2). Tous ces bassins vont être progressivement remplis d'eau (3). Lorsque l'eau de deux bassins se rejoint, on atteint une ligne de partage des eaux (4) : cela correspond à la zone où l'image va être segmentée. La segmentation sera terminée lorsque tous les bassins se seront rejoint, donc que toute l'image sera segmentée (5). On peut voir pourquoi il est nécessaire de lisser l'image : du fait du bruit, un très grand nombre de minima locaux supplémentaires va apparaître, et l'image va donc être segmentée en de très nombreux morceaux qui n'auront aucune signification physique. De plus, on inverse les zones sombres et claires de l'image pour que les bulles correspondent aux bassins et que la segmentation se fasse sur le contour des bulles. Si l'on reprend la figure 3.9, l'image a donc été segmentée par l'algorithme de "watershed" (4). Chaque bassin se voit affecter un niveau d'intensité uniforme et unique, tandis que les contours sont d'intensité 0. Ensuite, tous les bassins en contact avec le bord de l'image vont se voir affecter une intensité égale à 0 (5) ; en effet, si une bulle est coupée par le bord de l'image, on ne peut connaître sa taille totale et donc prendre cette bulle en compte fausserait les statistiques. On peut alors superposer toutes les zones d'intensité 0 sur l'image de base, ce qui permet de donner un aperçu de sa segmentation (6). Pour compter le nombre de bulles, le programme va compter le nombre de niveaux d'intensité différents sur l'image 5. Pour connaître la taille de chaque bulle, tous les pixels de même niveau d'intensité sont sommés. A partir de ce nombre de pixel et de la résolution de l'image, on peut calculer la taille de la bulle, ce qui permet ensuite de tracer la distribution de taille des bulles (7).

3.4 Code d'hydrodynamique radiative CHIC

Le code CHIC, développé au CELIA, est un code d'hydrodynamique radiative lagrangienne, bidimensionnel en géométrie axi-symétrique [81, 82]. Ce code comprend une modélisation fluide à deux températures, électronique et ionique, sur maillage non-structuré lagrangien avec reprojexion sur une grille eulérienne (ALE pour Arbitrary Lagrangian-Eulerian). Pour traiter le transport électronique, il inclut le modèle de diffusion de Spitzer-Härm avec limitation de flux. La propagation des faisceaux laser à travers la couronne de plasma est calculée par un algorithme de lancer de rayon tridimensionnel. Le transport de photons est décrit par un modèle de diffusion multigroupe. La loi d'Ohm généralisée permet de calculer les champs magnétiques auto-générés dus aux gradients croisés de densité et de température. Ce code est utilisé couramment dans le dimensionnement et l'interprétation d'expériences et pour le design de cibles de FCI. Les principaux intérêts de ce code dans l'étude des instabilités hydrodynamiques sont la richesse des modèles physiques ainsi que l'accès à la phase non-linéaire de l'instabilité.

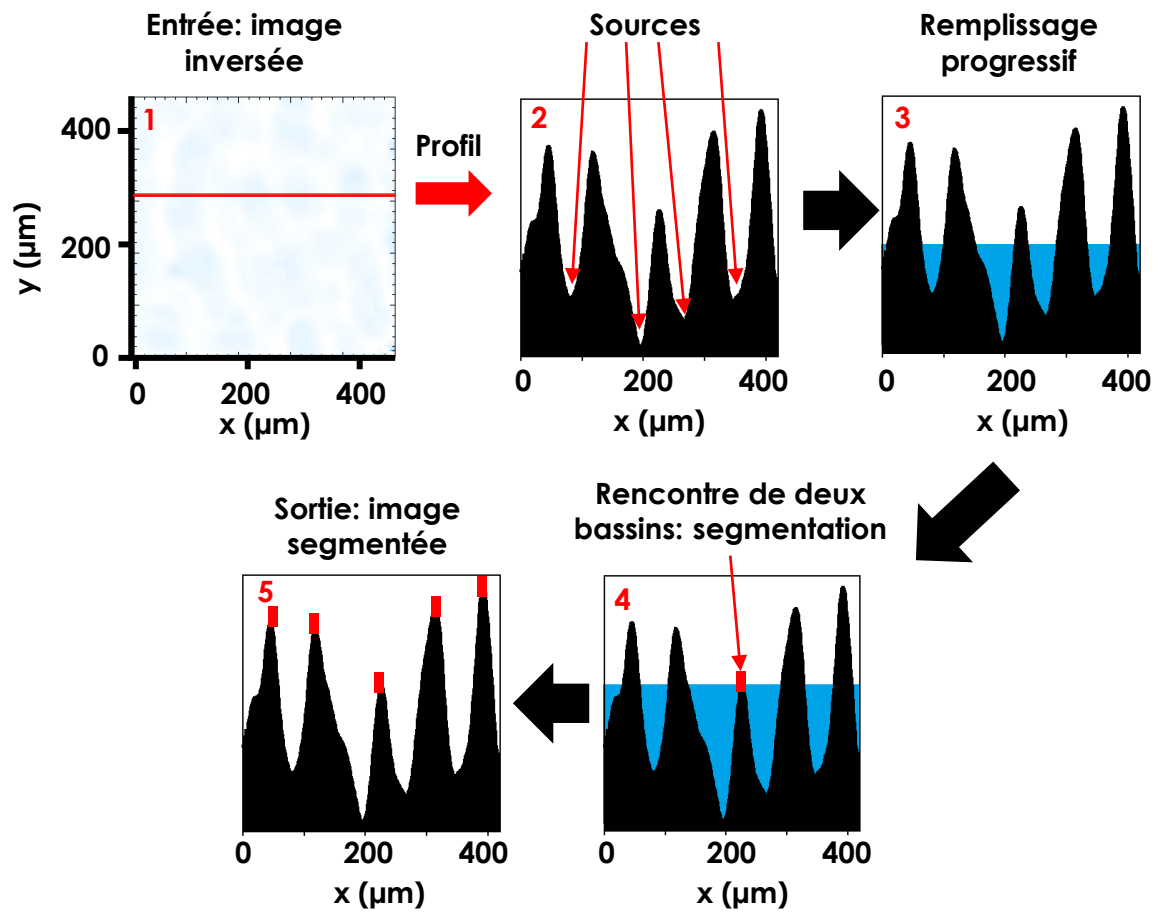


FIGURE 3.10 – Fonctionnement de l'algorithme de "watershed" utilisé dans le programme de segmentation d'image décrit en figure 3.9.

Nous avons effectué trois types de simulations avec CHIC : des simulations 2D utilisant une description réaliste des faisceaux et de la tache focale pour concevoir et interpréter les expériences, des simulations 1D pour étudier la dynamique de la cible (débouché de choc, accélération du front d'ablation, ...) et déterminer diverses quantités hydrodynamiques (densité, pression, température mais aussi vitesse acoustique ou vitesse d'ablation), et des simulations 2D monomodes de l'écoulement perturbé pour calculer la croissance de modulations du front d'ablation. Pour illustrer, la figure 3.11 présente un exemple de ces différents types de simulations. Dans les trois cas, il s'agit d'une cible de CH_2 de $15\ \mu\text{m}$ d'épaisseur illuminée à $5.10^{13}\ \text{W}/\text{cm}^2$. Pour la simulation de croissance de modulation, la longueur d'onde de la modulation imprimée est de $30\ \mu\text{m}$. Les profils de densité de ces trois types de simulations ont été représentés à 0, 400 et 1000 ps. Le laser se propage de la droite vers la gauche, selon un axe que l'on va appeler x . Dans le cas de la simulation 1D, sur l'exemple présenté en figure 3.11, il y a 150 mailles selon l'axe x et une seule selon l'axe y (axe vertical). La taille de la maille en y n'a pas d'importance ; à $t = 0$, une puissance laser uniforme est appliquée sur la première maille.

Dans le cas d'une simulation de croissance de modulations, qui est le type de calcul le plus fondamental dans le cadre de cette thèse, l'écoulement 1D est appliqué sur plusieurs mailles en y , avec une modulation de la puissance laser selon l'axe y . La modulation de la puissance laser est de type sinusoïdal, avec, sur la figure 3.11, le maximum de puissance vers le haut et le minimum vers le bas. La figure 3.12 schématise une cible maillée pour ce type de calcul, en faisant apparaître différents paramètres à définir. Entre 30 et 40 mailles par longueur d'onde ont été utilisées dans ces simulations. La convergence des calculs présentés ultérieurement a été vérifiée concernant le nombre de mailles, c'est-à-dire qu'ajouter des mailles ne modifie pas sensiblement le résultat des simulations. Le schéma de conduction utilisé est à 9 points, ce qui signifie que les électrons de conduction peuvent passer d'une maille à ses voisines horizontales, verticales et aussi en diagonale. Le limiteur de flux généralement utilisé est de type "sharp cut-off" à 7 %, ce qui est une valeur couramment utilisée pour des simulations dans la gamme d'intensité de nos expériences (quelques $10^{13}\ \text{W}/\text{cm}^2$ à quelques $10^{14}\ \text{W}/\text{cm}^2$). Selon l'axe x , l'épaisseur entière de la cible est modélisée tandis que selon l'axe y , on travaille uniquement sur une demie longueur d'onde de la modulation de l'intensité laser, ce qui permet de réduire le temps de calcul. Concernant les conditions aux limites hydrodynamiques, à droite et à gauche de la cible, c'est-à-dire les zones dans lequel le plasma et la cible se détendent, la pression est fixée à la pression initiale. En haut et en bas, des conditions de symétrie sont nécessaires pour se retrouver hydrodynamiquement dans un cas équivalent à une répétition infinie de période entière de modulations. Il a été vérifié que, dans nos configurations, pour un calcul initialement monomode, réaliser le calcul sur une demie longueur d'onde de modulation de

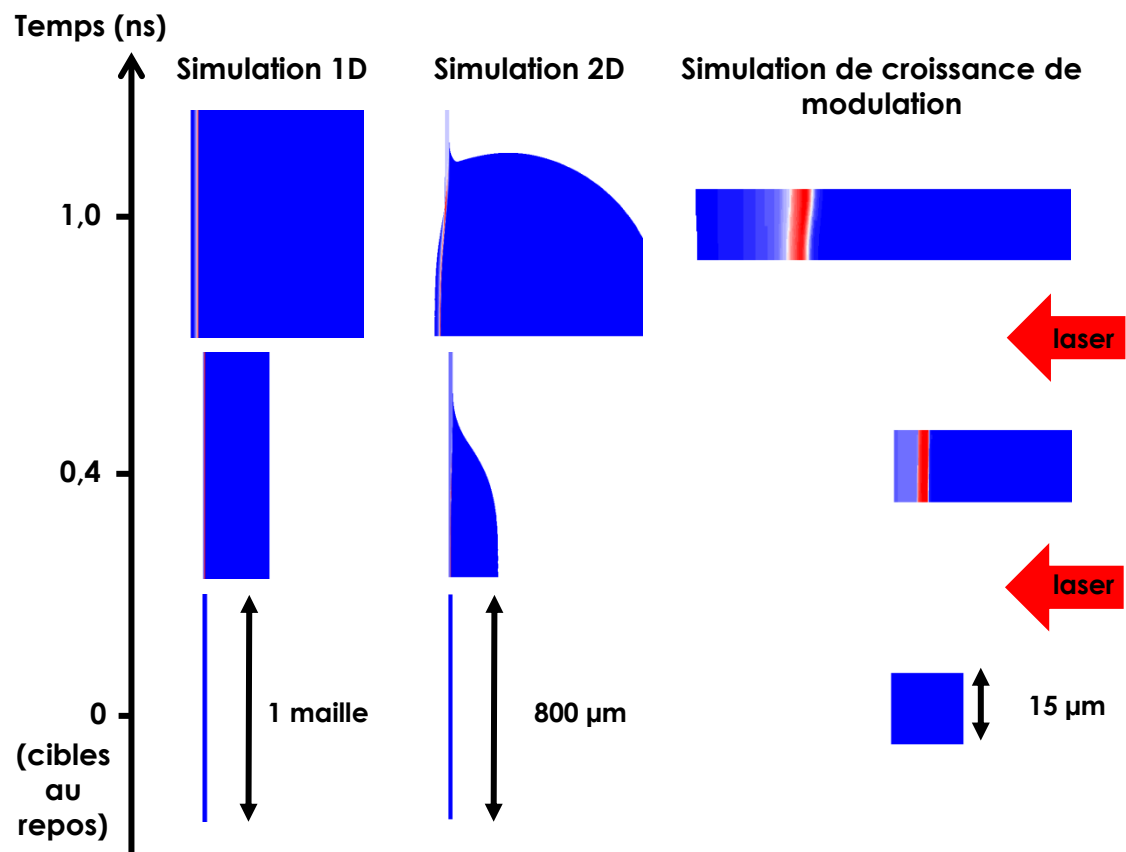


FIGURE 3.11 – Profils de densité issus de trois types de simulations : simulation 1D, simulation 2D et simulation de croissance de modulation de l'intensité laser imprimée. La cible, d'épaisseur $15 \mu\text{m}$, est constituée de CH_2 . L'intensité laser est de $5 \cdot 10^{13} \text{ W/cm}^2$ et la modulation de l'intensité laser est de longueur d'onde $30 \mu\text{m}$. Les simulations sont représentées à 0, 0,4 et 1,0 ns.

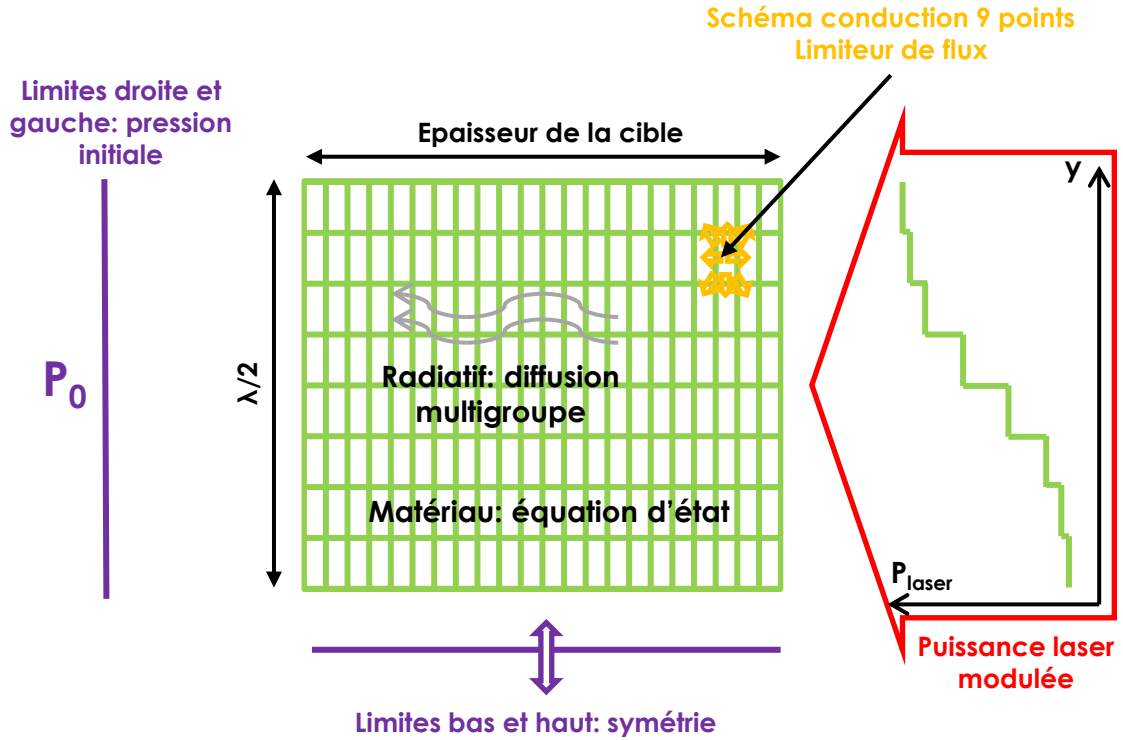


FIGURE 3.12 – Schéma des différents paramètres utilisés lors d'une simulation de croissance de modulation de l'intensité laser imprimée.

l'intensité laser était équivalent à effectuer le calcul sur une ou plusieurs longueurs d'onde.

Enfin, des simulations 2D de toute la cible en géométrie axi-symétrique ont aussi été réalisées, comme on peut le voir sur la figure 3.11. Dans ces simulations, les faisceaux laser sont modélisés un par un et chacun est pourvu de la loi de puissance mesurée expérimentalement. De plus, chaque faisceau est aussi pourvu de la tache focale induite par la lame de phase qui lui est associée. Un grand nombre de mailles est aussi nécessaire selon l'axe y (150 dans l'exemple présenté).

Deux modules de post-traitement des simulations CHIC ont principalement été utilisés. Le premier, utilisé avec les simulations de croissance des modulations, intègre la densité sur les différentes mailles selon l'axe x (axe perpendiculaire à la cible au repos) pour calculer les modulations de densité surfacique. Le repérage du front d'ablation, en calculant par exemple le minimum du gradient de densité, permet aussi de calculer l'amplitude des modulations du front d'ablation. Le deuxième module, associé aux simulations 1D, permet de calculer différentes quantités en se basant sur le repérage du front d'ablation de la même manière que pour le premier module. Ce module calcule la masse ablatée en fonction du temps en intégrant toute la masse de plasma éjectée à partir du front d'ablation. En dérivant cette masse ablatée par rapport au temps, le taux de masse ablaté

est calculé, ainsi que la vitesse d'ablation en divisant ce taux de masse ablatée par la densité du front d'ablation. Ce module calcule aussi, entre autres, la vitesse et l'accélération du front d'ablation. Ces différents paramètres sont utilisés pour calculer les taux de croissance et le modèle de Goncharov.

Dans ce chapitre, nous avons présenté le laser OMEGA : cette installation permet de réaliser des expériences nanosecondes à plusieurs dizaines de kJ dans des géométries variées et possède de très nombreux diagnostics. C'est donc une installation idéale pour réaliser des expériences d'instabilités hydrodynamiques. Nous avons développé de nombreux programmes pour dépouiller les données mesurées lors d'expériences sur cette installation, entre autre un programme de corrélation croisée et plusieurs autres d'analyse de Fourier. De plus, le code d'hydrodynamique radiative CHIC du CELIA permet de simuler les expériences d'instabilités hydrodynamiques réalisées sur OMEGA.

Chapitre 4

Mesure de la phase Richtmyer-Meshkov imprimée par laser

Comme on a pu le voir au cours du chapitre 2, l'IRM ablative a été mesurée sur le laser Nike KFr [10] et le laser OMEGA [66]. Cependant, dans les deux cas, les perturbations du front d'ablation sont issues de modulations usinées sur la cible. L'IRM ablative imprimée par des modulations de l'intensité laser n'a jamais été mesurée : réaliser cette mesure est l'objectif de ce chapitre. L'IRM ablative n'est pas instable et provoque l'oscillation des modulations du front d'ablation : ainsi, en comparaison avec l'IRT ablative, les modulations vont être de plus petite amplitude. D'autre part, plus les cibles sont épaisses, plus la durée de transit de choc est longue, et donc plus la phase Richtmyer-Meshkov dure longtemps. Mais plus la cible est épaisse, plus l'amplitude des modulations représente un faible pourcentage de l'épaisseur de la cible et donc plus la mesure de ces modulations par radiographie de face sera délicate. Ces contraintes nous ont donc poussé à utiliser plusieurs épaisseurs de cible et à optimiser les couples épaisseur de cible/source de radiographie. L'autre point important à définir dans cette expérience est les conditions d'empreinte, qui doivent suivre deux conditions majeures : le motif d'empreinte doit être fiable, pour ne pas varier d'un tir à l'autre, et ce motif doit imprimer des modulations faciles à analyser et à comparer à la théorie (comme dans la réf. [10] par exemple). Dans ce chapitre, nous allons donc présenter la configuration expérimentale issue de ces différentes contraintes et le résultat des mesures effectuées dans cette configuration.

4.1 Configuration expérimentale

Nous avons vu au cours du chapitre précédent que 7 domaines étaient à définir dans les SRF pour une campagne sur le laser OMEGA : "General", "Drivers", "Target", "Beams",

"TIM", "Fixed" et "Neutronics". Une expérience est donc décrite de manière complète par ces différents éléments. Après avoir présenté les contraintes et le choix de conception que nous avons effectués, nous allons définir la partie "Target" décrivant les diverses cibles utilisées au cours de ces campagnes. Ensuite, les différents faisceaux nécessaires seront définis, ce qui permet de remplir les parties "Drivers" et "Beams". Pour finir, les différents diagnostics utilisés (partie "TIM" et "Fixed") seront évoqués.

4.1.1 Contexte de conception de l'expérience

En septembre 2011, A. Casner et nos collègues américains ont réalisé une expérience pour mesurer l'IRM ablative imprimée par laser sur OMEGA EP. Un faisceau défocalisé ne portant pas de lame de phase illuminait la cible, imprimant des modulations 3D. Ces modulations apparaissaient clairement lors de la phase Rayleigh-Taylor ablative. Cependant, il n'y eu pas de mesure de qualité dans la phase Richtmyer-Meshkov ablative du fait d'un niveau d'empreinte trop faible. C'est donc pour obtenir un niveau d'empreinte suffisant que le choix fut fait lors des expériences sur OMEGA d'utiliser, en plus de certains tirs avec empreinte de modulations 3D (cf paragraphe 4.2.2), des lames de phase créant spécifiquement des modulations d'intensité marquées. De plus, ces lames de phases doivent créer des modulations 2D supposées a priori monomodes ; la comparaison des données expérimentales au modèle de Goncharov, qui exprime l'évolution par l'IRM ablative de modulations d'une longueur d'onde donnée, en sera plus aisée.

Lors de la conception de l'expérience, nous avons choisi d'utiliser une caméra à images intégrales plutôt qu'une caméra à balayage de fente pour effectuer la radiographie de face de la cible, et ce malgré le fait qu'une CBF permette une mesure continue en temps. Plusieurs raisons nous ont orienté vers ce choix. Tout d'abord, lors des journées de tir où nous avons aussi étudié l'empreinte de modulations 3D, il n'était pas possible d'utiliser une caméra à balayage de fente qui ne peut être utilisée que pour la radiographie de modulations 2D. Il en est de même pour les expériences où nous avons étudié l'empreinte des modulations croisées avec des modulations préimposées. L'autre problème qui se pose est le pointage de la CBF, qui doit être exactement à la perpendiculaire des modulations. Lorsqu'on utilise des modulations préimposées sur la cible, il est assez facile d'avoir un repère visuel pour l'alignement. Cependant, dans le cas de modulations imprimées, qui ne sont donc pas présentes au début de l'expérience, le pointage est bien plus complexe. En effet, même si l'orientation des modulations 2D est connue par le positionnement de la lame de phase, l'exacte perpendicularité de la fente avec celles-ci sera délicate à obtenir. Enfin, la caméra à balayage de fente utilisée comme imageur pour des radiographies de face d'instabilités hydrodynamiques par O. V. Gotchev et ses collaborateurs dans la réf. [66] n'est plus utilisée de manière standard sur OMEGA. Ainsi, lors de la seule journée où

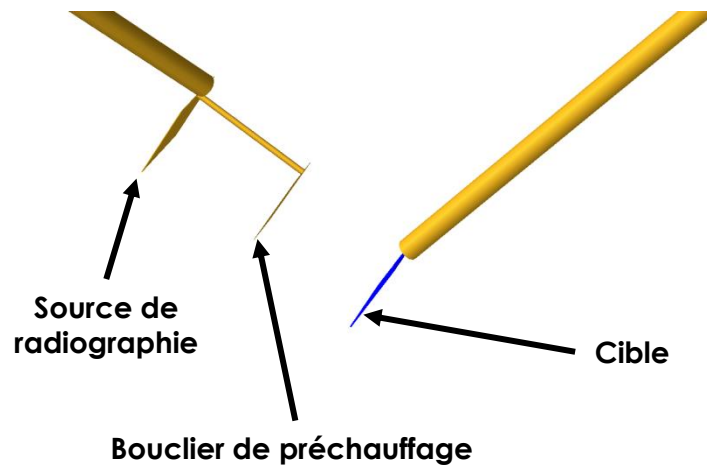


FIGURE 4.1 – Configuration des cibles utilisées lors des expériences d'IRM ablative imprimée par laser.

nous avons essayé d'utiliser une CBF pour la radiographie de face, après un début de journée prometteur (calibrations temporelle et spatiale réussies), nous n'avons pas mesuré de modulations imprimées suite à un mauvais alignement de la caméra. Nous sommes donc revenus en deuxième partie de journée à l'utilisation d'une caméra à images intégrales.

Enfin, il faut noter que chaque journée de tir a été précédée d'une série de simulations CHIC. En effet, celles-ci étaient nécessaires pour déterminer les instants de radiographie. On cherchait en effet à mesurer des modulations dans la phase Richtmyer-Meshkov ablative, donc avant le début d'accélération du front d'ablation. Il fallait néanmoins que l'amplitude des modulations soit suffisamment importante pour qu'elles puissent être mesurées. Ces conditions ont donc été déterminées à partir de simulations. De plus, lors des tirs avec les cibles les plus épaisses, on cherchait à mesurer une inversion de phase. L'instant d'inversion a lui aussi été déterminé avant l'expérience grâce à des simulations CHIC de croissance de modulations, pour pouvoir déclencher la mesure par l'imageur au bon moment.

4.1.2 Cibles

La configuration des différents éléments de type cibles est représentée en figure 4.1. Les cibles utilisées sont des feuilles de polystyrène (CH_2) de masse volumique $0,90 \text{ g/cm}^3$. Leurs dimensions sont de $3 \text{ mm} \times 3 \text{ mm} \times \text{épaisseur}$; l'épaisseur de la cible est de 15, 30, 50 ou $80 \mu\text{m}$ selon le tir. La cible est maintenue grâce à un porte-cible inséré dans le TIM 4 de la chambre. Pour réaliser une radiographie de face, des cibles pour générer des rayons X sont nécessaires. Ces cibles sont faites de feuilles de $3 \text{ mm} \times 3 \text{ mm} \times 25 \mu\text{m}$ en divers

matériau : uranium (U) avec émission à 1,3 keV, samarium (Sm) avec émission à 1,8 keV et tantale (Ta) avec émission à 2,5 keV. Plus une cible est fine et/ou plus l'énergie des photons émis par une source de radiographie est importante, plus la transmission de la cible est grande. Si l'on choisit un couple épaisseur de cible/énergie d'émission donnant une transmission trop faible, le signal récupéré par la XRFC sera trop faible, dans le niveau de bruit et la mesure sera ratée. Dans le cas contraire d'une transmission trop forte, la quantité de rayons X absorbée va peu varier entre une épaisseur de la cible au niveau d'un creux des modulations ou d'un pic : la sensibilité de la mesure sera donc faible. Il est donc nécessaire d'optimiser le couple épaisseur de cible/énergie d'émission. Un autre point doit être noté : l'IRM ablative ne crée pas de modulations de grande amplitude (quelques μm dans [10]). On ne peut donc pas utiliser une cible trop épaisse car l'énergie minimale des rayons X de radiographie nécessaire pour que ces derniers ne soient pas tous absorbés dans la cible serait alors trop importante pour avoir une bonne sensibilité pour des modulations de quelques μm . Lors de nos campagnes, des mesures réalisées avec des cibles de 100 μm d'épaisseur ont ainsi été infructueuses. D'un autre côté, une épaisseur de cible plus importante induit des temps de transit de choc et de remontée de l'onde de raréfaction plus grands et donc une période d'évolution de l'IRM ablative plus longue, ce qui est intéressant pour l'étudier. C'est pour cela que nous avons choisi différentes épaisseurs de cible. A partir d'éléments de la littérature [48], nous avons défini les couples épaisseur de cible/source de radiographie optimaux suivants pour les 2 premières journées de tir : U pour les cibles de 15 μm , Sm pour celles de 30 μm et Ta pour celles de 50 et 80 μm . L'utilisation de ces couples s'est avérée fructueuse lors de nos expériences. La transmission de ces différentes cibles est représentée en figure 4.2. Ainsi, $T = 0,30$ pour U/15 μm , $T = 0,39$ pour Sm/30 μm , $T = 0,55$ pour Ta/50 μm et $T = 0,38$ pour Ta/80 μm . Etant donné l'efficacité des radiographies malgré des transmissions différentes, nous avons utilisé lors des deux dernières journées de tir des sources en Sm pour les cibles de 15 et 50 μm , couples pour lesquels la transmission sera respectivement $T = 0,62$ et $T = 0,21$. Malgré ce changement, les radiographies obtenues ont encore été de bonne qualité.

Comme on peut le voir en figure 4.3 (a), les sources en U possèdent une bande assez intense de rayons X de quelques centaines d'eV. Pour éviter le préchauffage des cibles par ces rayons X, un "bouclier" d'aluminium (Al) de 1,5 mm * 1,5 mm * 6 μm est placé entre la source de radiographie et la cible. Ce filtre fait partie de l'ensemble du système d'imagerie. La courbe présentée en figure 4.3 (b) (issue, comme le spectre d'émission de l'uranium, de la thèse de V. A. Smalyuk [83]) montre la convolution du spectre d'émission de l'uranium par la réponse spectrale du système d'imagerie. Cette courbe correspond donc au spectre utilisé pour effectuer la radiographie de la cible. Si l'on excepte le petit pic aux environs de 3,5 keV, ce spectre est formé d'un pic de 600 eV de large centré

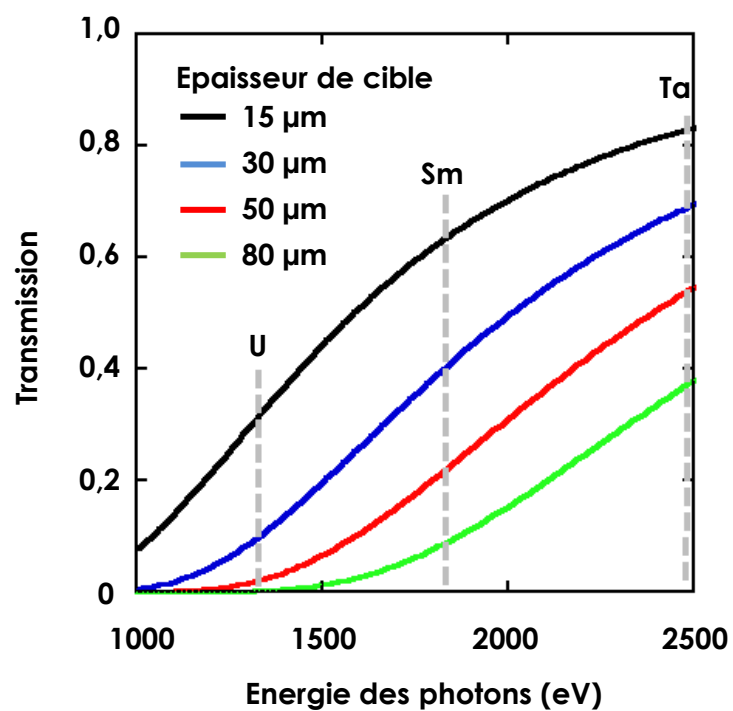


FIGURE 4.2 – Transmission en fonction de l'énergie des photons incident pour différentes épaisseurs de CH_2 . Les barres verticales représentent l'énergie d'émission des différentes sources de radiographie utilisées lors de nos expériences.

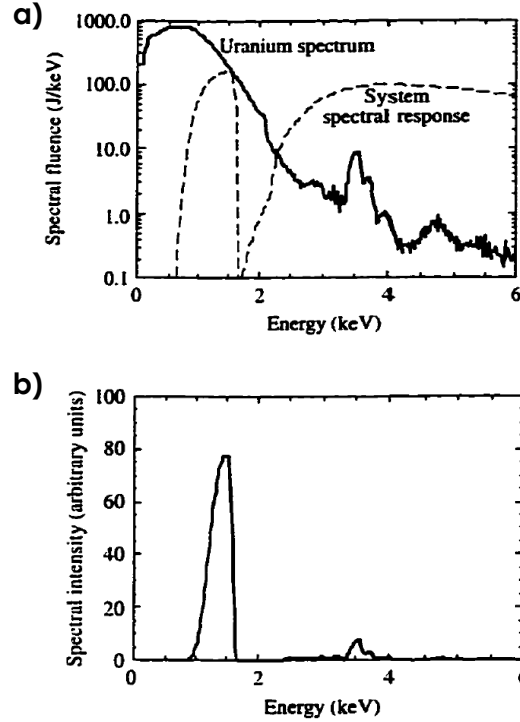


FIGURE 4.3 – a) Spectre d’émission de l’uranium (ligne pleine) et réponse spectrale du système d’imagerie (filtres et XRFC) (ligne tiretée). b) Spectre utilisé pour l’imagerie issue de la convolution du spectre de l’uranium et de la réponse spectrale du système d’imagerie. Ces deux graphes sont tirés de la thèse de V. A. Smalyuk [83].

autour de 1,3 keV. Nous étudierons brièvement dans le paragraphe 4.2.3 la validité de l’approximation faite lorsqu’on considère que l’émission de la source d’uranium peut être résumée à un pic monoénergétique de 1,3 keV.

4.1.3 Faisceaux laser sur cible

La configuration des faisceaux laser est représentée en figure 4.4. L’attaque de la cible est assurée par 3 faisceaux (faisceaux rouges et vert). Les deux faisceaux rouges sont utilisés avec des DPR et la DPP SG4 : le spot laser sur cible ainsi formé est de forme super-gaussienne d’ordre 4,2 avec un rayon de $352\ \mu\text{m}$. Ces faisceaux ont un angle d’incidence de 23° par rapport à la normale à la cible. Ils sont destinés à obtenir une intensité en rapport avec des conditions typiques de FCI. Le faisceau vert est nommé faisceau d’empreinte : il est destiné à créer les modulations du front d’ablation qui évolueront sous l’effet de l’IRM ablative. Pour cela, il est équipé d’une lame de phase spéciale, qui peut être soit la M30, la M60 ou la M120 en fonction du tir ; le faisceau portera alors des modulations de l’intensité 2D (cf figure 4.5 (a-c)) de longueur d’onde respectivement 30, 60 et $120\ \mu\text{m}$.

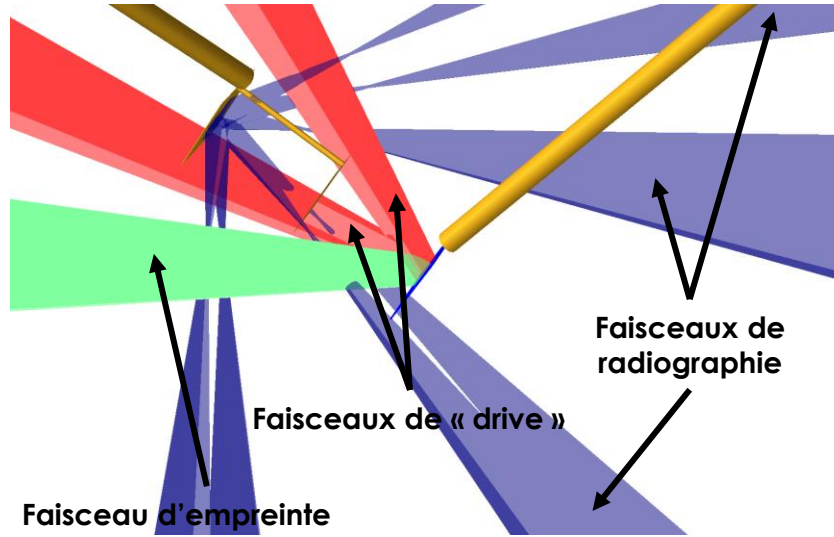


FIGURE 4.4 – Configuration expérimentale des expériences d’IRM ablative imprimée par laser. Les faisceaux de ”drive” (rouge), d’empreinte (vert) et de radiographie (bleu) sont représentés.

L’amplitude des modulations de l’intensité induites par ces lames de phase sera abordée à la fin de ce chapitre. Cependant, du fait de l’angle que forme le faisceau d’empreinte avec la cible, la longueur d’onde des modulations de l’intensité imprimées va être plus grande que la longueur d’onde nominale imposée par la lame de phase. Comme on peut le voir en figure 4.5 (d-e), la longueur d’onde imprimée va dépendre de deux angles : l’angle θ entre le faisceau et la normale à la cible, et l’angle ϕ entre les modulations et le plan perpendiculaire au plan de la cible dans lequel le faisceau est contenu, représenté par un axe en figure 4.5 (e). La longueur d’onde λ_i imprimée est donnée par

$$\lambda_i = \lambda_0 \sqrt{\frac{\sin^2 \phi}{\cos^2 \theta} + \cos^2 \phi} \quad (4.1)$$

avec λ_0 la longueur d’onde nominale de la lame de phase. On retrouve bien $\lambda_i = \lambda_0$ si le faisceau est perpendiculaire à la cible. Lors de nos expériences, $\theta = 23^\circ$ pour les faisceaux porteurs de la M30 et de la M120 et $\theta = 48^\circ$ pour celui porteur de la M60. De plus, $\phi = -32^\circ$ pour la M30 et la M60 et $\phi = 29^\circ$ pour la M120. On a donc $\lambda_i = 30,7\mu\text{m}$ pour la M30, $\lambda_i = 69,6\mu\text{m}$ pour la M60 et $\lambda_i = 122,5\mu\text{m}$ pour la M120. Pour certains des derniers tirs que nous avons réalisés avec la M30, la lame de phase avait dû être pivotée de 90° comme nous le verrons ultérieurement dans ce chapitre. Pour ces tirs, on avait donc $\phi = 58^\circ$ et donc $\lambda_i = 31,9\mu\text{m}$. Tous les faisceaux illuminant la cible sont issus du driver ”SSD” et l’impulsion portée par chaque faisceau est carrée de durée 3 ns, pour une énergie de 280 J/faisceau. Le faisceau d’empreinte est activé 200 ps avant les deux autres faisceaux pour que l’empreinte soit maximale ; en effet, la zone de conduction va se former moins vite si le faisceau d’empreinte est seul. L’avancer de plus de 200 ps ne va pas augmenter l’efficacité de l’empreinte d’après les auteurs de la réf. [61]. Le SSD n’a pas été utilisé dans

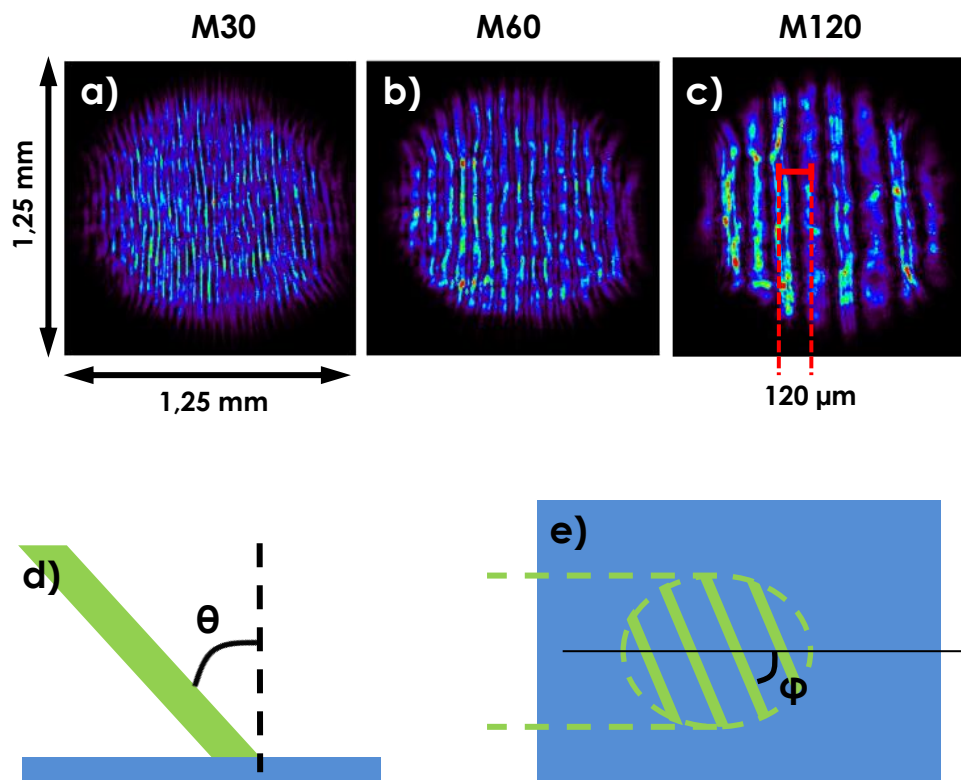


FIGURE 4.5 – Images obtenues par mesure ETP (Equivalent Target Plane) dans un plan perpendiculaire à la direction de propagation du faisceau pour les lames de phase a) M30, b) M60 et c) M120. Schéma représentant d) l'angle θ entre le faisceau et la normale à la cible et e) l'angle ϕ entre les modulations et le plan perpendiculaire au plan de la cible dans lequel le faisceau est contenu.

ces expériences pour maximiser l’empreinte. Sur la figure 4.4, les 7 faisceaux bleus sont les faisceaux de radiographie : ils illuminent la source de radiographie pour créer une émission de rayons X qui vont traverser la cible et être mesurés par une XRFC pour réaliser une radiographie de face de la cible. Ces faisceaux ne sont pas pourvus de DPP : le but est d’utiliser les points chauds non lissés des faisceaux pour créer un maximum de rayons X. Tous les faisceaux de radiographie ont une impulsion de 2 ns et une énergie de 330 J. Ces faisceaux sont défocalisés pour obtenir une tache focale d’environ $900\ \mu\text{m}$ de diamètre au niveau de la source de radiographie. Ils sont issus du "leg" 1 (cf chapitre précédent), qui est alimentée par le driver "Backlighter" ce qui permet d’utiliser une forme d’impulsion différente de celle des faisceaux qui illuminent la cible. Tous les faisceaux de radiographie sont concomitants, pour former une impulsion intense de 2 ns. Le début de cette impulsion est adapté à chaque tir, en fonction des instants auxquels on désire effectuer les mesures.

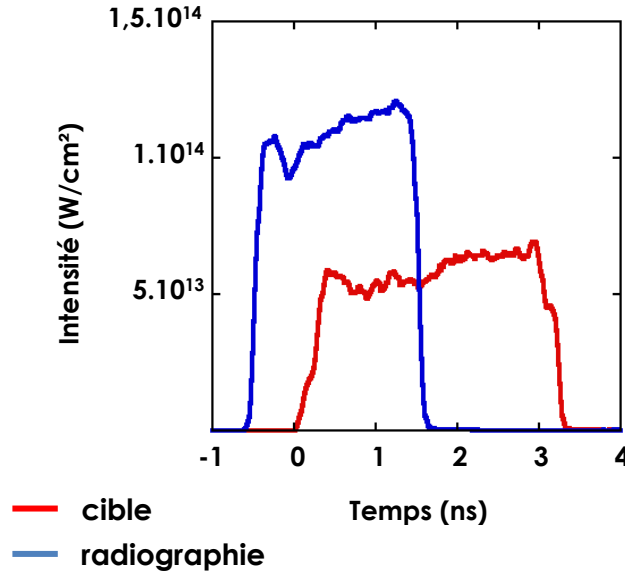


FIGURE 4.6 – Intensité totale sur cible (courbe rouge) et sur source de radiographie (courbe bleue) obtenue grâce aux diagnostics laser P510.

4.1.4 Radiographie de face et autres diagnostics

Lors de ces journées de tirs, notre diagnostic principal et primordial était une XRFC, placée dans le TIM 3 (port H18) et donc dans l'axe H3-H18, c'est-à-dire à la perpendiculaire de la cible, pour pouvoir réaliser des radiographies de face par l'utilisation couplée d'une source de radiographie. Ces radiographies ont été réalisées sur films Kodak T-Max 3200 [47]. Nous avons utilisé un grandissement de $12 \times$ et une grille de sténopés permettant de former 8 images par tir, deux par bande. Un filtre de $100 \mu\text{m}$ d'Al était utilisé avec les sources en U et un filtre de $25 \mu\text{m}$ de Be avec les autres sources. Les autres diagnostics utilisés n'avaient que des objectifs de contrôle : une autre XRFC ou une caméra à balayage de fente dans le TIM 5 pour vérifier la bonne émission de la source de radiographie, les différentes XRPHC pour vérifier aussi la présence de taches focales laser sur la cible et sur la source de radiographie. Enfin les diagnostics laser jouaient un rôle important car les P510 permettaient de connaître l'impulsion réellement délivrée par chaque faisceau, ce qui a permis d'utiliser l'impulsion laser réelle pour chaque tir dans les simulations numériques post-expérience. La figure 4.6 représente l'intensité totale sur cible et sur source de radiographie pour un tir type de la campagne expérimentale. L'allure temporelle de la puissance délivrée à la cible est mesurée par les caméras P510, puis nous avons calculé l'intensité à partir de la taille de la tache focale laser sur les cibles. On peut voir que l'impulsion sur cible ne varie que de quelques pour-cents durant les 3 ns ; de plus, l'impulsion est très stable d'un tir à l'autre.

4.2 Mesure de l'évolution des modulations de densité surfacique pendant la phase Richtmyer-Meshkov

4.2.1 Déroulement des journées de tir

Quatre journées de tirs ont eu lieu sur OMEGA pour étudier l'IRM ablative imprimée par laser : le 22/11/2011, le 02/05/2012, le 03/12/2013 et le 05/06/2014. Nous étions co-PI avec D. Martinez, avec qui j'ai personnellement dirigé les deux dernières journées de tir. Lors de la première journée, des cibles de 15, 30 et 50 μm d'épaisseur ont été imprimées à l'aide de la lame de phase M60 à une intensité d'environ 5.10^{13} W/cm^2 . Des tirs ont aussi été réalisés sans lame de phase pour effectuer une empreinte 3D de type speckle. Lors de la deuxième journée, les lames de phase M30 et M120 ont été utilisées. Lors de ces deux journées, la mesure de l'évolution temporelle des modulations de densité surfacique au front d'ablation se faisait par une XRFC. Lors des 3ème et 4ème journées de tirs, seule la M30 fut utilisée pour réaliser l'empreinte. En effet, une longueur d'onde plus petite induit une fréquence d'oscillation de l'IRM ablative plus petite et donne donc une plus grande probabilité d'observer une inversion de phase. Lors de la 3ème journée de tirs, le but était d'utiliser une caméra à balayage de fente pour effectuer la radiographie de face de la cible et donc pour mesurer de manière continue l'évolution des modulations. Cependant, une mauvaise orientation de la fente de la caméra n'a pas permis de mesure. Quelques tirs ont été fructueux en fin de journée grâce à un retour à une XRFC. Lors de la 4ème journée, des cibles de 30, 50 et 80 μm ont été utilisées, et l'intensité lors de certains tirs a été portée à environ 1.10^{14} W/cm^2 car la période d'oscillation de l'IRM ablative dépend des vitesses d'ablation et de "blow-off", et donc de l'intensité laser. Des modulations de 30 μm avaient aussi été usinées sur la cible, le but étant de réaliser l'empreinte par la M30 croisée à ces modulations préimposées, ce qui permet une analyse par FFT 1D dans chacune des deux directions une fois que le signal non désiré aura été retiré par un traitement d'image [61].

4.2.2 Exemples de radiographies de face

Les figures 4.7, 4.8, 4.9 et 4.10 présentent des exemples de radiographies de face effectuées durant ces journées de tir. On peut voir sur chaque radiographie les quatres bandes, correspondant à quatre instants différents, comportant chacune deux images de la cible. Le temps va du haut vers le bas et de la droite vers la gauche. Sur une même bande, les deux images sont séparées par environ 120 ps. Sur les figures 4.7, 4.8 et 4.9, on peut voir des radiographies de tirs avec empreinte réalisée respectivement grâce à la lame de phase M30, M60 et M120, pour des cibles respectivement de 30, 50 et 30 μm d'épaisseur. Dans les trois cas, on voit une cible qui n'est pas visiblement perturbée aux temps les plus courts, puis on voit les modulations apparaître. On remarque notamment que pour les cibles de 30 μm d'épaisseur, les modulations sont visibles à l'oeil nu autour de l'instant

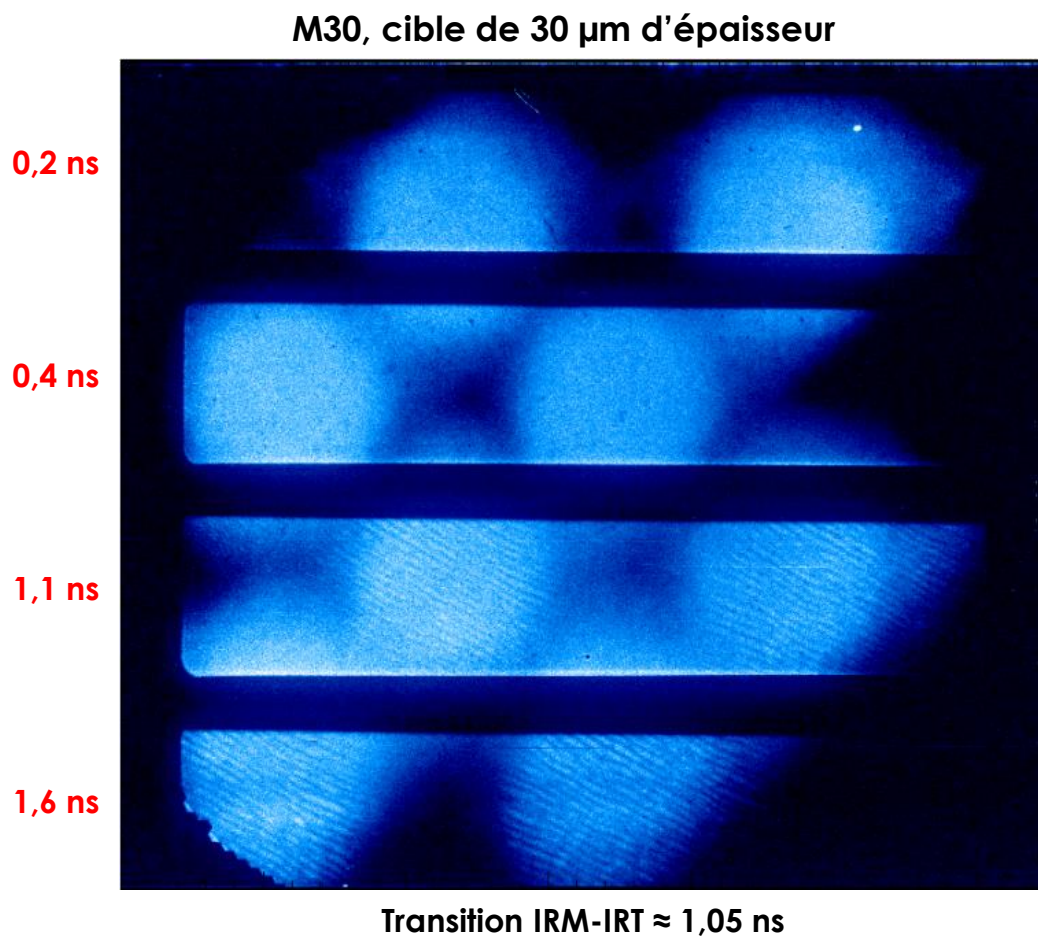


FIGURE 4.7 – Radiographie de face d'une cible de 30 μm d'épaisseur illuminée par deux faisceaux porteurs de la lame de phase SG4 et un faisceau d'empreinte porteur de la lame de phase M30. Les mesures débutent à partir de la 3ème piste.

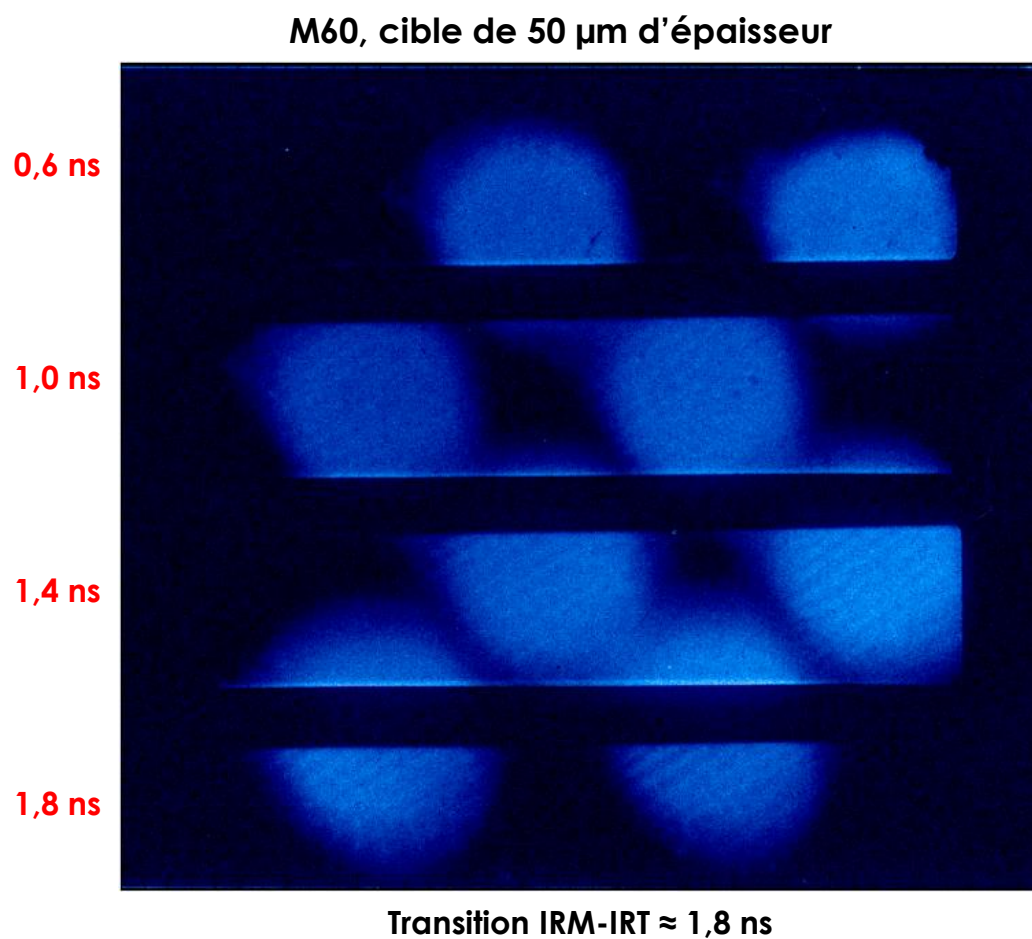


FIGURE 4.8 – Radiographie de face d'une cible de 50 μm d'épaisseur illuminée par deux faisceaux porteurs de la lame de phase SG4 et un faisceau d'empreinte porteur de la lame de phase M60. Les mesures débutent à partir de la 2ème piste.

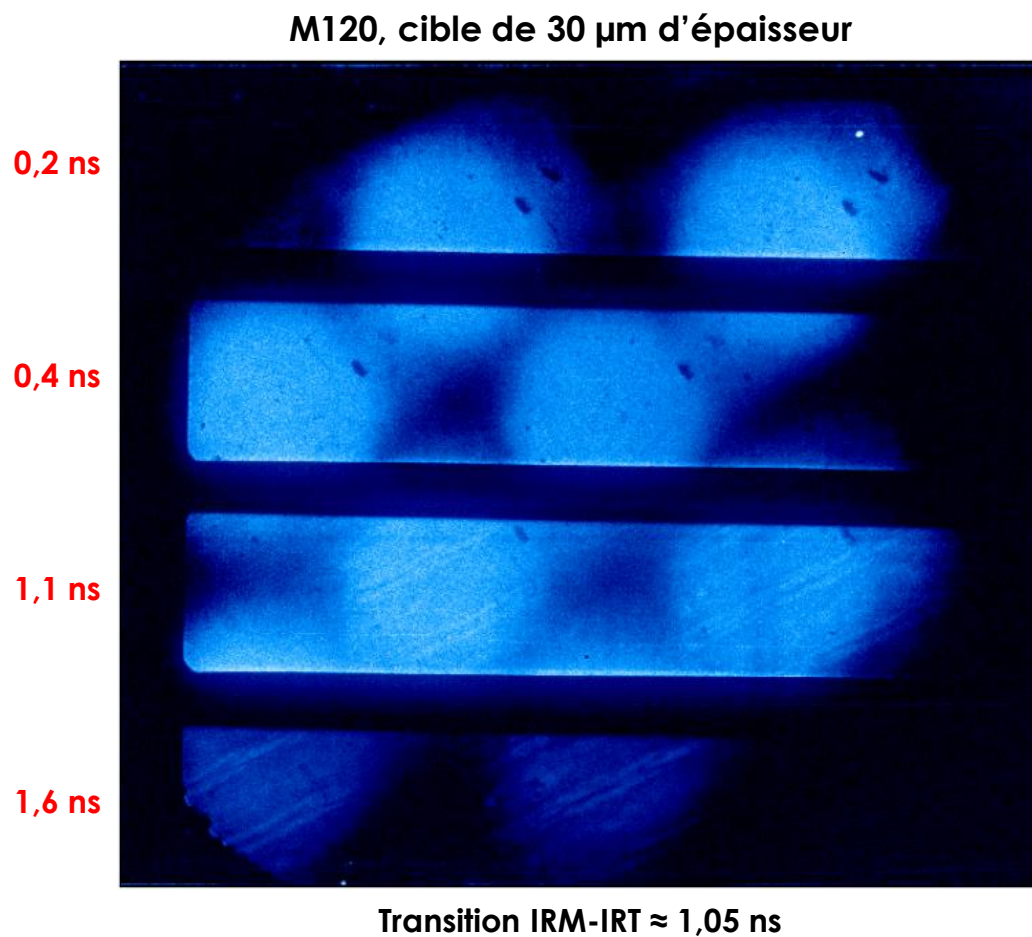


FIGURE 4.9 – Radiographie de face d'une cible de 30 μm d'épaisseur illuminée par deux faisceaux porteurs de la lame de phase SG4 et un faisceau d'empreinte porteur de la lame de phase M120. Les mesures débutent à partir de la 3ème piste.

de transition de l'IRM ablative à l'IRT ablative. Les modulations qui apparaissent sont 2D, et pour chaque lame de phase, elles sont orientées dans une direction différente qui dépend du faisceau d'empreinte et de l'orientation de la lame de phase. Dans tous les cas, l'amplitude des modulations en phase Richtmyer-Meshkov ablative ou à la transition IRM-IRT ablatives est faible, les modulations sortant à peine du niveau de bruit. Enfin, pour le tir avec la M120, on peut remarquer que les modulations qui apparaissent sur les deux dernières bandes sont composées des modulations de longueur d'onde $120\ \mu\text{m}$, ainsi que de sous structures, qui correspondent aux harmoniques à 60 et $30\ \mu\text{m}$ qui sont imprimés par la lame de phase, celle-ci n'étant pas monomode.

La figure 4.10 permet d'envisager l'efficacité de l'empreinte de modulations 3D par rapport à celle des lames de phase spéciales. Lors du tir dont la radiographie est présentée en figure 4.10 (a), le faisceau d'empreinte portant la lame de phase spéciale n'a pas été utilisé. Seuls deux faisceaux, porteurs de lames de phase SG4, illuminent donc la cible. La radiographie est effectuée jusqu'à $2,6\ \text{ns}$ après le début de l'illumination laser. Pourtant, aucune modulation n'apparaît sur la radiographie. Cela montre donc que l'empreinte des faisceaux porteurs des lames de phase SG4 est négligeable par rapport à l'empreinte 2D du faisceau porteur de la lame de phase spéciale. Ainsi, les modulations observées sur les radiographies sont donc bien issues de l'empreinte de ce dernier faisceau.

Lors de plusieurs tirs, le faisceau d'empreinte ne portait pas de lame de phase spéciale ; au contraire, le faisceau était dépourvu de lame de phase et défocalisé pour augmenter la taille des figures de speckle. Les modulations imprimées sur la cible étaient donc 3D. Les radiographies présentées en figure 4.10 (b-c) ont été effectuées lors de deux tirs de ce type, respectivement durant la phase Richtmyer-Meshkov ablative (b) et durant la phase Rayleigh-Taylor ablative (c). Sur la radiographie (b), on n'observe aucune modulation, et sur la radiographie (c), les modulations sont observables mais d'amplitude faible sur la première bande, soit environ $1\ \text{ns}$ après le début d'accélération du front d'ablation et donc après une croissance conséquente du fait de l'IRT ablative. Pour une cible de même épaisseur imprimée par la M60, en figure 4.10 (b), on voit que les modulations apparaissent au bout de $1,0\ \text{ns}$. Ainsi l'empreinte par les lames de phase spéciale est nettement plus forte que l'empreinte des figures de speckle. De plus, nous n'avons pas pu mesurer des modulations 3D dans la phase Richtmyer-Meshkov ablative lors de nos expériences.

4.2.3 Dépouillement des radiographies de face

Ces quatre journées de tirs représentent un nombre important de radiographies de face réalisées grâce à une XRFC [47]. Le début du dépouillement de ces radiographies est exposé en figure 4.11. Deux images sont présentes sur chaque film : la radiographie (2) et

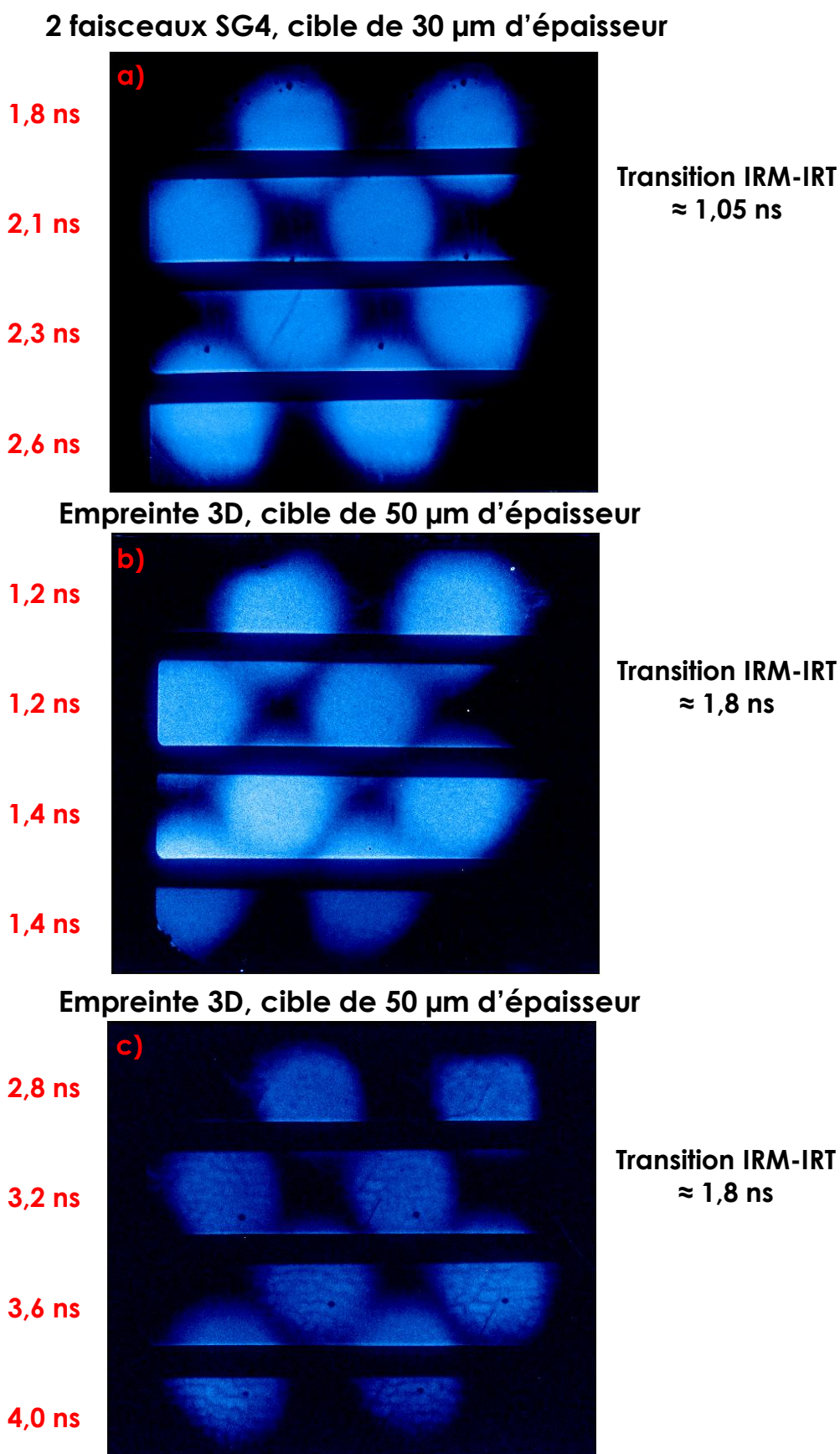


FIGURE 4.10 – Radiographies de face a) d'une cible illuminée par deux faisceaux porteurs de la lame de phase SG4 et b-c) d'une cible illuminée par deux faisceaux porteurs de la lame de phase SG4 et un faisceau d'empreinte défocalisé ne portant pas de lame de phase.

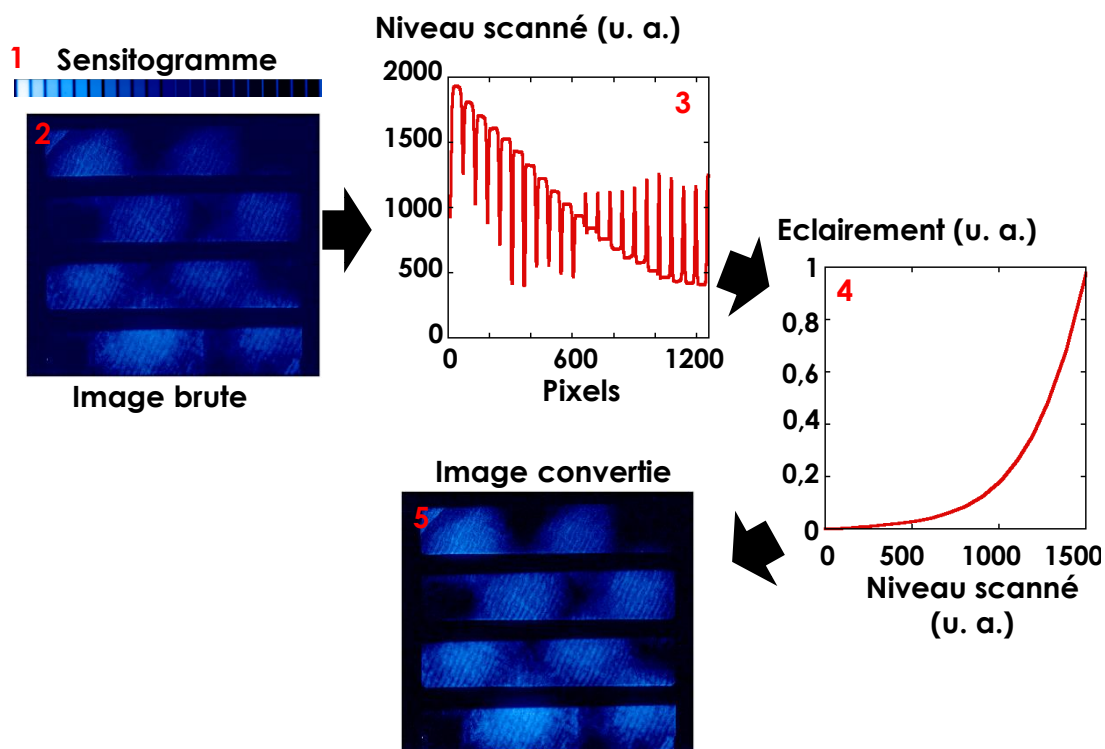


FIGURE 4.11 – Schéma de la méthode de dépouillement utilisée pour convertir les scans des radiographies de face en image d'éclairement : l'image brute (2) est accompagnée d'un sensitogramme, qui correspond à une échelle de conversion du niveau de clarté (1). Le profil extrait de ce sensitogramme (3) permet d'obtenir une courbe (4) qui, appliqué à l'image brute (2), va la convertir en image d'éclairement (5).

un sensitogramme (1). Ces images sont scannées avec un pas de 20 μm au LLE. Sachant que le grandissement du nez d'imagerie est de $12 \times$, le pas en espace de la radiographie est de $20/12 = 1,67 \mu\text{m}$. Comme on réalise un scan de la radiographie, on obtient un fichier image pour lequel le niveau de clarté/obscurité n'a pas de sens physique. Le sensitogramme a donc pour rôle de convertir ce niveau obtenu après scan (appelé niveau scanné sur la figure 4.11) en intensité X (ou éclairement) reçue par le film. Pour cela, on extrait un profil du sensitogramme (3) - dont la calibration est connue - à partir des 21 plages. Un fit polynomial de ces 21 points est réalisé pour obtenir une courbe de conversion continue (4). La radiographie initiale (2) peut ainsi être convertie en radiographie d'éclairement (5) par le truchement de cette courbe (4). On peut ensuite effectuer une corrélation croisée des 8 images de la radiographie puis appliquer la méthode de FFT 1D à chacune de ces images. Les images utilisées en entrée du programme de FFT 1D sont des images en éclairement. Cet éclairement noté I varie à travers l'image pour des rayons X monochromatiques selon [84]

$$I(x, y) = I_{BL}(x, y)e^{-\mu(E)\rho e(x, y)} \quad (4.2)$$

avec I_{BL} l'intensité incidente des rayons X de la source de radiographie, $\mu(E)$ le taux d'absorption massique de la cible à l'énergie E , ρ la densité de la cible et e l'épaisseur de la cible. Cependant on ne peut connaître la valeur absolue de I et de I_{BL} , et on ne pourra donc connaître la densité surfacique (produit de la densité et de l'épaisseur) absolue de la cible. En effet, le sensitogramme donne une calibration relative des niveaux d'intensité les uns par rapport aux autres et on ne peut mesurer absolument l'intensité des rayons X de radiographie en entrée et en sortie de la cible. Cependant, on peut mesurer les variations relatives de densité surfacique notées $\delta(\rho R)$ qui vont être définies par

$$\delta(\rho R) = \frac{\delta(PO)}{\mu(E)} \quad (4.3)$$

avec $\delta(PO)$ la variation de profondeur optique définie par

$$\delta(PO) = -\ln \frac{\delta I(x, y)}{I_{BL}(x, y)} \quad (4.4)$$

Les formules (4.3) et (4.4) sont obtenues à partir de l'équation (4.2). Dans un cas idéal, l'éclairement de la source de radiographie serait homogène et la partie I_{BL} serait une constante qui n'aurait pas d'influence sur les images (qui donnent comme expliqué des informations de variations relatives). Cependant cet éclairement n'est pas constant ; c'est pour cela qu'on divise l'image par un fit polynomial d'ordre 4 dans le programme de FFT 1D (cf chapitre précédent). Le logarithme de l'image résultante est aussi calculé dans ce programme. Ainsi, la FFT 1D donne des variations de profondeur optique à diverses longueurs d'onde. Pour calculer les variations de densité surfacique associées, il faut donc

connaître la valeur du taux d'absorption massique μ . Pour cela, nous avons utilisé le fait que la transmission d'une cible est définie comme

$$T = \frac{I}{I_{BL}} = e^{-\mu \rho e} \quad (4.5)$$

et donc que

$$\mu = -\frac{\ln T}{\rho e} \quad (4.6)$$

La transmission de chaque cible froide (c'est-à-dire avant le début de l'illumination par le laser) a été obtenue à partir des données du site internet www.cxro.lbl.gov. Ainsi pour une cible de CH_2 d'épaisseur $50 \mu\text{m}$ et de densité $0,90 \text{ g/cm}^3$, le taux de transmission des rayons X à $2,5 \text{ keV}$ (source de Ta) est $\mu_{\text{CH}_2}(2,5 \text{ keV}) = 134 \text{ cm}^2/\text{g}$. Il suffit donc de diviser la valeur du pic de profondeur optique à la longueur d'onde désirée sur le spectre de Fourier pour obtenir la valeur de densité surfacique correspondante. Il faut noter que cette méthode se base sur l'hypothèse que la transmission de la cible change peu au cours de la mesure et que les variations sont petites devant la valeur de la transmission de la cible froide. Un autre point à prendre en compte lors du dépouillement des radiographies est la fonction de transfert des modulations (FTM) qui traduit la conversion de l'objet (cible) en image par la XRFC. Par exemple, pour un objet de contraste quasiment infini comme un bord franc, l'image du bord va être étendue sur une certaine distance du fait de la résolution finie des différentes étapes de la mesure (XRFC, film, scanner). Ainsi les signaux de grandes longueurs d'onde vont être bien résolus tandis que l'amplitude des signaux de petites longueurs d'onde va être étalée et donc sous-évaluée [85]. Pour les XRFC que nous avons utilisées, la FTM s'écrit [86]

$$FTM(k) = \frac{1}{1 + \alpha} (e^{-\sigma_1^2 k^2 / 2} + \alpha e^{-\sigma_2^2 k^2 / 2}) \quad (4.7)$$

avec $k = \frac{2\pi}{\lambda}$, $\alpha = 5,25$, $\sigma_1 = 46,5 \mu\text{m}$ et $\sigma_2 = 4,2 \mu\text{m}$. On trouve donc pour les modulations imprimées par la M30 $FTM = 0,57$, pour celles imprimées par la M60 $FTM = 0,78$ et $FTM = 0,83$ pour la M120. Pour finir, le bruit B est calculé par FFT 1D dans la direction perpendiculaire aux modulations. On a donc

$$\delta(\rho R)(k) = \frac{\delta(OD) - B(k)}{\mu FTM(k)} \quad (4.8)$$

Cette formule permet donc de calculer la valeur de densité surfacique correspondant à chaque image de la cible sur les radiographies, et donc d'obtenir les points expérimentaux. L'incertitude sur la mesure est évaluée par le programme dédié présenté au chapitre précédent. Les données collectées sont présentées en figure 4.13 pour les 3 lames de phase. Le dernier aspect du dépouillement provient du fait que, lors de la 4ème journée de tir, des modulations préimposées étaient présentes sur la cible et ne se trouvaient pas à 90°

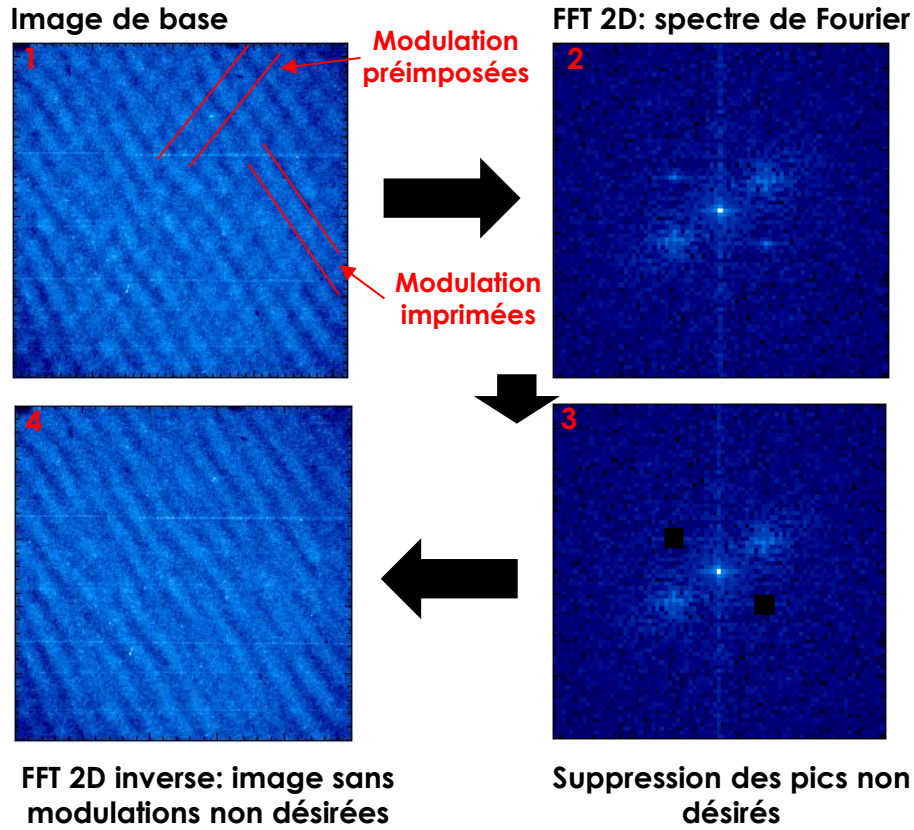


FIGURE 4.12 – Détail de la méthode utilisée pour supprimer les modulations préimposées non-désirées sur les images.

des modulations imprimées, pouvant modifier les valeurs trouvées lors de l'analyse par FFT 1D. Pour s'affranchir de ce problème, nous avons utilisé la méthode présentée en figure 4.12. On peut voir sur l'image à traiter (1) des modulations 2D dans deux directions différentes. Les modulations préimposées, qui vont en diagonale du coin bas gauche au coin haut droit, sont celles que l'on désire supprimer. Pour cela, on effectue la FFT 2D du signal (2). Les deux points nettement définis correspondent aux modulations à supprimer : on définit donc un niveau de 0 pour toute la zone associée (3). La FFT 2D inverse du spectre traité est alors réalisée et on retrouve la radiographie sans les modulations préimposées (4).

Pour une cible de CH_2 de $15 \mu\text{m}$ d'épaisseur dont la radiographie est réalisée à l'aide d'une source d'U, on trouve $\mu_{\text{CH}_2}(1,3 \text{ keV}) = 904 \text{ cm}^2/\text{g}$. Dans ce cas, on fait l'approximation d'une radiographie faite par des photons monoénergétiques à $1,3 \text{ keV}$. La figure 4.3 (b) donne le spectre réel $S_U(E)$. On peut donc calculer le taux d'absorption massique réel de la cible

$$\mu_{\text{CH}_2}(\text{uranium}) = \frac{\int S_U(E) \mu_{\text{CH}_2}(E) dE}{\int S_U(E) dE}. \quad (4.9)$$

On trouve alors que $\mu_{\text{CH}_2}(\text{uranium}) \approx 865 \text{ cm}^2/\text{g}$, ce qui donne un écart de seulement 5 % environ avec le taux d'absorption massique calculé en faisant l'approximation de photons monoénergétiques à 1,3 keV. Cette approximation est donc bien valable.

4.2.4 Données issues du dépouillement des radiographies de face

La figure 4.14 présente les données des 4 tirs à haute intensité ($I \approx 1.10^{14} \text{ W/cm}^2$) réalisés avec la M30. Les barres verticales pointillées représentent le temps de début d'accélération du front d'ablation c'est-à-dire la transition de l'IRM à l'IRT ablative. Ce temps est déterminé grâce à des simulations CHIC 1D. Les différentes couleurs correspondent à différentes épaisseurs de cible, et les différentes formes des points à différents tirs. Pour les deux intensités et les trois longueurs d'onde, l'IRM ablative imprimée par laser a été mesurée pour la première fois. Des mesures de l'IRM ablative ont été réalisées pour des cibles de 30, 50 et 80 μm d'épaisseur. Pour la M30 et avec les cibles de 50 et 80 μm d'épaisseur, à haute et basse intensité, une inversion de phase est observée -les données décroissent puis se remettent à croître. L'inversion de phase apparaît plus tôt à haute intensité : ceci est cohérent avec la théorie (cf équation 2.37) qui prévoit que la fréquence d'oscillation croît avec la vitesse d'ablation et la vitesse de "blow-off". Avec les cibles de 50 μm , on peut voir que les données décroissent plusieurs centaines de ps après le début d'accélération du front d'ablation : c'est le signe que la compétition entre IRM et IRT dure pendant un temps non négligeable après le début d'accélération du front d'ablation, avant que l'IRT prenne le pas sur l'IRM. Pour obtenir ces données, la taille des boîtes d'analyse était d'environ 400 $\mu\text{m} \times 400 \mu\text{m}$. Ainsi, l'analyse a été réalisée sur 13 longueurs d'onde pour les tirs avec M30, 6 pour ceux avec M60 et 3 pour ceux avec M120. Il semble qu'une limitation pour l'analyse des tirs avec la M120 est la faible statistique de l'analyse.

4.3 Détermination du niveau d'empreinte induit par les lames de phase

Un point fondamental pour utiliser le modèle de Goncharov ou réaliser des simulations CHIC de l'empreinte de modulations 2D est de connaître l'amplitude relative des modulations, c'est-à-dire le $\frac{\delta I}{I}$ du modèle de Goncharov (cf chapitre 2). Des images d'un faisceau portant les différents lames de phase d'empreinte sont exposées en figure 4.15. Ces images sont appelées images "Equivalent Target Plane" (ETP). Pour les obtenir, une caméra est placée au centre chambre et le faisceau à bas flux est focalisé dessus. La résolution de ces images est de 4,58 $\mu\text{m}/\text{pixel}$. Pour trouver l'amplitude relative des modulations associées aux différentes lames de phase, nous avons choisi la même taille de boîte (400 μm de côté) que pour l'analyse de données et la même méthode de traitement par FFT 1D. Ceci est nécessaire pour pouvoir comparer les simulations et modèles aux données

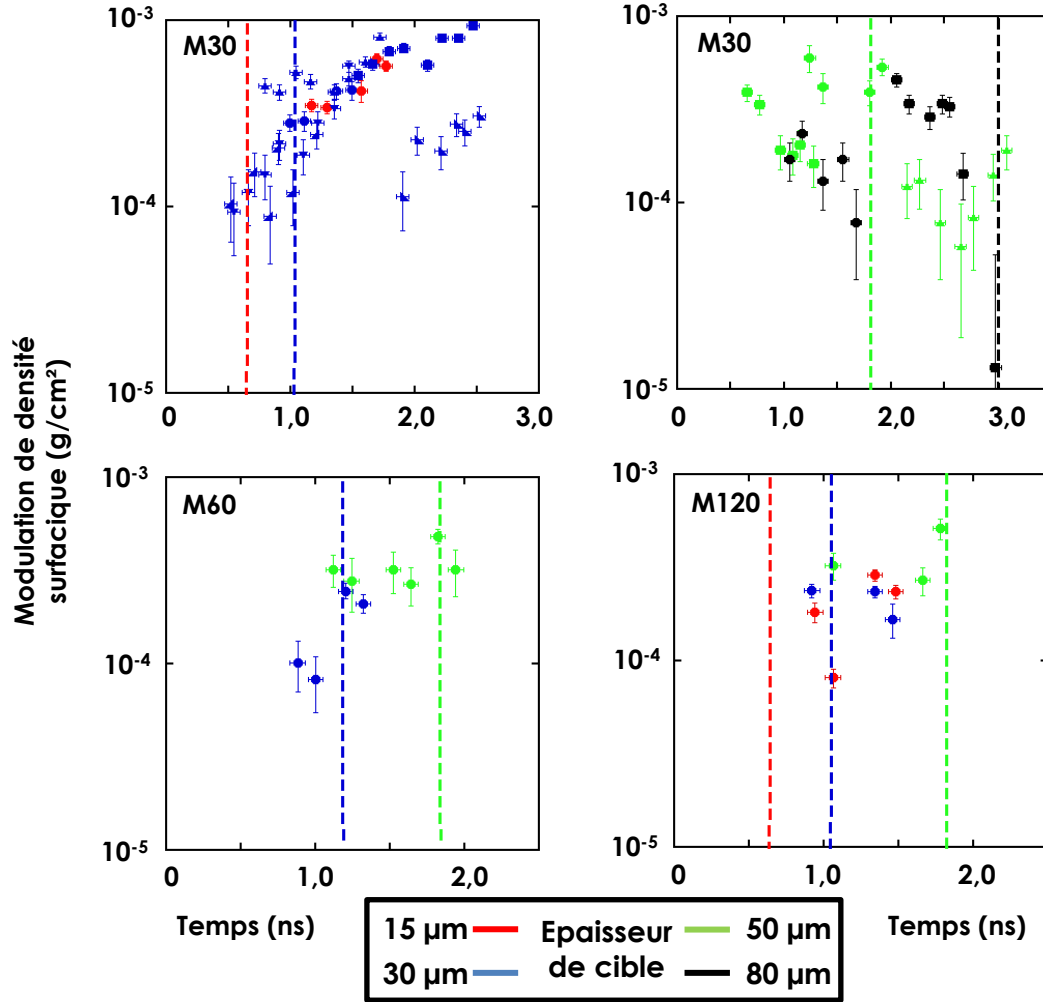


FIGURE 4.13 – Evolution des modulations de densité surfacique en fonction du temps pour les données mesurées lors des campagnes d'IRM ablative imprimée par laser à une intensité $I \approx 5.10^{13} \text{ W}/\text{cm}^2$. Les quatre figures correspondent aux trois lames de phase M30 (données séparées en deux figures pour plus de clarté), M60 et M120. Les barres pointillées verticales représentent le temps de début d'accélération du front d'ablation. Les différentes couleurs représentent différentes épaisseurs de cible : 15 μm en rouge, 30 μm en bleu, 50 μm en vert et 80 μm en noir. Les différents symboles (ronds, carrés, triangles...) correspondent à différents tirs.

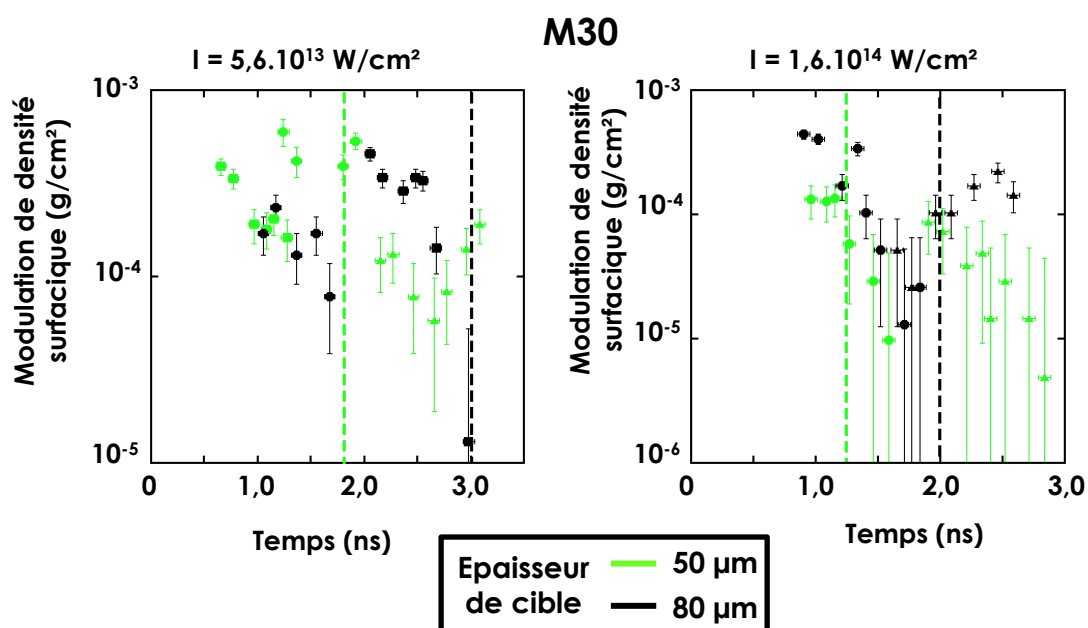


FIGURE 4.14 – Evolution des modulations de densité surfacique en fonction du temps pour les données obtenues avec la M30 à une intensité $I \approx 5,6 \cdot 10^{13} \text{ W/cm}^2$ (figure de gauche) et $I \approx 1,6 \cdot 10^{14} \text{ W/cm}^2$ (figure de droite). Les barres pointillées verticales représentent le temps de début d'accélération du front d'ablation. Les différentes couleurs représentent différentes épaisseurs de cible : 15 μm en rouge, 30 μm en bleu, 50 μm en vert et 80 μm en noir. Les différents symboles (ronds, carrés, triangles...) correspondent à différents tirs.

expérimentales ; si l'on choisit une autre taille de boîte ou méthode d'analyse, le résultat différera. Par exemple, si le profil est moyenné sur une zone plus large, le niveau de modulation mesuré va décroître car les modulations ont une forme un peu ondulée. Comme on ne sait pas quelle partie de la lame de phase a imprimé les modulations que l'on mesure avec les XRFC, nous avons choisi de prendre des boîtes carrées de $400\text{ }\mu\text{m}$ de côté en 9 endroits différents de l'image ETP. Le traitement par FFT 1D est ensuite réalisé pour ces 9 sous-images, et le niveau de modulation trouvé pour chacune est moyenné. On trouve ainsi que $\frac{\delta I}{I} = 18,6 \pm 5,9\%$ pour la M30, $55,4 \pm 12,8\%$ pour la M60 et $84,9 \pm 6,1\%$ pour la M120. L'écart-type relatif est le plus important pour la M30 : en fonction de la position de la boîte, on peut trouver un niveau de modulations de 10 ou 30 %. Un autre point à noter est qu'on ne connaît pas la FTM de modulation de la caméra utilisée pour mesurer les images ETP ; les niveaux de modulation n'ont donc pas pu être corrigés. Or cette FTM a probablement un effet important pour des modulations de $30\text{ }\mu\text{m}$ et nettement plus faible sur des modulations de $120\text{ }\mu\text{m}$.

Enfin, la figure 4.16 montre 2 images d'un faisceau porteur de la M30 : l'image 1 est l'image ETP montrée dans la figure 4.15 et l'image 3 est obtenue à partir d'une mesure de front d'onde transmis ("transmitted wavefront" en anglais). Cette mesure est réalisée avant le triplement de fréquence dans la chaîne laser OMEGA. Une FFT de cette mesure permet ensuite d'obtenir l'allure du faisceau focalisé (la FFT joue le rôle de la lentille de focalisation pour le passage de champ proche en champ lointain). Le niveau de modulation obtenu avec la mesure de front d'onde est nettement plus important qu'avec l'image ETP : $71,3 \pm 6,4\%$ au lieu de $18,6 \pm 5,9\%$. Cet important écart peut s'expliquer, au moins partiellement, de deux manières : tout d'abord, la résolution de l'image obtenue par mesure de front d'onde - $3\text{ }\mu\text{m}/\text{pixel}$ - est meilleure que celle des images ETP, et donc les pics et creux d'intensité vont être mieux résolus. De plus, la mesure du front d'onde transmis est effectuée avant le triplement de fréquence, donc avant le passage par des nombreuses optiques ; de nombreuses aberrations induites par ces optiques ne sont donc pas présentes sur l'image obtenue par mesure du front d'onde alors qu'on les trouve sur les images ETP. C'est d'ailleurs pour cela que lorsqu'on compare les images (3) et (4), les modulations sur l'image ETP (3) semblent nettement plus ondulées que sur l'image (4). Ces ondulations causent une diminution de l'amplitude des pics lorsqu'on calcule un profil perpendiculaire aux modulations moyenné sur une certaine épaisseur. Les aberrations induisent donc des ondulations transverses aux modulations ; ces ondulations semblent avoir une taille d'environ $20\text{ }\mu\text{m}$ que l'on peut estimer directement sur les images. Elle vont donc avoir une plus grande influence sur le niveau de modulation associé à la M30 que sur celui associé à la M60 et encore moins à la M120. Ces ondulations vont s'imprimer sur les cibles, et donc se retrouver dans l'analyse des radiographies. Cependant, il est compliqué d'évaluer dans quelle mesure elles vont s'imprimer par rapport aux modulations 2D car, d'après le

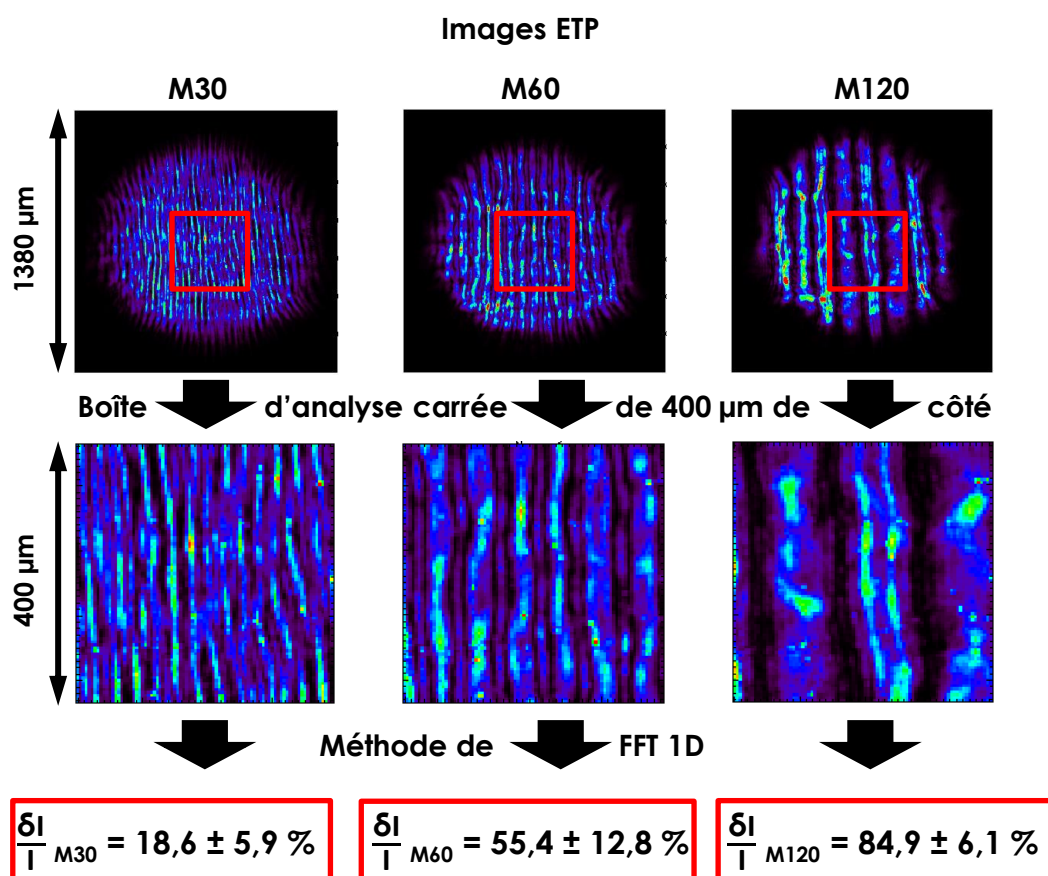


FIGURE 4.15 – Images obtenues par mesure ETP (Equivalent Target-Plane) dans un plan perpendiculaire à la direction de propagation d'un faisceau porteur des lames de phase M30, M60 et M120 et niveau de modulation de l'intensité laser associé à chacune de ces lames de phase.

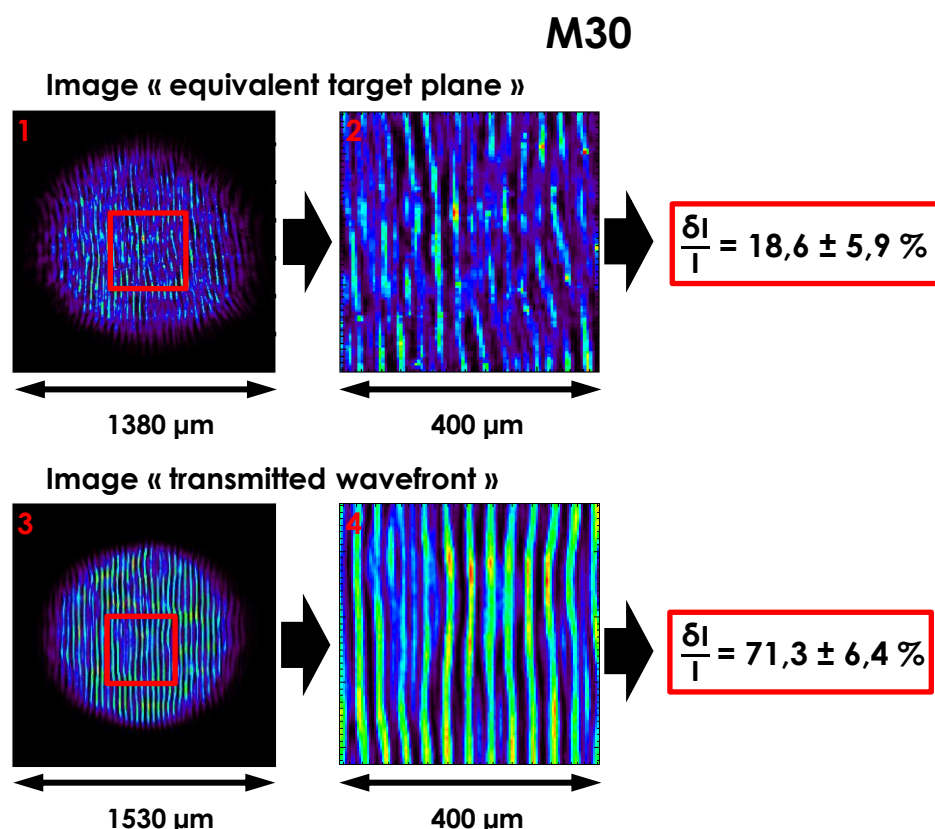


FIGURE 4.16 – Images dans un plan perpendiculaire à la direction de propagation d'un faisceau porteur de la M30 obtenues par deux méthodes : mesure ETP (1) et mesure de transmission du front d'onde (3). Les zones d'analyse extraites de ces images sont respectivement représentée en (2) et (4), ainsi que le niveau de modulation de l'intensité laser associé.

modèle de Goncharov, l'efficacité d'empreinte dépend de la longueur d'onde. Il est ainsi possible que le motif imprimé diffère légèrement du motif initial de la lame de phase, et qu'encore une fois cet effet soit plus important avec la M30 qu'avec les autres lames de phase. Quoiqu'il en soit, nous nous baserons sur les niveaux de modulation calculés à partir des images ETP car l'effet de toutes les optiques de fin de chaîne est pris en compte. Mais du fait de la résolution, de l'écart-type du niveau de modulation calculé, de la méconnaissance de la FTM de la caméra utilisée pour les mesures ETP et de la différence entre motif incident et motif imprimé, on doit considérer que le niveau de modulation imprimé peut fortement différer - et plus spécifiquement être sensiblement supérieur - du niveau de modulation mesuré pour la M30. Ces imprécisions doivent être atténuées pour la M60 et plus atténuées encore pour la M120.

Pour résumer, voici les principaux résultats présentés dans ce chapitre :

— Une configuration expérimentale a été définie pour mesurer l'IRM ablative impri-

mée par laser. Notamment, le choix des couples source de radiographie/épaisseur de cible a été optimisé. De plus, des lames de phase spéciales produisant des modulations 2D monomodes de l'intensité laser ont été choisies pour réaliser l'empreinte des modulations sur les cibles.

- L'amplitude des modulations de l'intensité créées par ces lames de phase a été déterminée, mais la valeur trouvée est sous-évaluée, particulièrement pour la M30, entre autres car on ne connaît pas la FTM de la caméra utilisée pour ces mesures.
- L'IRM ablative imprimée par laser a été mesurée pour la première fois au cours de ces expériences. Vingt tirs ont permis cette mesure, pour trois longueurs d'onde, trois épaisseurs de cible et deux intensités différentes.
- Pour la longueur d'onde de $30\text{ }\mu\text{m}$, une inversion de phase a été mesurée pour les deux intensités. Conformément à la théorie, cette inversion de phase apparaît plus tôt à plus haute intensité.

L'IRM ablative imprimée par laser ayant été mesurée, nous allons pouvoir, au cours du chapitre suivant, interpréter ces données à l'aide des simulations CHIC et du modèle développé par Goncharov, et confronter pour la première fois ce modèle à l'expérience.

Chapitre 5

Interprétation des mesures l'IRM imprimée par laser à partir du code CHIC et du modèle de Goncharov

Au cours du chapitre précédent, nous avons présenté des mesures de l'IRM ablative imprimée par laser. Il est pertinent de confronter ces résultats expérimentaux au modèle de Goncharov et aux simulations CHIC pour évaluer leur capacité à prédire l'évolution induite par cette instabilité et pour mieux comprendre les résultats expérimentaux. En effet, améliorer la fiabilité des codes et des modèles théoriques permettrait de concevoir les expériences de manière à limiter au maximum l'effet délétère des instabilités hydrodynamiques (cf chapitre 1 Introduction). Par exemple, on pourrait concevoir les cibles de manière à ce que la transition IRM-IRT ait lieu au moment où, du fait des oscillations de l'IRM, l'amplitude des modulations du front d'ablation s'annule ; l'amplitude des modulations du front d'ablation au moment du passage à l'IRT serait ainsi bien plus basse qu'en début d'expérience. En retour, les modèles théoriques peuvent être améliorés par leur comparaison avec les données expérimentales.

Nous présenterons donc dans ce chapitre en premier lieu les différents paramètres du modèle de Goncharov et la manière dont ils ont été évalués à partir des simulations 1D, puis la comparaison des données expérimentales avec ce modèle et les simulations bidimensionnelles de l'écoulement perturbé. Enfin, nous étudierons l'importance du choix des différents paramètres sur les simulations et sur leur accord avec les données expérimentales.

5.1 Réalisation de simulations CHIC de la croissance des modulations imprimées par laser et calcul du modèle de Goncharov

Pour réaliser les calculs, nous avons procédé en trois parties. Tout d'abord, nous avons réalisé des simulations 1D. Ces simulations nous ont permis d'obtenir l'instant de transition de l'IRM à l'IRT. De plus, nous en avons extrait différents paramètres nécessaires pour effectuer la deuxième partie : le calcul du modèle de Goncharov. Enfin, la troisième partie a consisté, pour chaque épaisseur de cible et chaque longueur d'onde imprimée, à réaliser des simulations de croissance de modulations.

5.1.1 Intensité équivalente pour les simulations CHIC 1D

Pour déterminer l'intensité laser à utiliser dans les simulations 1D, nous avons tiré parti des mesures issues des diagnostics laser P510. Ces mesures donnent l'évolution de la puissance délivrée par chaque faisceau au cours du temps. Pour des simulations 1D, Le code CHIC utilise en entrée une puissance moyenne uniforme délivrée sur la surface de la première maille, qui est un carré de 1 cm sur 1 cm. On cherche donc à calculer à partir des données P510 une intensité moyenne, sachant que dans le cas de l'expérience (ou d'une simulations 2D), l'intensité est répartie selon les caractéristiques de la tache focale et n'est donc pas uniforme spatialement.

Les impulsions issues des mesures P510 pour chacun des trois faisceaux illuminant la cible d'un tir avec M30 sont représentées en figure 5.1 (a). On peut d'ailleurs remarquer que le faisceau d'empreinte est bien avancé d'environ 200 ps sur les autres faisceaux (cf chapitre précédent). La puissance totale $W(t)$ correspond à la somme de ces trois faisceaux. On définit l'intensité $I(t)$ pour les simulations 1D de la manière suivante :

$$I(t) = \frac{W(t)\cos\theta}{\pi r_0^2} = \frac{0,85.\cos\theta}{\pi(352.10^{-4})^2}P(t), \quad (5.1)$$

avec θ l'angle d'incidence des trois faisceaux et r_0 le rayon de la tache focale. Pour les expériences avec les lames de phase M30 et la M120, les trois faisceaux ont un angle $\theta = 23^\circ$ et pour la M60, pour laquelle le faisceau d'empreinte a un angle d'incidence de 48° , $\cos\theta$ est pris comme étant la moyenne des cosinus des trois faisceaux. Le dénominateur (divisé par $\cos\theta$) correspond à l'aire de l'ellipse formé sur la cible par les faisceaux porteur de la SG4. On fait donc l'approximation suivante : la tache formée par le faisceau porteur de la lame de phase d'empreinte est de taille similaire à celles formées par les faisceaux porteurs des SG4. La lame de phase SG4 impose aux faisceaux qui en sont munis un profil d'intensité supergaussien d'ordre 4,2 et de rayon $r_0=352 \mu\text{m}$. Ainsi on peut écrire

$$I(r) = I_0 e^{-(\frac{r}{r_0})^{4,2}}, \quad (5.2)$$

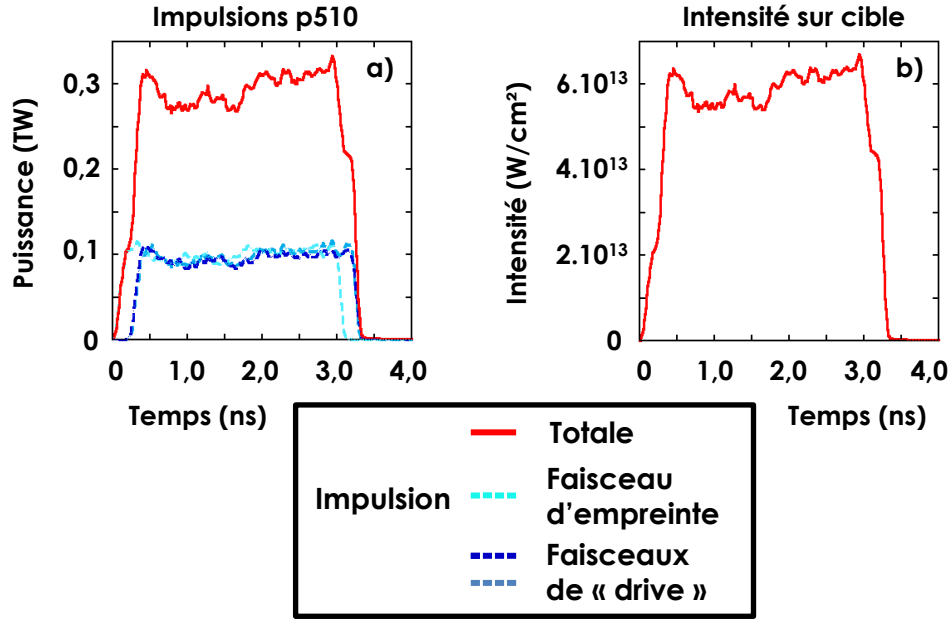


FIGURE 5.1 – a) Puissance délivrée en fonction du temps par chacun des trois faisceaux illuminant la cible et puissance totale pour le tir 65941 (M30, cible de 30 μm d'épaisseur). Ces courbes sont directement issues des données P510. b) Intensité totale sur cible en fonction du temps.

avec $r = 0$ au centre du faisceau. La puissance W_{r_0} contenue dans le disque de rayon r_0 et la puissance W_{tot} contenue dans l'ensemble du faisceau peuvent donc être définies par

$$W_{r_0} = 2\pi \int_0^{r_0} r I(r) dr, \quad (5.3)$$

$$W_{tot} = 2\pi \int_0^{+\infty} r I(r) dr. \quad (5.4)$$

Le pourcentage de la puissance totale du faisceau comprise dans le disque de rayon r_0 est donc $W_{r_0}/W_{tot} = 85\%$. Ce qui explique donc le coefficient 0,85 dans l'équation (5.1) : 15 % de l'énergie du faisceau est perdue dans les ailes de la supergaussienne. L'équation (5.1) permet donc d'obtenir $I(t)$. En multipliant $I(t)$ par 1 cm², on obtient la puissance utilisée dans les simulations 1D. Les variations de $I(t)$ sont représentées en figure 5.1 (b).

Pour nos simulations, nous avons utilisé un limiteur de flux de 7 % ("sharp cut-off") et l'équation d'état SESAME 7180.

5.1.2 Calcul analytique avec le modèle de Goncharov

Pour pouvoir utiliser le modèle d'IRM proposé par V. N. Goncharov et al [11] et présenté au chapitre 2 section 2.3.2, de nombreux paramètres doivent être évalués : la

Int. (W/cm ²)	P (Mbar)	c _s (cm/s)	V _c (cm/s)	V _a (cm/s)	ρ _a (g/cm ³)	ρ _{bl} (g/cm ³)	ρ (g/cm ³)	u (cm/s)
5.10 ¹³	8,8	2,5.10 ⁶	12,6.10 ⁵	2,3.10 ⁵	1,25	0,095 (M30) 0,050 (M60) 0,027 (M120)	3,6	3,8.10 ⁶
1,6.10 ¹⁴	21	3,4.10 ⁶	32,2.10 ⁵	3,3.10 ⁵	1,6	0,190 (M30)	3,9	5,4.10 ⁶

TABLE 5.1 – Résumé des différentes données hydrodynamiques nécessaires au calcul du modèle de Goncharov. Ces données ont été obtenues à partir de simulations CHIC 1D.

vitesse de formation de la zone de conduction V_c , la vitesse d'ablation V_a , la vitesse de "blow-off" ou vitesse d'éjection V_{bl} , la vitesse du choc u , les vitesses acoustiques pré- et post-choc c_{s0} et c_s et les densités pré- et post-choc ρ_0 et ρ . Nous allons calculer c_s en utilisant sa définition $c_s = \sqrt{\gamma P / \rho}$ avec P la pression, et V_{bl} en utilisant $V_{bl} = \rho_a V_a / \rho_{bl}$. Nous aurons également besoin de connaître la densité au front d'ablation ρ_a , la densité du plasma de "blow-off" ρ_{bl} et les pressions pré- et post-choc. Tous ces différents paramètres ont été obtenus à partir de simulations CHIC 1D. Le tableau 5.1 présente les valeurs de ces différents paramètres pour des simulations réalisées à des intensités de 5, 6.10¹³ W/cm² et 1, 6.10¹⁴ W/cm², que l'on appellera dans la suite de ce chapitre respectivement basse et haute intensités.

Pour calculer V_c , on mesure la taille de la zone de conduction D_c , entre la densité critique et le front d'ablation, à divers instants des simulations. L'évolution temporelle de D_c est représentée en figure 5.2. On peut remarquer que la croissance est nettement plus importante à haute intensité, et donc que le découplage des modulations de l'intensité et du front d'ablation va avoir lieu plus rapidement. On calcule V_c comme la vitesse moyenne sur la première nanoseconde. Pour obtenir la position du front d'ablation, on repère les extrema du gradient de densité : l'un correspond à la position du front de choc, l'autre à celle du front d'ablation. Repérer cette position permet aussi de relever la densité au front d'ablation ρ_a . On peut alors calculer ρ_{bl} en relevant la densité à une distance $1/2k$ du front d'ablation dans la direction du plasma en expansion. Pour obtenir V_{bl} , il reste à calculer V_a . Pour cela, on évalue la dérivée temporelle de la masse ablatée, ce qui donne le taux de masse ablatée $\dot{m}_a$. Sachant que $V_a = \dot{m}_a / \rho$, on relève la valeur de ρ dans la matière non-ablatée après le passage du choc et on obtient donc V_a . La pression P , donnée directement en sortie par le code CHIC, est relevée entre le front de choc et le front d'ablation et dans la matière non perturbée, ce qui permet de calculer les vitesses acoustiques.

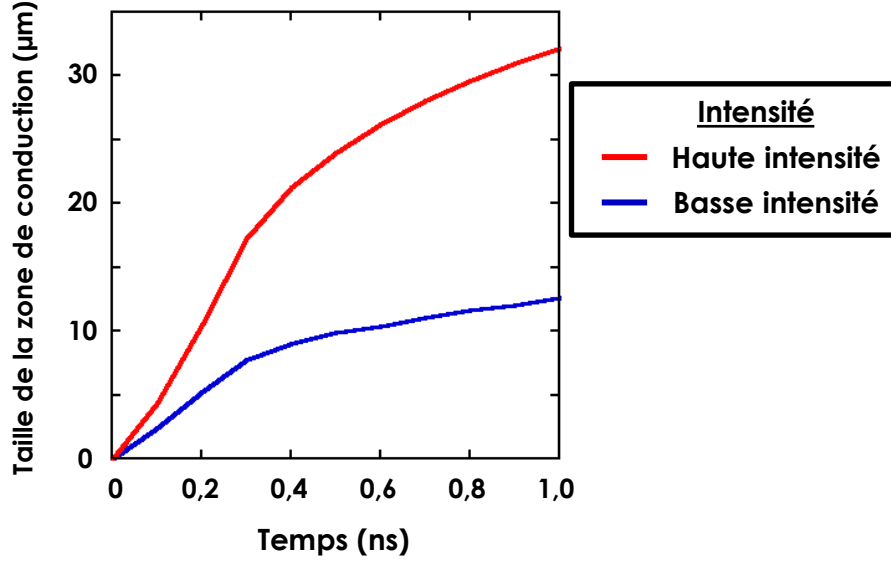


FIGURE 5.2 – Evolution temporelle de la taille de la zone de conduction, extraite de simulations CHIC 1D, pour des tirs à des intensités de $5,6 \cdot 10^{13} \text{ W/cm}^2$ et $1,6 \cdot 10^{14} \text{ W/cm}^2$.

5.1.3 Simulation de la croissance des modulations imprimée par laser

Pour calculer la croissance d'une modulation de l'intensité laser imprimée sur la cible, on effectue des simulations CHIC du type de celle présentée sur la figure 3.12 du chapitre 3. Ainsi, on effectue la simulation sur une demie longueur d'onde pour des raisons de réduction du temps de calcul. Il a été vérifié précédemment que le résultat de la simulation n'était pas affecté par ce choix. Ces simulations sont initialement monomodes : seule la modulation correspondant à la longueur d'onde principale imprimée par la lame de phase est considérée. Une modulation sinusoïdale est donc appliquée à la loi de puissance laser. Cette puissance laser est calculée en multipliant $I(t)$ (cf équation (5.1)) par la surface $\frac{\lambda_{\text{empreinte}}}{2} \times 1 \text{ cm}^2$. Enfin, nous avons là aussi utilisé un limiteur de flux de 7 % ("sharp cut-off") et l'équation d'état SESAME 7180. L'influence de ces deux paramètres sur le résultat des simulations est discutée en fin de chapitre.

5.2 Comparaison des simulations CHIC et des données expérimentales

La Figure 5.3 présente la comparaison entre les données expérimentales et les simulations réalisées avec CHIC. Les courbes pleines sont obtenues en utilisant le niveau de modulation moyen pour chaque lame de phase ; ces courbes sont encadrées par 2 courbes pointillées, qui correspondent aux simulations réalisées en utilisant les barres d'erreurs minimales et maximales du niveau de modulation. On peut noter plusieurs points :

- Pour les lames de phase M60 et M120 (figure 5.3 (d-e)), les simulations présentent un bon accord en terme d'allure de l'évolution temporelle avec les données. Quantitativement, les simulations semblent surestimer systématiquement l'amplitude des modulations, les écarts allant de quelques pour-cents à un facteur 2 pour les plus importants. Cependant, des écarts de ce niveau sont raisonnables si l'on compare avec la littérature (cf par exemple [84, 87, 10]).
- Pour la M30 (figure 5.3 (a-b)), les simulations concordent avec les données expérimentales en terme d'évolution temporelle. Pour les cibles de 15 et 30 μm d'épaisseur (a), la transition de l'IRM à l'IRT a lieu rapidement, l'IRM n'a pas le temps d'osciller et la densité surfacique ne cesse de croître. Pour les cibles de 50 et 80 μm (b), on observe expérimentalement et numériquement la décroissance des modulations liée au phénomène d'inversion de phase. En revanche, les simulations donnent une inversion de phase plus rapide que les expériences ; elles sous-estiment donc la période d'oscillation de l'IRM ablative. Cela explique aussi pourquoi on observe une sorte de palier dans les simulations pour la cible de 30 μm d'épaisseur et pas dans l'expérience : les oscillations sont plus rapides dans les simulations, le palier correspond donc à un début d'oscillation avant que l'IRT ne prenne le pas sur l'IRM. Ce palier n'est pas observé dans l'expérience car les oscillations y sont plus lentes.
- En revanche, en terme qualitatif, les données expérimentales sont systématiquement supérieures aux simulations, à l'exception du tir avec une cible de 15 μm et d'un tir avec une cible de 30 μm . Les données de ce dernier tir ne sont pas dans le nuage de points formé par les 5 autres tirs avec des cibles de 30 μm . On peut donc considérer que ce tir n'est pas fiable. Pour la cible de 15 μm , on remarque que la croissance est plus faible que dans la simulation. Cette réduction de croissance s'observe aussi pour les amplitudes de modulation les plus élevées des tirs avec cibles de 30 μm , qui semblent saturer et leur pente s'incliner sans suivre la forte croissance des simulations correspondantes. Ce phénomène apparaît quand les données se rapprochent du mg/cm^2 . Or la densité surfacique des cibles froides est de 1,35 mg/cm^2 pour la cible de 15 μm d'épaisseur et de 2,7 mg/cm^2 pour celle de 30 μm . Ainsi, quand les modulations tendent vers le mg/cm^2 , elles se rapprochent de la densité surfacique totale de la cible. Sachant qu'une partie de celle-ci a été ablatée, on peut supposer que l'affaissement des données expérimentales pour les cibles de 15 et 30 μm correspond au fait que les modulations ont percé la cible et donc qu'elles ne peuvent plus croître. Par rapport à l'expérience ou aux simulations 2D utilisant la description réaliste de la tache focale et des faisceaux laser, la cible ne peut pas se détendre latéralement dans les simulations de croissance de modulations utilisant l'écoulement perturbé. Le phénomène de chute de la densité surfacique peut donc apparaître plus tardivement, ce qui pour-

rait expliquer pourquoi cette tendance n'est pas observée dans ces simulations.

Ainsi, on voit que les données des tirs avec une cible de $30\ \mu\text{m}$ d'épaisseur sont supérieures aux simulations entre 1 et 2 ns, puis qu'elles se rejoignent. De la même manière, les points du tir avec la cible de $15\ \mu\text{m}$ se situent sous la courbe de simulation après 1,3 ns. On peut donc supposer que si les mesures avaient été effectuées plus tôt pour ce tir, les points obtenus auraient été au-dessus de la simulation. Ce résultat n'est donc pas incompatible avec le tableau global de données expérimentales au-dessus des simulations pour la M30.

- Pour les tirs avec M30 à haute intensité (figure 5.3 (c)), on observe à nouveau un bon accord simulation-expérience en terme d'allure de l'évolution temporelle des modulations. On peut aussi voir que les inversions de phase ont encore lieu plus tôt dans les simulations que dans l'expérience, et que le niveau des données expérimentales est supérieur à celui des simulations. En revanche, un tir avec une cible de $50\ \mu\text{m}$ ne suit pas ces tendances (triangles verts) : l'inversion de phase a lieu nettement plus tardivement pour ce tir que pour les autres.

Il faut aussi noter que les simulations réalisées sont initialement monomodes (initialement car des harmoniques peuvent ensuite se développer), ce qui n'est pas le cas du motif imprimé par les lames de phase (cf chapitre précédent). Cependant, les modulations de densité surfacique ne dépassent jamais $1\ \text{mg}/\text{cm}^2$, ce qui correspond à une amplitude d'environ $3\ \mu\text{m}$. Ainsi, l'amplitude des modulations n'excède jamais le dixième de leur longueur d'onde, et reste même pour la majorité des données très en-deçà de cette amplitude. On peut alors supposer que l'on reste toujours en phase linéaire ou en début de phase faiblement non-linéaire et que les modes croissent indépendamment les uns des autres. La description des expériences par des simulations initialement monomodes doit donc être représentative de l'expérience.

5.3 Interprétation avec le modèle de Goncharov

5.3.1 Comparaison des simulations CHIC et du modèle de Goncharov

Pour interpréter les différentes observations faites dans la partie précédente, nous allons utiliser le modèle de Goncharov. Ce modèle n'est valable que dans la limite $kC_s t \gg 1$. Nous considérons donc qu'il est valide pour $kC_s t \geq \pi$ comme dans la réf. [11]. On trouve alors comme temps de début de validité $t_{\text{valid}}=0,75\ \text{ns}$ pour la M30 à basse intensité et $0,5\ \text{ns}$ à haute intensité, $1,5\ \text{ns}$ pour la M60 et $3,0\ \text{ns}$ pour la M120. Toutes les données des tirs utilisant la M120 ayant été collectées avant $2,0\ \text{ns}$, nous ne les confronterons donc pas avec le modèle de Goncharov. Pour les autres tirs, la comparaison modèle-simulations

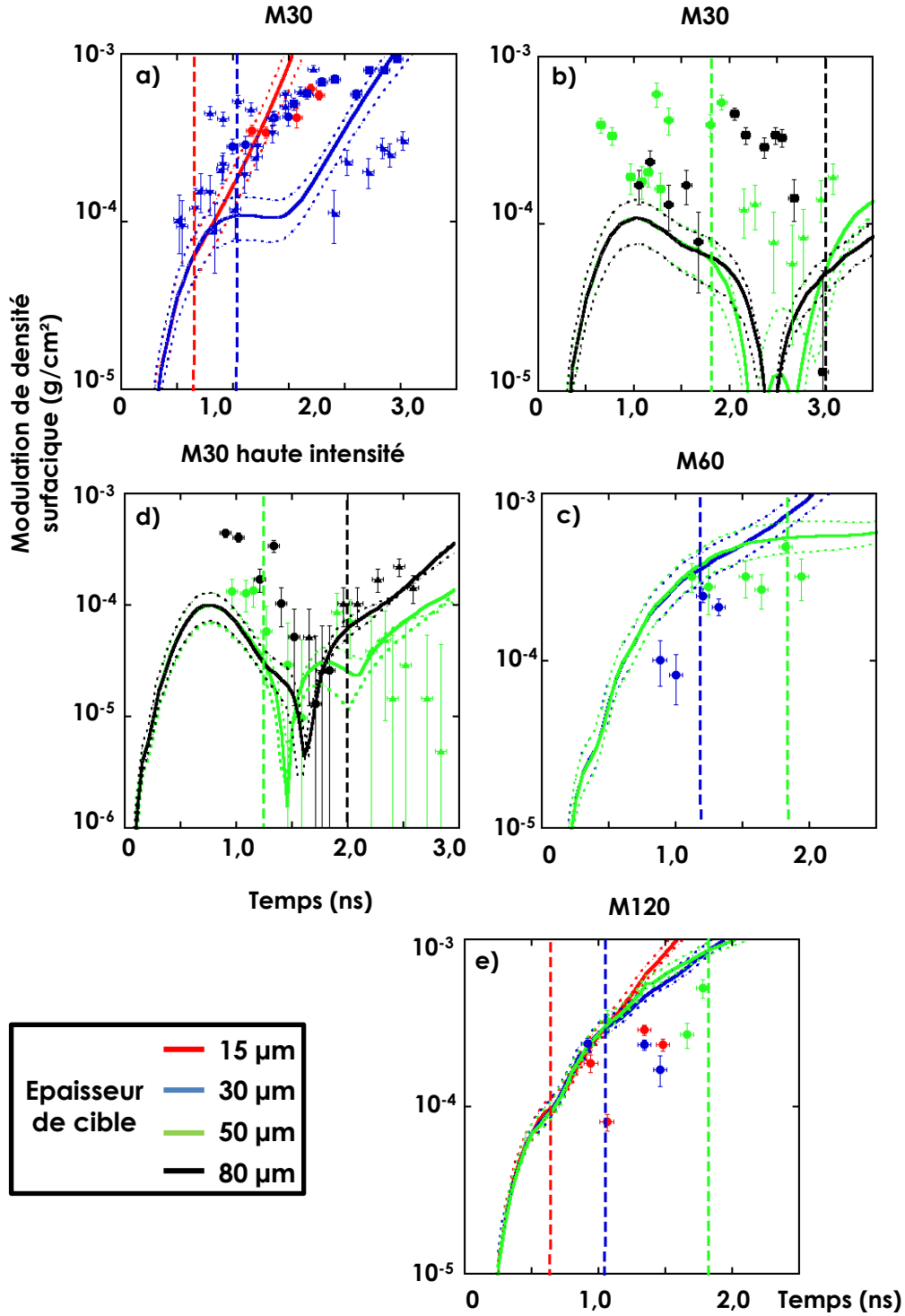


FIGURE 5.3 – Comparaison de l'évolution temporelle des modulations de densité surfacique extraites des expériences, et simulations correspondantes réalisées avec CHIC pour la lame de phase M30 à basse intensité (a-b), à haute intensité (c), pour la M60 (d) et pour la M120 (e). Les différentes couleurs correspondent à différentes épaisseurs de cible. Les barres verticales représentent la transition de l'IRM à l'IRT. Les simulations en trait plein ont été réalisées avec les valeurs moyennes de modulation de l'intensité laser ; les courbes pointillées les encadrant l'ont été avec les valeurs minimales et maximales de modulation de l'intensité laser.

est présentée sur la figure 5.4. Le modèle correspond à la courbe grise ; la ligne pointillée représente sa partie non valide ($t < t_{valid}$) et la ligne tireté sa partie valide ($t > t_{valid}$). Le modèle de Goncharov décrit l’empreinte des modulations et leur évolution par l’IRM ablative. Il ne prend pas compte le fait que la cible va finir par accélérer et que l’IRM ablative va laisser place à l’IRT ablative - la cible est considérée d’épaisseur infinie. Ainsi, l’épaisseur de cible n’entre pas en compte dans ce modèle ; sur chaque graphe, les simulations faites pour différentes épaisseurs sont à comparer à la même courbe qui représente le modèle de Goncharov. De plus, à partir du début d’accélération (barres verticales), le modèle n’est plus valable. Le modèle peut cependant continuer à montrer un accord avec les simulations comme avec les données expérimentales quelques centaines de picosecondes après le début de l’accélération du front d’ablation ; ceci est dû au fait que l’IRM et l’IRT ablatives sont en compétition durant une période avant que l’IRT ne prenne le pas sur l’IRM.

Les comparaisons entre le modèle et les simulations pour la M30 (basse et haute intensité) et pour la M60, représentées en figure 5.4, montrent une même tendance sur les trois graphes : l’amplitude maximale des modulations et la fréquence des oscillations sont systématiquement supérieures de quelques dizaines de pourcents avec le modèle. Elle restent néanmoins du même ordre de grandeur dans les deux cas. Ces différences peuvent s’expliquer par le fait que l’on utilise des paramètres constants pour calculer la solution de Goncharov alors que ces paramètres varient au cours du temps dans les simulations.

5.3.2 Interprétation des observations expérimentales par le modèle de Goncharov

La figure 5.5 présente la comparaison entre le modèle de Goncharov et les données expérimentales ; c’est la première fois que ce modèle est confronté à l’expérience. On remarque le même type de désaccords qu’entre expérience et simulations : amplitude légèrement surestimée pour la M60 et plus nettement sous-estimée pour la M30, période d’oscillation plus courte. Cette cohérence des désaccords semble logique car les paramètres utilisés pour calculer le modèle sont issus de simulations. Pour interpréter ces observations, les deux paramètres fondamentaux issus de l’équation (2.39) et de la réf. [11] sont la période d’oscillation T_{RM} et l’amplitude maximale des modulations η_m

$$T_{RM} = \frac{\lambda}{\sqrt{V_a V_{bl}}} \quad (5.5)$$

$$\eta_m \propto \frac{c_s^2 \delta I / I}{k V_c \sqrt{V_a V_{bl}}} \quad (5.6)$$

La fréquence des oscillations semblant sur-évaluée par le modèle et les simulations, nous avons donc recalculé la solution du modèle en réalisant un abattement de 35 % sur

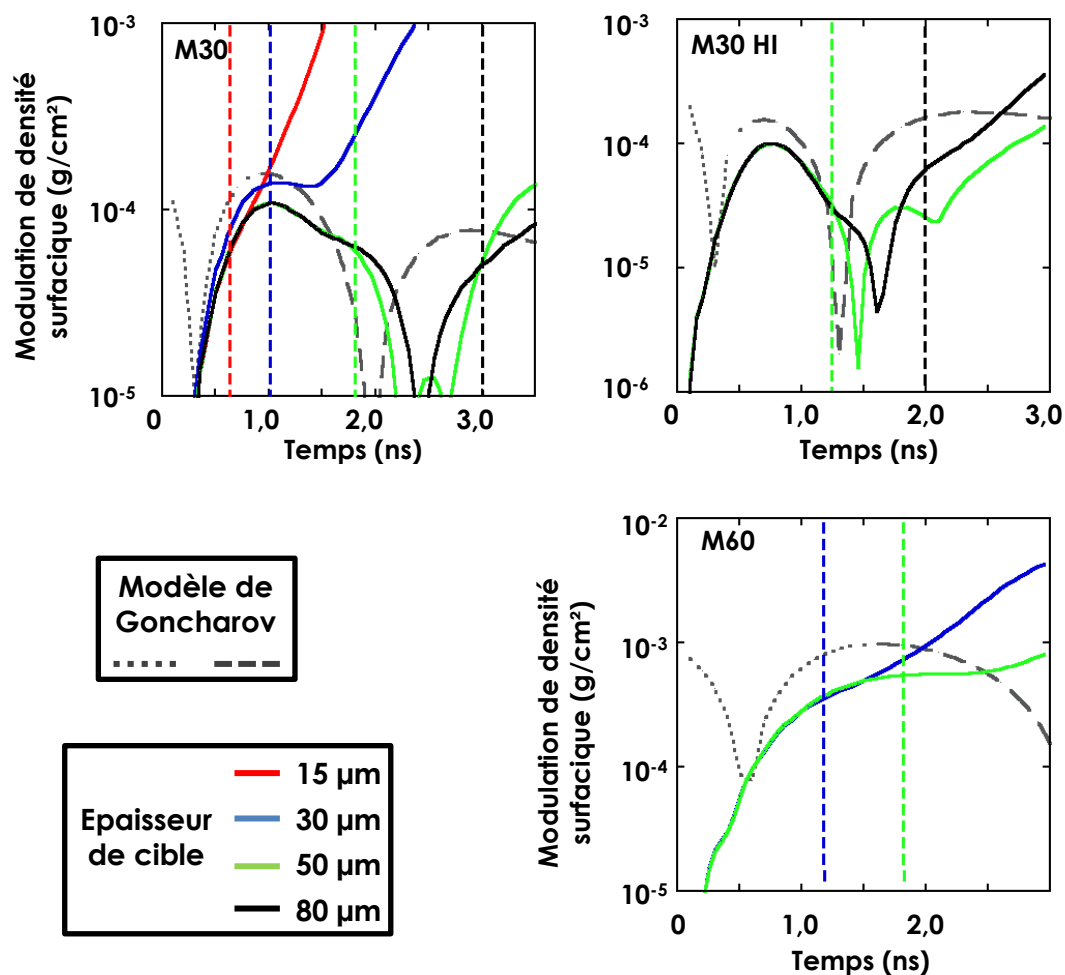


FIGURE 5.4 – Comparaison des données expérimentales, des simulations utilisant la valeur de modulation de l'intensité moyenne et du modèle de Goncharov (courbes grises) pour la M30 (a), M30 à haute intensité (b) et M60 (c). La partie pointillée de la courbe issue du modèle correspond à la zone où le modèle n'est pas applicable, et la partie tiretée là où le modèle est applicable.

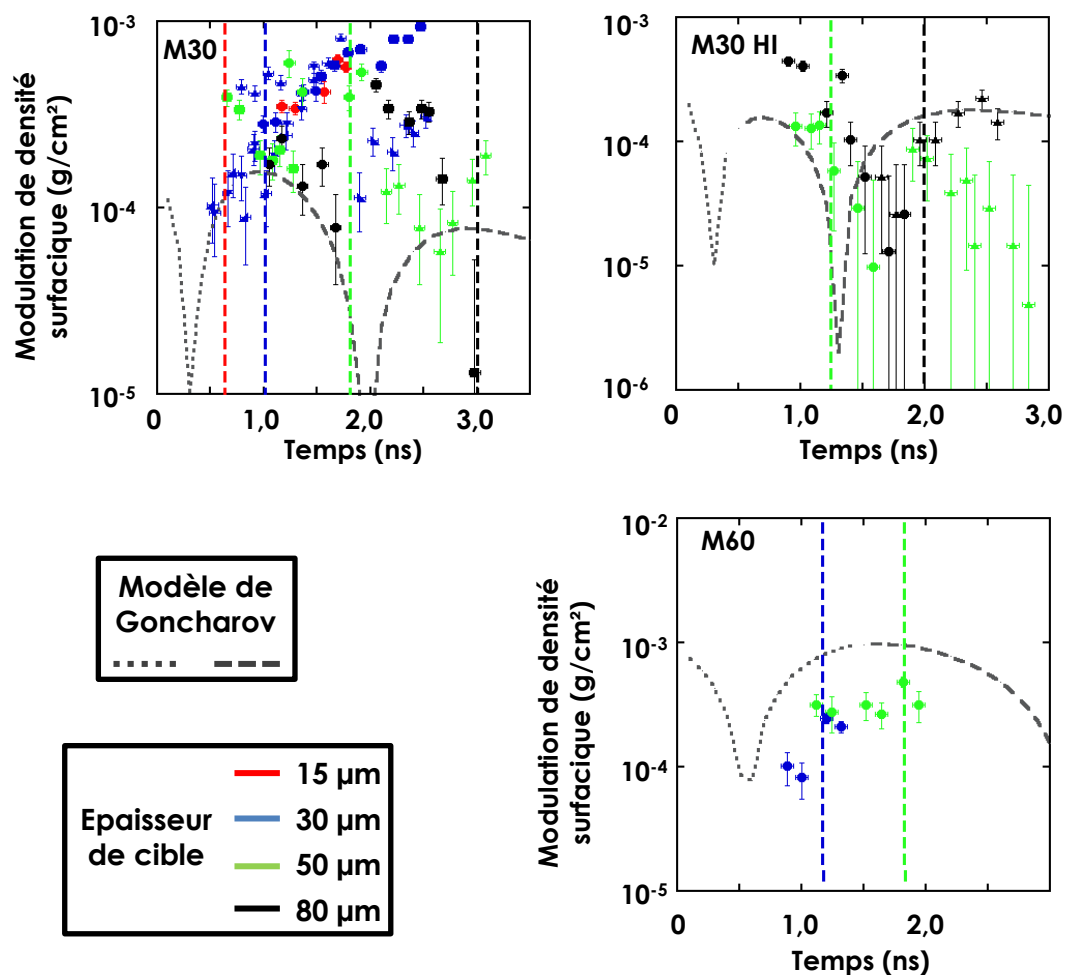


FIGURE 5.5 – Comparaison des données expérimentales, des simulations utilisant la valeur de modulation de l'intensité moyenne et du modèle de Goncharov (courbes grises) pour la M30 (a), M30 à haute intensité (b) et M60 (c). La partie pointillée de la courbe issue du modèle correspond à la zone où le modèle n'est pas applicable, et la partie tiretée là où le modèle est applicable.

V_a et de 25 % sur c_s . Ce type d'abattement est justifié : en effet, la comparaison d'une simulation CHIC 1D et d'une simulation CHIC 2D équivalente incluant la configuration réaliste de l'expérience montre un débouché de choc plus tardif et une accélération moins importante dans le cas de la simulation 2D. Ainsi, pour se recalculer sur l'écoulement 2D, qui est plus représentatif de l'expérience, il faudrait diminuer l'intensité utilisée dans les simulations 1D de quelques dizaines de pourcents. Une telle réduction de l'intensité ne suffirait néanmoins pas à retrouver les abattements utilisés pour V_a et c_s .

Le résultat de ce calcul avec abattement est présenté en figure 5.6. Le modèle avec l'abattement sur les paramètres est représenté en rouge, et celui avec les paramètres initiaux en gris. Pour la M60, l'accord en terme d'allure comme d'amplitude est presque parfait. Pour la M30 à basse intensité, l'inversion de phase du modèle se produit au même instant que dans l'expérience. A haute intensité, la période d'oscillation est même un peu surestimée. En effet, en diminuant la vitesse d'ablation (et donc aussi la vitesse de "blow-off"), d'après l'équation (5.5), la période d'oscillation augmente. D'autre part, pour toutes les données obtenues avec la M30, l'amplitude des modulations est toujours nettement sous-estimée avec le modèle. Or, d'après le chapitre précédent, le $\frac{\delta I}{I}$ de la M30 est probablement sous-évalué. D'après la formule (5.6), augmenter $\frac{\delta I}{I}$ accroît l'amplitude maximale des modulations. En résumé, il semble que les simulations ou leur post-traitement surestiment V_a et c_s , et que le $\frac{\delta I}{I}$ de la M30 calculé à partir des images ETP est sous-estimé. Cette dernière hypothèse va être testée dans la section suivante.

Un autre point reste à expliquer : les données d'un tir M30 à la plus haute intensité pour une épaisseur de cible de 50 μm (triangles verts) présentent une inversion de phase vers 3,0 ns, tandis que les autres données, ainsi que le modèle et les simulations prédisent l'inversion vers 1,5 ns. Deux images de la cible extraites de la radiographie XRFC de ce tir problématique sont présentées en figure 5.7. L'image (a) correspond à un temps de mesure de 2,52 ns et l'image (b) à un temps de 2,83 ns. Sur l'image (a), la partie de droite présente des modulations, ce qui n'est pas le cas de la partie de gauche, tandis que sur l'image (b), les modulations ne sont observées nulle part. Ainsi les modulations de la partie droite s'inversent entre 2,52 et 2,83 ns tandis que sur la partie gauche de l'image, il n'y a déjà plus de modulations à 2,52 ns. C'est le signe que les modulations se sont inversées plus tôt. La partie droite correspond donc à une zone où T_{RM} est plus grand, d'où une vitesse d'ablation et une intensité plus faibles. On peut donc expliquer le désaccord entre les données du tir représenté par les triangles verts et les données des autres tirs à haute intensité par le fait que les images de ce tir problématique ont été obtenues sur le bord du spot laser, la partie de droite correspondant à la zone la plus décentrée.

A travers le modèle de Goncharov, nous avons vu la sensibilité de l'IRM ablative, et particulièrement de l'inversion de phase, à des petites variations des paramètres. Dans la section suivante, nous allons donc étudier l'effet des variations des paramètres expérimen-

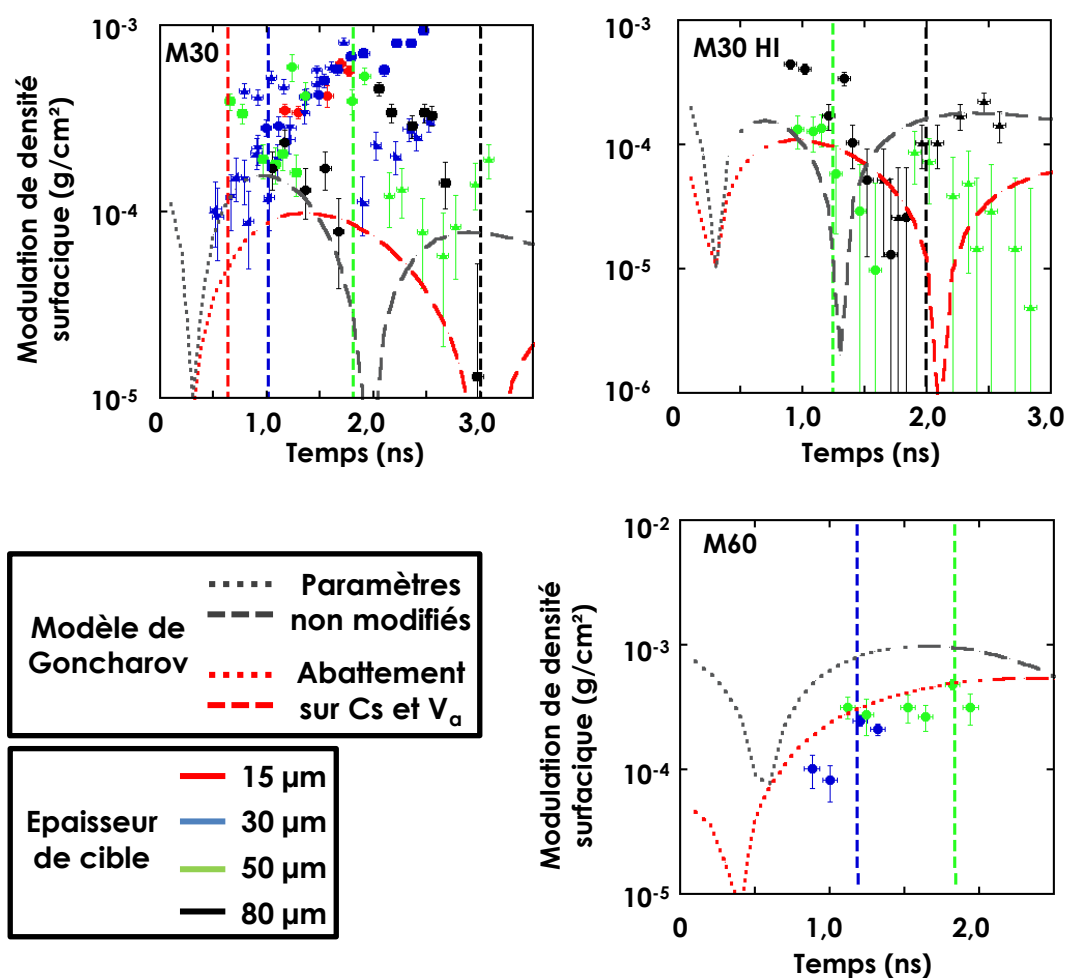


FIGURE 5.6 – Comparaison des données expérimentales, du modèle de Goncharov utilisant les paramètres extraits des simulations (courbes grises) et du modèle de Goncharov calculé avec un abattement de 25 % sur c_s et de 35 % sur V_a .

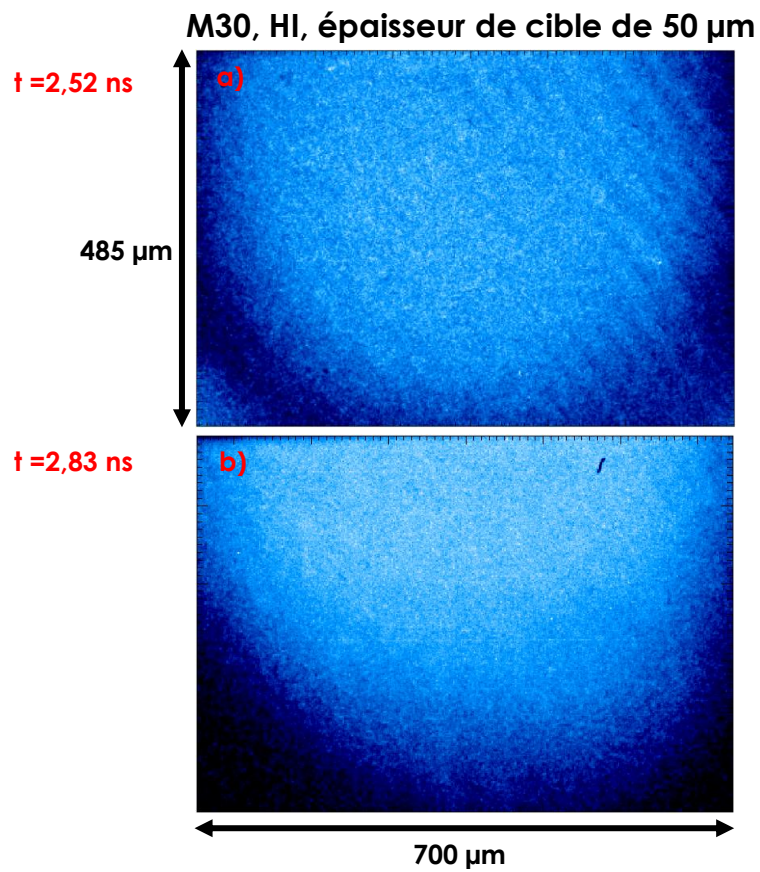


FIGURE 5.7 – Images de la cible extraites de la radiographie de face du tir avec M30 et cible de 50 μm à haute intensité présentant un fort désaccord avec les simulations (cf figure 5.5 (b)) à $t=2,52 \text{ ns}$ (a) et $t=2,83 \text{ ns}$ (b).

taux et numériques sur les simulations.

5.4 Etude des variations des simulations CHIC en fonction des paramètres numériques et expérimentaux

5.4.1 Variation du niveau d’empreinte

Pour vérifier l’hypothèse de la modulation de l’intensité laser $\frac{\delta I}{I}$ sous-estimée pour la M30, nous avons réalisé des simulations et calculé le modèle de Goncharov en utilisant des niveaux de modulation de l’intensité laser de $2 \times \frac{\delta I}{I}$ et de $3 \times \frac{\delta I}{I}$. Ces résultats sont représentés en figure 5.8 (a-b) pour le niveau de modulation de $2 \times \frac{\delta I}{I}$ et (c-d) pour $3 \times \frac{\delta I}{I}$, à basse et haute intensités. La période d’oscillation est toujours sous-estimée car le niveau de modulation de l’intensité n’intervient pas dans l’équation (5.5). En revanche, l’amplitude des oscillations de l’IRM ablative est plus importante quand on augmente $\frac{\delta I}{I}$. Certaines données expérimentales sont d’amplitude similaire aux calculs à $2 \times \frac{\delta I}{I}$, d’autres à $3 \times \frac{\delta I}{I}$; dans tous les cas, augmenter le niveau de modulation de l’intensité permet de faire disparaître la sous-estimation systématique de l’amplitude des modulations de densité surfacique dans les calculs. Le fait qu’un facteur 2 convienne bien pour certains tirs et un facteur 3 mieux pour d’autres traduit probablement le fait que d’un tir à l’autre, l’image analysée peut correspondre à l’impression d’une zone différente de la lame de phase M30 et qu’elle peut être plus ou moins décentrée par rapport au spot du laser.

5.4.2 Variation de la forme de l’impulsion

Les données expérimentales sont issues de différents tirs; or les comparaisons effectuées sur les figures précédentes qui englobent tous les tirs reviennent à les considérer comme équivalents d’un point de vue hydrodynamique. Les conditions expérimentales requises pour les différents tirs sont similaires; cependant, la forme de l’impulsion peut varier d’un tir à un autre comme on peut le voir sur la figure 5.9. Sur cette figure, les impulsions de trois tirs sont représentées. Bien que l’allure générale des impulsions soit très stable, on peut observer des écarts de quelques pour-cents à certains moments. Nous avons donc réalisé des simulations avec ces différentes impulsions pour deux conditions expérimentales : cible de $80 \mu\text{m}$ imprimée par M30 et cible de $50 \mu\text{m}$ imprimée par M60. Les résultats de ces simulations sont représentés en figure 5.10. On peut voir que pour la M60, pour laquelle il n’y a pas d’inversion de phase, il n’y a aucune différence entre les différentes simulations. En revanche, pour la M30, on remarque que l’instant d’inversion de phase varie sensiblement d’une simulation à l’autre, avec un écart d’environ 100 ps entre celle utilisant l’impulsion 65941 et celle utilisant l’impulsion 65944. D’un autre côté,

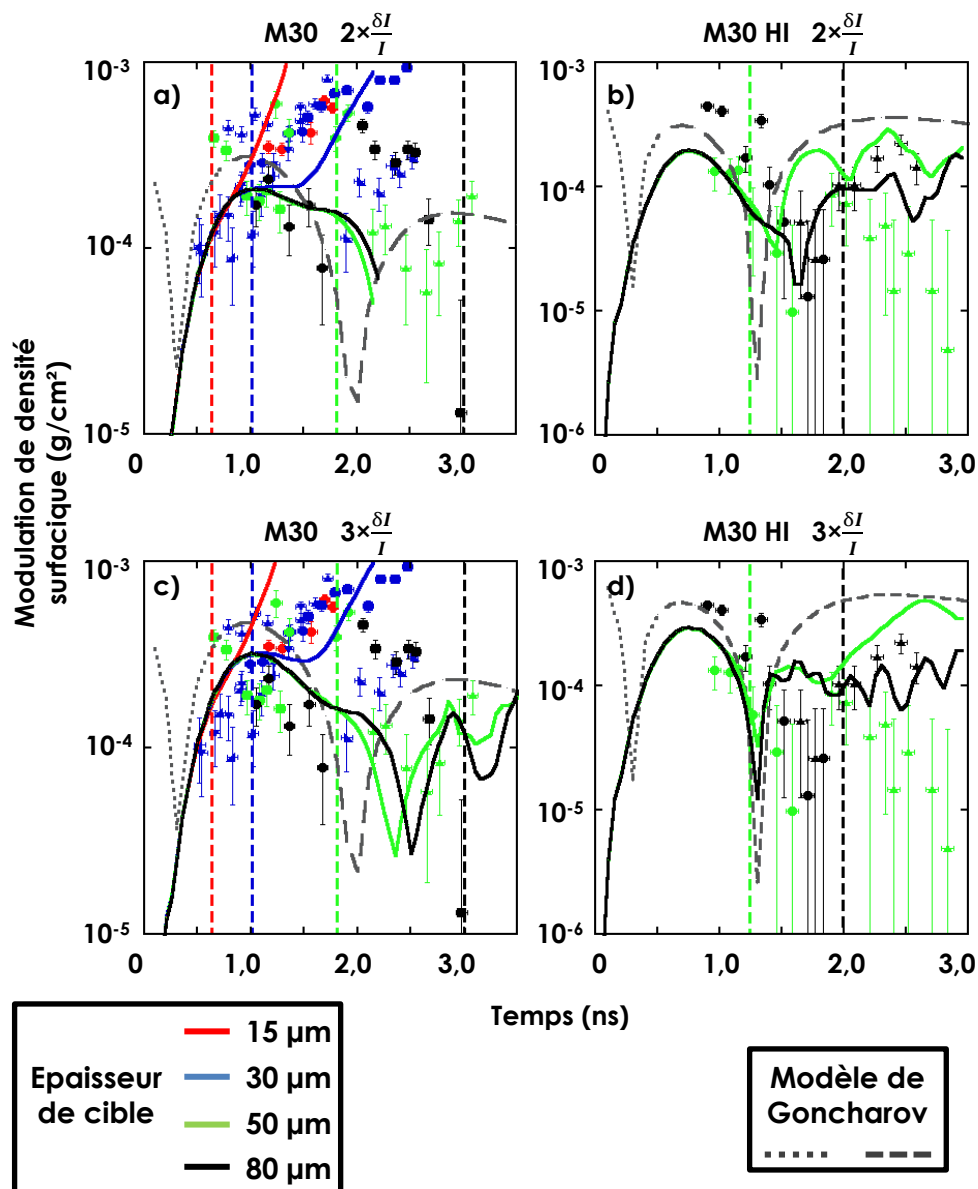


FIGURE 5.8 – Comparaison des données expérimentales avec le modèle de Goncharov et des simulations CHIC pour la M30 en utilisant un niveau de modulation de l'intensité multiplié d'un facteur 2, à basse (a) et haute (b) intensité, et multiplié d'un facteur 3, à basse (c) et haute (d) intensité.

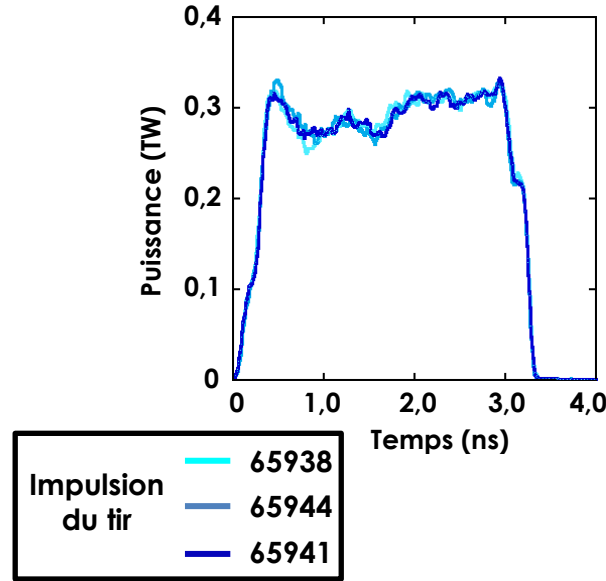


FIGURE 5.9 – Comparaison des puissances délivrées dans les impulsions de trois tirs différents en fonction du temps.

l'amplitude maximale des modulations est la même pour toutes les simulations.

En résumé, un point fondamental à retenir de ces variations sur l'impulsion est que l'instant d'inversion de phase est extrêmement sensible aux variations des paramètres et des conditions expérimentales, bien plus que le reste de la phase Richtmyer-Meshkov ou que la croissance par l'IRT ablative. Ainsi, la prédiction de cet instant préalablement à une expérience ne peut se faire qu'avec une précision toute relative. D'autre part, englober les différents tirs dans un même graphe lors des comparaisons expériences-calculs semble raisonnable, en gardant néanmoins à l'esprit la sensibilité sur l'instant d'inversion de phase.

5.4.3 Variation de l'équation d'état et du limiteur de flux

Les équations d'état permettent de relier les différentes quantités thermodynamiques entre elles. Dans la réf. [88], des mesures de vitesse de choc et de température de cibles de CH et de CH₂ sont comparées à des simulations réalisées en utilisant différentes équations d'état. Pour le CH₂, les auteurs montrent que les équations d'état SESAME 7171 comme 7180 permettent d'obtenir un bon accord entre les simulations et les expériences. Nous avons donc réalisé des simulations avec différentes équations d'état pour déterminer si le choix de la SESAME 7180 était optimal. Quatre équations d'état ont été étudiées : SESAME 7171, SESAME 7180, QEoS développée par Moore et collaborateurs [89] et modélisation par un gaz parfait (GP). Les résultats, pour des simulations avec une cible de 80 μm imprimée par M30 et cible de 50 μm imprimée par M60, sont présentés en figure

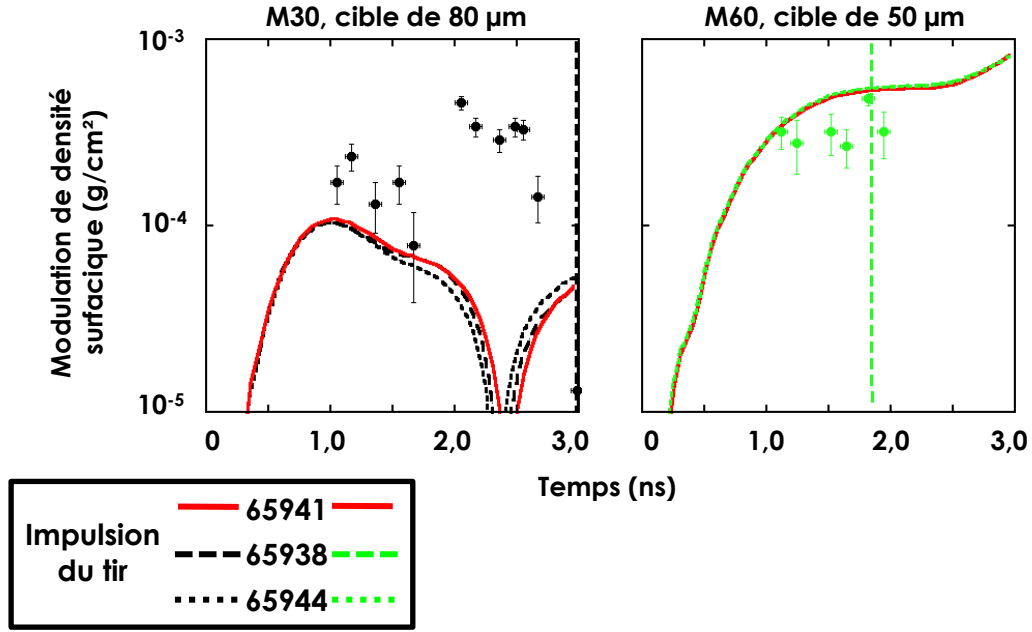


FIGURE 5.10 – Simulations réalisées en utilisant les différents impulsions présentées en figure 5.9 pour une cible de $80\ \mu\text{m}$ imprimée par M30 et une cible de $50\ \mu\text{m}$ imprimée par M60.

5.11. On peut noter deux points : tout d'abord, par rapport aux simulations réalisées avec les équations SESAME, les simulations utilisant les modèles QEoS et GP présentent des inversions de phase arrivant plus tôt. Or nous avons vu précédemment que les simulations sous-estimaient la période d'oscillation des modulations ; ainsi, il semble que les simulations réalisées avec les équations d'état SESAME soient optimales. D'autre part, avec les équations d'état SESAME 7171 et 7180, les simulations sont extrêmement semblables. Ceci vient corroborer les résultats obtenus dans la réf. [88]. On peut donc penser que le choix de la SESAME 7180 pour réaliser les simulations convenait.

Comme on a pu le voir dans le chapitre 2 et la section 3.4 du chapitre 3, un autre paramètre à fixer dans les simulations est le limiteur de flux, qui représente généralement quelques pour-cents du flux de chaleur maximal et a été introduit dans les codes de FCI pour retrouver les résultats expérimentaux. Pour les simulations présentées précédemment, nous avons utilisé un limiteur de flux de 7 % de type "sharp cut-off" car c'est une valeur couramment utilisée dans les codes d'hydrodynamique radiative dans la gamme d'intensité utilisée dans ces expériences. Pour étudier l'influence de ce paramètre, nous avons donc réalisé des simulations utilisant divers valeurs de limiteurs de flux : 3, 7 et 11 %. Les résultats de diverses simulations effectuées sont représentés en figure 5.12 pour les couples épaisseur de cible/lame de phase suivants : $80\ \mu\text{m}$ /M30 (a), $50\ \mu\text{m}$ /M60 (b), $50\ \mu\text{m}$ /M30 (haute intensité) (c) et $80\ \mu\text{m}$ /M30 (haute intensité) (d). Pour les simulations

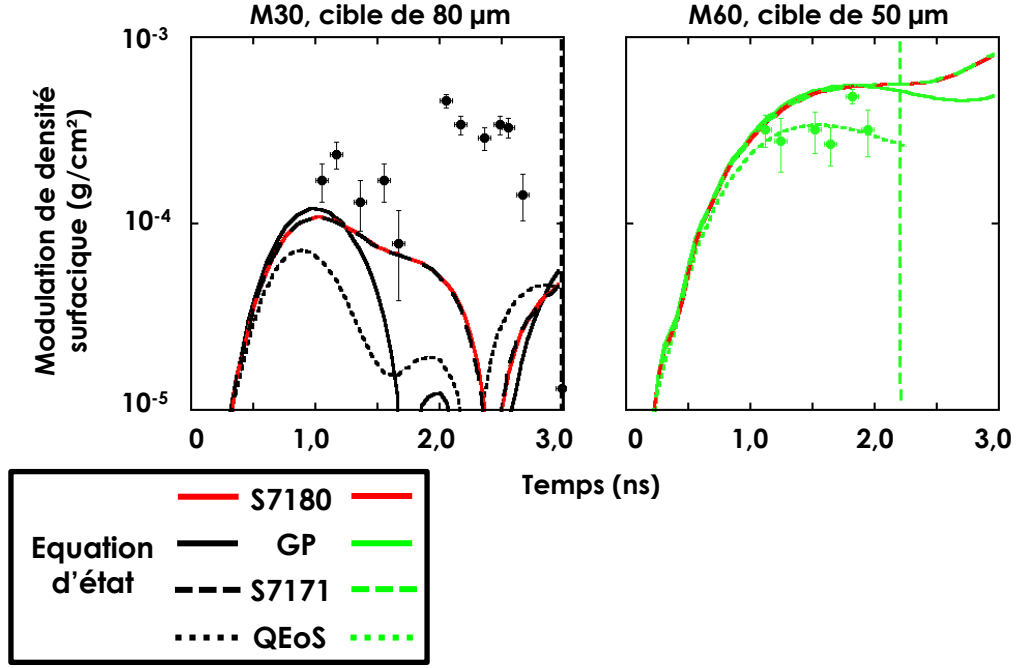


FIGURE 5.11 – Simulations réalisées en utilisant les différentes équations d'état SESAME 7171 (courbes tiretées vertes et noires), SESAME 7180 (courbes pleines rouges), gaz parfait (courbes pleines vertes et noires) et QEOs (courbes pointillées vertes et noires) pour une cible de 80 µm imprimée par M30 et une cible de 50 µm imprimée par M60.

à basse intensité (a-b), on peut voir que les simulations avec différents limiteurs de flux sont semblables, celles à 7 et 11 % sont mêmes confondues, ce qui sous-tend que le flux de chaleur n'atteint pas 7 % du flux maximal. A haute intensité, les simulations avec des limiteurs de flux de 7 et 11 % sont encore très proches. La simulation à 3 % donne une période d'oscillation plus longue. Ceci pourrait être expliqué par le fait que diminuer la valeur du limiteur de flux réduit l'ablation, de même pour la vitesse d'ablation et donc la fréquence des oscillations. L'amplitude des modulations n'est de son côté pas affectée par les variations de limiteur de flux. Cependant, si on s'intéresse par exemple à la réf. [66], on peut voir que pour reproduire l'hydrodynamique des cibles et le comportement de l'IRM ablative, les auteurs ne descendent jamais leur limiteur de flux en dessous de 4 %, même pour un limiteur de flux variable en temps. Une piste étudiée dans la réf. [66] est l'utilisation d'un modèle de transport non-local. Mais dans le cas de cet article, les expériences sont réalisées à une intensité supérieure à 4.10^{14} W/cm². Aux intensités auxquelles nous travaillons, le flux d'électrons ne devrait pas être délocalisé et donc l'utilisation d'un modèle de transport non-local ne devrait pas influencer le résultat des simulations (cf réf. [90] par exemple). En revanche, une mesure de trajectoire de la feuille pourrait permettre, par croisement avec les simulations de croissance des modulations imprimées, de caler la valeur du limiteur de flux.

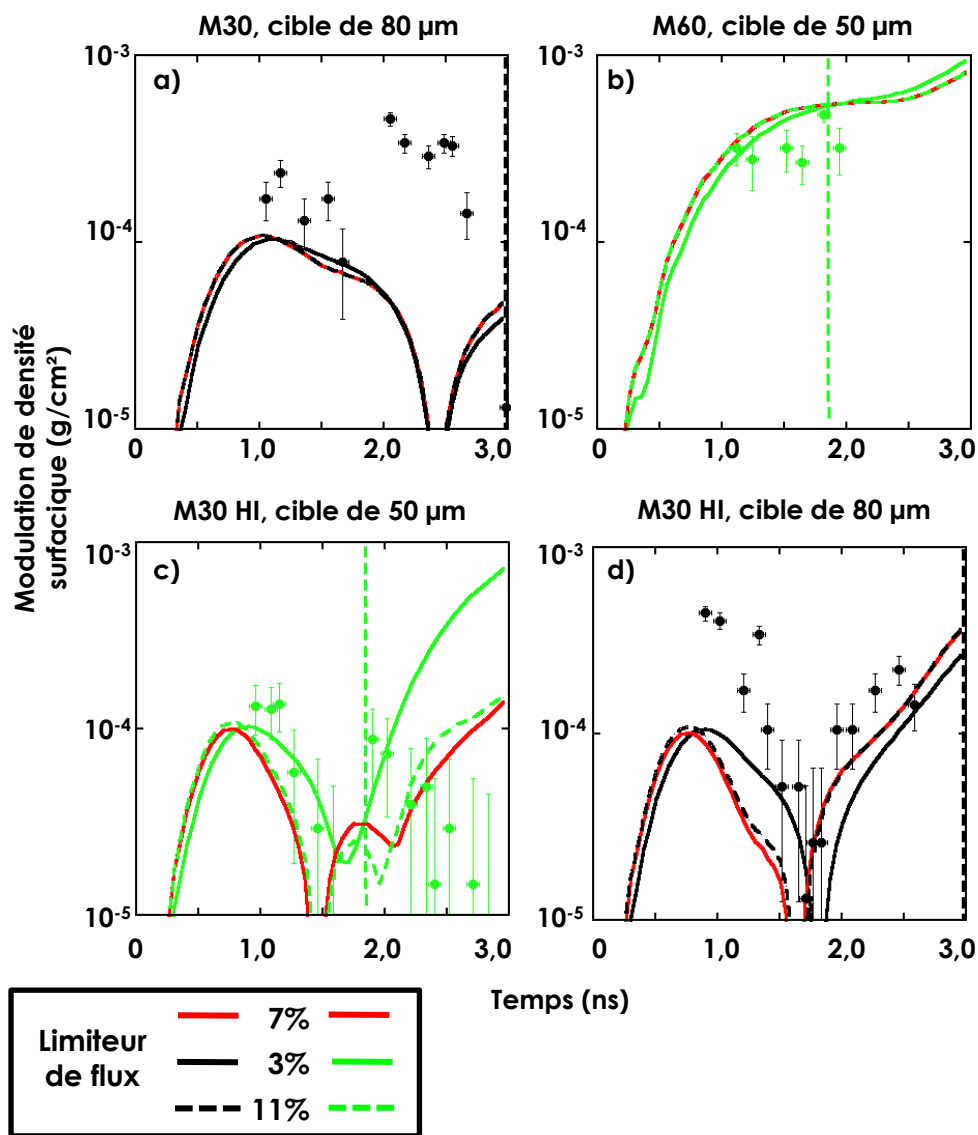


FIGURE 5.12 – Simulations réalisées en utilisant les valeurs de limiteur de flux à 11 % (courbes tiretées vertes et noires), 7 % (courbes pleines rouges) et 3 % (courbes pleines vertes et noires) pour une cible de 80 μm imprimée par M30 (a), une cible de 50 μm imprimée par M60 (b), et les tirs à haute intensité avec des cibles de 50 μm (c) et 80 μm (d).

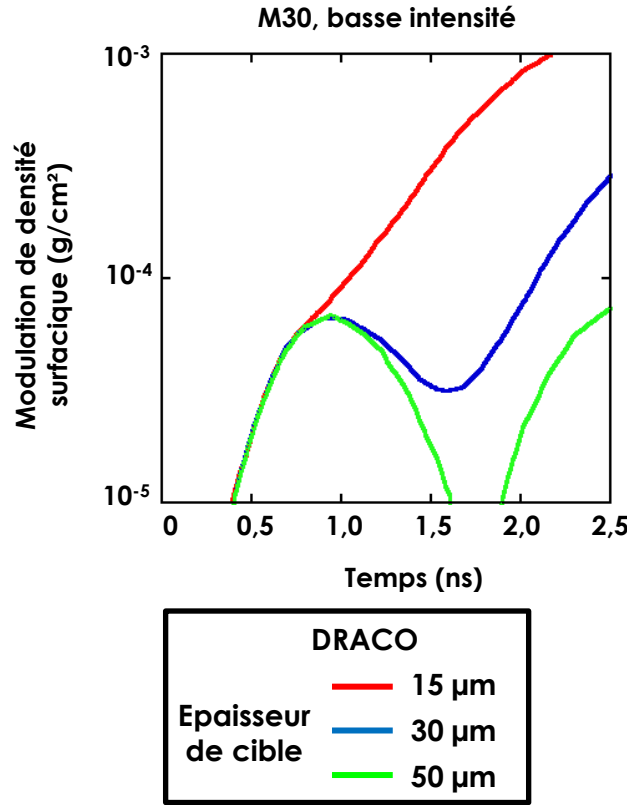


FIGURE 5.13 – Simulations DRACO réalisées avec une intensité d'environ $5 \cdot 10^{13} \text{ W/cm}^2$ pour des cibles de 15 μm , 30 μm et 50 μm d'épaisseur imprimées par M30.

5.5 Conséquences de l'étude au LLE

Notre collaborateur américain I. Igumenshev (LLE) a lui aussi réalisé certaines simulations de croissance des modulations par en utilisant le code DRACO du LLE. Le résultats des simulations DRACO pour des cibles imprimées par la lame de phase M30 à basse intensité est présentée en figure 5.13. Malgré l'utilisation de deux codes différents, les désaccords observés entre les simulations et les expériences sont identiques : l'amplitude des modulations et la période d'oscillation sont largement sous-estimées avec la lame de phase M30. D'autre part, comme pour CHIC, un bon accord est trouvé entre les simulations DRACO et les expériences pour les tirs avec la lame de phase M60.

Il est intéressant de noter ici qu'au début de nos expériences conjointes avec V. Sma-lyuk et D. Martinez (LLNL), la mesure de la phase Richtmyer-Meshkov imprimée par laser suscitait un intérêt modéré de la part de certains chercheurs au sein du LLE, du fait des mesures passées de l'IRM issue de modulations préimposées sur la cible [66] et de l'existence du modèle pour l'IRM imprimée par laser [11]. Cependant, plusieurs événements ont suscité un regain d'intérêt pour cette campagne expérimentale. Premièrement, les difficultés du code DRACO à simuler les expériences pour la longueur d'onde 30 μm .

De plus, les récents résultats d'implosions cryogéniques en AD sur OMEGA ont montré que les instabilités hydrodynamiques avaient un effet plus important que prévu [91]. Enfin, des mesures réalisées sur OMEGA, effectuées par T. Boehly et ses collaborateurs et non publiées, ont montré, par l'utilisation d'un VISAR 2D [92], que les modulations d'un choc créé par la M30 étaient plus importantes que prévues. Ces différents points posent question : les différents désaccords observés viennent-ils seulement d'une mauvaise qualité de la lame de phase M30, ou sont-ils aussi liés à des questions de physique relatives à l'IRM ablative pour les plus petites longueurs d'onde ? Pour essayer de répondre à ces questions, il est prévu une conception d'une nouvelle lame de phase M30 de meilleure qualité, ainsi qu'une campagne interne de mesure de l'empreinte laser au sein du LLE.

5.6 Conclusion sur les interprétations

Pour conclure, les divers éléments à retenir de ce chapitre sont les suivants :

- En s'appuyant sur la puissance laser délivré mesurée lors des tirs et sur des paramètres issus de simulations CHIC 1D, le modèle de Goncharov a été calculé et des simulations 2D monomodes de croissance des modulations ont été réalisées pour les différentes longueurs d'onde imprimées et les différentes épaisseurs de cible.
- Les simulations CHIC montrent un bon accord avec les données expérimentales pour les tirs avec les lames de phase M60 et M120. Avec la lame de phase M30, des désaccords apparaissent tant en terme d'amplitude des modulations que d'instant d'inversion de phase. Nos collègues américains ont constaté les mêmes comportements lors de la simulations des expériences avec le code DRACO du LLE.
- Les simulations ont été aussi comparés au modèle de Goncharov ; il en ressort un bon accord global et de petites différences qui peuvent être expliquées par le fait que le modèle a été calculé avec des paramètres qui ne varient pas au cours du temps.
- Le modèle de Goncharov pour l'empreinte de défauts de l'intensité laser a été comparé à des données expérimentales pour la première fois. Lorsqu'on compare ce modèle avec les expériences, les mêmes désaccords que ceux obtenus entre les simulations et les expériences apparaissent. L'interprétation à l'aide du modèle montre que ces écarts peuvent être imputés à une sur-estimation de la vitesse d'ablation et de la vitesse acoustique d'environ 20 %.
- L'utilisation du modèle a aussi permis de confirmer les observations réalisées à partir des simulations : le niveau initial des modulations de l'intensité laser est sous-estimé d'un facteur 2 à 3 pour la lame de phase M30, comme supposé à la fin du chapitre 4.
- Plus généralement, le modèle de Goncharov s'avère être un outil d'interprétation

puissant, très utile notamment pour regarder la sensibilité de l'inversion de phase par rapport aux paramètres hydrodynamiques.

- Le modèle, ainsi que les simulations réalisées avec différentes impulsions font apparaître un point notable : l'instant d'inversion de phase est complexe à déterminer précisément et très sensible aux variations de paramètres numériques et expérimentaux.
- Si l'on désire contrôler l'IRM ablative imprimée par laser, on dispose de deux possibilités : concevoir la cible et l'impulsion de manière à ce que les longueurs d'onde les plus dangereuses soient autour de l'inversion de phase au moment de la transition avec l'IRT ablative, ou diminuer l'amplitude maximale des oscillations donnée par (5.6). La première méthode peut être intéressante mais complexe à mettre en place précisément : cela signifie adapter la cible et/ou la vitesse d'ablation (dont dépend la période d'oscillation (5.5)) sur des designs d'implosions dont les paramètres sont déjà optimisés, pour un bénéfice qui reste incertain. En effet, l'instant d'inversion est particulièrement délicat à déterminer, comme on a pu le voir, et très sensible aux paramètres expérimentaux, comme la variation des impulsions d'un tir à l'autre (cf figure 5.10). De plus, si certaines longueurs d'onde seront en train de s'inverser au moment de la transition IRM-IRT, d'autres seront plus proches de leur maximum d'oscillation. On peut par exemple, comme dans la réf. [14] pour des modulations issues de la rugosité de la cible, chercher à minimiser la croissance par l'IRM ablative des modes qui sont les plus nombreux à être imprimés. Mais dans ce cas, les modes qui se situent au pic de la courbe de croissance de l'IRT ablative pourraient se retrouver à un maximum d'oscillation au moment de la transition IRM-IRT ablative.
- La méthode la plus directe pour contrôler l'IRM ablative imprimée par laser est donc de réduire l'amplitude maximale (5.6). Pour cela, on peut par exemple penser à augmenter la vitesse d'ablation au début des implosions en augmentant l'intensité. Cependant, cela peut compliquer la compression en augmentant l'entropie et surtout, cela ne diminuera pas forcément η_m car la vitesse acoustique va alors augmenter. On voit d'ailleurs sur les simulations présentées en figure 5.3 que l'amplitude maximale d'oscillation pour la cible de $80 \mu\text{m}$ est la même à $5,6 \cdot 10^{13} \text{ W/cm}^2$ et à $1,6 \cdot 10^{14} \text{ W/cm}^2$. Le meilleur moyen de contrôler l'ensemencement de l'IRT ablative par l'IRM ablative est donc de diminuer le niveau de modulation de l'intensité laser $\frac{\delta I}{I}$. Dans le chapitre suivant, nous allons donc étudier la réduction des instabilités hydrodynamiques par l'utilisation de mousses sous-denses.

Chapitre 6

Démonstration de la réduction par des mousses sous-denses des instabilités hydrodynamiques 2D imprimées par un motif de référence

Dans le chapitre précédent, nous avons montré la bonne capacité des simulations et d'un modèle à reproduire le comportement de l'IRM ablative imprimée par laser. Cependant, les simulations peinent à prévoir précisément l'instant d'inversion de phase des modulations du front d'ablation. Ainsi, il semble qu'on puisse utiliser les simulations pour éviter d'être à un pic d'oscillation lors du passage à l'IRT, voire pour se trouver à des amplitudes faibles, mais il reste nécessaire de réduire le niveau initial des modulations imprimées par le laser. Au cours de ce chapitre, nous nous proposons de démontrer l'efficacité d'une nouvelle méthode de réduction de l'empreinte des défauts de l'intensité laser par l'utilisation de mousses de densité sous-critique par rapport à la longueur d'onde laser.

6.1 Configuration expérimentale

6.1.1 Contexte de l'expérience

Depuis plus de 20 ans, l'utilisation de mousses dans les expériences de FCI a suscité un fort intérêt et donc des études variées sur le sujet [20, 93, 94, 95, 96, 97]. Les réfs. [20, 93] présentent notamment des expériences dans lesquelles le lissage des inhomogénéités de l'intensité laser par conduction thermique a été montré. Cependant, dans ces diverses publications, l'objet d'étude était des mousses sur-denses, donc de densité supérieure à la densité critique à la longueur d'onde du laser. De plus, dans les études sur le lissage des défauts laser, la mousse sur-dense est ionisée par un flash de rayons X qui peut causer le préchauffage de la cible.

Un voie alternative pour le lissage des faisceaux est l'utilisation de plasmas sous-dense (dont la densité est inférieure à la densité critique du laser), donc dans lesquels le laser peut se propager et effectuer l'ionisation lui-même, sans passer par un flash de rayons X. Dans la réf. [98], les auteurs étudient la propagation d'un faisceau laser dans un gaz sous-dense et montrent le lissage des défauts d'intensité laser. Cependant, en attaque directe, des contraintes techniques peuvent empêcher l'utilisation d'un gaz autour de la cible pour permettre le lissage. Ainsi, un autre voie intéressante est l'utilisation de mousses sous-denses. Des expériences s'intéressant uniquement à l'aspect lissage ont été réalisées sur la LIL [22] avec des intensités de l'ordre de quelques 10^{14} W/cm². Les mousses étaient placées dans des supports ouverts des deux côtés opposés. Le faisceau laser se propageait à travers la mousse et la répartition d'intensité du faisceau en sortie était mesurée. Ces expériences ont permis de démontrer des effets de lissage des faisceaux laser par les instabilités paramétriques. Dans la continuité de ces travaux et en utilisant le même type de mousses, nous cherchons donc à démontrer pour la première fois l'effet de lissage de ces mousses sur le développement l'IRT ablative sur des feuilles de CH recouvertes de mousses sous-denses.

Comme nous cherchons à mesurer l'effet des mousses sur la croissance de l'IRT ablative, il faut que les défauts d'intensité laser puissent s'imprimer à la surface d'un ablateur. Ainsi, nous avons placé une feuille de CH en sortie de mousse. Cependant, cette nécessité empêche une mesure des faisceaux laser après propagation dans la mousse et l'on ne peut donc pas étudier les instabilités paramétriques responsables du lissage. Ce type d'étude a déjà été réalisé, notamment dans la réf. [22]. Au cours de la journée de tirs, nous n'avons donc pas fait de mesure de ce type et la journée a été consacrée uniquement à la mesure des instabilités hydrodynamiques par radiographie de face. La conception de l'expérience, que j'ai assurée, s'appuie sur la configuration utilisée lors du chapitre 4.

6.1.2 Cibles

Les cibles sont représentées en figure 6.1 : ce sont des feuilles de CH de 15 μ m d'épaisseur surmontées d'un support en cuivre. Dans le cas des cibles de CH seul, le support est vide ; dans le cas des cibles avec mousse, il est rempli d'une mousse sous-dense de densité 5, 7 ou 10 mg/cm³ et d'épaisseur 300 ou 500 μ m. Le support en Cu est de l'épaisseur de la mousse, de diamètre externe 8 mm et de diamètre interne 2,5 mm ; il comporte une fente de 1 mm de large pour permettre de mesurer l'émission propre de la mousse par caméra à balayage de fente. Les feuilles de CH ont été fournies par S. Fujioka de l'Institute of Laser Engineering (ILE) d'Osaka (Japon). Pour obtenir une rugosité suffisamment faible (< 20 nm), requise pour des expériences d'instabilités hydrodynamiques, les feuilles ont été comprimées entre deux plaques de verre de qualité optique et placées dans un four, dont elles ressortent avec la même rugosité que les plaques de verre. Une mesure de l'état

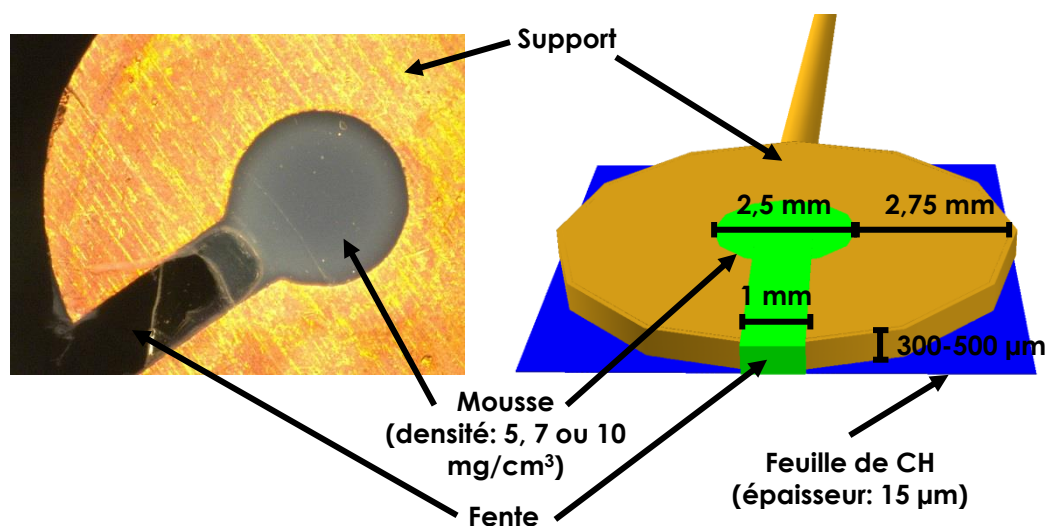


FIGURE 6.1 – Photographie (gauche) et schéma (droite) d'une cible avec mousse.

de surface d'une feuille de CH avant collage au support est montrée en figure 6.2 (a), ainsi qu'un profil associé en (b). On peut noter que la feuille semble courbée du coin bas à droite de l'image vers le coin haut à gauche. Cette courbure apparaît sur le profil. Les mousses, des aérogels de $C_{15}H_{20}O_6$ dont la taille des pores est de l'ordre du micron, sont fabriquées par N. Borisenko du Lebedev Institute de Moscou (Russie) ; une image de ces pores est présentée en figure 6.2 (c). Les mousses sont déposées dans les supports en Cu qui sont ensuite collés aux feuilles de CH. Enfin la pose des capillaires a lieu au LLE. Les feuilles de CH, d'épaisseur $15\ \mu\text{m}$ et de densité $1,05\ \text{g}/\text{cm}^3$ ont une densité surfacique de $1,575\ \text{mg}/\text{cm}^2$. La densité surfacique la plus élevée pour une mousse, celle de $7\ \text{mg}\cdot\text{cm}^{-3}/500\ \mu\text{m}$, est de $0,35\ \text{mg}/\text{cm}^2$, donc petite devant la densité surfacique de la feuille de CH. Nous choisirons donc la même source de radiographie pour effectuer des radiographies de face, soit de l'uranium, comme pour les cibles de $15\ \mu\text{m}$ d'épaisseur lors des expériences d'IRM ablative. La détermination de la densité de la mousse est faite au Lebedev Institute à l'aide de plusieurs méthodes : radiographie par rayons X mous, microscopie de fines épaisseurs de mousses, mesure de la masse de la mousse et de son volume. Grâce à la corrélation des résultats obtenus avec ces différentes méthodes, la densité des mousses est connue avec une barre d'erreur inférieure à 15 %.

6.1.3 Faisceaux laser

La figure 6.3 présente la configuration de l'expérience. Comme dans les expériences d'IRM ablative, un faisceau d'empreinte est utilisé. Il est porteur d'une lame de phase M30 ou M60 selon le tir. L'objectif est d'évaluer la capacité des mousses à réduire l'empreinte et l'IRT ablative subséquente ; l'utilisation des lames de phase M30 et M60 permet d'avoir un motif d'empreinte mesurable et contrôlé. Les autres faisceaux d'accélération (drive)

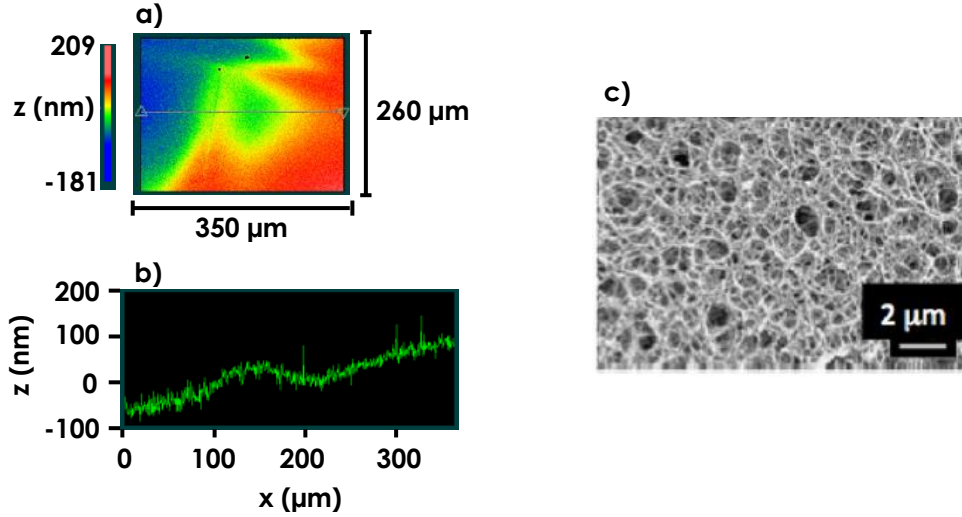


FIGURE 6.2 – a) Etat de surface d’une feuille de CH avant collage au support de Cu et création de la mousse. b) Profil horizontal extrait de l’image (a). Ces données nous ont été communiquées par N. Borisenko. c) Image des pores d’une mousse de 10 mg/cm^3 . Cette image est extraite de la réf. [99].

sont munis de lames de phase SG4. L’intensité laser requise est de quelques 10^{14} W/cm^2 , car les phénomènes d’instabilités paramétriques générant le lissage dans les mousses n’apparaissent qu’à partir d’une intensité seuil. Ensuite, la durée d’impulsion laser doit être supérieure à 1 ns : l’ionisation de la mousse peut prendre jusqu’à 500 ps et il est nécessaire d’attendre ensuite le développement des instabilités hydrodynamiques pour être sûr de mesurer les modulations imprimées sur la feuille de CH. On choisit donc pour chaque faisceau une impulsion carrée de 1 ns ; l’impulsion totale est donc obtenue par addition de plusieurs faisceaux. Ces impulsions permettent d’extraire 500 J/faisceau contre 330 J/faisceau pour celles de 2 ns et 180 J/faisceau pour celles de 3 ns. Les impulsions totales utilisées sont représentées en figure 6.4. Pour les cibles de CH seul, on utilise des impulsions carrées de 2 ns à environ $3 \cdot 10^{14} \text{ W/cm}^2$, formées par la superposition de 3 faisceaux à 23° pendant la première nanoseconde et 2 faisceaux à 23° et un à 48° pendant la deuxième nanoseconde (ce qui explique la petite baisse d’intensité pendant la deuxième nanoseconde). Pour les cibles avec mousse, une marche d’intensité plus élevée à environ $5 \cdot 10^{14} \text{ W/cm}^2$ est formée pendant la première nanoseconde par 3 faisceaux supplémentaires à 48° . Cette surintensité a pour but de compenser l’énergie laser absorbée par la mousse pour son ionisation et son chauffage. Ainsi, d’après des simulations CHIC, la puissance délivrée à la feuille de CH devrait être approximativement similaire avec et sans mousse, et donc les accélérations semblables. Ceci est fondamental pour que le taux de croissance de l’IRT soit le même dans les deux cas et donc que les différences observées entre les modulations des feuilles de CH avec et sans mousse ne proviennent que de l’effet de ces mousses. Le moment où le laser atteindra la feuille et donc le début de son accéléra-

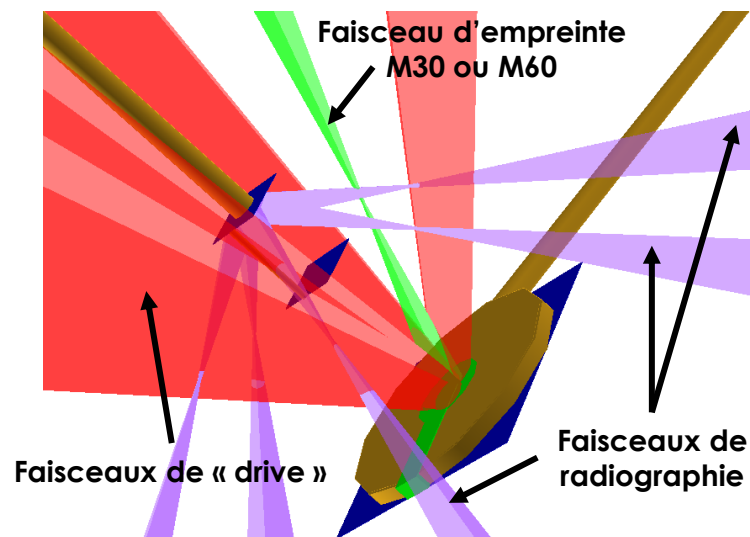


FIGURE 6.3 – Configuration expérimentale des expériences de réduction de l'empreinte laser par utilisation de mousses sous-denses. Les faisceaux de "drive" (rouge), d'empreinte (vert) et de radiographie (bleu) sont représentés.

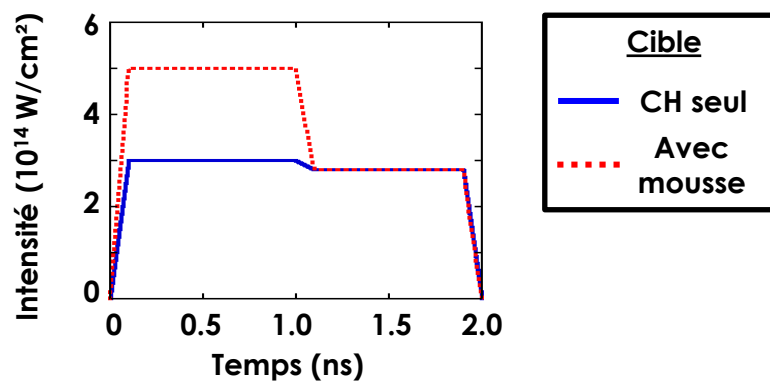


FIGURE 6.4 – Impulsions totales sur cible dans le cas de CH seul (courbe pleine bleue) et recouvert de mousse (courbe pointillée rouge).

tion seront simplement retardés du temps d'ionisation de la mousse. La portion d'énergie absorbée dans la mousse a été évaluée avant les expériences par des simulations CHIC. Dans l'idéal, il aurait fallu adapter la hauteur de la marche en fonction de la densité surfacique de la mousse, mais cela aurait nécessité des ajustements de réglage de conversion 3ω des faisceaux. Ces réglages sont réalisables mais nécessitent a minima 2 tirs pour être répétitifs. Ce n'était pas compatible avec notre plan de tir.

6.1.4 Diagnostics

Comme lors des expériences d'IRM ablative imprimée par laser, le diagnostic principal de cette expérience était un imageur X multi-sténopés (XRFC) pour effectuer des radiographies de face de la cible et mesurer ainsi l'évolution des modulations imprimées. Il était placé en TIM 2 (port H3), à l'opposé de l'expérience précédente, pour pouvoir illuminer la cible avec les faisceaux 25 et 30 qui sont les seuls pourvus des diagnostics de rétrodiffusion FABS. Le grandissement choisi était aussi de 12, la grille de sténopé permettant de former 16 images (4 par bande) pour les 3 premiers tirs puis 8 images pour les suivants, car avec la première grille, les images d'une bande se superposaient un peu sur celles des autres bandes et réduisaient donc la zone d'intérêt des mesures. De la même manière que pour les expériences d'IRM ablative, nous avons utilisé un autre imageur X multi-sténopé et les différents imageurs fixes à sténopé (XRPHC) pour contrôler les spots laser et l'émission de radiographie ; les diagnostics laser ont aussi permis de mesurer les impulsions réelles pour réaliser les simulations post-expérience. En revanche, deux diagnostics importants ont été ajoutés : une caméra à balayage de fente appelée SSCA (pour SIM Streaked Camera) dans le TIM 4 et les FABS. La SSCA a pour but de mesurer l'émission propre de la mousse pour obtenir la vitesse d'ionisation de celle-ci ainsi que la mise en vitesse de la feuille de CH. Les cibles ont donc été montées de façon à ce que la fente se situe dans la ligne de visée de la SSCA. Cependant, comme on peut le voir sur la figure 6.5 qui représente la vue des différents diagnostics, la SSCA n'est pas exactement orientée dans le plan de la feuille de CH mais forme un angle de $10,2^\circ$ avec celui-ci car il n'existe pas d'inserteur de diagnostic (TIM) qui soit perpendiculaire à l'axe H3-H18 (ni à aucun autre axe permettant de faire une radiographie de face d'ailleurs). Il faudra donc tenir compte de cet angle lors de l'interprétation des données. La résolution des images SSCA est de $1 \mu\text{m}/\text{pixel}$ en ordonnée et $2,3 \text{ ps}/\text{pixel}$ en abscisse. Les FABS vont pouvoir permettre de récolter la lumière rétrodiffusée pour un angle de 23° (FABS 30) et un angle de 48° (FABS 25) dans l'axe des faisceaux concernés.

Lors de la journée d'expérience, nous avons cherché à étudier l'effet des variations de densité des mousses sur le lissage. Ainsi, nous avons réalisé des tirs avec une mousse de $5 \text{ mg.cm}^{-3}/500 \mu\text{m}$, une de $7 \text{ mg.cm}^{-3}/500 \mu\text{m}$ et une de $10 \text{ mg.cm}^{-3}/300 \mu\text{m}$ pour les deux conditions d'empreinte (M30 et M60). Les mousses de 10 mg.cm^{-3} étaient d'une épaisseur $300 \mu\text{m}$ pour minimiser l'énergie absorbée par la mousse pour son chauffage. Nous disposions aussi d'une mousse de $7 \text{ mg.cm}^{-3}/300 \mu\text{m}$, mais celle-ci étant abîmée et en un seul exemplaire, nous n'avons pu obtenir de données significatives avec, et nous n'avons donc pu étudier l'effet de la longueur de la mousse à densité donnée.

Dans cette section, nous avons montré comment la conception de l'expérience a été optimisée pour obtenir un maximum d'information : radiographie de face des modulations de densité surfacique de la cible, trajectoire des zones d'auto-émission et énergie rétrodiffusée. Dans la section suivante, nous présentons les mesures d'auto-émission effectuées par

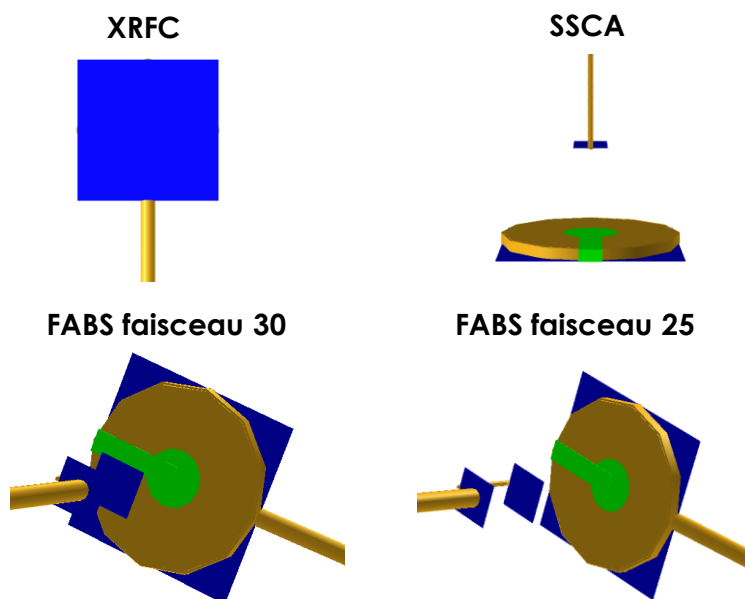


FIGURE 6.5 – Vues depuis les différents diagnostics : XRFC, SSCA et les deux FABS.

la caméra à balayage de fente, qui vont nous permettre de déterminer si les conditions d'accélération étaient similaires pour les différents types de cibles.

6.2 Analyse des données de la caméra à balayage de fente

6.2.1 Interprétation des images d'auto-émission

La figure 6.6 présente une image mesurée par la SSCA pour une cible avec mousse de 5 mg.cm^{-3} et de longueur $500 \mu\text{m}$. C'est une image d'auto-émission : la SSCA n'est pas couplée à une source de radiographie de côté, elle récupère seulement l'émission du plasma de la mousse et de la feuille. Les images (a-e) permettent d'interpréter l'image SSCA. La fente de la caméra, représentée en pointillés roses, est orientée dans la direction de la fente du support de cuivre (a). Lorsque les faisceaux laser sont allumés, un spot se forme à la surface de la mousse (b) qui va être ionisée et commencer à émettre. Sur l'image SSCA, on observe pour les différents tirs une zone de 150 à $200 \mu\text{m}$ au démarrage de l'émission, qui correspond à la projection du spot laser d'environ 1 mm sur l'angle de 10° de la SSCA. Ensuite, le plasma de la mousse progressivement formé va s'étendre dans la direction opposée à la feuille de CH ; cela correspond à la pente croissante sur la partie supérieure de l'image SSCA. D'autre part, le laser va se propager dans la mousse en l'ionisant (c), la transformant donc en plasma qui émet ; cette onde d'ionisation se retrouve dans la pente décroissant très fortement juste après le début d'émission. Le laser atteint ensuite la cible (d). Pendant plusieurs centaines de ps, le temps du transit du choc

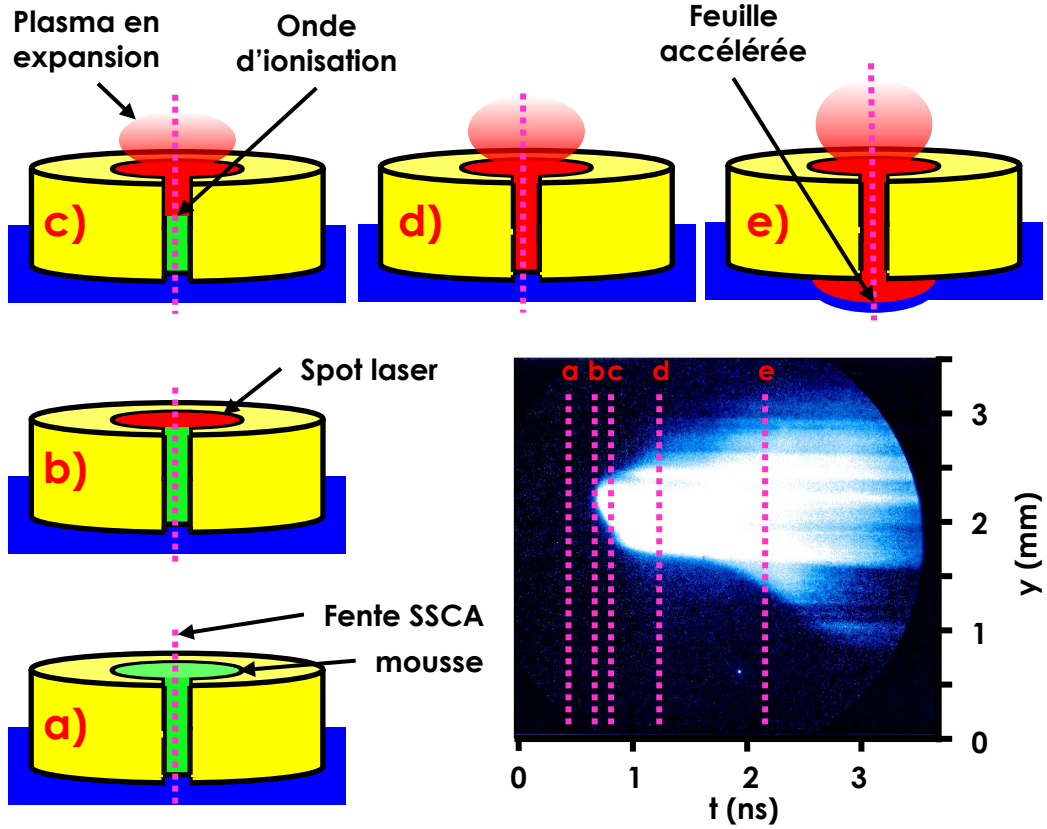


FIGURE 6.6 – Image SSCA extraite d'un tir avec mousse de $5 \text{ mg.cm}^{-3}/500 \mu\text{m}$ (en bas à droite) et schéma de la cible à différents instants : a) avant le début de l'illumination laser, b) au moment où le laser touche la surface de la mousse, c) quand l'onde d'ionisation se propage, d) quand le laser atteint la surface de la feuille de CH et e) quand la feuille a été mise en vitesse. Les barres verticales pointillées roses représentent la visée de la fente de la SSCA. Les zones en rouge sur les schémas (a-e) correspondent aux zones d'émission.

et de l'onde de raréfaction, le front d'ablation ne se déplace que du fait de l'ablation et donc à la vitesse V_a . La zone d'émission de la partie inférieure de l'image est alors stable, quasiment horizontale. Le front d'ablation va ensuite accélérer (e), on voit alors cette zone d'émission se déplacer vers le bas de l'image avec une pente de plus en plus forte. On a vu dans le chapitre précédent que pour des intensités de l'ordre de celles rencontrées ici, la vitesse d'ablation était de l'ordre de 10^5 cm/s . Or en mesurant la pente de la zone des images SSCA correspondant à la mise en vitesse de la feuille de CH vers 1,5 ns, on trouve pour tous les tirs une vitesse comprise entre 3 et 5.10^7 cm/s . Cela explique que la pente de la partie avant mise en vitesse, où le seul déplacement du front d'ablation est dû à la vitesse d'ablation, est négligeable devant le pente de la partie où la feuille est accélérée.

Sur ces images SSCA, nous pouvons extraire deux données : la vitesse de l'onde d'ionisation et la trajectoire de la zone émissive, qui d'après les simulations se trouve être

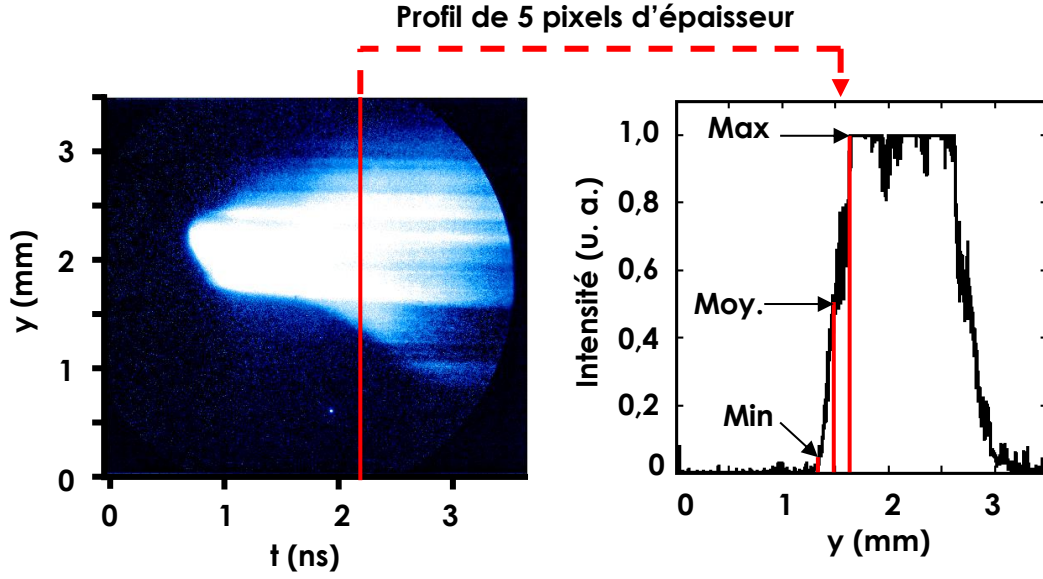


FIGURE 6.7 – Détail de la méthode de mesure de la mise en vitesse du front d'ablation.

très proche (quelques dizaines de μm) du front d'ablation. Cependant, comme on peut le voir sur la figure 6.6, la pente qui correspond à l'ionisation de la mousse est très raide et la résolution n'est pas très adaptée car on cherche à mesurer la mise en vitesse de la feuille sur un temps supérieur à une nanoseconde tandis que l'ionisation se fait sur une durée de l'ordre de la centaine de ps. De plus, les images sont saturées en intensité. La pente correspondant à l'ionisation de la mousse n'est donc observable que sur 2 tirs parmi les 8 où des images SSCA ont été obtenues. On trouve des vitesses d'ionisation d'environ $1,2 \mu\text{m}/\text{ps}$ pour l'un et $0,7 \mu\text{m}/\text{ps}$ pour l'autre, sachant que ces deux tirs ont été réalisés avec des mousses de $5 \text{ mg}/\text{cm}^3$. Ces valeurs sont donc sujettes à caution mais l'ordre de grandeur correspond à ce que l'on peut trouver dans la littérature [99]. La mesure de la mise en vitesse de la feuille a lieu sur des temps plus longs et est donc plus fiable. La méthode de dépouillement des données est exposée en figure 6.7. On prend des profils verticaux de largeur $5 \mu\text{m}$ (pour atténuer un peu le bruit) tous les $30 \mu\text{m}$ sur les images SSCA. La brusque montée en intensité correspond au plasma en expansion à partir du front d'ablation. Cette montée doit donc se situer à quelques dizaines de μm du front d'ablation, distance qui doit rester constante au cours du temps : l'évolution temporelle de la position de ce front d'émission et de la position du front d'ablation doivent donc être similaires. Le milieu de la montée est relevé comme position du front d'émission, le début et la fin comme barres d'erreurs. On peut remarquer, comme évoqué précédemment, que les données sont saturées, ce qui fausse un peu la détermination de la position de la fin du front de montée. En effectuant ces relevés tous les $30 \mu\text{m}$, on obtient la variation temporelle du front d'émission.

6.2.2 Profils extraits des données d’auto-émission mesurées par la caméra à balayage de fente

La figure 6.8 présente les résultats des mesures de trajectoire du front d’ablation. En (a), les trajectoires pour 3 types de cibles (CH seul, mousse de $5 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$ et de $7 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$) sont comparées. Les trois autres images comparent des simulations CHIC de l’émission du plasma avec ces trajectoires. Pour la figure 6.8 (a), il n’y avait pas de calage temporel absolu des données car le fiducial était cassé le jour de l’expérience. Les trajectoires ont donc été recalées temporellement pour trouver le meilleur accord avec la trajectoire de la cible de CH seul. On peut voir que les trajectoires sont très proches les unes des autres. Une conséquence importante est qu’à un écart temporel donné près - écart inconnu par absence de calage temporel absolu, les feuilles de CH avec et sans mousse vont subir des accélérations semblables dans les différents tirs. Les différentes cibles seront donc dans les mêmes conditions vis-à-vis de la croissance de l’IRT ablative : si l’on tient compte du retard dans le début d’accélération dû à l’ionisation de la mousse, les différences observées au niveau des modulations du front d’ablation entre les tirs seront imputables à l’effet des mousses. Cela signifie aussi que la surintensité pour les cibles avec mousse a bien permis de compenser l’absorption d’énergie pour l’ionisation et le chauffage des mousses. D’autre part, les images observées en figure 6.8 (b-d) montrent toutes un bon accord entre les simulations CHIC et les trajectoires mesurées. Ainsi, le code CHIC simule correctement la mise en vitesse de la feuille ; cela permet de se fier à sa représentativité pour interpréter les données des radiographies de face.

L’interprétation des données d’auto-émission mesurées par les caméras à balayage de fente a deux conséquences notables. Premièrement, cela permet de confirmer la fiabilité du code CHIC pour interpréter les tirs, ce qui avait déjà été le cas pour d’autres expériences avec des mousses sous-denses [100, 99]. D’autre part, on a montré qu’à un décalage temporel inconnu près, les différentes cibles subissent des accélérations similaires, et se trouvent donc dans les mêmes conditions vis-à-vis de l’IRT ablative. Dans la section suivante, nous allons utiliser les données de rétrodiffusion mesurées par les FABS pour déterminer l’ordre de grandeur du décalage temporel entre l’accélération des cibles avec et sans mousses.

6.3 Détermination du temps d’ionisation

6.3.1 Simulations de l’ionisation des mousses

Pour pouvoir comparer les données issues des radiographies de face, il reste à déterminer le temps nécessaire pour ioniser la mousse. Deux méthodes ont été utilisées pour cela. Tout d’abord, des simulations 2D ont été réalisées. Ici se pose donc la question de la modélisation de la mousse. Comme on l’a vu, la mousse a une structure poreuse : elle est

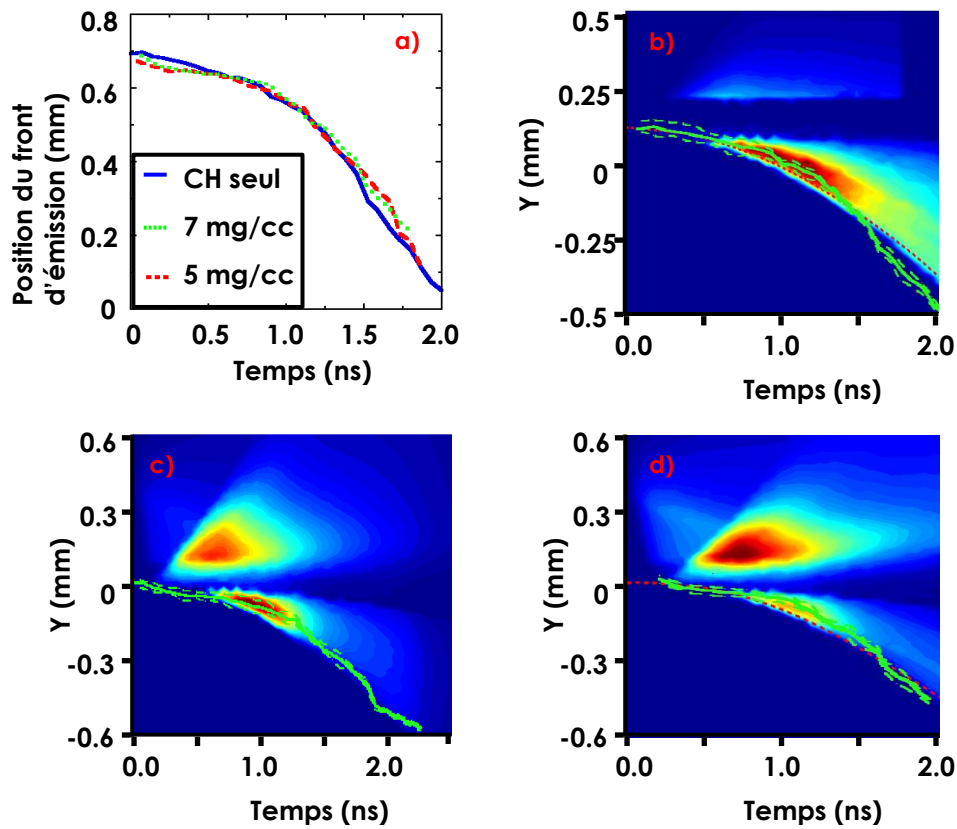


FIGURE 6.8 – a) Comparaison des trajectoires du front d’ablation mesurées pour une cible de CH seul (bleu), une cible mousse de $5 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$ (vert) et de $7 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$ (rouge) et comparaison des trajectoires mesurées et des simulations CHIC d’émission pour b) une cible de CH seul (bleu), c) une cible mousse de $5 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$ (vert) et d) de $7 \text{ mg.cm}^{-3}/500 \text{ }\mu\text{m}$

constituée de pores d'une taille de l'ordre du μm séparés par des parois denses de quelques dizaines de nm. Or prendre cette structure stochastique en compte dans les simulations, pour des mousses de centaines de μm d'épaisseur, impliquerait un temps de calcul très important. Ainsi la mousse est modélisée par un matériau homogène de densité équivalente constante dans les simulations. Un modèle pour la propagation d'un laser et d'une onde d'ionisation dans des mousses sous-denses est développé dans la réf. [101]. Les cas d'un matériau homogène et d'un matériau poreux, tous deux sous-denses, sont abordés. Dans le cas d'un matériau homogène, le laser se propage dans un milieu transparent. Cependant, le laser excite des électrons du matériau, qui vont transmettre leur énergie par collisions ; le matériau va donc chauffer et s'ioniser et une partie de l'énergie laser va être absorbée. Au fur et à mesure de sa propagation à travers le matériau, une partie de l'impulsion laser va donc être absorbée. Ainsi, il va se créer une onde d'ionisation, dont la vitesse dépend du rapport entre le flux d'énergie délivré par le laser et l'absorption du plasma pour son chauffage et son ionisation. Dans le cas d'un matériau poreux, la propagation du laser est différente. Le laser traverse un pore vide d'une taille de l'ordre du μm , puis atteint une paroi sur-dense de quelques dizaines de nm. Il ionise et chauffe le matériau de la paroi, qui se détend dans les pores voisins. C'est seulement quand la densité du matériau en détente sera passée sous la densité critique que le laser reprendra sa propagation. Ainsi, l'onde d'ionisation va plus lentement dans un matériau poreux que dans un matériau homogène. Lorsqu'on modélise la mousse par un matériau homogène dans les simulations, on sous-estime donc le temps de traversée de la mousse par le laser. Les auteurs de la réf. [99] montrent que la sous-estimation dans les simulations CHIC est d'un facteur 2 environ par rapport à des expériences réalisées sur GEKKO XII avec les mêmes types de mousse et d'intensité laser que nous avons utilisés. Dans notre cas, le temps de traversée de l'onde d'ionisation donné par les simulations est de 150 ps pour les différentes mousses, les variations de densité surfacique induisant des écarts peu importants comparés aux incertitudes temporelles expérimentales. Ainsi, nous corrigeons ce temps d'ionisation issu des simulations d'un facteur 2 et considérons donc que le temps d'ionisation est d'environ 300 ps. Un dernier point que l'on peut noter est que l'équation d'état de la mousse utilisée dans les simulations ne joue pas sur la propagation du laser. En effet, au passage de l'onde d'ionisation, la mousse est très rapidement transformée en plasma. Des simulations effectuées avec différentes équation d'état (SESAME, QEoS, gaz parfait) ont montré que la description de la mousse par un gaz parfait était suffisante et qu'il n'y avait pas de différences entre les résultats des différentes simulations.

6.3.2 Données FABS

Une deuxième méthode a consisté à utiliser les données issues des FABS. Ces données pour les spectres Brillouin d'un tir avec mousse de $10 \text{ mg.cm}^{-3}/300 \text{ }\mu\text{m}$ sont représentées en figure 6.9. Il faut bien comprendre que les conditions d'interaction laser-plasma sont très complexes : 6 faisceaux pendant la première nanoseconde et 3 pendant la deuxième arrivent de directions différentes avec des angles d'incidences différents dans un plasma de mousse tout d'abord puis de CH se détendant dans le plasma de mousse. Les figures 6.9 (a) et (b) représentent respectivement les spectres Brillouin pour les faisceaux 25 et 30. Pendant la première nanoseconde, seul le faisceau 30 est allumé. Pourtant, une énergie 5 fois plus importante est récupérée par le FABS 25 comparée à celle collectée par le FABS 30. C'est donc le signe que l'énergie récupérée dans le FABS 25 n'est pas rétrodiffusée par le faisceau 25 (celui-ci n'étant pas allumé) mais qu'on a plutôt affaire à une énergie diffusée dans toutes les directions du fait de l'interaction des différents faisceaux laser et du plasma. Sur les figures 6.9 (a) et (b), on peut voir un signal intense décalé vers le rouge puis une décroissance de la longueur d'onde qui commence vers 550 ps pour le FABS 25 et vers 450 ps pour le FABS 30, accompagnée d'une baisse importante de l'intensité du signal reçue par les FABS. Des simulations CHIC montrent que lorsque le laser atteint la feuille de CH, le plasma de CH se détend dans le plasma de mousse, créant une onde de choc qui se propage dedans dans le sens opposé à la propagation des faisceaux laser. Cette onde de choc augmente la température des ions du plasma de mousse à la température des électrons, initialement plus chauds. Or un écart de température entre ions et électrons est une condition nécessaire au développement de l'instabilité Brillouin. Ainsi, quand on voit la longueur d'onde et l'énergie du signal rétrodiffusé décroître, c'est le signe que le laser a déjà atteint la feuille de CH. Il faut cependant un temps de 100 à 200 ps pour que le choc se forme et thermalise les ions. Les temps de 450 et 550 ps trouvés sont donc des bornes supérieures du temps d'ionisation, le temps réel se situant plus autour de 300-350 ps, comparable au temps évalué par les simulations. Nous ne pouvons cependant pas expliquer l'écart entre le temps trouvé avec le FABS 25 et celui avec le FABS 30 à l'heure actuelle.

Le recouplement des données d'énergie rétrodiffusée mesurées par les FABS, des simulations CHIC de notre expérience mais aussi de résultats précédemment obtenus avec les mêmes mousses permet d'obtenir une valeur de la durée d'ionisation des mousses, qui peut être évaluée à environ 300 ps. Cependant, les imprécisions intrinsèques aux méthodes utilisées ne permettent pas de connaître cette durée mieux qu'à plus ou moins 100 ps près au minimum. La valeur de 300 ps doit donc plus être considérée comme un ordre de grandeur. Dans les deux sections précédentes, nous avons étudié la dynamique des cibles. Nous pouvons désormais nous intéresser aux données des radiographies de face.

Spectres Brillouin pour cible avec mousse de 10 mg/cm³ et 300 μm

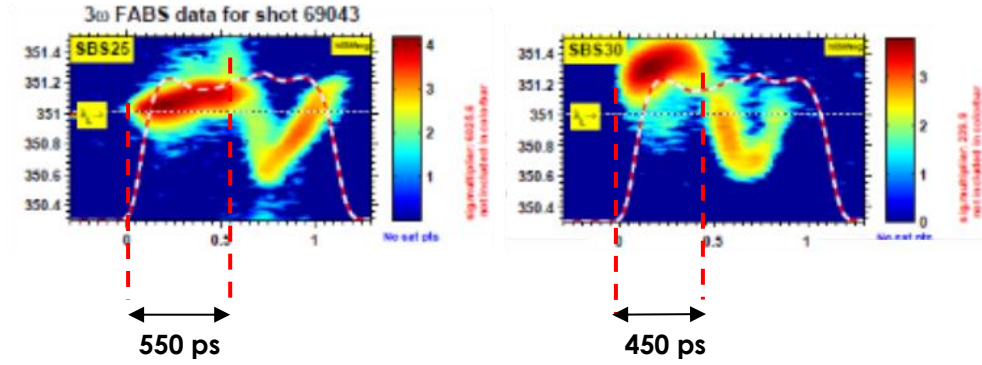


FIGURE 6.9 – Spectres Brillouin issus des FABS 25 et 30 pour un tir avec mousse de 10 mg/cm³/300 μm.

6.4 Dépouillement des données de radiographie de face

6.4.1 Comparaison des radiographies de face avec différents types de cibles et de motifs d'empreinte

Nous allons à présent comparer les cibles de CH seul à trois types de cibles : mousse de 10 mg.cm⁻³/300 μm, de 7 mg.cm⁻³/500 μm et de 5 mg.cm⁻³/500 μm. Comme chaque densité de cible correspond à une épaisseur donnée, nous ne qualifierons plus les cibles que par leur densité, sans ajouter l'épaisseur. La figure 6.10 présente des images issues des radiographies de face pour les différentes cibles : de (a) à (d) pour la M30 et de (i) à (k) pour la M60. Les spectres de Fourier 2D associés à ces images sont représentés respectivement de (e) à (h) et de (l) à (n). Aucune donnée de qualité n'a pu être extraite avec la mousse de 10 mg/cm³ pour la M60 du fait d'un changement de source de radiographie. Les images de cibles avec mousse ont été prises plus tardivement que celles de CH seul pour tenir compte, au moins partiellement, de la durée d'ionisation de la mousse. Sur l'image (a), pour la cible de CH seul (faisceau d'empreinte muni de la M30), on voit clairement les modulations 2D imprimées par la M30 qui correspondent aux deux spots dans la direction perpendiculaire sur le spectre de Fourier 2D associé (e). Avec la mousse de 10 mg/cm³, on observe toujours ce motif d'empreinte 2D auquel viennent s'ajouter des modulations 3D sous la forme de bulles, incluses entre les traits des modulations 2D. Ce phénomène est illustré sur le spectre (f) par l'apparition d'un anneau d'une certaine largeur fréquentielle autour du centre ; les deux spots du motif initial sont néanmoins toujours présents. Avec la mousse de 7 mg/cm³ (c), le motif initial est déjà plus difficile à identifier, et les bulles sont plus grandes et plus marquées. Sur le spectre associé (g), les spots correspondants au

motif d’empreinte initial sont ténus et sont presque fondus dans l’anneau. Enfin, pour la mousse de 5 mg/cm^3 (d), le motif initial a complètement disparu au profit de structures 3D de diverses tailles. Cette impression est corroborée par le spectre associé à cette image : on voit toujours un anneau, plus resserré sur le centre du spectre car les bulles sont plus grosses, mais pas de trace des spots correspondant au motif initial. Les observations sont similaires pour les images de cibles imprimées par M60. On voit donc que plus la densité des mousses est petite, plus le lissage du motif initial est important (pour les densités concernées par l’expérience ; il doit exister un optimum), mais ce lissage s’accompagne de l’apparition inattendue de structures 3D, particulièrement pour les mousses de 5 mg/cm^3 . Les causes possibles de l’apparition de ces structures sont discutées ultérieurement.

6.4.2 Evolution temporelle des modulations de densité surfacique

Nous avons étudié l’évolution temporelle de la densité surfacique de la longueur d’onde nominale du motif 2D imprimé pour les cibles dont les radiographies montraient des modulations les plus proches du cas du CH seul, c’est-à-dire la mousse de 10 mg/cm^3 pour la M30 et celle de 7 mg/cm^3 pour la M60. Pour cela, nous avons utilisé la variante de la méthode de FFT 1D présentée en annexe B. Les résultats sont représentés sur la figure 6.11 (a) pour la M30 et (b) pour la M60. Pour la M30, on peut voir que les données décroissent pour le CH seul comme pour la mousse de 10 mg/cm^3 . Ce phénomène est probablement dû au fait que les modulations avaient déjà percé la cible quand les mesures ont commencé. Ainsi il n’est pas possible de s’appuyer sur ces données pour effectuer une analyse quantitative des radiographies. Cependant, si l’on compare les données des deux cibles, il semble difficile à imaginer que les modulations de la cible avec mousse de 10 mg/cm^3 aient pu être supérieures à celles de la cible de CH seul. Pour la M60 en revanche, on observe une croissance continue pour la mousse de 7 mg/cm^3 ; pour la cible de CH seul, on observe une croissance sur les premiers points, puis une stabilisation suivie d’une décroissance comme pour la M30. Comme on observe une croissance pour les 2 types de cible, il sera possible d’utiliser ces données pour une analyse quantitative, au moins au début des mesures avec la cible de CH seul. On peut remarquer que les mesures ont été effectuées aux mêmes instants avec les deux lames de phase et que pourtant la croissance n’a été observée qu’avec la M60 : ceci est dû au fait que la longueur d’onde $30 \mu\text{m}$ croît plus rapidement que celle de $60 \mu\text{m}$ et donc que les modulations de $30 \mu\text{m}$ vont atteindre plus vite l’épaisseur de la cible.

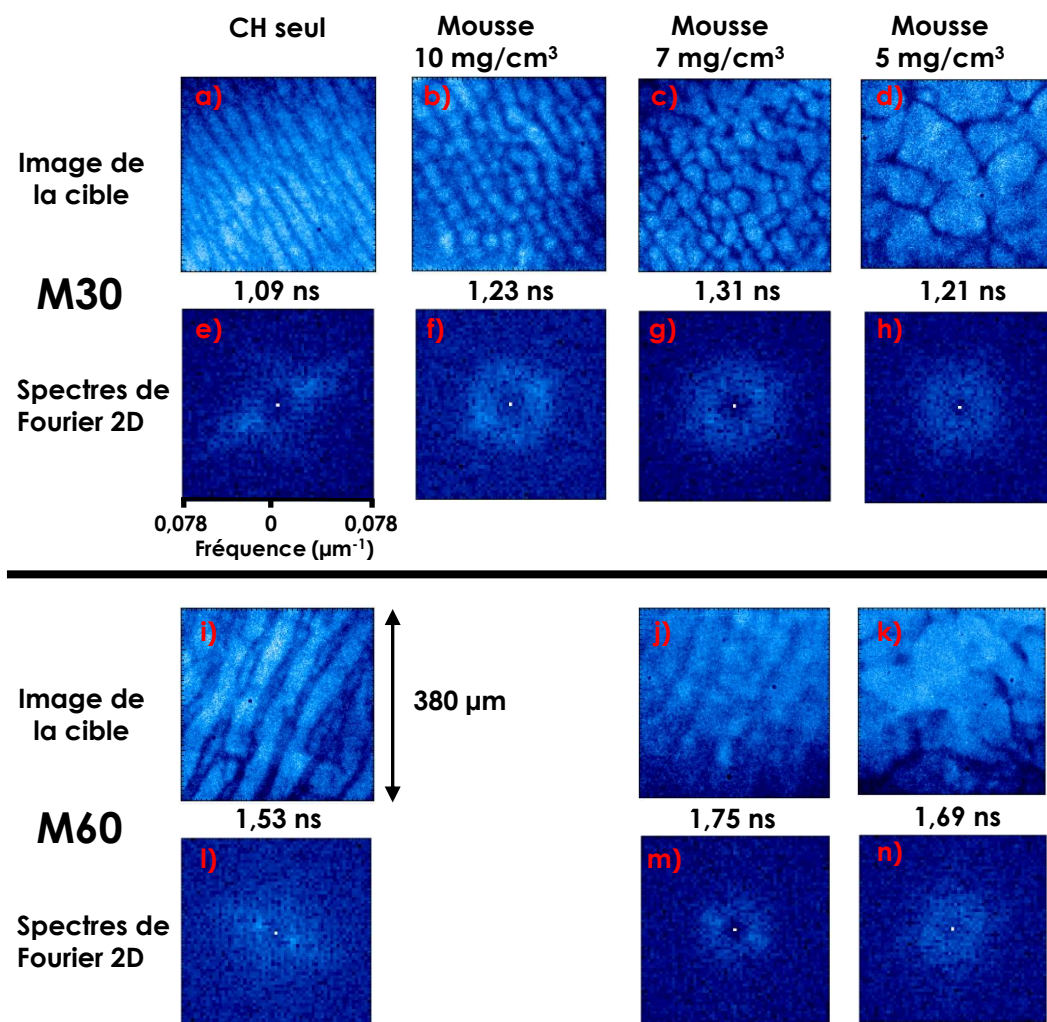


FIGURE 6.10 – Images issues des radiographies de face réalisées avec une XRFC pour différents types de cibles -a) et i) CH seul, b) mousse de 10 mg/cm³, c) et j) mousse de 7 mg/cm³ et d) et k) mousse de 5 mg/cm³- et de conditions d’empreinte - (a-d) M30 et (i-k) M60. Les spectres de Fourier 2D associés aux images (a-d) et (i-k) sont représentés respectivement en (e-h) et (l-n).

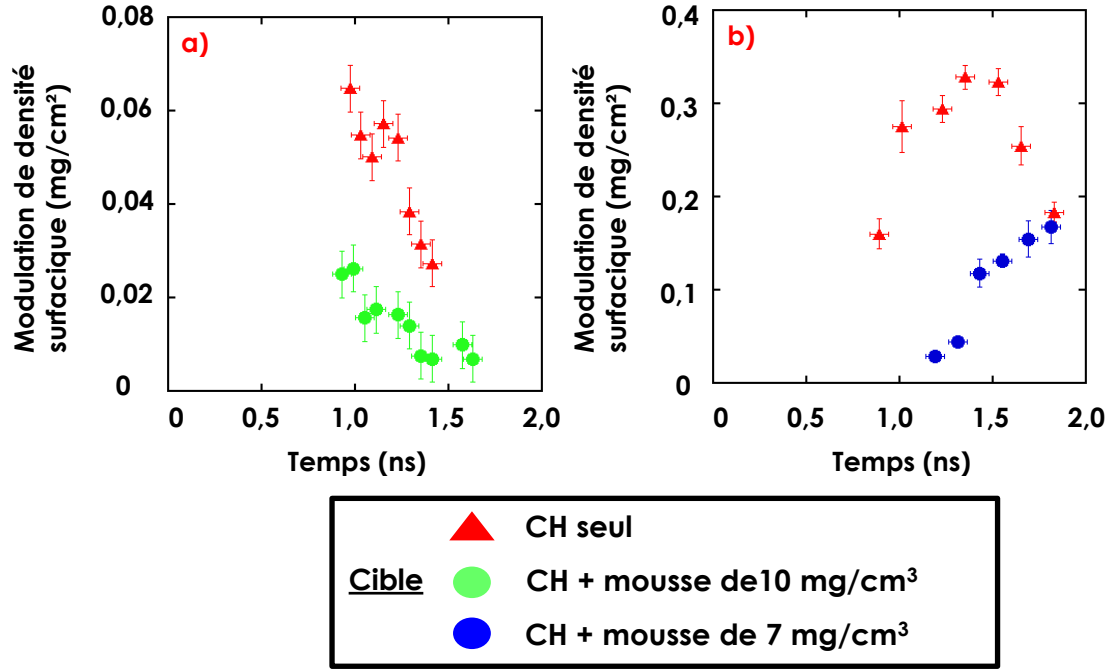


FIGURE 6.11 – Evolution temporelle des modulations de densité surfacique mesurées expérimentalement a) pour une cible de CH seul (triangles rouges) et une cible avec mousse de 10 mg/cm^3 (ronds verts) imprimées par M30 et b) pour une cible de CH seul (triangles rouges) et une cible avec mousse de 7 mg/cm^3 (ronds bleus) imprimées par M60.

6.4.3 Analyse des distributions de bulles obtenues avec les cibles avec mousses

La taille des bulles observées sur les images (c) et (d) de la figure 6.10 semble différer. Pour mesurer la distribution de la taille des bulles, nous avons utilisé le programme de segmentation d'image sur deux radiographies des cibles avec mousses de 7 mg/cm^3 à $1,25 \text{ ns}$ et de 5 mg/cm^3 à $1,27 \text{ ns}$ imprimées par M30. Ces images sont représentées respectivement en figure 6.12 (a) et (b). L'image (c) montre les distributions pour ces deux images. La distribution pour la cible avec mousse de 7 mg/cm^3 est très piquée, quasiment toutes les bulles ayant une taille comprise entre 20 et $50 \mu\text{m}$. Pour la mousse de 5 mg/cm^3 , la distribution est beaucoup plus étalée, la taille des bulles allant de 30 à $120 \mu\text{m}$. On voit donc que les distributions diffèrent grandement à un instant donné. La figure 6.13 présente l'évolution des distributions pour une cible avec mousse de 5 mg/cm^3 imprimée par M60. Des images XRFC à $0,81 \text{ ns}$ (a), $1,13 \text{ ns}$ (b) et $1,41 \text{ ns}$ (c) peuvent être observées. Ces images ont été extraites respectivement de la 1ère, la 2ème et la 3ème piste de la radiographie associée. Pour obtenir les 3 distributions représentées en (d), on a moyenné la distribution extraite de ces images avec la distribution extraite de l'autre image située sur la même piste. Aucune distribution n'a été calculée pour la 4ème piste car, sur les images qui en sont extraites, une grosse bulle occupe la quasi-totalité

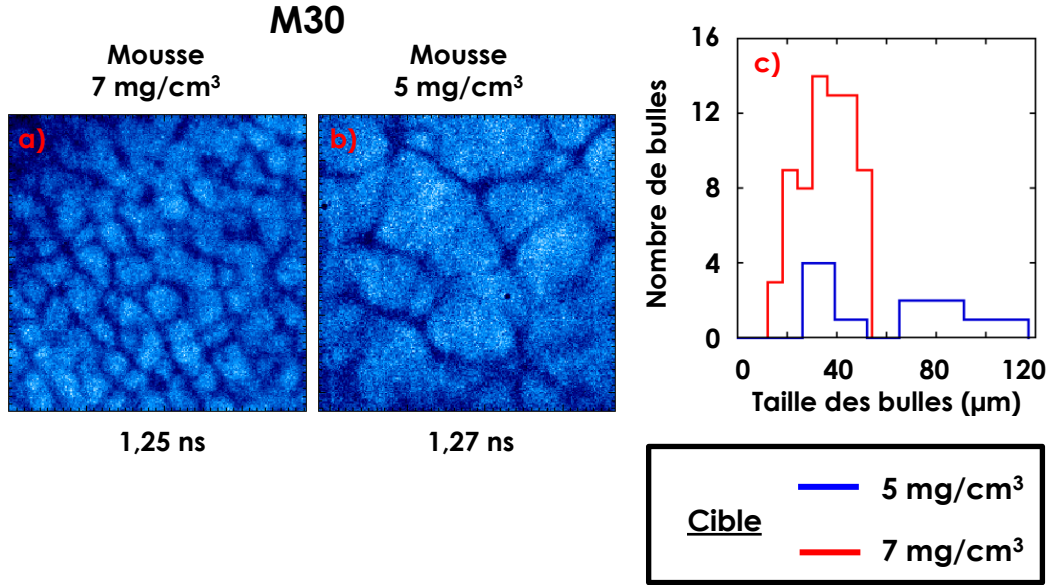


FIGURE 6.12 – Radiographies de face d’une cible a) de 7 mg/cm³ à 1,25 ns et b) de 5 mg/cm³ à 1,27 ns imprimées par M30. c) Distribution de taille des bulles obtenues par segmentation d’image pour la radiographie a) (ligne rouge) et b) (ligne bleue).

de l’image (cf figure 6.10 (k)) et le programme de segmentation n’a alors plus d’utilité. L’observation des images comme des distributions nous apporte la même information : aux premiers instants de mesure, les bulles sont assez petites (majoritairement entre 20 et 60 μm) et la distribution est piquée. Puis, du fait du phénomène de coalescence des bulles, la taille moyenne des bulles grossit et la distribution s’étale, avec des tailles allant de 20 à 120 μm pour la 3ème piste. Ceci permet de rappeler que ce qui est observé sur les radiographies n’est pas le motif d’inhomogénéités de l’intensité du laser directement mais les modulations imprimées ayant évolué sous l’effet des instabilités hydrodynamiques. Si l’on fait la même analyse temporelle pour la mousse de 7 mg/cm³, on voit que les distributions sont beaucoup plus stables. De plus, la distribution piquée observée en figure 6.12 (c) pour la mousse de 7 mg/cm³ est assez semblable à celle de la figure 6.13 (d) pour la 1ère piste d’un tir avec mousse de 5 mg/cm³, bien que les lames de phase d’empreinte soient différentes. On peut supposer que les modulations 3D sont initialement de même taille avec les deux types de mousse. Cependant, la mousse de 5 mg/cm³ semble lisser parfaitement le motif 2D tandis que ce motif est toujours présent pour la mousse de 7 mg/cm³. Ainsi, l’évolution des bulles pourrait être différente en fonction de la persistance ou non du motif initial 2D. Cette possibilité ne reste cependant qu’au stade d’hypothèse, faute de justifications théoriques ou numériques.

6.4.4 Interprétation de l'origine des structures 3D

Concernant l'origine des modulations 3D, plusieurs hypothèses peuvent être émises mais aucune ne semble pour l'instant complètement cohérente. On sait tout d'abord que ces bulles ne peuvent provenir directement de la radiographie de structures présentes dans la mousse : pour atteindre des niveaux de modulation de densité surfacique non négligeables, il faudrait que les structures de la mousse soient d'une épaisseur de l'ordre de la centaine de μm . Or les modulations sont toujours observées aux temps longs, quand la mousse est complètement ionisée et n'est plus qu'un plasma homogène en expansion et ne peut donc contenir les structures observées. On peut aussi noter que de telles structures ne sont pas observées sur la figure 6.2 (c). Ainsi, une première hypothèse est que les modulations 3D proviennent de défauts de surface de la feuille de CH, malgré le traitement censé apporter une rugosité inférieure à 20 nm. De telles structures n'apparaissent pas sur la figure 6.2 (a-b), mais ces mesures d'état de surface ont été réalisées avant collage de la feuille au support contenant la mousse. Cette phase de collage pourrait avoir occasionné des défauts de surface supplémentaires. Cependant, il semble peu probable dans ce cas que l'on n'observe pas de modulations avec la cible de CH seul. Une autre possibilité serait liée à la présence de petites structures de quelques dizaines de μm dans la mousse, par exemple à des bulles de vide, régulièrement réparties sur l'épaisseur et la longueur de la mousse. La vitesse de propagation du laser est de l'ordre du $\mu\text{m}/\text{ps}$ dans la mousse et de 300 $\mu\text{m}/\text{ps}$ dans le vide. Ainsi, entre une épaisseur de mousse possédant une structure de quelques dizaines de μm et une épaisseur sans structure, un retard de propagation du laser de quelques dizaines de ps va apparaître. Or, l'empreinte laser est sensible à un décalage temporel de quelques dizaines de ps [61]. Un argument s'oppose néanmoins à cette hypothèse : lors de la vérification de la taille des pores au Lebedev Institute, aucune structure de dizaines de μm n'a été observée (cf figure 6.2 (c)). On pourra cependant pondérer cet argument car la vérification de la taille des pores ne peut se faire qu'en surface de la mousse. Une autre cause de l'apparition de ces structures 3D pourrait être la création de défauts d'intensité de ce type par des phénomènes d'instabilités paramétriques, comme de l'autofocalisation. Néanmoins, des phénomènes produisant des modulations d'intensité qui imprimeraient les structures 3D observées sur les radiographies ne sont pas observés dans les simulations d'interaction laser-plasma ou dans la littérature sur ces mousses sous-denses.

En résumé, les radiographies de face font apparaître un phénomène de lissage du motif initial 2D. Ce lissage semble d'autant plus efficace que la densité de la mousse est faible. Cependant, ce lissage s'accompagne de l'apparition de structures 3D de quelques dizaines de microns, dont la source n'a pas été identifiée de manière certaine. L'analyse des radiographies par transformée de Fourier montre que seules les données des tirs utilisant la lame de phase M60 pourront être utilisées pour une analyse quantitative. En effet, la

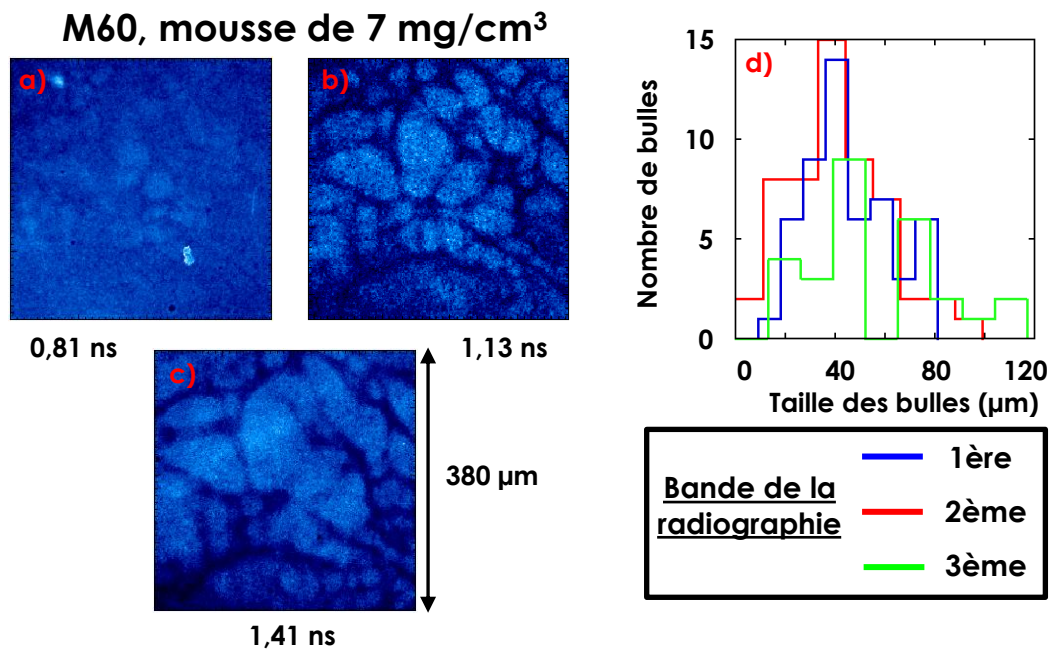


FIGURE 6.13 – Radiographies de la cible avec mousse de 7 mg/cm^3 imprimée par M60 à a) 0,81 ns, b) 1,13 ns et c) 1,41 ns. d) Distribution de taille des bulles obtenues par segmentation d'image à ces différents instants; les distributions sont moyennées sur les deux images obtenues sur chaque piste de la radiographie XRFC : 1ère piste (ligne bleue, image (a)), 2ème piste (ligne rouge, image (b)) et 3ème piste (ligne verte, image (c)). Les données issues de la 4ème piste ne sont pas représentées car le programme de segmentation d'image n'a pu leur être appliqué.

feuille de CH a été percée par les modulations avant le début des mesures lors de tirs avec la M30. Dans la section suivante, nous présentons donc la comparaison des simulations CHIC et des simulations d'interaction laser-plasma PARAX avec les données des tirs utilisant la M60.

6.5 Interprétation des données expérimentales par les simulations CHIC et PARAX

6.5.1 Simulations d'interaction laser-plasma PARAX

Le code CHIC est un code d'hydrodynamique radiative ; il n'a donc pas la capacité de simuler les interactions laser-plasma telles que les instabilités paramétriques. Ainsi, G. Riazuelo, un chercheur du CEA, a effectué des simulations de la traversée par le laser du plasma de mousse en utilisant le code PARAX [102, 103]. C'est un code électromagnétique paraxial 3D qui simule la propagation des faisceaux laser à travers le plasma en prenant en compte l'absorption Bremsstrahlung, la filamentation thermique et pondéromotrice ainsi que la diffusion Brillouin stimulée vers l'avant. La réf. [100] présente la chaîne de simulations utilisée ; tout d'abord, une simulation CHIC 2D utilisant la configuration réaliste de l'expérience (faisceaux lasers, tache focale) est réalisée. On extrait ensuite les profils de densité et de température à l'instant où l'on veut démarrer la simulation PARAX. Ces profils sont utilisés en entrée pour le code PARAX, qui va simuler la propagation des faisceaux lasers dans le plasma issu des simulations CHIC. Les simulations PARAX durent 100 ps jusqu'à l'atteinte de la quasi-stationnarité. Cette durée de 100 ps est suffisante pour que les instabilités paramétriques se développent. Les résultats présentés en figure 6.14 sont issus d'une simulation PARAX utilisant les profils hydrodynamiques CHIC à 150 ps (moment où le laser atteint la feuille dans les simulations) pour une mousse de 5 mg/cm^3 imprimée par M60. Des coupes perpendiculaires à la direction de propagation du laser sont ensuite réalisées à l'entrée de la mousse (faisceau initial non lissé), au milieu de la mousse et en sortie de mousse (à l'arrivée sur la feuille de CH). De la même manière que pour les données expérimentales, on calcule ensuite un profil perpendiculaire aux modulations 2D de la M60 moyenné sur la largeur du faisceau. Les spectres de Fourier sont ensuite obtenus par calcul de la FFT 1D de ces profils. En comparant les différents spectres, il apparaît clairement qu'entre l'entrée et la sortie de la mousse, le pic principal à $60 \mu\text{m}$, le pic du second harmonique et les petites longueurs d'onde sont lissés par la mousse. Le pic principal notamment passe d'une amplitude de 0,031 à 0,012, diminuant donc d'un facteur 2,6. Cependant, ce pic ne présente aucun effet de lissage en milieu de mousse. En effet, les speckles laser qui portent entre autres les modulations de l'intensité laser ont une forme de cigares de rayon $3 \mu\text{m}$ et de longueur environ $100 \mu\text{m}$. Or le milieu de la mousse est à $250 \mu\text{m}$ et la fin à $500 \mu\text{m}$. On peut donc comprendre que le plasma doit être d'une

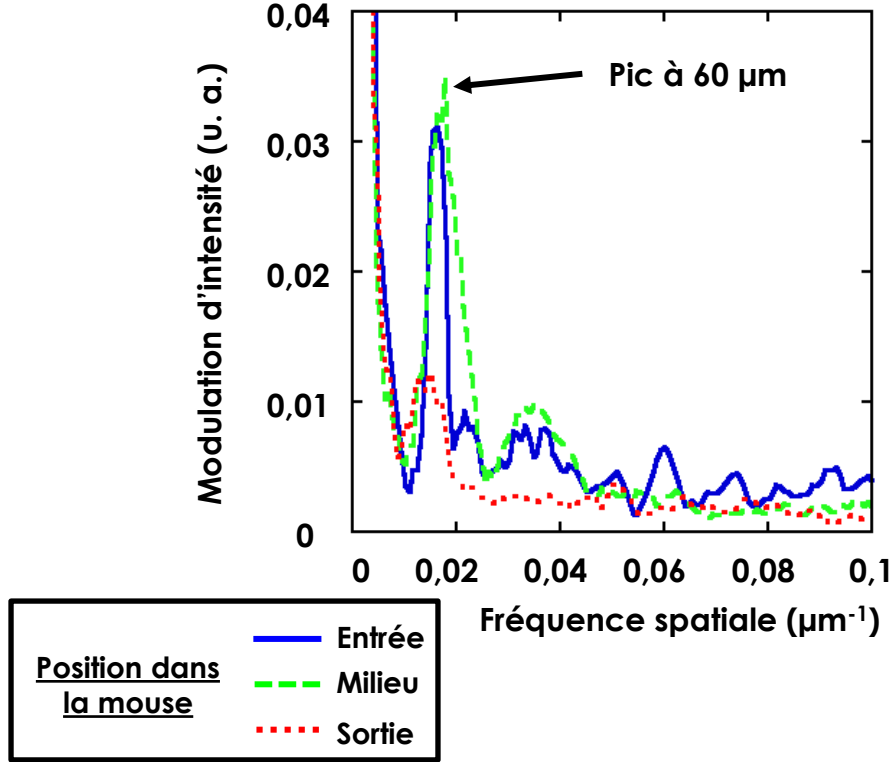


FIGURE 6.14 – Simulation PARAX pour une cible avec mousse de 5 mg/cm^3 imprimée par M60 : spectres de Fourier de l'intensité laser normalisée calculés dans la direction transverse aux modulations 2D de la M60 dans des plans perpendiculaires à la direction de propagation des faisceaux laser. Ces plans sont sélectionnés avant l'entrée dans la mousse (ligne pleine bleue), au milieu de la mousse (ligne tireté verte) et à la fin de la propagation à travers le plasma (ligne pointillée rouge).

longueur de quelques points chauds minimum pour que les instabilités paramétriques se développent et soient effectives.

Comme cela a été démontré dans des publications précédentes [22, 100, 104], le mécanisme physique de lissage est lié à l'autofocalisation des speckles et la diffusion Brillouin stimulée dans les speckles. Ces deux processus créent des fluctuations de la densité de mousse stochastiques pour les longueurs d'onde de l'ordre de la taille de speckle et celles plus petites. Cette turbulence en petite échelle produit une décroissance efficace des modulations de l'intensité du faisceau laser en grande échelle imposées par la lame de phase.

6.5.2 Comparaison des courbes de croissance expérimentales et numériques

Les simulations PARAX confirment l'effet lissant des instabilités paramétriques opérant dans les mousses. Pour montrer à l'aide de CHIC l'effet de lissage laser par le plasma

sur l'IRT ablative, nous avons repris les résultats expérimentaux présentés en figure 6.11 (b) pour l'empreinte par M60 et nous avons réalisé des simulations CHIC des tirs correspondants pour la feuille de CH derrière la mousse. Il est important de noter que toutes les simulations CHIC de croissance des modulations ont été recalées par rapport à la dynamique de l'écoulement donné par des simulations 2D utilisant la configuration réaliste de l'expérience. La figure 6.15 montre la comparaison des données expérimentales et des simulations CHIC. Le niveau de modulation de l'intensité absorbée par la feuille a été fixé à 16 % pour la simulation du tir avec CH seul, niveau qui permet de trouver le meilleur accord avec les données expérimentales en phase de croissance. Ce niveau correspond à la superposition du faisceau porteur de la M60 avec deux faisceaux d'angle d'incidence de 23° porteurs de la SG4. Pour la mousse, trois faisceaux d'angle d'incidence 48° porteurs de la SG4 sont ajoutés. Le niveau de modulation de l'intensité absorbée devient alors 9,3 %. Or, pour trouver le meilleur accord entre les données expérimentales et la simulation pour la mousse de 7 mg/cm^3 , le niveau de modulation de l'intensité doit être fixé à 5 %, soit une diminution quasiment d'un facteur 2. Il est important de rappeler ici que dans ces simulations, nous avons simulé seulement la dynamique de la feuille, les instabilités paramétriques ne se développent pas dans les simulations CHIC et donc leur effet de lissage est pris en compte par le profil de l'intensité laser absorbée sur la feuille. On peut aussi rappeler qu'on a montré dans la section 6.2 la fiabilité du code CHIC pour simuler la dynamique des cibles, à un recalage temporel près. Cela signifie donc que pour le tir avec mousse de 7 mg/cm^3 imprimé par M60, les instabilités paramétriques dans la mousse font passer le niveau de modulation de l'intensité de 9,3 % initialement à 5 % en sortie de mousse. On trouve donc le même ordre de grandeur de facteur de réduction du pic principal que celui trouvé dans la simulation PARAX pour une cible avec mousse de 5 mg/cm^3 présentée en figure 6.14. Ce résultat est assez remarquable quand on considère d'un côté que le facteur déterminé par les simulations a été obtenu par une chaîne de codes [100], et de l'autre côté les imprécisions expérimentales, que ce soit par exemple sur la valeur de la sur-intensité pendant la première nanoseconde ou du fait que les modulations ont percé rapidement la cible après le début de la mesure dans le cas du CH seul. Il faut aussi noter que l'on n'a pas encore comparé les données expérimentales avec une simulation PARAX utilisant une cible de mousse à 7 mg/cm^3 . En effet, ces simulations durent longtemps et nous n'avons pas pour l'instant toutes les données de la simulations avec une mousse de 7 mg/cm^3 . Néanmoins, des simulations avec une mousse de 10 mg/cm^3 permettent de retrouver le même ordre de grandeur de réduction du pic principal qu'avec la mousse de 5 mg/cm^3 . Le résultat pour la mousse de 7 mg/cm^3 ne devrait donc pas trop différer.

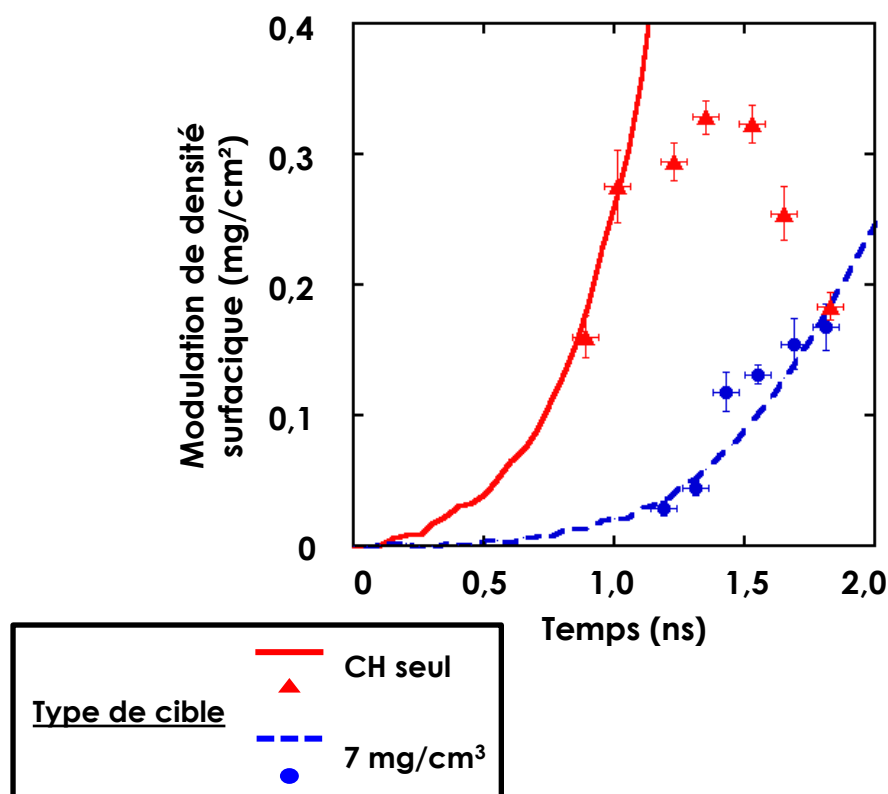


FIGURE 6.15 – Evolution temporelle des modulations de densité surfacique à la longueur d’onde $60\ \mu\text{m}$ pour une cible de CH seul (ligne pleine rouge pour la simulation CHIC et triangles rouges pour les données expérimentales) et pour une cible avec mousse de $7\ \text{mg}/\text{cm}^3$ (ligne tireté bleue pour la simulation CHIC et ronds bleus pour les données expérimentales). Le niveau de modulation de l’intensité laser est de 16 % pour la simulation avec cible de CH seul et de 5 % pour celle de la cible avec mousse de $7\ \text{mg}/\text{cm}^3$.

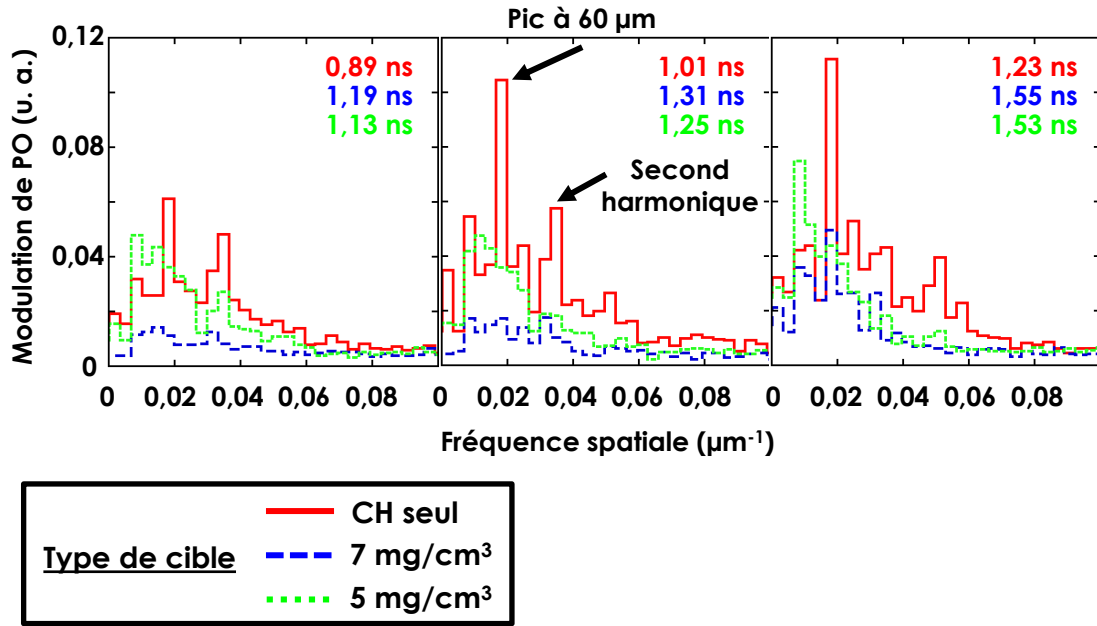


FIGURE 6.16 – Spectres de Fourier obtenus par l’analyse des radiographies de cibles de CH seul (ligne pleine rouge), avec mousse de 7 mg/cm^3 (ligne tireté bleue) et avec mousse de 5 mg/cm^3 (ligne pointillée verte). Les spectres des différentes cibles sont calculés respectivement à a) 0,89 ns, 1,19 ns et 1,13 ns, b) 1,01 ns, 1,31 ns et 1,25 ns, et c) 1,23 ns, 1,55 ns et 1,53 ns.

6.5.3 Effet des mousses sur l’ensemble du spectre spatial

Le paragraphe précédent présentait l’évolution des modulations de densité surfacique associées à la longueur d’onde principale de la M60. Pour avoir une idée de l’effet des mousses sur l’ensemble des longueurs d’ondes, la comparaison des spectres de Fourier de cibles de CH seul et avec mousses de 7 mg/cm^3 et 5 mg/cm^3 imprimées par M60 sont montrées en figure 6.16 à différents instants dans la phase de croissance de la modulation principale du CH seul. Les spectres avec mousses représentés ont été calculés environ 250-300 ps après les spectres du CH seul pour tenir compte de la durée d’ionisation des mousses et donc comparer les spectres après des durées similaires de croissance des modulations du fait de l’IRT. Tout d’abord, on peut voir que les deux mousses lissent les petites longueurs d’onde (grandes fréquences spatiales) inférieures à la trentaine de μm , ce qui est cohérent avec le résultat des simulations PARAX. Ces mousses lissent aussi le pic principal imprimé par la M60 et son second harmonique. On peut noter un très bon lissage global de la mousse de 7 mg/cm^3 . En revanche, un pic de grandes longueurs d’onde apparaît avec la mousse de 5 mg/cm^3 . Ce pic correspond aux larges structures 3D observées sur les radiographies.

Pour résumer, voici les points les plus importants abordés dans ce chapitre :

— Nous avons mesuré pour la première fois la modification de l’empreinte laser et

des instabilités hydrodynamiques subséquentes sur des feuilles de CH recouvertes par des mousses sous-denses.

- Les données d’auto-émission mesurées par caméra à balayage de fente permettent d’estimer que les cibles recouvertes par des mousses différentes subissent une accélération similaire, décalée du temps d’ionisation de la mousse par rapport à la feuille de CH seule.
- Le temps d’ionisation des mousses, par le croisement des informations données par des simulations CHIC et des diagnostics de rétrodiffusion, a été évalué à environ 300 ps.
- Les radiographies de face montrent une disparition progressive du motif initial d’empreinte 2D lorsqu’on passe à des mousses de plus petites densités, mais aussi l’apparition de structures 3D. L’origine de ces structures n’est pas déterminée parfaitement à l’heure actuelle.
- L’analyse comparative des courbes de croissance issues de simulations CHIC et des données expérimentales obtenues avec la M60 montre une réduction de l’empreinte de la longueur d’onde principale d’un facteur quasiment égal à 2 avec la mousse de 7 mg/cm^3 .
- La comparaison des spectres de Fourier calculés à partir des radiographies en plusieurs instants pour les différents types de cible montre un lissage de la longueur d’onde principale, de son second harmonique ainsi que des petites longueurs d’onde avec les mousses de 5 et 7 mg/cm^3 imprimées par M60.

Chapitre 7

Conclusion et perspectives

Au cours de cette thèse, nous avons étudié les conditions initiales de l'instabilité de Rayleigh-Taylor ablative en attaque directe. Ce travail a été composé d'une étude de l'instabilité de Richtmyer-Meshkov ablative imprimée par laser et de la mesure de l'effet de mousses sous-denses sur l'empreinte d'inhomogénéités de l'intensité laser.

7.1 Etude de l'instabilité de Richtmyer-Meshkov imprimée par des inhomogénéités de l'intensité laser

L'instabilité de Richtmyer-Meshkov ablative imprimée par des défauts de l'intensité laser n'avait jamais été mesurée. Nous avons mis en place une configuration expérimentale sur le laser OMEGA pour la mesurer sur des cibles planes de CH_2 . Les couples épaisseur de cible/matériau de la source de radiographie ont été optimisés et le choix de la méthode d'empreinte s'est porté sur l'utilisation d'un faisceau porteur d'une lame de phase spéciale, créant des modulations 2D de l'intensité laser de longueur d'onde $30\text{ }\mu\text{m}$ (M30), $60\text{ }\mu\text{m}$ (M60) ou $120\text{ }\mu\text{m}$ (M120). Quatre journées de tirs ont été réalisées, au cours desquelles l'instabilité de Richtmyer-Meshkov ablative imprimée par laser a été mesurée pour la première fois, et ce pour les trois longueurs d'onde. Pour la longueur d'onde de $30\text{ }\mu\text{m}$, deux intensités avaient été utilisées pour illuminer les cibles : $5,6 \cdot 10^{13}\text{ W/cm}^2$ et $1,6 \cdot 10^{14}\text{ W/cm}^2$. Dans les deux cas, une inversion de phase a été observée (cf chapitre 4). Des simulations monomodes de croissance des modulations imprimées par laser ont été réalisées avec le code d'hydrodynamique radiative CHIC du CELIA, en utilisant les impulsions laser mesurées pendant les expériences. D'autre part, le modèle de Goncharov qui décrit l'instabilité de Richtmyer-Meshkov ablative imprimée par laser [11] a été appliqué avec les paramètres extraits des simulations CHIC 1D (cf chapitre 5). Les simulations et le modèle montrent un bon accord ; les légères différences observées peuvent être imputées au

fait que le modèle est calculé avec des paramètres (vitesse d'ablation, vitesse de formation de la zone de conduction, ...) constants, alors que ces paramètres varient temporellement dans les simulations.

La comparaison des simulations et des données expérimentales d'une part, et du modèle avec ces mêmes données expérimentales d'autre part, fait apparaître des tendances similaires dans les deux cas. L'accord est assez bon, sur l'allure comme sur l'amplitude, à quelques dizaines de pourcents près, pour la lame de phase M60 (et pour la lame de phase M120 avec les simulations). Avec la lame de phase M30 en revanche, l'accord est moins bon, particulièrement concernant l'amplitude des modulations. A la lumière de l'interprétation par le modèle et les simulations, et des résultats d'autres expériences sur OMEGA, ce désaccord peut s'expliquer par un défaut de qualité de la lame de phase M30. Il est envisagé actuellement au Laboratory for Laser Energetics de fabriquer une nouvelle lame de phase induisant des modulations 2D de $30\ \mu\text{m}$ possédant moins d'imperfections.

Le modèle de Goncharov de l'instabilité de Richtmyer-Meshkov imprimée par des défauts de l'intensité laser a été confronté pour la première fois à des expériences. Les bons accords obtenus montrent sa grande utilité pour l'interprétation des données et le design des cibles. Il permet notamment de mettre en lumière le lien de l'amplitude maximale de l'instabilité de Richtmyer-Meshkov ablative (5.6) et de la fréquence des oscillations (5.5) avec des paramètres contrôlables expérimentalement. Pour contrôler l'instabilité de Richtmyer-Meshkov ablative imprimée par laser, deux méthodes peuvent être utilisées : concevoir cibles et impulsions pour que les longueurs d'onde les plus néfastes soient en train de s'inverser au moment de la transition avec l'instabilité de Rayleigh-Taylor ablative, ou diminuer l'amplitude maximale des oscillations. Concernant la première méthode, les simulations et le modèle font apparaître que l'instant d'inversion de phase est très sensible aux paramètres du laser et de la cible. Il est compliqué de mettre en oeuvre cette méthode. Un autre moyen de contrôler l'instabilité de Richtmyer-Meshkov ablative est de diminuer le niveau d'amplitude des modulations de l'intensité laser, ce qui peut être réalisé en recouvrant la cible d'une mousse sous-dense.

Bien que les résultats soient très encourageants, il est souhaitable de faire de nouvelles mesures avec la future lame de phase à $30\ \mu\text{m}$ et des lames de phase de plus petites longueurs d'onde pour obtenir des modulations plus proches du maximum de croissance de l'instabilité de Richtmyer-Meshkov et de l'instabilité de Rayleigh-Taylor ablatives. L'expérience et son interprétation peuvent être aussi améliorées. Du côté des expériences, il serait souhaitable de mesurer l'accélération de la feuille, d'utiliser des intensités plus élevées et de réaliser des tirs avec empreinte de défauts 3D et avec des cibles de deutérium solide cryogénique pour se rapprocher des conditions des expériences de fusion. En attaque indirecte, des tirs avec une couche de deutérium-tritium cryogénique perturbée sont d'ailleurs prévus sur le National Ignition Facility dans les mois qui viennent. Du côté de l'interprétation, il est souhaitable d'enrichir les simulations avec un modèle non-local

complet qui pourrait être étalonné en utilisant les résultats des expériences à plus haute intensité.

7.2 Etude du lissage de l’empreinte laser par des mousses sous-denses

Nous avons étudié le développement des instabilités hydrodynamiques dans des cibles recouvertes de mousses sous-denses dans le but de mettre en évidence la possible réduction de l’empreinte de défauts de l’intensité laser par l’action d’instabilités paramétriques. L’utilisation d’un faisceau d’empreinte porteur de modulations de l’intensité 2D et de cibles composées d’une feuille de CH recouverte ou non par une mousse sous-dense nous ont permis de mesurer pour la première fois la diminution du niveau des modulations et le retard du développement de l’instabilité de Rayleigh-Taylor ablative.

L’analyse des données d’auto-émission mesurées de côté par une caméra à balayage de fente montre que le front d’émission, se situant à quelques dizaines de microns du front d’ablation d’après les simulations CHIC, suit une trajectoire similaire avec les différentes cibles, à la durée d’ionisation de la mousse près. Les mesures de rétrodiffusion et les simulations CHIC ont permis d’évaluer un temps d’ionisation de la mousse de l’ordre de 300 ps.

L’observation des radiographies de face fait apparaître des tendances similaires avec la M30 et la M60 : plus la densité de la mousse est petite, plus le lissage du motif d’empreinte 2D est efficace, mais ce lissage s’accompagne de l’apparition de structures 3D d’origine inconnue. La comparaison des données expérimentales de modulation de densité surfacique et des simulations CHIC de croissance des modulations pour la cible de CH seul et celle avec mousse de 7 mg.cm^{-3} imprimées par M60 fait apparaître une réduction du niveau de modulation de l’intensité laser d’un facteur quasiment égal à 2 pour la longueur d’onde nominale. Les simulations effectuées avec le code d’interaction laser-plasma PARAX par G. Riazuelo (CEA/DAM/DIF) pour une mousse de 5 mg.cm^{-3} montrent aussi le lissage de la longueur d’onde principale d’un facteur du même ordre de grandeur. Enfin, les spectres de Fourier expérimentaux montrent qu’en plus de la longueur d’onde à $60 \mu\text{m}$, le second harmonique et les longueurs d’onde inférieures à $30 \mu\text{m}$ sont lissés par les mousses.

La diminution du niveau d’empreinte d’un facteur 2 et le retard de croissance de l’instabilité de Rayleigh-Taylor d’environ 1 ns (cf figure 6.15) sont des résultats nouveaux et encourageants. Mais il reste à expliquer l’origine des perturbations 3D, surtout pour les mousses de basse densité. Ces structures peuvent être néfastes car le front de bulle croît plus vite et sature à une amplitude supérieure en géométrie 3D qu’en géométrie 2D [50]. Des simulations récentes montrent même qu’en 3D, le front de bulle ne saturerait pas du

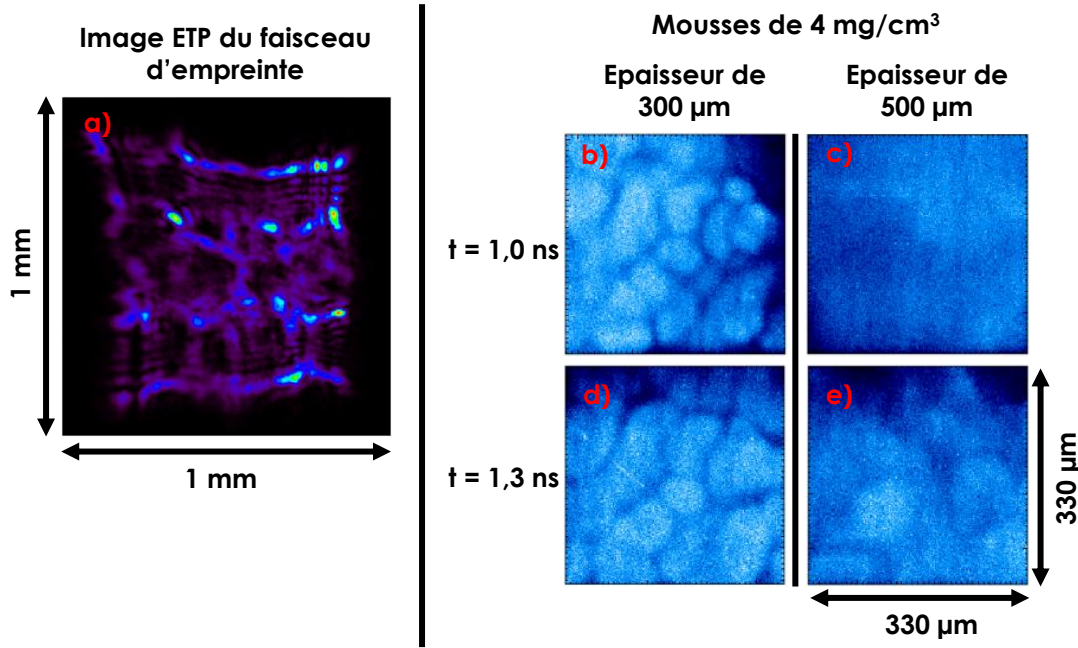


FIGURE 7.1 – a) Image "Equivalent Target Plane" du faisceau 4 d'OMEGA EP défocalisé de $700\ \mu\text{m}$. Radiographies de face de cibles, imprimées par ce faisceau défocalisé, avec mousse de $4\ \text{mg}\cdot\text{cm}^{-3}$ d'épaisseur $300\ \mu\text{m}$ b) à $t=1,0\ \text{ns}$ et d) à $t=1,3\ \text{ns}$, et d'épaisseur $500\ \mu\text{m}$ c) à $t=1,0\ \text{ns}$ et e) à $t=1,3\ \text{ns}$.

fait de l'accumulation de la vorticit  dans les bulles [105]. Mais de toute mani re, dans le cas d'implosions cryog niques en attaque directe, les d fauts d'intensit  laser sont des speckles, donc de g om trie 3D. Il faudrait donc r aliser des exp riences avec empreinte de points chauds pour pouvoir affirmer ou infirmer que les mousses introduisent des d fauts 3D n fastes.

Une nouvelle exp rience r alis e r cemment, en juillet 2014, sur OMEGA EP, avait pour objectif d' tudier l'effet de l'empreinte g n rique du laser et de l' paisseur des mousses. Le faisceau d'empreinte a  t  d focalis  pour cr er les structures 3D apparaissant en figure 7.1 (a), d'une taille allant de quelques dizaines de microns   environ $200\ \mu\text{m}$. Les radiographies pr sent es en figure 7.1 (b-e) montrent la comparaison, aux m mes instants, de deux mousses de densit  $4\ \text{mg}\cdot\text{cm}^{-3}$ et d' paisseur $300\ \mu\text{m}$ et $500\ \mu\text{m}$. Il appara t clairement que les structures sont nettement moins d velopp es avec la mousse de $500\ \mu\text{m}$ d' paisseur, m me en consid rant la dur e d'ionisation l g rement sup rieure dans ce cas. Cela montre qu'une certaine  paisseur de mousse est n cessaire pour effectuer le lissage, confirmant notre r sultat pr c dent. Il reste encore   analyser ces tirs pour  valuer la possible r duction de l'empreinte pour les cibles avec mousse par rapport   celles de CH seul.

Plusieurs pistes sont ouvertes pour des futures exp riences. Il serait int ressant de comparer les diff rentes m thodes de lissage existantes, dans les m mes conditions, et de pond rer ainsi leurs avantages et inconv nients respectifs - par exemple, les mousses lissent

le faisceau sans préchauffage mais consomment une partie de l'énergie laser pour leur ionisation et leur chauffage. De plus, pour connaître le potentiel des mousses dans le cas d'implosions cryogéniques, il faudrait effectuer des expériences avec empreinte de défauts de taille comprise entre 3 et 30 μm , qui sont les longueurs d'onde les plus dangereuses pour le développement des instabilités hydrodynamiques. La réalisation de telles expériences est actuellement limitée par la résolution spatiale des diagnostics de la radiographie. Les futurs imageurs à incidence rasante qui seront installés sur le National Ignition Facility et le Laser Méga-Joule auront une résolution spatiale approchant 5 μm , ce qui permettrait de mesurer directement l'empreinte de ces petits défauts. Mais on peut dès à présent s'affranchir de ces problèmes de résolution en étudiant la phase non-linéaire de l'instabilité de Rayleigh-Taylor, quand des bulles d'une plus grande taille apparaîtront du fait de la coalescence de bulles. Il conviendra néanmoins d'optimiser la taille des feuilles de CH, suffisamment épaisses pour qu'elles ne soient pas percées trop vite mais suffisamment fines pour que les modulations de densité surfaciques puissent être mesurées. Actuellement, il existe aussi une limite technologique car la faisabilité de la fabrication de cibles sphériques recouvertes de mousse n'a pas été démontrée, et seuls un ou deux laboratoires dans le monde ont la capacité de fabriquer ces mousses sous-denses.

Dans l'optique de réduire l'effet des instabilités hydrodynamiques, d'autres pistes de recherche existent. On a vu qu'en attaque indirecte, des ablateurs laminés alternant couches de CH d'un micron dopées et non dopées au germanium ont permis d'obtenir une réduction de la croissance de l'instabilité de Rayleigh-Taylor [16]. Les simulations des réf. [106, 107] montrent que les ablateurs laminés peuvent aussi réduire les modulations du front d'ablation en attaque directe, dans la phase Richtmyer-Meshkov comme dans la phase Rayleigh-Taylor. Réaliser une expérience avec des ablateurs laminés en attaque directe pourrait permettre de vérifier ces prédictions.

Annexe A

Annexe A : Détail du modèle de Goncharov

Le modèle présenté dans la réf. [11] définit l'évolution des perturbations du front d'ablation de longueur d'onde $2\pi/k$ sous l'effet de l'IRM ablative imprimée par laser, pour $D_c = V_c t$ dans la limite $kc_s t \gg 1$ avec c_s la vitesse acoustique de la matière post-choc, par

$$\eta(t) = \frac{2}{3\gamma k} \frac{\delta I}{I} \left[\frac{c_s^2}{V_c^2 + V_a V_{bl}} e^{-kV_c t} + (A \cos \omega_{RM} t + B \sin \omega_{RM} t) e^{-2kV_a t} \right] + \eta_v(t) \quad (\text{A.1})$$

avec les constantes

$$A = -\frac{c_s^2}{V_a V_{bl} + V_c^2} + C_0 + 0,5 - \eta_v(0), \quad (\text{A.2})$$

$$B = \frac{c_s}{\sqrt{V_a V_{bl}}} \left[\frac{M_s^2 + 1}{2M_s^2 M_1} - 1 - C_1 + \frac{c_s(V_c - 2V_a)}{V_a V_{bl} + V_c^2} \right], \quad (\text{A.3})$$

$$C_0 = 2 \sum_{i=1}^3 i \Lambda_i, \quad (\text{A.4})$$

$$C_1 = \sum_{i=1}^3 \Lambda_i, \quad (\text{A.5})$$

$$\Lambda_1 = 2 \left[1 - \frac{(2M_1^2 + 1)M_s^2 + M_1^2}{M_1((2 + M_1^2)M_s^2 + 1)} \right], \quad (\text{A.6})$$

$$\Lambda_2 = 2 \left[1 + (\Lambda_1 - 2) \frac{M_s^2(M_1^4 + 4M_1^2 + 1) + 2}{M_s^2(M_1^4 + 8M_1^2 + 3) + 2(M_1^2 + 1)} \right], \quad (\text{A.7})$$

$$\Lambda_3 = -(\coth \theta_s + \Lambda_1 \sinh 2\theta_s + \Lambda_2 \sinh 4\theta_s) / \sinh 6\theta_s, \quad (\text{A.8})$$

$$\theta_s = (u/c_s). \quad (\text{A.9})$$

On a $\frac{\delta I}{I}$ le niveau de modulation de l'intensité laser, γ le ratio des chaleurs spécifiques, V_c la vitesse d'expansion de la zone de conduction, V_a la vitesse d'ablation, V_{bl} la vitesse d'éjection du plasma dite de "blow-off", ω_{RM} la fréquence d'oscillation de l'IRM ablative, η_v le terme de convection de vorticit , M_s le nombre de Mach du choc d fini par $M_s = u/c_0$ avec u la vitesse du choc et c_0 la vitesse acoustique de la mati re avant le choc, et $M_1 = u/c_s$ avec u la vitesse de la mati re post-choc dans le r f rentiel du choc. La fr quence d'oscillation de l'IRM ablative ω_{RM} est donn e par

$$\omega_{RM} = k\sqrt{V_a V_{bl}} \quad (\text{A.10})$$

et le terme de convection de vorticit  η_v par

$$\eta_v(t) = (2c_s/V_{bl})e^{kV_a t} \int_{\infty}^{kV_a t} \Omega(\eta)e^{-\eta} d\eta \quad (\text{A.11})$$

avec

$$\Omega(\eta) = \frac{1 - M_s^2}{[(\gamma + 1)M_s^2 M_1 \sinh^2 \theta_s]} \tilde{w}_s(\eta / \sinh \theta_s), \quad (\text{A.12})$$

o 

$$\tilde{w}_s(x) \approx \frac{2}{3\gamma k} \frac{\delta I}{I} \sum_{i=1}^3 (2e^{-2i\theta_s} + \Lambda_i \sinh 2i\theta_s) J_{2i}(x) \quad (\text{A.13})$$

avec J la fonction de Bessel.

Annexe B

Annexe B : Méthode alternative de dépouillement des radiographies de face

Dans le chapitre 6, on cherche à comparer des images qui possèdent uniquement des structures 2D, d'autres des modulations 3D, et enfin d'autres les deux types de modulations. Comme le montre la figure 6.10, les FFT 2D sont un outil intéressant pour une analyse qualitative des radiographies. Pour l'analyse quantitative, la situation est plus compliquée. La figure B.1 représente l'analyse par FFT 2D de deux images générées artificiellement : les modulations de l'image (a) varient en $\cos(x)$ et celles de l'image (d) en $\cos(r)$. L'amplitude et la fréquence des modulations sont les mêmes sur les deux images, seule la géométrie diffère : modulations 2D pour l'image (a) et 3D pour l'image (d). Sur le spectre, dans le 1er cas, on voit les deux spots dans la direction perpendiculaire aux modulations. Dans le 2ème cas, un cercle, centré sur le milieu du spectre, à même distance du centre que les deux spots. Les moyennes azimutales des deux spectres, qui permettent d'obtenir un spectre 1D équivalent, sont représentées sur les courbes (c) et (f). Le pic correspondant aux modulations se situe à la même fréquence pour les deux spectres ; cependant, son amplitude est 20 fois plus importante dans le cas des modulations 3D. Le fait de moyenner azimutalement atténue nettement les deux pics associés aux modulations 2D ; dans le cas des images expérimentales des tirs avec cibles de CH seul, les pics disparaissent même complètement dans le bruit lors de la moyenne azimutale. Dans le cas de la figure B.1, les modulations représentent le même danger dans les deux cas à cet instant donné pour une cible car elles sont de même amplitude mais l'analyse par moyenne azimutale du spectre de Fourier 2D donne des résultats différents. Ainsi, cette méthode est pertinente lorsqu'on compare deux images ne comportant que des modulations 3D, comme dans la réf. [54], mais ne convient pas pour comparer du 2D et du 3D ensemble. De plus, on cherche dans tous les cas à comparer l'effet des mousses sur le motif

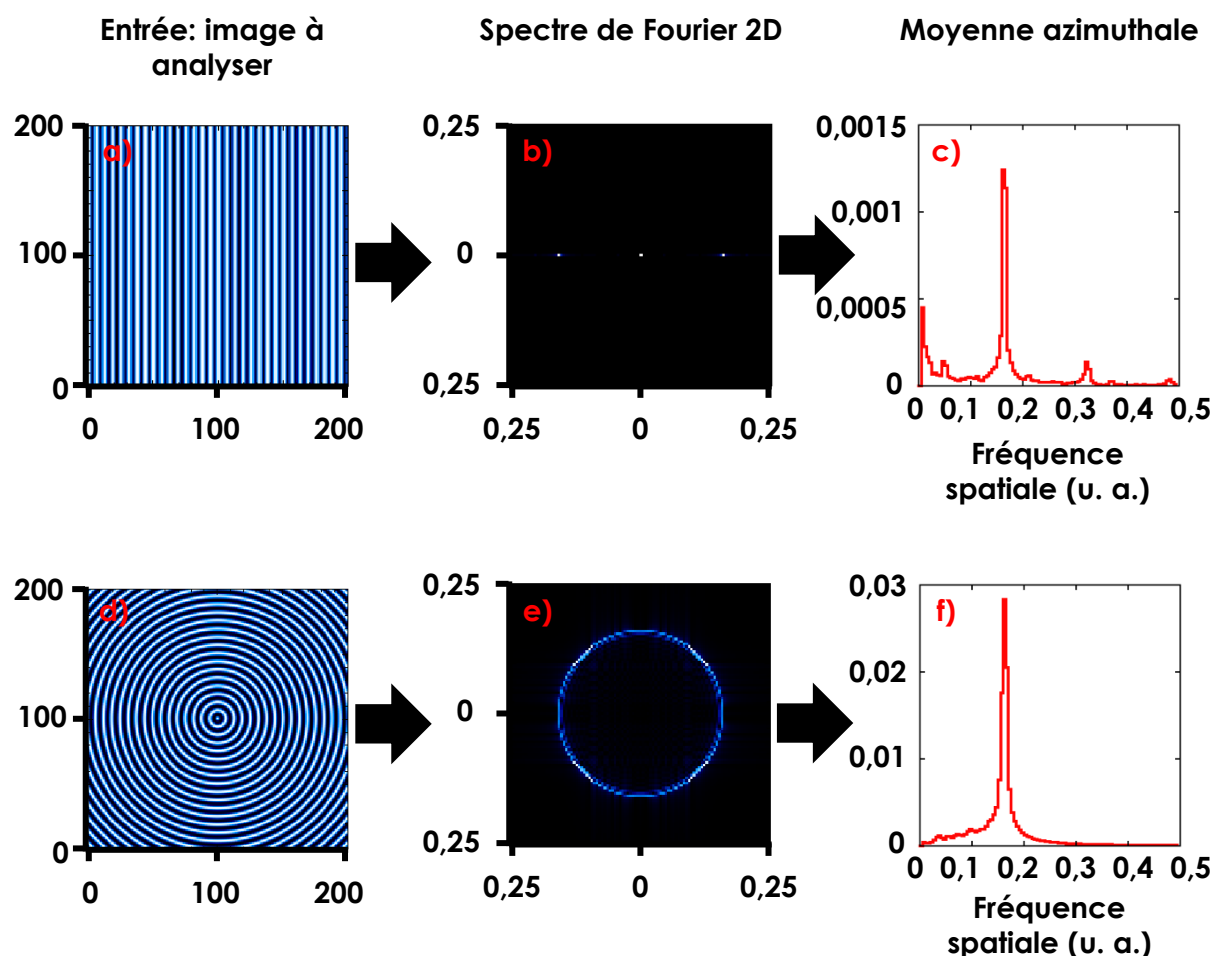


FIGURE B.1 – Images générées par un a) cosinus horizontal et d) un cosinus radial de mêmes fréquences et même amplitude. Les spectres de Fourier 2D et leur moyenne azimuthale associée sont représentés en (b-c) pour le cosinus horizontal et (e-f) pour le cosinus radial.

2D originel, donc dans une direction privilégiée. Une FFT 1D et une FFT 2D ne sont pas comparables : si l'on veut obtenir des valeurs de modulation de densité surfacique, il va dans tous les cas falloir passer par une analyse par FFT 1D. Un problème se pose néanmoins pour l'analyse par FFT 1D : la méthode décrite dans le chapitre 3 utilise un profil moyenné sur la moitié de la largeur de l'image, c'est-à-dire $190\ \mu\text{m}$ dans notre cas. Or on a vu que la taille des bulles pouvait descendre jusqu'à $20\ \mu\text{m}$. Moyenner sur $190\ \mu\text{m}$ va alors complètement noyer l'information de ces petites bulles qui n'apparaîtra pas dans le spectres de Fourier. La figure B.2 (a) présente une radiographie à $1,27\ \text{ns}$ d'un tir avec mousse de $7\ \text{mg}/\text{cm}^3$ imprimée par M30 ; sur cette image, qui a été préalablement tournée pour que le motif 2D original soit horizontal, sont représentées trois zones, de largeur $1,67\ \mu\text{m}$ (1 pixel), $30\ \mu\text{m}$ (18 pixels) et $190\ \mu\text{m}$ (115 pixels). Si l'on regarde en (b) les profils associés, on peut noter que les profils sur $1,67$ et $30\ \mu\text{m}$ sont très similaires, à ceci près que

le 2ème nommé est beaucoup moins bruité. Le profil calculé sur $190\ \mu\text{m}$ est semblable en certains endroits, comme sur la partie du profil entre $x=0$ et $x=100\ \mu\text{m}$, mais présente de nettes différences ailleurs, comme au niveau du pic autour de $x=230\ \mu\text{m}$. Ces différences sont dues au fait que moyenner sur une si grande largeur cause la disparition du signal des bulles dans le profil, seule l'information des modulations 2D est conservée. Les différences entre le profil moyenné sur $30\ \mu\text{m}$ et celui moyenné sur $190\ \mu\text{m}$ sont encore plus fortes pour les images des tirs avec mousses de $5\ \text{mg}/\text{cm}^3$ car ces images ne portent que du signal 3D. On se propose donc, pour garder l'information des modulations 3D tout en analysant les modulations 2D, de réaliser les FFT 1D de 6 profils de $30\ \mu\text{m}$ calculés les uns à côté des autres (zone totale analysée de $180\ \mu\text{m}$), et de moyenner ces 6 spectres pour obtenir le spectre final, au lieu de calculer le spectre d'un seul profil de $190\ \mu\text{m}$ de large. On ne descend pas en dessous de $30\ \mu\text{m}$ de largeur car le bruit devient alors rédhibitoire. Il faut néanmoins vérifier si cette méthode est fiable vis-à-vis de l'analyse des modulations 2D. On applique donc les deux méthodes à une radiographie d'un tir de M60 sur du CH seul à $1,23\ \text{ns}$ (c). Les deux spectres correspondants sont représentés en figure B.2 (d). Le point principal est qu'avec les deux méthodes, on trouve exactement la même amplitude pour le pic correspondant à la longueur d'onde nominale du motif 2D imprimé. Cette cohérence indique donc que la nouvelle méthode peut être appliquée à l'analyse des données. On remarque aussi quelques différences entre les deux spectres et entre autres une amplitude moyenne plus importante avec la nouvelle méthode pour les grandes fréquences spatiales, ce qui indique probablement un niveau de bruit supérieur. Malgré ce désavantage, nous utiliserons la nouvelle méthode pour l'analyse des données pour prendre en compte les structures 3D.

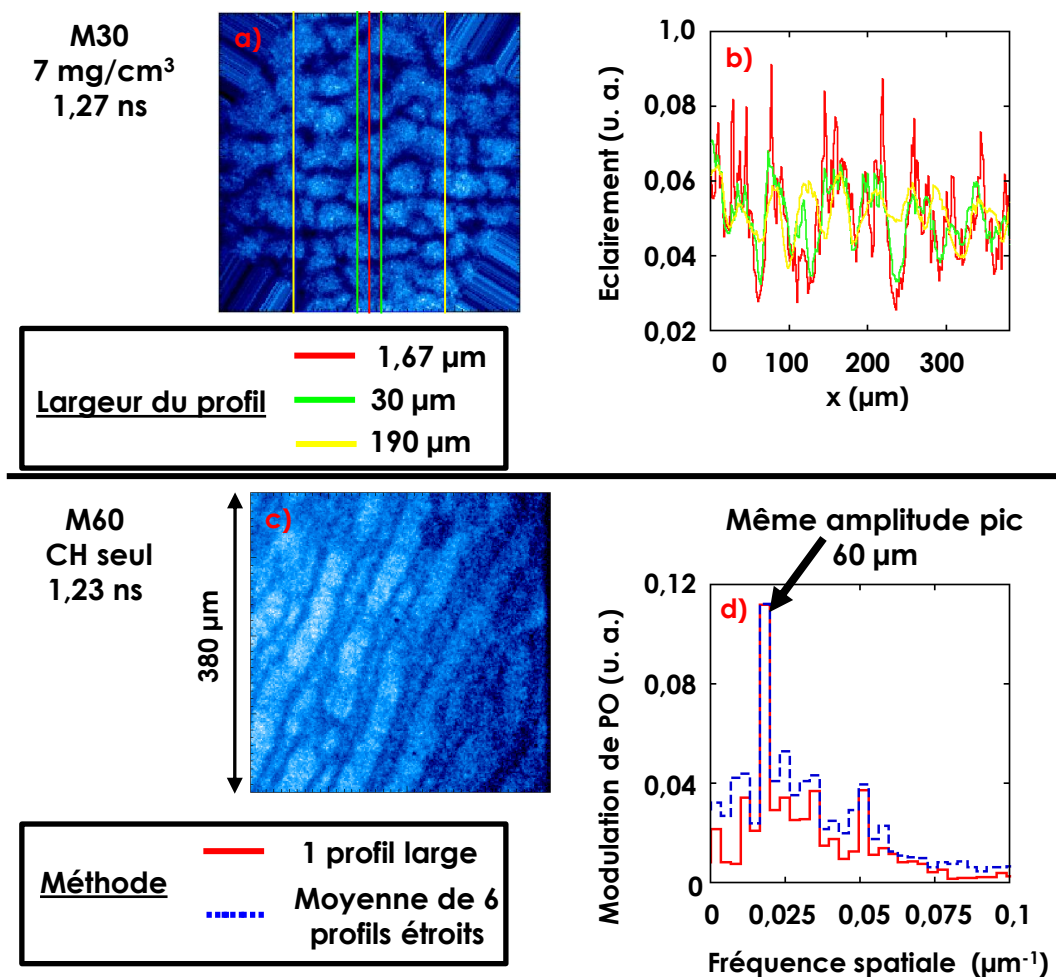


FIGURE B.2 – a) Radiographie à 1,27 ns d’une cible avec mousse de 7 mg/cm³ imprimée par M30, tournée pour que les modulations issues du motif d’empreinte initial soient horizontales. b) Courbes issues des profils verticaux réalisés sur l’image (a), moyennés sur une largeur de 1,67 μm (rouge), 30 μm (vert) et 190 μm (jaune). c) Radiographie à 1,23 ns d’une cible de CH seul imprimée par M60. d) Spectres de Fourier issus de l’analyse de l’image (c) par FFT 1D d’un profil moyenné sur 190 μm (courbe pleine rouge) et par moyenne des spectres obtenus par FFT 1D de 6 profils accolés moyennés sur 30 μm (courbe tireté bleue).

Annexe C

Annexe C : Laser OMEGA EP

Dans l'annexe du bâtiment du LLE se trouve une autre installation laser nommée OMEGA Extended Performance (OMEGA EP). Elle possède 4 faisceaux du type de ceux utilisés sur le NIF, capables de délivrer au total presque 20 kJ d'énergie sur la cible. Les faisceaux 3 et 4 portent des impulsions ns à 3ω . Les faisceaux 1 et 2 peuvent aussi être utilisés comme les faisceaux 3 et 4 à 3ω ou bien à 1ω avec des impulsions courtes ps. Ces deux faisceaux peuvent être dirigés soit vers la TC d'OMEGA EP soit vers la TC d'OMEGA pour réaliser des tirs combinés avec les 60 faisceaux d'OMEGA. Un schéma d'une chaîne laser d'OMEGA EP pour les faisceaux 1 et 2 est représenté en figure C.1. Les chaînes laser des faisceaux 3 et 4 sont similaires, à ceci près qu'elles ne possèdent pas de lien vers la chambre de compression. Contrairement à OMEGA, chaque faisceau possède sa propre chaîne laser sans éléments communs. La partie "driver" des faisceaux longs est assez semblable à celle d'OMEGA, avec quelques changements apportés aux amplificateurs régénératifs pour permettre la création d'impulsions allant jusqu'à 10 ns. L'impulsion, après avoir été formée temporellement, passe par des amplificateurs (amplificateur régénératif et verre dopé au Nd), des apodiseurs pour donner une forme carrée au faisceau pour correspondre aux amplificateurs, un isolateur qui sélectionne une impulsion, des filtres, puis est dirigée vers la partie amplification de la chaîne.

La partie génératrice d'impulsions courtes pour les faisceaux 1 et 2 est représentée en figure C.2. L'amplification de l'impulsion courtes est basée sur la méthode nommée "Optical Parametric Chirped-Pulse Amplification" (OPCPA), qui est en fait la combinaison des méthodes de "Chirped-Pulse Amplification" (CPA) [108] et de "Optical Parametric Amplification" (OPA). Au début de la chaîne, un laser produit des impulsions de 2 nJ, 200 fs avec une largeur de bande de 8 nm. Cette largeur de bande est nécessaire : en effet, il n'est pas possible d'amplifier une impulsion courte (fs ou ps) à de hautes énergies. Pour cela, on utilise donc la méthode CPA : l'impulsion courte est étirée par une série de réseau. Le chemin optique parcouru dépend de la longueur d'onde. Ainsi, si l'impulsion est strictement monochromatique, il n'y aura pas d'étirement ce qui sera en revanche le

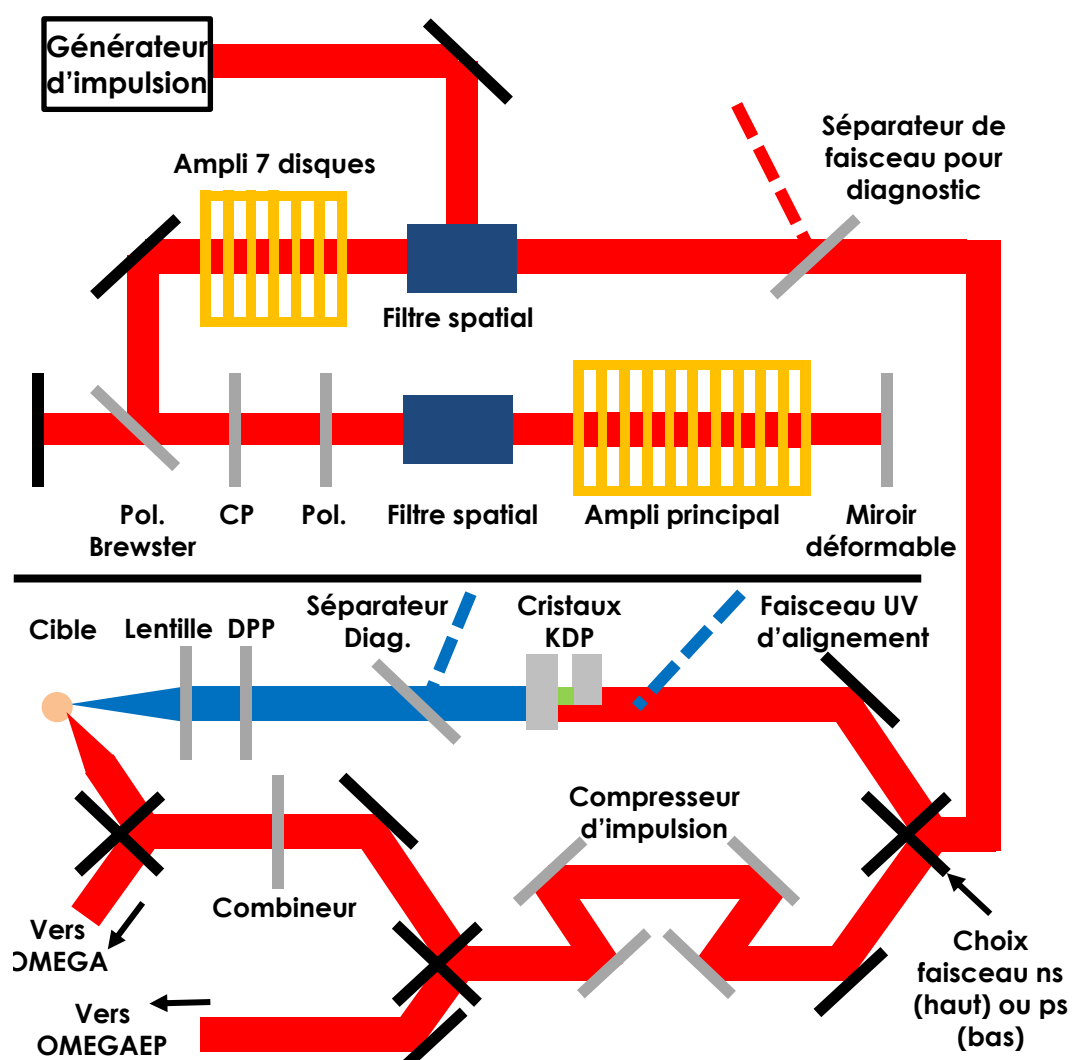


FIGURE C.1 – Schéma d'une chaîne laser OMEGA EP. Tous les composants ne sont pas représentés.

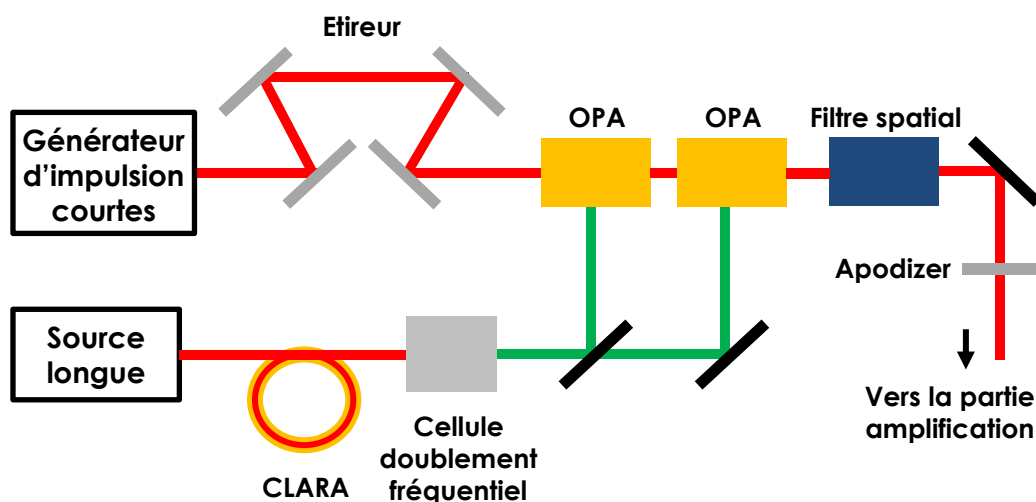


FIGURE C.2 – Schéma de la partie génératrice d'impulsions courtes sur OMEGA EP pour les faisceaux 1 et 2. Tous les composants ne sont pas représentés.

cas pour une bande spectrale de largeur finie. Ensuite, le faisceau est amplifié puis recompressé. Ici, la technique d'amplification est l'OPA : la méthode utilisée est donc qualifiée d'OPCPA. Pour en revenir au cas particulier d'OMEGA EP, l'impulsion initiale est étirée de 200 fs à 2,4 ns. Puis elle passe dans les deux OPA. Les OPA sont des cristaux de triborate de lithium (notés LBO). Ils agissent d'une manière similaire aux cristaux de KDP : ils sont pompés par un laser à 2ω , le faisceau à 1ω étiré est incident et en sortie on trouve le faisceau étiré plus un faisceau appelé oisif à 1ω . La somme des fréquences du faisceau incident et du faisceau oisif est égale à la fréquence du faisceau pompe et on a donc en sortie un faisceau amplifié (somme du faisceau incident et du faisceau "oisif"). A la sortie des deux OPA, l'énergie de l'impulsion de 2,4 ns est d'environ 250 mJ. La méthode OPA a été choisie pour l'amplification car l'amplification se fait sur toute la largeur spectrale, qui est donc préservée, ce qui est nécessaire pour pouvoir compresser l'impulsion ultérieurement. Le faisceau pompe à 2ω est issu du générateur d'impulsion longue amplifié par un "Crystal Large-Aperture Ring Amplifier" (CLARA) puis doublé en fréquence. Les OPA sont pompés par des impulsions de 2,4 ns à 0,2 J pour le premier et 1,2 J pour le deuxième. Après cette amplification, le faisceau est filtré une première fois et mis en forme spatiale par un apodiseur, puis suit la fin de la partie "driver" commune avec les impulsions longues (isolation de l'impulsion, filtres, apodiseurs, préamplificateurs, ...) avant d'être injecté dans la partie "amplification". Lorsque le faisceau est injecté dans la partie "amplification", il passe d'abord par un filtre spatial, puis par un 1er amplificateur à 7 disques. Les amplificateurs sont des disques de forme carrée de 40 cm de côté, en verre dopé au Nd, pompés par des lampes flash au Xe. Le faisceau est ensuite dirigé par un miroir vers la ligne d'amplification principale. Il rencontre un polariseur de Brewster ; sa polarisation ayant au préalable été fixée horizontalement par rapport à ce polariseur, le

faisceau est réfléchi. Il traverse une cellule de Pockels désactivée puis un polariseur placé là pour empêcher une impulsion courte réfléchie dans la TC lors d'un tir, de pouvoir atteindre l'amplificateur principal et de risquer de le détériorer. Puis le faisceau passe par un filtre et atteint l'amplificateur principal composé de 11 disques. Il est ensuite réfléchi par un miroir déformable qui permet de corriger le front du faisceau d'inhomogénéités de taille supérieure à 33 mm. Le faisceau repasse par l'amplificateur et le filtre, et atteint la cellule de Pockels qui a alors été activée, et provoque donc une rotation de 90° de la polarisation. Le faisceau n'est alors plus réfléchi (vers l'amplificateur à 7 disques) mais transmis par le polariseur de Brewster car de polarisation verticale par rapport à celui-ci. Il est réfléchi au fond de la ligne par un miroir puis repasse le long de la ligne (ainsi que par la cellule de Pockels toujours active qui va le re-polariser horizontalement) pour être à nouveau réfléchi par le miroir déformable. Le faisceau va donc repasser 2 fois supplémentaires par l'amplificateur principal. La cellule de Pockels est alors désactivée ; le faisceau, qui reste polarisé horizontalement, va être réfléchi par le polariseur de Brewster et repartir vers l'amplificateur à 7 disques et le filtre spatial associé. L'impulsion est donc amplifiée 4 fois par l'amplificateur principal et 2 fois par l'amplificateur à 7 disques. Il est à noter que les filtres spatiaux permettent entre autres de mesurer et transmettre le front d'onde du faisceau au miroir déformable qui en corrigera les aberrations. Le faisceau se dirige ensuite vers un séparateur qui en prélève 0,2 % pour être analysé par divers diagnostics laser.

C'est alors la fin de la partie amplification. Les faisceaux 3 et 4, ainsi que les faisceaux 1 et 2 s'ils sont en configuration impulsion longue, sont triplés en fréquence par deux cristaux de KDP, puis dirigés vers la TC et la cible par deux miroirs. Ils traversent une DPP (si requis) et sont focalisés par une lentille de focale 3,4 m. Ils entrent ensuite dans la TC par un hublot qui maintient le vide et un bouclier à débris puis illuminent la cible. L'efficacité de conversion est mesurée grâce à un séparateur qui prélève 4 % du faisceau après les cristaux de KDP et transmet cette lumière à un ensemble de diagnostics laser. Un de ces diagnostics est un capteur qui permet de détecter les faisceaux d'alignement IR et UV. Le faisceau d'alignement UV est inséré juste avant les cristaux de KDP tandis que le faisceau d'alignement IR est injecté au début de la partie amplification.

Les faisceaux 1 et 2, en configuration impulsion courte, sont dirigés vers la chambre de compression. Le faisceau 2 rencontre un ensemble de 4 réseaux qui compressent l'impulsion de 100 ps jusqu'à 1 ps selon le choix, puis une partie du faisceau (0,5 à 1 %) est prélevée pour les diagnostics d'impulsions courtes. Le faisceau atteint alors le combineur de faisceaux, puis est dirigé vers la TC d'OMEGA ou d'OMEGA EP. Dans le cas d'OMEGA EP, deux miroirs dirigent le faisceau vers la TC ; la focalisation se fait par un miroir parabolique de focale 1 m. On ne peut utiliser une lentille qui serait détériorée par le passage du faisceau très intense. Le faisceau 1 suit un chemin similaire, à ceci près que

l'on peut choisir de le diriger vers le combineur de faisceaux ou non. S'il passe par le combineur, les faisceaux 1 et 2 arriveront dans la TC depuis une localisation très proche, et l'angle entre les deux faisceaux sera faible. Si le faisceau 1 ne passe pas par le combineur, il arrivera dans la TC à 90° du faisceau 2, ce qui permet par exemple d'éclairer à la fois des sources pour une radiographie de face et de côté lors d'un même tir. Le faisceau 1 sera alors qualifié de faisceau "sidelighter" et le faisceau 2 de faisceau "backlighter". L'alignement de la chambre de compression se fait grâce à deux faisceaux IR issus du même lieu que l'ensemble de diagnostics laser situé après les réseaux de compression. Pour finir, la TC d'OMEGA EP est très semblable à celle d'OMEGA. De même taille, elle possède aussi 6 inserteurs de diagnostics (TIMs) et est supportée par une structure comprenant 3 étages.

Les SRF sur OMEGA EP sont très similaires à celles d'OMEGA. Elles possèdent aussi les parties "General", "Target", "TIM", "Fixed" et "Neutronics". Cependant, beaucoup moins de diagnostics neutroniques sont installés sur OMEGA EP. Les diagnostics fixes sont aussi différents. De plus, il faut aussi noter que bien que de nombreux diagnostics insérables dans les TIM soient communs à OMEGA et OMEGA EP, certains diagnostics ne peuvent être utilisés que sur une seule installation. Les deux autres parties, "Select Beam/Source" et "Set up Beam/Source", permettent de définir les faisceaux. Les PI choisissent dans la partie "Select Beam/Source" quels faisceaux doivent être allumés, si les faisceaux 1 et 2 doivent être utilisés en impulsion ns ou ps, et dans ce dernier cas s'ils doivent être combinés ou injectés dans la TC à 90° l'un de l'autre. La partie "Set up Beam/Source" permet de définir les caractéristiques de chaque faisceau utilisé : forme de l'impulsion, décalage du faisceau par rapport au t_0 d'OMEGA EP, énergie du faisceau, focalisation (meilleure focalisation ou agrandissement du spot), s'il faut une DPP et laquelle, ainsi que le pointage du faisceau.

Concernant les diagnostics laser, on en trouve principalement en 3 endroits : à la sortie de l'amplification (le séparateur de faisceau correspondant est représenté en figure C.1), après le triplement en fréquence pour les impulsions longues et dans la chambre de compression pour les impulsions courtes. Pour l'ensemble en sortie d'amplification, 0,2 % de l'énergie du faisceau IR est prélevée. L'énergie du faisceau est mesurée, ainsi que l'allure spatiale du faisceau en champ proche et champ lointain (avant et après focalisation). Une mesure du front d'onde est aussi effectuée par un détecteur. Une CBF appelée ROSS, l'équivalent de la P510 sur OMEGA, et un spectromètre permettent de mesurer la forme temporelle de l'impulsion. Pour les impulsions longues UV, l'ensemble de diagnostics comprend des capteurs pour l'alignement du faisceau UV. On trouve aussi des appareils de mesure de la forme du spot laser en champ proche et lointain, une CBF ROSS pour mesurer la forme de l'impulsion et un calorimètre. Pour les impulsions courtes IR, les séparateurs prélèvent 0,5 à 1 % de l'énergie de chaque faisceau pour analyse par les

diagnostics laser. Comme pour celui en sortie d'amplification, l'ensemble de diagnostics comprend des caméras en champ proche et champ lointain, un appareil de mesure du front d'onde et de l'énergie de l'impulsion et des capteurs pour l'alignement. De plus, un diagnostic de spot laser permet de caractériser la distribution de l'intensité du faisceau sur la cible, et un diagnostic temporel ultrarapide, composé entre autres d'une caméra à balayage de fente rapide, permet de mesurer la durée et la forme temporelle de l'impulsion.

Références

- [1] J Nuckolls, L Wood, A Thiessen, and G Zimmerman. *Nature*, 239(15368) :139, 1972.
- [2] ITER Physics Basis Editors, ITER Physics Expert Group Chairs, Co-Chairs, ITER Joint Central Team, and Physics Integration Unit. *Nuclear Fusion*, 39(12) :2137, 1999.
- [3] S. Atzeni and J. Meyer ter Vehn. *The Physics of Inertial Fusion*. Oxford Science Publications, 2004.
- [4] Stephen E. Bodner. *Phys. Rev. Lett.*, 33 :761–764, Sep 1974.
- [5] C. A. Haynam, P. J. Wegner, J. M. Auerbach, M. W. Bowers, S. N. Dixit, G. V. Erbert, G. M. Heestand, M. A. Henesian, M. R. Hermann, K. S. Jancaitis, K. R. Manes, C. D. Marshall, N. C. Mehta, J. Menapace, E. Moses, J. R. Murray, M. C. Nostrand, C. D. Orth, R. Patterson, R. A. Sacks, M. J. Shaw, M. Spaeth, S. B. Sutton, W. H. Williams, C. C. Widmayer, R. K. White, S. T. Yang, , and B. M. Van Wonterghem. *Appl. Opt.*, 46 :3276, 2007.
- [6] John Lindl, Otto Landen, John Edwards, Ed Moses, and NIC Team. *Physics of Plasmas (1994-present)*, 21(2) :–, 2014.
- [7] S. P. Regan, R. Epstein, B. A. Hammel, L. J. Suter, H. A. Scott, M. A. Barrios, D. K. Bradley, D. A. Callahan, C. Cerjan, G. W. Collins, S. N. Dixit, T. Döppner, M. J. Edwards, D. R. Farley, K. B. Fournier, S. Glenn, S. H. Glenzer, I. E. Golovkin, S. W. Haan, A. Hamza, D. G. Hicks, N. Izumi, O. S. Jones, J. D. Kilkenny, J. L. Kline, G. A. Kyrala, O. L. Landen, T. Ma, J. J. MacFarlane, A. J. MacKinnon, R. C. Mancini, R. L. McCrory, N. B. Meezan, D. D. Meyerhofer, A. Nikroo, H.-S. Park, J. Ralph, B. A. Remington, T. C. Sangster, V. A. Smalyuk, P. T. Springer, and R. P. J. Town. *Phys. Rev. Lett.*, 111 :045001, Jul 2013.
- [8] V. N. Goncharov, T. C. Sangster, R. Betti, T. R. Boehly, M. J. Bonino, T. J. B. Collins, R. S. Craxton, J. A. Delettrez, D. H. Edgell, R. Epstein, R. K. Follett, C. J. Forrest, D. H. Froula, V. Yu. Glebov, D. R. Harding, R. J. Henchen, S. X. Hu, I. V. Igumenshchev, R. Janezic, J. H. Kelly, T. J. Kessler, T. Z. Kosc, S. J. Loucks, J. A. Marozas, F. J. Marshall, A. V. Maximov, R. L. McCrory, P. W. McKenty, D. D. Meyerhofer, D. T. Michel, J. F. Myatt, R. Nora, P. B. Radha, S. P. Regan, W. Seka, W. T. Shmayda, R. W. Short, A. Shvydky, S. Skupsky, C. Stoeckl, B. Yaakobi, J. A.

- Frenje, M. Gatu-Johnson, R. D. Petrasso, and D. T. Casey. *Physics of Plasmas (1994-present)*, 21(5) :–, 2014.
- [9] V. Goncharov. *Phys. Rev. Lett.*, 82 :2091–2094, 1999.
- [10] Y. Aglitskiy, A. L. Velikovich, M. Karasik, V. Serlin, C. J. Pawley, A. J. Schmitt, S. P. Obenschain, A. N. Mostovych, J. H. Gardner, and N. Metzler. *Physics of Plasmas*, 9(5) :2264–2276, 2002.
- [11] V. N. Goncharov, S. Skupsky, T. R. Boehly, J. P. Knauer, P. McKenty, V. A. Smalyuk, R. P. J. Town, O. V. Gotchev, R. Betti, and D. D. Meyerhofer. *Physics of Plasmas (1994-present)*, 7(5) :2062–2068, 2000.
- [12] E. N. Loomis, D. Braun, S. H. Batha, C. Sorce, and O. L. Landen. *Physics of Plasmas*, 18(9) :092702, 2011.
- [13] E. N. Loomis, D. Braun, S. H. Batha, and O. L. Landen. *Physics of Plasmas (1994-present)*, 19(12) :–, 2012.
- [14] J. L. Peterson, D. S. Clark, L. P. Masse, and L. J. Suter. The effects of early time laser drive on hydrodynamic instability growth in national ignition facility implosions. *Physics of Plasmas (1994-present)*, 21(9) :–, 2014.
- [15] L. Masse, A. Casner, D. Galmiche, G. Huser, S. Liberatore, and M. Theobald. *Phys. Rev. E*, 83(5) :055401, 2011.
- [16] G. Huser, A. Casner, L. Masse, S. Liberatore, D. Galmiche, L. Jacquet, and M. Theobald. *Physics of Plasmas (1994-present)*, 18(1) :–, 2011.
- [17] Shinsuke Fujioka, Atsushi Sunahara, Katsunobu Nishihara, Naofumi Ohnishi, Tomoyuki Johzaki, Hiroyuki Shiraga, Keisuke Shigemori, Mitsuo Nakai, Tadashi Ikegawa, Masakatsu Murakami, Keiji Nagai, Takayoshi Norimatsu, Hiroshi Azechi, and Tatsuhiko Yamanaka. *Phys. Rev. Lett.*, 92 :195001, May 2004.
- [18] C. Yanez, J. Sanz, M. Olazabal-Loume, and L. F. Ibanez. *Physics of Plasmas*, 18(5) :052701, 2011.
- [19] S. Skupsky, R. W. Short, T. Kessler, R. S. Craxton, S. Letzring, and J. M. Soures. *Journal of Applied Physics*, 66 :3456–3462, October 1989.
- [20] R. G. Watt, J. Duke, C. J. Fontes, P. L. Gobby, R. V. Hollis, R. A. Kopp, R. J. Mason, D. C. Wilson, C. P. Verdon, T. R. Boehly, J. P. Knauer, D. D. Meyerhofer, V. Smalyuk, R. P. Town, A. Iwase, and O. Willi. *Physical Review Letters*, 81 :4644–4647, November 1998.
- [21] T. R. Boehly, D. L. Brown, R. S. Craxton, R. L. Keck, J. P. Knauer, J. H. Kelly, T. J. Kessler, S. A. Kumpan, S. J. Loucks, S. A. Letzring, F. J. Marshall, R. L. McCrory, S. F. B. Morse, W. Seka, J. M. Soures, and C. P. Verdon. *Optics Communications*, 133 :495–506, February 1997.

- [22] S. Depierreux, C. Labaune, D. T. Michel, C. Stenz, P. Nicolaï, M. Grech, G. Riazuelo, S. Weber, C. Riconda, V. T. Tikhonchuk, P. Loiseau, N. G. Borisenko, W. Nazarov, S. Hüller, D. Pesme, M. Casanova, J. Limpouch, C. Meyer, P. Di-Nicola, R. Wrobel, E. Alozy, P. Romary, G. Thiell, G. Soullié, C. Reverdin, and B. Villette. *Phys. Rev. Lett.*, 102 :195005, May 2009.
- [23] R. Hansen, LLE, and LLNL. *Polar drive ignition campaign - Conceptual design*. 2012.
- [24] W. L. Kruer. *The Physics of Laser Plasma Interactions*. Addison-Wesley, 1988.
- [25] Lyman Spitzer and Richard Härm. *Phys. Rev.*, 89 :977–981, Mar 1953.
- [26] R. Dautray and J. P. Watteau. *La fusion thermonucléaire inertielle par laser*. Eyrolles, 1993.
- [27] Lord Rayleigh. *Proc. London Math. Soc.*, XIV :170, 1883.
- [28] G.I. Taylor. *Proc. R. Soc. London Ser. A*, 201 :192, 1950.
- [29] D. H. Sharp. *Physica 12D*, page 3, 1984.
- [30] Hermann Ludwig Ferdinand von Helmholtz. *Monatsberichte der Königlich Preussische Akademie der Wissenschaften zu Berlin*, 23 :215, 1868.
- [31] Lord William Thomson Kelvin. *Philosophical Magazine*, 42 :362, 1871.
- [32] J. Sanz. *Phys. Rev. Lett.*, 73 :2700–2703, Nov 1994.
- [33] H. J. Kull. *Physics of Fluids B : Plasma Physics (1989-1993)*, 1(1) :170–182, 1989.
- [34] A. R. Piriz, J. Sanz, and L. F. Ibañez. *Physics of Plasmas (1994-present)*, 4(4) :1117–1126, 1997.
- [35] R. Betti, V. N. Goncharov, R. L. McCrory, P. Sorokin, and C. P. Verdon. *Physics of Plasmas*, 3(5) :2122–2128, 1996.
- [36] J. Sanz, J. Ram´, R. Ramis, R. Betti, and R. P. J. Town. *Phys. Rev. Lett.*, 89 :195002, Oct 2002.
- [37] J. Sanz, R. Betti, R. Ramis, and J. Ramírez. *Plasma Physics and Controlled Fusion*, 46 :B367–B380, December 2004.
- [38] J. Garnier, P.-A. Raviart, C. Cherfils-Cléroutin, and L. Masse. *Phys. Rev. Lett.*, 90 :185003, May 2003.
- [39] H. Takabe, K. Mima, L. Montierth, and R. L. Morse. *Physics of Fluids (1958-1988)*, 28(12) :3676–3682, 1985.
- [40] R. Betti, V. N. Goncharov, R. L. McCrory, and C. P. Verdon. *Physics of Plasmas (1994-present)*, 2(10) :3844–3851, 1995.
- [41] V. N. Goncharov, R. Betti, R. L. McCrory, P. Sorokin, and C. P. Verdon. *Physics of Plasmas (1994-present)*, 3(4) :1402–1414, 1996.

- [42] V. N. Goncharov, R. Betti, R. L. McCrory, and C. P. Verdon. *Physics of Plasmas (1994-present)*, 3(12) :4665–4676, 1996.
- [43] J. Sanz. *Phys. Rev. E*, 53 :4026–4045, Apr 1996.
- [44] L. Masse. *Phys. Rev. Lett.*, 98 :245001, Jun 2007.
- [45] M J deC Henshaw, G J Pert, and D L Youngs. *Plasma Physics and Controlled Fusion*, 29(3) :405, 1987.
- [46] Steven W. Haan. *Phys. Rev. A*, 39(11) :5812–5825, 1989.
- [47] V. A. Smalyuk, T. R. Boehly, D. K. Bradley, V. N. Goncharov, J. A. Delettrez, J. P. Knauer, D. D. Meyerhofer, D. Oron, and D. Shvarts. *Phys. Rev. Lett.*, 81 :5342–5345, Dec 1998.
- [48] V. A. Smalyuk, O. Sadot, J. A. Delettrez, D. D. Meyerhofer, S. P. Regan, and T. C. Sangster. *Phys. Rev. Lett.*, 95(21) :215001, 2005.
- [49] J. Garnier and L. Masse. *Physics of Plasmas (1994-present)*, 12(6) :–, 2005.
- [50] D. Oron, L. Arazi, D. Kartoon, A. Rikanati, U. Alon, and D. Shvarts. *Phys. Plasmas*, 8(6) :2883–2889, 2001.
- [51] Uri Alon, Jacob Hecht, David Mukamel, and Dov Shvarts. *Phys. Rev. Lett.*, 72 :2867–2870, May 1994.
- [52] Uri Alon, Dov Shvarts, and David Mukamel. *Phys. Rev. E*, 48 :1008–1014, Aug 1993.
- [53] Guy Dimonte, P. Ramaprabhu, D. L. Youngs, M. J. Andrews, and R. Rosner. *Physics of Plasmas (1994-present)*, 12(5) :–, 2005.
- [54] V. A. Smalyuk, O. Sadot, R. Betti, V. N. Goncharov, J. A. Delettrez, D. D. Meyerhofer, S. P. Regan, T. C. Sangster, and D. Shvarts. *Physics of Plasmas (1994-present)*, 13(5) :–, 2006.
- [55] O. Sadot, V. A. Smalyuk, J. A. Delettrez, D. D. Meyerhofer, T. C. Sangster, R. Betti, V. N. Goncharov, and D. Shvarts. *Phys. Rev. Lett.*, 95(26) :265001, 2005.
- [56] R. D. Richtmyer. *Commun. Pure Appl. Math.*, 13 :297, 1960.
- [57] E. E. Meshkov. *Fluid Dyn.*, 4 :101, 1969.
- [58] R. Ishizaki and K. Nishihara. *Phys. Rev. Lett.*, 78 :1920–1923, Mar 1997.
- [59] Alexander L. Velikovich, Jill P. Dahlburg, John H. Gardner, and Robert J. Taylor. Saturation of perturbation growth in ablatively driven planar laser targets. *Physics of Plasmas (1994-present)*, 5(5) :1491–1505, 1998.
- [60] R. Taylor, A. Velikovich, J. Dahlburg, and J. Gardner. *Phys. Rev. Lett.*, 79 :1861–1864, Sep 1997.

- [61] V. A. Smalyuk, V. N. Goncharov, T. R. Boehly, J. A. Delettrez, D. Y. Li, J. A. Marozas, A. V. Maximov, D. D. Meyerhofer, S. P. Regan, and T. C. Sangster. *Phys. Plasmas*, 12(7) :072703, 2005.
- [62] Hideaki Takabe, Katsunobu Nishihara, and Tosiya Taniuti. *Journal of the Physical Society of Japan*, 45(6) :2001–2008, 1978.
- [63] V. N. Goncharov, O. V. Gotchev, E. Vianello, T. R. Boehly, J. P. Knauer, P. W. McKenty, P. B. Radha, S. P. Regan, T. C. Sangster, S. Skupsky, V. A. Smalyuk, R. Betti, R. L. McCrory, D. D. Meyerhofer, and C. Cherfils-Clerouin. *Physics of Plasmas*, 13(1) :012702, 2006.
- [64] P. B. Radha, V. N. Goncharov, T. J. B. Collins, J. A. Delettrez, Y. Elbaz, V. Yu. Glebov, R. L. Keck, D. E. Keller, J. P. Knauer, J. A. Marozas, F. J. Marshall, P. W. McKenty, D. D. Meyerhofer, S. P. Regan, T. C. Sangster, D. Shvarts, S. Skupsky, Y. Srebro, R. P. J. Town, and C. Stoeckl. *Physics of Plasmas (1994-present)*, 12(3) :–, 2005.
- [65] S. P. Obenschain, S. E. Bodner, D. Colombant, K. Gerber, R. H. Lehmberg, E. A. McLean, A. N. Mostovych, M. S. Pronko, C. J. Pawley, A. J. Schmitt, J. D. Sethian, V. Serlin, J. A. Stamper, C. A. Sullivan, J. P. Dahlburg, J. H. Gardner, Y. Chan, A. V. Deniz, J. Hardgrove, T. Lehecka, and M. Klapisch. *Physics of Plasmas (1994-present)*, 3(5) :2098–2107, 1996.
- [66] O. V. Gotchev, V. N. Goncharov, J. P. Knauer, T. R. Boehly, T. J. B. Collins, R. Epstein, P. A. Jaanimagi, and D. D. Meyerhofer. *Physical Review Letters*, 96(11) :115005, 2006.
- [67] P. M. Celliers, G. W. Collins, L. B. Da Silva, D. M. Gold, and R. Cauble. *Applied Physics Letters*, 73(10) :1320–1322, 1998.
- [68] L. M. Barker and R. E. Hollenbach. *Journal of Applied Physics*, 43(11) :4669–4675, 1972.
- [69] Y. Kato, K. Mima, N. Miyanaga, S. Arinaga, Y. Kitagawa, M. Nakatsuka, and C. Yamanaka. *Phys. Rev. Lett.*, 53 :1057–1060, Sep 1984.
- [70] T. J. B. Collins, J. A. Marozas, K. S. Anderson, R. Betti, R. S. Craxton, J. A. Delettrez, V. N. Goncharov, D. R. Harding, F. J. Marshall, R. L. McCrory, D. D. Meyerhofer, P. W. McKenty, P. B. Radha, A. Shvydky, S. Skupsky, and J. D. Zuegel. *Physics of Plasmas (1994-present)*, 19(5) :–, 2012.
- [71] Laboratory for Laser Energetics. *LLE review*, volume 114. 2008.
- [72] S. P. Obenschain, D. G. Colombant, M. Karasik, C. J. Pawley, V. Serlin, A. J. Schmitt, J. L. Weaver, J. H. Gardner, L. Phillips, Y. Aglitskiy, Y. Chan, J. P. Dahlburg, and M. Klapisch. *Physics of Plasmas (1994-present)*, 9(5) :2234–2243, 2002.

- [73] C. Yamanaka, Yoshiaki Kato, Yasukazu Izawa, Kunio Yoshida, Tatsuhiko Yamanaka, Takatomo Sasaki, M. Nakatsuka, T. Mochizuki, J. Kuroda, and Sadao Nakai. *Quantum Electronics, IEEE Journal of*, 17(9) :1639–1649, Sep 1981.
- [74] Shinsuke Fujioka, Hiroyuki Shiraga, Masaharu Nishikino, Manabu Heya, Keisuke Shigemori, Mitsuo Nakai, Hiroshi Azechi, Sadao Nakai, and Tatsuhiko Yamanaka. *Review of Scientific Instruments*, 73(7) :2588–2596, 2002.
- [75] K. Otani, K. Shigemori, T. Sakaiya, S. Fujioka, A. Sunahara, M. Nakai, H. Shiraga, H. Azechi, and K. Mima. *Physics of Plasmas*, 14(12) :–, 2007.
- [76] Documentation sur idl. http://www.exelisvis.com/docs/using_idl_home.html.
- [77] M. J. Eckart, R. L. Hanks, J. D. Kilkenny, R. Pasha, J. D. Wiedwald, and J. D. Hares. *Review of Scientific Instruments*, 57(8) :2046–2048, 1986.
- [78] J. D. Kilkenny, P. Bell, R. Hanks, G. Power, R. E. Turner, and J. Wiedwald. *Review of Scientific Instruments*, 59(8) :1793–1796, 1988.
- [79] J. P. Knauer, R. Betti, D. K. Bradley, T. R. Boehly, T. J. B. Collins, V. N. Goncharov, P. W. McKenty, D. D. Meyerhofer, V. A. Smalyuk, C. P. Verdon, S. G. Glendinning, D. H. Kalantar, and R. G. Watt. *Physics of Plasmas (1994-present)*, 7(1) :338–345, 2000.
- [80] L. Vincent and P. Soille. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 13 :583, 1991.
- [81] J. Breil and P.-H. Maire. *Journal of Computational Physics*, 224(2) :785 – 823, 2007.
- [82] J. Breil, S. Galera, and P.-H. Maire. *Computers Fluids*, 46(1) :161 – 167, 2011. 10th {ICFD} Conference Series on Numerical Methods for Fluid Dynamics (ICFD 2010).
- [83] V. A. Smalyuk. *Thèse : Experimental investigation of the nonlinear Rayleigh-Taylor instability in CH foils irradiated by UV light*. 1999.
- [84] G. Fiksel, S. X. Hu, V. A. Goncharov, D. D. Meyerhofer, T. C. Sangster, V. A. Smalyuk, B. Yaakobi, M. J. Bonino, and R. Jungquist. *Physics of Plasmas*, 19(6) :062704, 2012.
- [85] V. A. Smalyuk, T. R. Boehly, D. K. Bradley, J. P. Knauer, and D. D. Meyerhofer. *Review of Scientific Instruments*, 70 :647–650, January 1999.
- [86] A. Casner, D. Galmiche, G. Huser, J-P. Jadaud, S. Liberatore, and M. Vandenboomgaerde. *Phys. Plasmas*, 16(9) :092701, 2009.
- [87] J. P. Knauer, R. Betti, D. K. Bradley, T. R. Boehly, T. J. B. Collins, V. N. Goncharov, P. W. McKenty, D. D. Meyerhofer, V. A. Smalyuk, C. P. Verdon, S. G. Glendinning, D. H. Kalantar, and R. G. Watt. *Physics of Plasmas*, 7 :338–345, January 2000.

- [88] M. A. Barrios, D. G. Hicks, T. R. Boehly, D. E. Fratanduono, J. H. Eggert, P. M. Celliers, G. W. Collins, and D. D. Meyerhofer. *Physics of Plasmas*, 17(5) :056307, 2010.
- [89] R. M. More, K. H. Warren, D. A. Young, and G. B. Zimmerman. *Physics of Fluids (1958-1988)*, 31(10) :3059–3078, 1988.
- [90] S. Hu, V. Smalyuk, V. Goncharov, S. Skupsky, T. Sangster, D. Meyerhofer, and D. Shvarts. *Phys. Rev. Lett.*, 101 :055002, Jul 2008.
- [91] V. N. Goncharov, T. C. Sangster, R. Epstein, S. X. Hu, I. V. Igumenshchev, C. J. Forrest, D. H. Froula, F. J. Marshall, R. L. McCrory, D. D. Meyerhofer, D. T. Michel, P. B. Radha, S. P. Regan, W. Seka, and C. Stoeckl. Presented at the 56th Annual Meeting of the APS Division of Plasma Physics, 2014.
- [92] T. R. Boehly, V. N. Goncharov, S. X. Hu, G. Fiksel, and P. M. Celliers. Presented at the 56th Annual Meeting of the APS Division of Plasma Physics, 2014.
- [93] M. Dunne, M. Borghesi, A. Iwase, M. Jones, R. Taylor, O. Willi, R. Gibson, S. Goldman, J. Mack, and R. Watt. *Phys. Rev. Lett.*, 75 :3858–3861, Nov 1995.
- [94] Michel Koenig, Alessandra Benuzzi, Franck Philippe, Dimitri Batani, Tom Hall, Nicolas Grandjouan, and Wigen Nazarov. *Physics of Plasmas (1994-present)*, 6(8) :3296–3301, 1999.
- [95] Giora Hazak, Alexander L. Velikovich, John H. Gardner, and Jill P. Dahlburg. *Physics of Plasmas (1994-present)*, 5(12) :4357–4365, 1998.
- [96] O. Willi, L. Barringer, A. Bell, M. Borghesi, J. Davies, R. Gaillard, A. Iwase, A. MacKinnon, G. Malka, C. Meyer, S. Nuruzzaman, R. Taylor, C. Vickers, D. Hoarty, P. Gobby, R. Johnson, R.G. Watt, N. Blanchot, B. Canaud, H. Croso, B. Meyer, J.L. Miquel, C. Reverdin, A. Pukhov, and J. Meyer ter Vehn. *Nuclear Fusion*, 40(3Y) :537, 2000.
- [97] D. Hoarty, O. Willi, L. Barringer, C. Vickers, R. Watt, and W. Nazarov. *Physics of Plasmas (1994-present)*, 6(5) :2171–2177, 1999.
- [98] V. Malka, J. Faure, S. Hüller, V. Tikhonchuk, S. Weber, and F. Amiranoff. *Phys. Rev. Lett.*, 90 :075002, Feb 2003.
- [99] Ph. Nicolaï, M. Olazabal-Loumé, S. Fujioka, A. Sunahara, N. Borisenko, S. Gusakov, A. Orekov, M. Grech, G. Riazuelo, C. Labaune, J. Velechowski, and V. Tikhonchuk. *Physics of Plasmas (1994-present)*, 19(11) :–, 2012.
- [100] M Olazabal-Loumé, Ph Nicolaï, G Riazuelo, M Grech, J Breil, S Fujioka, A Sunahara, N Borisenko, and V T Tikhonchuk. *New Journal of Physics*, 15(8) :085033, 2013.
- [101] S. Yu. GusŤkov, J. Limpouch, Ph. Nicolaï, and V. T. Tikhonchuk. *Physics of Plasmas (1994-present)*, 18(10) :–, 2011.

- [102] F. Walraet, G. Riazuelo, and G. Bonnaud. *Physics of Plasmas*, 10 :811–819, March 2003.
- [103] G. Riazuelo and G. Bonnaud. *Physics of Plasmas*, 7 :3841–3844, October 2000.
- [104] M. Grech, G. Riazuelo, D. Pesme, S. Weber, and V. T. Tikhonchuk. *Phys. Rev. Lett.*, 102 :155001, Apr 2009.
- [105] R. Yan and R. Betti. Presented at the 56th Annual Meeting of the APS Division of Plasma Physics, 2014.
- [106] M. Olazabal-Loumé and L. Masse. *EPJ Web of Conferences*, 59, Nov 2013.
- [107] M. Lafon, R. Betti, K. S. Anderson, T. J. B. Collins, S. Skupsky, and P. W. McKenty. Presented at the 56th Annual Meeting of the APS Division of Plasma Physics, 2014.
- [108] D. Strickland and G. Mourou. *Optics Communications*, 56(3) :219, 1985.

Publications et communications

Publications

- A. Casner, V. Smalyuk, L. Masse, A. Moore, **B. Delorme**, D. Martinez, I. Igumenshev, L. Jacquet, S. Liberatore, R. Seugling, C. Chicanne, H.S. Park and B.A. Remington, *Design and implementation plan for indirect-drive highly non-linear ablative Rayleigh Taylor instability experiments on the National Ignition Facility*, High Energy Density Physics 9 (1), 32 (2013)
- A. Casner, L. Masse, **B. Delorme**, D. Martinez, G. Huser, D. Galmiche, S. Liberatore, I. Igumenshchev, M. Olazabal-Loumé, Ph. Nicolaï, J. Breil, D. T. Michel, D. Froula, W. Seka, G. Riazuelo, S. Fujioka, A. Sunahara, M. Grech, C. Chicanne, M. Theobald, N. Borisenko, A. Orekhov, V. T. Tikhonchuk, B. Remington, V. N. Goncharov and V. A. Smalyuk, *Progress in indirect and direct-drive planar experiments on hydrodynamic instabilities at the ablation front*, Phys. Plasmas 21, 122702 (2014)
- **B. Delorme**, M. Olazabal-Loumé, Ph. Nicolaï, A. Casner, N. Borisenko, J. Breil, D. Froula, S. Fujioka, V. Goncharov, M. Grech, D. T. Michel, A. Orekhov, G. Riazuelo, A. Sunahara and V. T. Tikhonchuk, *Experimental Demonstration of Laser Imprint Reduction Using Underdense Foams*, soumis à Phys. Plasmas

Communications

- **Poster** : *Preliminary analysis of ablative Richtmyer-Meshkov experiments using 1D and 2D simulations*, NIF User Group Meeting, Livermore, Février 2012
- **Poster** : *Preliminary analysis of ablative Richtmyer-Meshkov experiments using 1D and 2D simulations*, Forum ILP, Sainte-Marie-en-Ré, Septembre 2012
- **Poster** : *Analysis and simulations of ablative Richtmyer-Meshkov direct drive experiments*, OMEGA Laser User Group Meeting, Rochester, Avril 2013
- **Oral** : *Analysis and simulations of ablative Richtmyer-Meshkov direct drive experiments at the OMEGA laser*, Direct Drive and Fast Ignition Workshop, Rome, Mai 2013
- **Oral** : *Analysis of ablative Richtmyer-Meshkov direct-drive experiments performed on the OMEGA laser facility*, 40th European Physics Society Conference on Plasma Physics, Espoo, Juillet 2013
- **Oral** : *Etude expérimentale des conditions initiales de l'instabilité Rayleigh-Taylor en Fusion par Confinement Inertiel*, Forum ILP, Orcières, Février 2014
- **Oral** : *Experimental demonstration of laser imprint reduction using underdense foams*, 13th International Workshop on the Fast Ignition of Fusion Targets, Oxford, Septembre 2014